

VOLUME 4,NUMBER 1 JANUARY 2006

ISSN:1548-5390 PRINT,1559-176X ONLINE



**JOURNAL
OF CONCRETE
AND APPLICABLE
MATHEMATICS**

EUDOXUS PRESS,LLC

SCOPE AND PRICES OF THE JOURNAL
Journal of Concrete and Applicable Mathematics

A quartely international publication of **Eudoxus Press, LLC**

Editor in Chief: George Anastassiou

Department of Mathematical Sciences,
University of Memphis
Memphis, TN 38152, U.S.A.
ganastss@memphis.edu

The main purpose of the "Journal of Concrete and Applicable Mathematics" is to publish high quality original research articles from all subareas of Non-Pure and/or Applicable Mathematics and its many real life applications, as well connections to other areas of Mathematical Sciences, as long as they are presented in a Concrete way. It welcomes also related research survey articles and book reviews. A sample list of connected mathematical areas with this publication includes and is not restricted to: Applied Analysis, Applied Functional Analysis, Probability theory, Stochastic Processes, Approximation Theory, O.D.E, P.D.E, Wavelet, Neural Networks, Difference Equations, Summability, Fractals, Special Functions, Splines, Asymptotic Analysis, Fractional Analysis, Inequalities, Moment Theory, Numerical Functional Analysis, Tomography, Asymptotic Expansions, Fourier Analysis, Applied Harmonic Analysis, Integral Equations, Signal Analysis, Numerical Analysis, Optimization, Operations Research, Linear Programming, Fuzzyness, Mathematical Finance, Stochastic Analysis, Game Theory, Math. Physics aspects, Applied Real and Complex Analysis, Computational Number Theory, Graph Theory, Combinatorics, Computer Science Math. related topics, combinations of the above, etc. In general any kind of Concretely presented Mathematics which is Applicable fits to the scope of this journal. Working Concretely and in Applicable Mathematics has become a main trend in many recent years, so we can understand better and deeper and solve the important problems of our real and scientific world. "Journal of Concrete and Applicable Mathematics" is a peer-reviewed International Quarterly Journal. We are calling for papers for possible publication. The contributor should send three copies of the contribution to the editor in-Chief typed in TEX, LATEX double spaced. [See: Instructions to Contributors]

Journal of Concrete and Applicable Mathematics(JCAAM)

ISSN:1548-5390 PRINT, 1559-176X ONLINE.

is published in January, April, July and October of each year by

EUDOXUS PRESS, LLC,

1424 Beaver Trail Drive, Cordova, TN 38016, USA,

Tel. 001-901-751-3553

anastassioug@yahoo.com

<http://www.EudoxusPress.com>.

Annual Subscription Current Prices: For USA and Canada, Institutional: Print \$250, Electronic \$220, Print and Electronic \$310. Individual: Print \$77, Electronic \$60, Print & Electronic \$110. For any other part of the world add \$25 more to the above prices for Print.

Single article PDF file for individual \$8. Single issue in PDF form for individual \$25.
The journal carries page charges \$8 per page of the pdf file of an article, payable upon acceptance of the article within one month and before publication.
No credit card payments. Only certified check, money order or international check in US dollars are acceptable.
Combination orders of any two from JoCAAA, JCAAM, JAFA receive 25% discount, all three receive 30% discount.

Copyright©2004 by Eudoxus Press, LLC all rights reserved. JCAAM is printed in USA.

JCAAM is reviewed and abstracted by AMS Mathematical Reviews, MATHSCI, and Zentralblatt MATH.

It is strictly prohibited the reproduction and transmission of any part of JCAAM and in any form and by any means without the written permission of the publisher. It is only allowed to educators to Xerox articles for educational purposes. The publisher assumes no responsibility for the content of published papers.

JCAAM IS A JOURNAL OF RAPID PUBLICATION

Editorial Board Associate Editors

Editor in -Chief:
George Anastassiou
Department of Mathematical Sciences
The University Of Memphis
Memphis, TN 38152, USA
tel. 901-678-3144, fax 901-678-2480
e-mail ganastss@memphis.edu
www.msci.memphis.edu/~ganastss/jcaam
Areas: Approximation Theory,
Probability, Moments, Wavelet,
Neural Networks, Inequalities, Fuzzyness.

Associate Editors:

1) Ravi Agarwal
Florida Institute of Technology
Applied Mathematics Program
150 W. University Blvd.
Melbourne, FL 32901, USA
agarwal@fit.edu
Differential Equations, Difference
Equations, inequalities

2) Shair Ahmad
University of Texas at San Antonio
Division of Math. & Stat.
San Antonio, TX 78249-0664, USA
SAHMAD@utsa.edu
Differential Equations, Mathematical
Biology

3) Drumi D. Bainov
Medical University of Sofia
P.O. Box 45, 1504 Sofia, Bulgaria
drumibainov@yahoo.com
Differential Equations, Optimal Control,
Numerical Analysis, Approximation Theory

4) Carlo Bardaro
Dipartimento di Matematica & Informatica
Universita' di Perugia
Via Vanvitelli 1
06123 Perugia, ITALY
tel. +390755855034, +390755853822,
fax +390755855024
bardaro@unipg.it ,
bardaro@dipmat.unipg.it
Functional Analysis and Approximation Th.,

19) Rupert Lasser
Institut fur Biomathematik & Biomertie, GSF
-National Research Center for environment and
health
Ingolstaedter landstr.1
D-85764 Neuherberg, Germany
lasser@gsf.de
Orthogonal Polynomials, Fourier Analysis,
Mathematical Biology

20) Alexandru Lupas
University of Sibiu
Faculty of Sciences
Department of Mathematics
Str. I. Ratiu nr. 7
2400-Sibiu, Romania
lupas@ulbsibiu.ro
Classical Analysis, Inequalities,
Special Functions, Umbral Calculus,
Approximation Th., Numerical Analysis
and Methods

21) Ram N. Mohapatra
Department of Mathematics
University of Central Florida
Orlando, FL 32816-1364
tel. 407-823-5080
ramm@mail.ucf.edu
Real and Complex analysis, Approximation Th.,
Fourier Analysis, Fuzzy Sets and Systems

22) Rainer Nagel
Arbeitsbereich Funktionalanalysis
Mathematisches Institut
Auf der Morgenstelle 10
D-72076 Tuebingen
Germany
tel. 49-7071-2973242
fax 49-7071-294322
rana@fa.uni-tuebingen.de
evolution equations, semigroups, spectral th.,
positivity

23) Panos M. Pardalos
Center for Appl. Optimization
University of Florida
303 Weil Hall
P.O. Box 116595
Gainesville, FL 32611-6595

Summability, Signal Analysis, Integral
Equations,
Measure Th., Real Analysis

5) Francoise Bastin
Institute of Mathematics
University of Liege
4000 Liege
BELGIUM
f.bastin@ulg.ac.be
Functional Analysis, Wavelets

6) Paul L. Butzer
RWTH Aachen
Lehrstuhl A für Mathematik
D-52056 Aachen
Germany
tel. 0049/241/80-94627 office,
0049/241/72833 home,
fax 0049/241/80-92212
Butzer@rwth-aachen.de
Approximation Th., Sampling Th., Signals,
Semigroups of Operators, Fourier Analysis

7) Yeol Je Cho
Department of Mathematics Education
College of Education
Gyeongsang National University
Chinju 660-701
KOREA
tel. 055-751-5673 Office,
055-755-3644 home,
fax 055-751-6117
yjcho@nongae.gsnu.ac.kr
Nonlinear operator Th., Inequalities,
Geometry of Banach Spaces

8) Sever S. Dragomir
School of Communications and Informatics
Victoria University of Technology
PO Box 14428
Melbourne City M.C
Victoria 8001, Australia
tel 61 3 9688 4437, fax 61 3 9688 4050
sever.dragomir@vu.edu.au,
sever@sci.vu.edu.au
Math. Analysis, Inequalities, Approximation
Th.,
Numerical Analysis, Geometry of Banach
Spaces,
Information Th. and Coding

9) A.M. Fink
Department of Mathematics
Iowa State University
Ames, IA 50011-0001, USA

tel. 352-392-9011
pardalos@ufl.edu
Optimization, Operations Research

24) Svetlozar T. Rachev
Dept. of Statistics and Applied Probability
Program
University of California, Santa Barbara
CA 93106-3110, USA
tel. 805-893-4869
rachev@pstat.ucsb.edu
AND
Chair of Econometrics and Statistics
School of Economics and Business Engineering
University of Karlsruhe
Kollegium am Schloss, Bau II, 20.12, R210
Postfach 6980, D-76128, Karlsruhe, Germany
tel. 011-49-721-608-7535
rachev@lsoe.uni-karlsruhe.de
Mathematical and Empirical Finance,
Applied Probability, Statistics and Econometrics

25) Paolo Emilio Ricci
Universita' degli Studi di Roma "La Sapienza"
Dipartimento di Matematica-Istituto
"G. Castelnuovo"
P.le A. Moro, 2-00185 Roma, ITALY
tel. ++39 0649913201, fax ++39 0644701007
riccip@uniroma1.it, Paoloemilio.Ricci@uniroma1.it
Orthogonal Polynomials and Special functions,
Numerical Analysis, Transforms, Operational
Calculus,
Differential and Difference equations

26) Cecil C. Rousseau
Department of Mathematical Sciences
The University of Memphis
Memphis, TN 38152, USA
tel. 901-678-2490, fax 901-678-2480
ccrousse@memphis.edu
Combinatorics, Graph Th.,
Asymptotic Approximations,
Applications to Physics

27) Tomasz Rychlik
Institute of Mathematics
Polish Academy of Sciences
Chopina 12, 87100 Torun, Poland
T.Rychlik@impan.gov.pl
Mathematical Statistics, Probabilistic
Inequalities

28) Bl. Sendov
Institute of Mathematics and Informatics
Bulgarian Academy of Sciences
Sofia 1090, Bulgaria

tel.515-294-8150
fink@math.iastate.edu
Inequalities, Ordinary Differential
Equations

10) Sorin Gal
Department of Mathematics
University of Oradea
Str.Armatei Romane 5
3700 Oradea, Romania
galso@uoradea.ro
Approximation Th., Fuzzyness, Complex
Analysis

11) Jerome A. Goldstein
Department of Mathematical Sciences
The University of Memphis,
Memphis, TN 38152, USA
tel.901-678-2484
jgoldste@memphis.edu
Partial Differential Equations,
Semigroups of Operators

12) Heiner H. Gonska
Department of Mathematics
University of Duisburg
Duisburg, D-47048
Germany
tel.0049-203-379-3542 office
gonska@informatik.uni-duisburg.de
Approximation Th., Computer Aided
Geometric Design

13) Dmitry Khavinson
Department of Mathematical Sciences
University of Arkansas
Fayetteville, AR 72701, USA
tel.(479)575-6331, fax(479)575-8630
dmitry@uark.edu
Potential Th., Complex Analysis, Holomorphic
PDE, Approximation Th., Function Th.

14) Virginia S. Kiryakova
Institute of Mathematics and Informatics
Bulgarian Academy of Sciences
Sofia 1090, Bulgaria
virginia@diogenes.bg
Special Functions, Integral Transforms,
Fractional Calculus

15) Hans-Bernd Knoop
Institute of Mathematics
Gerhard Mercator University
D-47048 Duisburg
Germany
tel.0049-203-379-2676

bSENDov@bas.bg
Approximation Th., Geometry of Polynomials,
Image Compression

29) Igor Shevchuk
Faculty of Mathematics and Mechanics
National Taras Shevchenko
University of Kyiv
252017 Kyiv
UKRAINE
shevchuk@univ.kiev.ua
Approximation Theory

30) H.M. Srivastava
Department of Mathematics and Statistics
University of Victoria
Victoria, British Columbia V8W 3P4
Canada
tel.250-721-7455 office, 250-477-6960 home,
fax 250-721-8962
harimsri@math.uvic.ca
Real and Complex Analysis, Fractional Calculus
and Appl.,
Integral Equations and Transforms, Higher
Transcendental
Functions and Appl., q-Series and q-Polynomials,
Analytic Number Th.

31) Ferenc Szidarovszky
Dept. Systems and Industrial Engineering
The University of Arizona
Engineering Building, 111
PO. Box 210020
Tucson, AZ 85721-0020, USA
szidar@sie.arizona.edu
Numerical Methods, Game Th., Dynamic Systems,
Multicriteria Decision making,
Conflict Resolution, Applications
in Economics and Natural Resources
Management

32) Gancho Tachev
Dept. of Mathematics
Univ. of Architecture, Civil Eng. and Geodesy
1 Hr. Smirnenski blvd
BG-1421 Sofia, Bulgaria
Approximation Theory

33) Manfred Tasche
Department of Mathematics
University of Rostock
D-18051 Rostock
Germany
manfred.tasche@mathematik.uni-rostock.de
Approximation Th., Wavelet, Fourier Analysis,
Numerical Methods, Signal Processing,

knoop@math.uni-duisburg.de
Approximation Theory, Interpolation

16) Jerry Koliha
Dept. of Mathematics & Statistics
University of Melbourne
VIC 3010, Melbourne
Australia
koliha@unimelb.edu.au
Inequalities, Operator Theory,
Matrix Analysis, Generalized Inverses

17) Mustafa Kulenovic
Department of Mathematics
University of Rhode Island
Kingston, RI 02881, USA
kulenm@math.uri.edu
Differential and Difference Equations

18) Gerassimos Ladas
Department of Mathematics
University of Rhode Island
Kingston, RI 02881, USA
gladas@math.uri.edu
Differential and Difference Equations

Image Processing, Harmonic Analysis

34) Chris P. Tsokos
Department of Mathematics
University of South Florida
4202 E. Fowler Ave., PHY 114
Tampa, FL 33620-5700, USA
profcpt@math.usf.edu, profcpt@chuma1.cas.usf.edu
Stochastic Systems, Biomathematics,
Environmental Systems, Reliability Th.

35) Lutz Volkmann
Lehrstuhl II fuer Mathematik
RWTH-Aachen
Templergraben 55
D-52062 Aachen
Germany
volkm@math2.rwth-aachen.de
Complex Analysis, Combinatorics, Graph Theory

From Generalized De Giorgi Estimates to Upper Gaussian Bounds for the Heat Kernel Associated to Higher Order Elliptic Operators

Mahmoud QAFSAOUI

*Université d'Orléans,
I.U.T. d'Orléans - Département Informatique,
Rue d'Issoudun, B.P. 6729, 45067 - Orléans cedex 2, France.
e-mail: Mahmoud.Qafsaoui@iut.univ-orleans.fr*

Abstract

In this paper, we derive upper bounds and regularity for the heat kernel from an elliptic regularity property. The operators studied here are of order $2m$, with complex coefficients and are elliptic in the sense of the Gårding inequality.

1991 Mathematics Subject Classification. 35J30, 35J45, 35A08, 35K25, 35K40.

Key words. Higher order elliptic operators – Elliptic regularity – Generalized Morrey-Campanato spaces – Heat kernel – Parabolic regularity – Gaussian bounds.

Contents

1	Introduction.	3
2	From elliptic regularity to upper bounds of the heat kernel.	6
3	Proof of the main result.	10
3.1	Regularity for solutions of inhomogeneous elliptic equations. . .	10
3.2	Regularity for inhomogeneous elliptic operators.	14
3.3	Iteration	15
3.4	Regularity of the semigroup e^{-zL}	16
3.5	Perturbation techniques	19
3.6	Scaling argument	20
3.7	Hölder regularity	21
4	Remarks	21

1 Introduction.

The aim of this paper is to establish a connection between an elliptic regularity property in terms of generalized Morrey-Campanato spaces and a parabolic one characterized by upper gaussian bounds for the heat kernel of a class of higher order operators in divergence form defined via sesquilinear forms.

Let us start by recalling known results to put our work in perspective. Regarding second order uniformly elliptic operators with real-valued measurable bounded coefficients, in 1958 Nash [28] established his fundamental work on the Hölder continuity of solutions of parabolic equations with measurable coefficients. As a consequence of this result, he obtained the Hölder regularity for solutions of the associated elliptic problem. In 1967, Aronson [3] established the link with the gaussian estimates for the heat kernel. His idea relies on the parabolic Harnack inequality of Moser [27]. He gave gaussian bounds for the fundamental solution of the heat equation $\frac{\partial u}{\partial t} - \operatorname{div}(A(x, t) \nabla u) = 0$, $x \in \mathbb{R}^n$, $0 \leq t \leq T$. More precisely, that solution, denoted by $\Gamma_t(x, y, \tau)$, verifies the following estimates

$$\frac{c_0}{(t - \tau)^{n/2}} e^{-a_0 \frac{|x-y|^2}{t-\tau}} \leq \Gamma_t(x, y, \tau) \leq \frac{c_1}{(t - \tau)^{n/2}} e^{-a_1 \frac{|x-y|^2}{t-\tau}}, \quad (1)$$

for all $0 \leq \tau < t$, $x, y \in \mathbb{R}^n$ and where c_0, c_1, a_0 and a_1 are nonnegative constants depending only on dimension and the ellipticity. Fabes and Strook [12] used another approach relying on the original ideas of Nash which allowed them to obtain the Harnack inequality. In fact, upper and lower bounds are equivalent to Harnack's inequality.

One may wonder what happens to (1) when the coefficients are complex-valued functions regardless of time. In this case, $\Gamma_t(x, y, \tau) = K_{t-\tau}(x, y)$ and taking into account what we said before, we expect the estimates

$$|K_t(x, y)| \leq \frac{C}{t^{n/2}} e^{-a \frac{|x-y|^2}{t}}, \quad (2)$$

$$|K_t(x + h, y) - K_t(x, y)| + |K_t(x, y + h) - K_t(x, y)| \leq \frac{C}{t^{n/2}} \left(\frac{|h|}{t^{1/2}} \right)^\nu, \quad (3)$$

where $\nu \in (0, 1)$. The arguments of Nash-Aronson seem unusable in this case.

In order to overcome such problem, other arguments were developed. Recently, Auscher, McIntosh and Tchamitchian [6] showed that the heat kernel has gaussian estimates of type (2)-(3) if $n = 1$ or 2 . On the other hand,

Auscher, Coulhon and Tchamitchian showed in [5] that for $n \geq 5$, there exists an elliptic operator $T = -\operatorname{div}(A\nabla)$ with complex coefficients such that the semigroup e^{-tT} generated by $(-T)$ is not bounded on $L^\infty(\mathbb{R}^n)$ for all $t > 0$. In particular, the associated kernel does not verify the gaussian estimates. Note that cases $n = 3$ and 4 are still open.

Considering these counter-examples, it may be of interest to have a simple and efficient criterion to decide if the heat kernel has upper gaussian bounds. The criterion found by Auscher in [4] is an elliptic regularity condition which we denote by (D) and concerns weak solutions u of the homogeneous equation $Tu = -\operatorname{div}(A\nabla u) = 0$ on euclidean balls. That condition is an estimation of type of De Giorgi:

$$\forall x \in \mathbb{R}^n, \forall r, R, 0 \leq r \leq R, \int_{B_r(x)} |\nabla u|^2 \leq c \left(\frac{r}{R}\right)^{n-2+2\mu_0} \int_{B_R(x)} |\nabla u|^2,$$

for all weak solutions u on a ball $B_R(x)$, with constants $c > 0$ and $\mu_0 \in (0, 1]$ depending neither on u, x nor on R .

Recall that De Giorgi's theorem [11] stipulates that, in the case of real coefficients, weak solutions of $Tu = 0$ verify estimates of type (D) . According to Morrey theorem ([24], Theorem 3.5.2.) the latter estimates imply the Hölder regularity of the solutions. However, De Giorgi's result is not always true when the coefficients are complex-valued functions. Indeed, Maz'ya, Nazarov et Plamenevskii [21] gave examples of equations $Tu = 0$ whose solutions are not locally bounded when the dimension $n \geq 5$. Note that these examples were used in [5].

In [4], Auscher established the equivalence between (D) (*i.e.* De Giorgi estimates) and $((2)-(3))$ (*i.e.* Nash-Aronson estimates) when $n > 2$. The literal idea of the implication $(D) \Rightarrow ((2) - (3))$ is to consider the parabolic equation $\frac{\partial u_t}{\partial t} + Tu_t = 0$ as an elliptic equation with second member $Tu_t = -\frac{\partial u_t}{\partial t}$ for every $t > 0$. Suppose that an elliptic regularity theorem stipulates that the solution of the equation $Tu = f$ verifies $\|u\|_{E_{g(\gamma)}} \leq C_\gamma \|f\|_{E_\gamma}$ for all $\gamma \in I$, where $(E_\gamma)_{\gamma \in I}$ is a family of spaces defined for γ in an interval I and $g : I \rightarrow I$. Then we get

$$\|t \frac{\partial u_t}{\partial t}\|_{E_{g(\gamma)}} \leq C \|u_t\|_{E_{g(\gamma)}} \leq C C_\gamma \left\| \frac{\partial u_t}{\partial t} \right\|_{E_\gamma},$$

the first inequality being a consequence of the analyticity of the semigroup e^{-tT} . Therefore, we only have to iterate such inequality. In [4], E_γ is the family of Morrey-Campanato spaces and we know that it is perfectly adapted

to elliptic operators.

Now, let us come to our class of higher order operators. They are homogeneous of order $2m$ of the form

$$L_0 = \sum_{|\alpha|=m, |\beta|=m} (-1)^m D^\alpha (a_{\alpha\beta} D^\beta),$$

where $a_{\alpha\beta} \in L^\infty(\mathbb{R}^n, \mathbb{C}^n)$ and independent on time. In addition, we assume that these operators are elliptic in the sense of the strong Gårding inequality: there exists a constant $\delta_0 > 0$ such that for all $u \in H^m(\mathbb{R}^n)$,

$$\Re \int_{\mathbb{R}^n} \sum_{|\alpha|=|\beta|=m} a_{\alpha\beta}(x) D^\beta u(x) \overline{D^\alpha v(x)} dx \geq \delta_0 \int_{\mathbb{R}^n} |\nabla^m u(x)|^2 dx. \quad (4)$$

Davies showed in [10] that the kernel $K_t(x, y)$ of the semigroup e^{-tL_0} always has upper gaussian bounds when $n < 2m$. Auscher and Tchamitchian [8] extended this result to the case $n = 2m$. For the remaining case, when the dimension is bigger than the order of the operator, Davies gave counter-examples in which upper bounds are not verified by the kernel [10].

In [7], we extended the equivalence between (D) and $((2)-(3))$ to the class of higher order operators L_0 elliptic in the sense of the Gårding inequality without any condition on the dimension n and the order of the operator $2m$. More precisely, we showed that the following are equivalent:

- (i) There exists a constant $c_0 > 0$ such that for all $R > 0, x_0 \in \mathbb{R}^n$ and for all weak solution u of $T_0 u = 0$ on the ball $B_R(x_0)$, one has for all $0 < r \leq R$,

$$\int_{B_r(x_0)} |\nabla^m u|^2 \leq c_0 \left(\frac{r}{R}\right)^{n-2m+2\mu} \int_{B_R(x_0)} |\nabla^m u|^2,$$

where $\mu \in (\max(0, m - n/2), m]$ and $T_0 = L_0$ or L_0^* .

- (ii) There exists $l \in \{0, 1, \dots, m-1\}$, $\nu \in (0, 1)$ and constants c and $a > 0$ such that for all $t > 0$ and all $x, y, h \in \mathbb{R}^n$, one has

$$|D_x^\gamma K_t(x, y)| + |D_y^\gamma K_t(x, y)| \leq \frac{c}{t^{\frac{n+|\gamma|}{2m}}} \exp\left(-a\left(\frac{|x-y|}{t^{1/2m}}\right)^{\frac{2m}{2m-1}}\right),$$

for all multi-indices $\gamma \in \mathbb{N}^n$ of length $|\gamma| \leq l$ and

$$\begin{aligned} |D_x^\gamma K_t(x+h, y) - D_x^\gamma K_t(x, y)| &\leq \frac{c}{t^{\frac{n+|\gamma|}{2m}}} \left(\frac{|h|}{t^{1/2m}}\right)^\nu, \\ |D_y^\gamma K_t(x, y+h) - D_y^\gamma K_t(x, y)| &\leq \frac{c}{t^{\frac{n+|\gamma|}{2m}}} \left(\frac{|h|}{t^{1/2m}}\right)^\nu, \end{aligned}$$

for all multi-indices $\gamma \in \mathbb{N}^n$ of length $|\gamma| = l$.

This extension to higher order operators is far from being obvious.

In this paper, our objective is to generalize the elliptic condition (i) by introducing a function ψ and from it to derive upper gaussian bounds and the Hölder regularity for the heat kernel according to the function ψ which is essentially nonnegative, almost increasing, satisfying the Dini condition and the doubling property. This is the core of the following section.

2 From elliptic regularity to upper bounds of the heat kernel.

Before stating the main result, let us give a few notations and definitions.

2.1. Notation.

For a multi-index $\alpha = (\alpha_1, \dots, \alpha_n) \in \mathbb{N}^n$, we set $|\alpha| = \alpha_1 + \dots + \alpha_n$. For any $x = (x_1, \dots, x_n) \in \mathbb{R}^n$ and any multi-index $\alpha = (\alpha_1, \dots, \alpha_n)$,

$$D_x^\alpha = \frac{\partial^{\alpha_1}}{\partial x_1^{\alpha_1}} \cdots \frac{\partial^{\alpha_n}}{\partial x_n^{\alpha_n}}.$$

By $\nabla^m u$ and $|\nabla^m u|$, we denote respectively the vector $(D^\alpha u)_{|\alpha|=m}$ and its length $|\nabla^m u| = \left(\sum_{|\alpha|=m} |D^\alpha u|^2 \right)^{1/2}$. The notation H^m stands for the Sobolev space $W^{m,2}$, $m \in \mathbb{Z}$. Norms in L^p -spaces will be denoted by $\|\cdot\|_p$.

Now, let us introduce the class of operators used here and mention some of their properties. Further details can be found in [1, 7, 29].

2.2. Higher order elliptic operators.

Let $m \in \mathbb{N}^*$. Let $a_{\alpha\beta}(x)$ be bounded measurable functions on $\mathbb{R}^n$ where α, β are multi-indices such that $|\alpha| = |\beta| = m$. Set

$$\mathcal{Q}_0(u, v) = \int_{\mathbb{R}^n} \sum_{|\alpha|=|\beta|=m} a_{\alpha\beta}(x) D^\beta u(x) D^\alpha v(x) dx \quad (5)$$

for all $u, v \in H^m(\mathbb{R}^n)$. The form $\mathcal{Q}_0$ is continuous on $H^m(\mathbb{R}^n)$ and if one defines $M = \|\|(a_{\alpha\beta}(x))\|\|_\infty$, where $\|(a_{\alpha\beta}(x))\|$ is the norm of the matrix $(a_{\alpha\beta}(x))$, then for all $u, v \in H^m(\mathbb{R}^n)$,

$$|\mathcal{Q}_0(u, v)| \leq M_0 \|\nabla^m u\|_2 \|\nabla^m v\|_2. \quad (6)$$

Under these assumptions, by a variation on the Lax-Milgram lemma, there exists a unique operator in divergence form $L_0 : H^m(\mathbb{R}^n) \longrightarrow H^{-m}(\mathbb{R}^n)$, linear and continuous, such that for all $u, v \in H^m(\mathbb{R}^n)$, $\langle L_0 u, v \rangle = \mathcal{Q}_0(u, v)$. We write this operator as $(-1)^m \sum_{|\alpha|=|\beta|=m} D^\alpha (a_{\alpha\beta} D^\beta)$. Its adjoint L_0^* is associated with the coefficients $\overline{a_{\beta\alpha}}$. Note that $\langle \cdot, \cdot \rangle$ stands for the usual scalar product on L^2 .

We suppose that the class of operators L_0 is elliptic in the sense of the Gårding inequality: there exist constants $\delta_0 > 0$ and $\lambda_0 \geq 0$ such that for all $u \in H^m(\mathbb{R}^n)$,

$$\mathcal{Q}_0(u, u) \geq \delta_0 \|\nabla^m u\|_2^2 - \lambda_0 \|u\|_2^2. \quad (7)$$

Set $\mathcal{D}(L_0) = \{u \in H^m(\mathbb{R}^n) \mid L_0 u \in L^2(\mathbb{R}^n)\}$. As a consequence of (7), the operator L_0 , restricted to $\mathcal{D}(L_0)$, is maximal accretive of type $\omega < \frac{\pi}{2}$ and $(-L_0)$ is the generator of a contraction semigroup e^{-tL_0} on $L^2(\mathbb{R}^n)$.

Definitions. 1. For $m \geq 1$, $\delta_0 \geq 0$, $\lambda_0 \geq 0$ and $M_0 \geq 0$, we denote by $\mathcal{E}_{2m}(\delta_0, \lambda_0, M_0)$ the class of the operators associated with a sesquilinear form satisfying (5), (6) and (7). The operators in this class are called homogeneous elliptic operators of order $2m$.

2. We denote by $\mathcal{L}_{2m}(\delta_0, \lambda_0, M_0, M)$ the class of inhomogeneous elliptic operators of order $2m$ of the form

$$L = \sum_{|\alpha| \leq m, |\beta| \leq m} (-1)^{|\alpha|} D^\alpha (a_{\alpha\beta} D^\beta)$$

whose leading part L_0 belongs to $\mathcal{E}_{2m}(\delta_0, \lambda_0, M_0)$ and whose other coefficients are bounded with $M = \sum_{k=1}^{2m} M_k^{2m/k}$, where

$$M_k = \sup\{\|a_{\alpha\beta}\|_\infty \mid |\alpha|, |\beta| \leq m, |\alpha| + |\beta| = 2m - k\}.$$

The operators L also satisfy a Gårding inequality as a consequence of the one for their leading part L_0 and there exists $\lambda \geq 0$ such that the operator $L + \lambda$, restricted to $\mathfrak{D}(L)$, is maximal accretive and $-(L + \lambda)$ is the generator of a contraction semigroup $e^{-t(L+\lambda)}$ on $L^2(\mathbb{R}^n)$ that is analytic. In addition, $\|e^{-tL}\|_{L^2 \rightarrow L^2} \leq e^{\lambda t}$.

2.3. Generalized Morrey-Campanato and Hölder space.

Let $0 \leq \lambda \leq n$ and ϕ a nonnegative continuous function on $[0, R_0]$. Define the generalized Morrey space $L_\phi^{2,\lambda}(\mathbb{R}^n) = L_\phi^{2,\lambda}$ as the space of functions $u \in L_{loc}^2$ such that

$$\|u\|_{L_\phi^{2,\lambda}} := \sup_{x \in \mathbb{R}^n, 0 < \rho \leq 1} \frac{1}{\phi(\rho)} \left(\rho^{-\lambda} \int_{B_\rho(x)} |u|^2 \right)^{1/2} < +\infty.$$

Note that $L_\phi^{2,\lambda_1} \subset L_\phi^{2,\lambda_2}$ if $0 \leq \lambda_2 < \lambda_1 \leq n$ and $L^{2,\lambda}$ will correspond to L_ϕ^{2,λ_2} when $\phi \equiv 1$.

Let $0 \leq \lambda \leq n + 2k$, where $k \in \mathbb{N}^*$. Define the generalized Campanato space $\mathcal{L}_{k,\phi}^{2,\lambda}(\mathbb{R}^n) = \mathcal{L}_{k,\phi}^{2,\lambda}$ as the space of functions $u \in L_{loc}^2$ such that

$$\|u\|_{\mathcal{L}_{k,\phi}^{2,\lambda}} := \sup_{x \in \mathbb{R}^n, 0 < \rho \leq 1} \frac{1}{\phi(\rho)} \left(\rho^{-\lambda} \inf_{P \in \mathcal{P}_{k-1}} \int_{B_\rho(x)} |u(y) - P(y)|^2 dy \right)^{1/2} < +\infty,$$

where $\mathcal{P}_l$ is the class of polynomials of degree less than or equal to l . The spaces $L_\phi^{2,\lambda}$ and $L_\phi^{2,0} \cap \mathcal{L}_{k,\phi}^{2,\lambda}$, equipped with the norm $\|u\|_{L_\phi^{2,\lambda}}$ and $\|u\|_{L_\phi^{2,0}} + \|u\|_{\mathcal{L}_{k,\phi}^{2,\lambda}}$ respectively, are Banach spaces. We give two useful properties for later use:

- (a) Poincaré inequality (see Lemma 18 of [7]): There exists a constant $c > 0$ such that for all $u \in H_{loc}^m$, $\lambda \geq 0$ and $0 \leq s \leq m$, we have

$$\|\nabla^s u\|_{\mathcal{L}_{m-s,\phi}^{2,\lambda+2(m-s)}} \leq c \|\nabla^m u\|_{L_\phi^{2,\lambda}}. \quad (8)$$

- (b) If $0 \leq \lambda < n + 2k$ and $0 \leq k \leq m$ then $L_\phi^{2,0} \cap \mathcal{L}_{m,\phi}^{2,\lambda} \simeq L_\phi^{2,0} \cap \mathcal{L}_{k,\phi}^{2,\lambda}$. The isomorphism remains true if one replaces $L_\phi^{2,0}$ by L^2 . Note that $\mathcal{L}_{0,\phi}^{2,\lambda} = L_\phi^{2,\lambda}$ and the symbol $\simeq$ means that the spaces are isomorphic as normed spaces.

Let ϕ be a nonnegative function such that $\phi(0) = 0$. By $\mathcal{C}_\phi(\Omega)$ we denote the generalized homogeneous Hölder space of continuous functions u such that $|u(x) - u(y)| \leq C\phi(|x - y|)$ for a constant C and for $x, y \in \Omega$, $x \neq y$.

2.4. Statement of the main result.

Let us give the elliptic and the parabolic properties.

Definition (Dirichlet property). Let $L_0 \in \mathcal{E}_{2m}(\delta_0, 0, M_0)$ and ψ be a non-negative function on $(0, +\infty)$. We say that L_0 verifies the property (D_ψ) if there exists a constant $c_0 > 0$ such that for all $R > 0, x_0 \in \mathbb{R}^n$ and for all v L_0 -harmonic function (i.e. weak solution of $L_0 u = 0$) on $B_R(x_0)$,

$$\int_{B_r(x_0)} |\nabla^m v(x)|^2 dx \leq c_0 \left(\frac{r}{R} \right)^{n-2m+\epsilon} \frac{\psi^2(r)}{\psi^2(R)} \int_{B_R(x_0)} |\nabla^m v(x)|^2 dx, \quad (D_\psi)$$

for all $0 < r \leq R$. The parameter $\epsilon > 0$ is assumed to be small enough.

In terms of regularity in the Morrey-Campanato spaces, if v verifies the inequality (D_ψ) for all balls, then $\nabla^m v$ belongs to the generalized Morrey space $L_\psi^{2,n-2m+\epsilon}$. We will often omit the centre of balls which will be assumed as co-centred.

Definition (Gaussian property). Let $L_0 \in \mathcal{E}_{2m}(\delta_0, 0, M_0)$ and $\tilde{\psi}$ be a nonnegative function on $(0, +\infty)$ satisfying $\tilde{\psi}(0) = 0$. We say that L_0 verifies the property $(G_{\tilde{\psi}})$ if there exist constants c and $a > 0$ such that for all $t > 0$ and all $x, y, h \in \mathbb{R}^n$,

$$|K_t(x, y)| \leq \frac{c}{t^{n/2m}} \exp\left(-a\left(\frac{|x-y|}{t^{1/2m}}\right)^{2m/2m-1}\right),$$

$$|K_t(x+h, y) - K_t(x, y)| + |K_t(x, y+h) - K_t(x, y)| \leq \frac{c}{t^{n/2m}} \tilde{\psi}\left(\frac{|h|}{t^{1/2m}}\right).$$

The main result of this paper is the following.

Theorem 2.1 (Main result). *Let $L_0 \in \mathcal{E}_{2m}(\delta_0, 0, M_0)$, $R_0 \in (0, +\infty)$ and let ψ be a nonnegative function on $(0, R_0]$ such that*

- (i) *ψ is almost increasing: there exists a constant $c_\psi > 0$ such that if $r \leq r'$ then $\psi(r) \leq c_\psi \psi(r')$,*
- (ii) *ψ verifies the doubling condition: there exist two constants c_ψ and c'_ψ such that $c_\psi \psi(r) \leq \psi(2r) \leq c'_\psi \psi(r)$ for all $0 < r \leq R_0$,*
- (iii) *there exist $\varepsilon > 0$ and $c > 0$ such that for all $0 < r \leq R_0$, $\psi(r) \geq cr^\varepsilon$,*
- (iv) *$\lim_{t \rightarrow 0} \psi(t) = 0$.*
- (v) *ψ verifies the Dini condition:*

$$\int_0^{R_0} \frac{\psi(t)}{t} dt < +\infty.$$

Under these assumptions, if both L_0 and L_0^ verify the Dirichlet condition (D_ψ) then L_0 verifies the Gaussian property $(G_{\tilde{\psi}})$ with*

$$\tilde{\psi}(r) = \psi(r) + \int_0^r \frac{\psi(t)}{t} dt.$$

It is worth noting that in the case of power functions one gets exactly the result in [7].

3 Proof of the main result.

In this section, we intend to prove the main theorem 2.1. The argument will be divided into six steps: 1- The Morrey-Campanato method consisting in improving the regularity of weak solutions to the equation $L_0 u = f$ with initial data in generalized Morrey-Campanato spaces, 2- A gain in regularity, in terms of Morrey-Campanato spaces, for inhomogeneous elliptic operators L whose leading part is L_0 . This step is a consequence of the previous one, 3- Iteration of L via an elliptic equation, 4- The continuity of the semigroup e^{-zL} from $L^1(\mathbb{R}^n)$ into $\mathcal{C}_{\tilde{\psi}}$ which involves estimates on kernels, 5- A variant of perturbation method due to Davies to obtain the exponential decay, 6- A scaling argument to get estimates on the heat kernel of e^{-tL_0} .

3.1 Regularity for solutions of inhomogeneous elliptic equations.

Before giving a regularity result which is useful for our argument, let us state a technical lemma which can be considered as a generalization of a result due to Campanato.

Lemma 3.1. *Let $R_0 \in \mathbb{R}^+$ and $H : [0, R_0] \rightarrow \mathbb{R}^+$ be an almost increasing function (i.e., $\exists C_H > 0$ such that $H(r) \leq C_H H(R)$ for $r \leq R$). Let $\tilde{g}$ and ω be two nonnegative functions on $[0, R_0]$. Assume that*

1. *there exist $A \geq 1$ and $B > 0$ such that for all $0 < r \leq R \leq R_0$,*

$$H(r) \leq A \frac{g(r)}{g(R)} H(R) + B \tilde{g}(R),$$

where $g(r) = \tilde{g}(r)\omega(r)$,

2. *$\tilde{g}$ verifies the doubling property on $[0, R_0]$,*
3. *there exist C_1, C_2, γ and $\theta > 0$ such that for all $0 < r \leq R \leq R_0$,*

$$(i) \sup_{0 \leq r \leq R} \left(\frac{r}{R} \right)^{-\gamma} \frac{\omega(r)}{\omega(R)} \leq C_1,$$

$$(ii) \int_r^{+\infty} \frac{dt}{t^{1-\theta}\omega(t)} \leq \frac{C_2}{r^{-\theta}\omega(r)}.$$

Then, there exists a constant C depending only on A, C_H, C_1 and C_2 such that for all $0 < r \leq R \leq R_0$,

$$H(r) \leq C \frac{\tilde{g}(r)}{\tilde{g}(R)} H(R) + C B \tilde{g}(r). \quad (9)$$

Proof. Let $0 < \tau < 1$ and $0 < R \leq R_0$. By induction we show that for all $k \in \mathbb{N}^*$,

$$H(\tau^k R) \leq A^k \frac{g(\tau^k R)}{g(R)} H(R) + B \sum_{i=0}^{k-1} A^{k-1-i} \frac{g(\tau^k R)}{g(\tau^{i+1} R)} \tilde{g}(\tau^i R). \quad (10)$$

Now, using successively the monotonicity of H , the inequality (10) and the doubling property of $\tilde{g}$ we get for $\tau^{k+1} R < r < \tau^k R$ ($k \geq 0$),

$$\begin{aligned} H(r) &\leq C_H H(\tau^k R) \\ &\leq C_H \left(A^k \frac{g(\tau^k R)}{g(R)} H(R) + B \sum_{i=0}^{k-1} A^{k-1-i} \frac{g(\tau^k R)}{g(\tau^{i+1} R)} \tilde{g}(\tau^i R) \right) \\ &= C_H \left(A^k \frac{\omega(\tau^k R)}{\omega(R)} \frac{\tilde{g}(\tau^k R)}{\tilde{g}(R)} H(R) + B \sum_{i=0}^{k-1} \frac{A^{k-1-i}}{\omega(\tau^{i+1} R)} \frac{\tilde{g}(\tau^i R)}{\tilde{g}(\tau^{i+1} R)} \omega(\tau^k R) \tilde{g}(\tau^k R) \right) \\ &\leq C_H \left(\left(A^k \frac{\omega(\tau^k R)}{\omega(R)} \right) \frac{\tilde{g}(\tau^k R)}{\tilde{g}(R)} H(R) + B C_{\tilde{g}} \left(\sum_{i=0}^{k-1} \frac{A^{k-1-i}}{\omega(\tau^{i+1} R)} \right) \omega(\tau^k R) \tilde{g}(\tau^k R) \right). \end{aligned}$$

Now, to obtain (9) it suffices to have

$$\begin{cases} \sup_{k,R} A^k \frac{\omega(\tau^k R)}{\omega(R)} \leq C_1 & (*) \\ \sum_{i=0}^{k-1} \frac{A^{k-1-i}}{\omega(\tau^{i+1} R)} \leq \frac{C_2}{\omega(\tau^k R)}. & (**) \end{cases}$$

Choose $A^k = \tau^{k \log_{\tau} A}$ ($r \approx \tau^k R$) and $t = \tau^{i+1} R$ ($\frac{r}{t} = \tau^{k-1-i}$) then (*) and (**) become respectively

$$\sup_{0 < r \leq R} \left(\frac{r}{R} \right)^{\log_{\tau} A} \frac{\omega(r)}{\omega(R)} \leq C_1,$$

i.e. there exists $\gamma > 0$ such that $\sup_{0 < r \leq R} \left(\frac{r}{R} \right)^{-\gamma} \frac{\omega(r)}{\omega(R)} \leq C_1$ and

$$\int_r^R \frac{dt}{t^{1+\log_{\tau} A} \omega(t)} \leq \frac{C_2}{t^{\log_{\tau} A} \omega(r)} \left(\text{i.e. } \exists \theta > 0 / \int_r^{+\infty} \frac{dt}{t^{1-\theta} \omega(t)} \leq \frac{C_2}{r^{-\theta} \omega(r)} \right)$$

$$\text{since } \sum_{i=0}^{k-1} \frac{A^{k-1-i}}{\omega(\tau^{i+1} R)} \approx \int_r^R \frac{(r/t)^{\log_{\tau} A}}{\omega(t)} \frac{dt}{t}.$$

Eventually, we conclude by choosing $0 < \tau < 1$ such that $\log_{\tau} A \geq -\gamma$ (resp. $\log_{\tau} A \leq -\theta$) for (*) (resp. (**)).

□

Here and thereafter, we assume that L_0 verifies the property (D_ψ) with the constant c_0 .

Theorem 3.1. *Let $R \leq 1$ and $u \in H^m(B_R)$ a weak solution on B_R to the equation*

$$L_0 u = g_0 + \sum_{0 < |\alpha| < m} D^\alpha g_\alpha + \sum_{|\alpha|=m} D^\alpha h_\alpha, \quad (11)$$

where $g_0 \in L_\psi^{2,\eta}$, $g_\alpha \in L_\psi^{2,\lambda}$ ($0 < |\alpha| < m$), $h_\alpha \in L_\psi^{2,\mu}$ ($|\alpha| = m$) and $0 \leq \eta, \lambda, \mu < n$. Then there exists a constant $c > 0$ depending on $c_0, \delta_0, M_0, \eta, \lambda$ and μ such that if $\nu = \inf(\eta + 2m, \lambda + 2, \mu, n - 2m)$ then for all $0 < r \leq R$,

$$\int_{B_r} |\nabla^m u|^2 \leq c \left(\frac{r}{R} \right)^\nu \frac{\psi^2(r)}{\psi^2(R)} \int_{B_R} |\nabla^m u|^2 + c r^\nu \psi^2(r) \mathcal{X},$$

$$\text{where } \mathcal{X} = \|g_0\|_{L_\psi^{2,\eta}}^2 + \sum_{0 < |\alpha| < m} \|g_\alpha\|_{L_\psi^{2,\lambda}}^2 + \sum_{|\alpha|=m} \|h_\alpha\|_{L_\psi^{2,\mu}}^2.$$

Proof. The method applied here is a variation on [7], Theorem 22. In this argument, $\|\cdot\|_2$ stands for $\left(\int_{B_R} |\cdot|^2 \right)^{1/2}$.

Let $0 < r \leq R$ and $u \in H^m(B_R)$ solution to (11). By Lax-Milgram theorem, for all $u \in H^m(B_R)$ there exists a unique function $w \in H_0^m(B_R)$ such that $v = w + u$ is L_0 -harmonic on B_R and

$$\int_{B_R} |\nabla^m v|^2 \leq c \int_{B_R} |\nabla^m u|^2, \quad (12)$$

where $c > 0$ is a constant depending only on the ellipticity constant δ_0 and M_0 . Applying (D_ψ) we get

$$\begin{aligned} \int_{B_r} |\nabla^m u|^2 &\leq 2 \left(\int_{B_r} |\nabla^m v|^2 + \int_{B_r} |\nabla^m(u-v)|^2 \right) \\ &\leq 2c_0 \left(\frac{r}{R} \right)^{n-2m+\epsilon} \frac{\psi^2(r)}{\psi^2(R)} \int_{B_R} |\nabla^m v|^2 + 2 \int_{B_R} |\nabla^m(u-v)|^2, \end{aligned}$$

and hence by virtue of (12)

$$\int_{B_r} |\nabla^m u|^2 \leq 2c_1 \left(\frac{r}{R} \right)^{n-2m+\epsilon} \frac{\psi^2(r)}{\psi^2(R)} \int_{B_R} |\nabla^m u|^2 + 2 \int_{B_R} |\nabla^m z|^2, \quad (13)$$

where $z = u - v$. One has $L_0 z = L_0 u = g_0 + \sum_{0 < |\alpha| < m} D^\alpha g_\alpha + \sum_{|\alpha|=m} D^\alpha h_\alpha$ and

$$\mathcal{Q}_0(z, z) = \langle g_0, z \rangle + \sum_{0 < |\alpha| < m} (-1)^{|\alpha|} \langle g_\alpha, D^\alpha z \rangle + (-1)^m \sum_{|\alpha|=m} \langle h_\alpha, D^\alpha z \rangle.$$

By applying the strong version of Gårding inequality (7) (i.e. $\lambda_0 = 0$) and Cauchy-Schwarz inequality, we obtain

$$\delta_0 \int_{B_R} |\nabla^m z|^2 \leq \|g_0\|_2 \|z\|_2 + \sum_{0 < |\alpha| < m} \|g_\alpha\|_2 \|D^\alpha z\|_2 + \sum_{|\alpha|=m} \|h_\alpha\|_2 \|D^\alpha z\|_2.$$

By Poincaré inequality (in $H_0^m(B_R)$), there exists a constant $c > 0$ independent on R and u such that for all $\alpha \in \mathbb{N}^n$, $0 \leq |\alpha| < m$,

$$\|D^\alpha z\|_2 \leq c R^{m-|\alpha|} \|\nabla^m z\|_2.$$

Therefore

$$\begin{aligned} & \int_{B_R} |\nabla^m z|^2 \leq \\ & c \left(R^m \|g_0\|_2 \|\nabla^m z\|_2 + \sum_{0 < |\alpha| < m} R^{(m-|\alpha|)} \|g_\alpha\|_2 \|\nabla^m z\|_2 + \sum_{|\alpha|=m} \|h_\alpha\|_2 \|D^\alpha z\|_2 \right) \\ & \leq c \left(R^m \|\nabla^m z\|_2 \|g_0\|_2 + R \|\nabla^m z\|_2 \sum_{0 < |\alpha| < m} \|g_\alpha\|_2 + \|\nabla^m z\|_2 \left(\sum_{|\alpha|=m} \|h_\alpha\|_2^2 \right)^{1/2} \right), \end{aligned}$$

where we have used the fact that $R \leq 1$. Simplifying by $\|\nabla^m z\|_2$, it follows

$$\begin{aligned} \int_{B_R} |\nabla^m z|^2 & \leq c^2 \left(R^m \|g_0\|_2 + R \sum_{0 < |\alpha| < m} \|g_\alpha\|_2 + \left(\sum_{|\alpha|=m} \|h_\alpha\|_2^2 \right)^{1/2} \right)^2 \\ & \leq c^2 \left(R^{2m} \|g_0\|_2^2 + R^2 \sum_{0 < |\alpha| < m} \|g_\alpha\|_2^2 + \sum_{|\alpha|=m} \|h_\alpha\|_2^2 \right). \end{aligned}$$

Since (recall that $\|\cdot\|_2 = (\int_{B_R} |\cdot|^2)^{1/2}$), $\|g_0\|_2 \leq R^{\eta/2} \psi(R) \|g_0\|_{L_\psi^{2,\eta}}$, $\|g_\alpha\|_2 \leq R^{\lambda/2} \psi(R) \|g_\alpha\|_{L_\psi^{2,\lambda}}$ and $\|h_\alpha\|_2 \leq R^{\mu/2} \psi(R) \|h_\alpha\|_{L_\psi^{2,\mu}}$. Then

$$\int_{B_R} |\nabla^m z|^2 \leq c^2 R^{\inf(\eta+2m, \lambda+2, \mu)} \psi^2(R) \mathcal{X}, \quad (14)$$

where

$$\mathcal{X} = \|g_0\|_{L_\psi^{2,\eta}}^2 + \sum_{0 < |\alpha| < m} \|g_\alpha\|_{L_\psi^{2,\lambda}}^2 + \sum_{|\alpha|=m} \|h_\alpha\|_{L_\psi^{2,\mu}}^2.$$

It then follows from (13) and (14) that

$$\int_{B_r} |\nabla^m u|^2 \leq C \left(\frac{r}{R} \right)^{n-2m+\epsilon} \frac{\psi^2(r)}{\psi^2(R)} \int_{B_R} |\nabla^m u|^2 + CR^{\inf(\eta+2m, \lambda+2, \mu)} \psi^2(R) \mathcal{X}.$$

Eventually, we conclude by applying Lemma 3.1 to the inequality above to obtain

$$\int_{B_r} |\nabla^m u|^2 \leq C \left(\frac{r}{R} \right)^\nu \frac{\psi^2(r)}{\psi^2(R)} \int_{B_R} |\nabla^m u|^2 + Cr^\nu \psi^2(r) \mathcal{X},$$

where $\nu = \inf(\eta + 2m, \lambda + 2, \mu, n - 2m)$.

3.2 Regularity for inhomogeneous elliptic operators.

Let us consider the class $\mathcal{L}_{2m}(\delta_0, 0, M_0, M)$ of inhomogeneous operators

$$L = \sum_{|\alpha| \leq m, |\beta| \leq m} (-1)^{|\alpha|} D^\alpha (a_{\alpha\beta} D^\beta),$$

whose leading part L_0 belongs to $\mathcal{E}_{2m}(\delta_0, 0, M_0)$. We can consider L as $L = L_0 +$ perturbation.

In the rest of our demonstration, we assume that L_0 and L_0^* verify (D_ψ) with the constant c_0 .

Theorem 3.2. *Let $u \in H_{loc, unif}^m$ solution to $Lu = h$. If $h \in L^2 \cap L_\psi^{2, \kappa}$ ($0 \leq \kappa < n$) then $\nabla^m u \in L_\psi^{2, \tau}$ where $\tau = \inf(\kappa + 2, n - 2m)$ and one has for all θ verifying $2 \leq \theta \leq \tau$,*

$$\|\nabla^m u\|_{L_\psi^{2, \theta}} \leq c \left(\|u\|_{H_{loc, unif}^m} + \|h\|_{L_\psi^{2, \kappa}} \right),$$

where c depends on c_0, δ_0, M and κ .

Proof. The argument is similar to the one used to prove Theorem 25 [7]. The adaptation is easy and left to the reader. $\square$

Now, to state the following result, suppose in addition that $\Re a_{00} > \lambda$ for $\lambda > 0$ large enough so that

$$L : H^m(\mathbb{R}^n) \longrightarrow H^{-m}(\mathbb{R}^n)$$

is an isomorphism. Theorem 3.2 yields

Corollary 3.1. *Under the assumptions above, for all real δ such that $0 \leq \delta < n$,*

$$\nabla^m L^{-1} : L^2 \cap L_\psi^{2, \delta} \longrightarrow L^2 \cap L_\psi^{2, \tau},$$

where $\tau = \inf(\delta + 2, n - 2m)$ and $T : \mathcal{A} \longrightarrow \mathcal{B}$ means that the operator T extends to a bounded operator from $\mathcal{A}$ into $\mathcal{B}$.

3.3 Iteration

Now, we iterate the inhomogeneous operator L via the equation $Lu_z = -\frac{du_z}{dz}$ where $u_z = e^{-zL}f$ for $f \in L^2$. For this purpose, we need the following lemma.

Lemma 3.2. *Let $\Gamma_\eta = \{z \in \mathbb{C}, \Re z > 0, |\arg z| < \eta, \eta \in (0, \frac{\pi}{2})\}$ a sector on which e^{-zL} is a contraction semigroup on L^2 . Let $z \mapsto u_z$ be an $L^2(\mathbb{R}^n)$ -valued holomorphic function on Γ_η and let E be a Banach space included in $L^2_{loc}(\mathbb{R}^n)$.*

Assume that for all $z \in \Gamma_\eta$, $u_z \in E$ with $\|u_z\|_E \leq \frac{C_1}{|z|^\alpha}$ where C_1 and α are nonnegative constants. Then for all sector $\Gamma_{\eta'}$ strictly included in Γ_η ($\eta' < \eta$), there exists a constant $C_2 = C_2(\eta, \eta', \alpha)$ such that for all $z \in \Gamma_{\eta'}$, $\frac{du_z}{dz} \in E$ and $\left\| \frac{du_z}{dz} \right\|_E \leq \frac{C_2}{|z|^{\alpha+1}}$.

Proof. Let $\mathcal{C}_z$ be the circle of equation $|w - z| = R_z$ where $R_z = \frac{1}{2}d(z, \Gamma_\eta^c)$. One has

$$\frac{du_z}{dz} = \frac{1}{2\pi i} \int_{\mathcal{C}_z} \frac{u_w}{(w - z)^2} dw, \quad z \in \Gamma_{\eta'}.$$

The integral converges a priori in L^2 , thus in L^2_{loc} . On the other hand

$$\int_{\mathcal{C}_z} \left\| \frac{u_w}{(w - z)^2} \right\|_E |dw| \leq \int_{\mathcal{C}_z} \frac{\|u_w\|_E}{|w - z|^2} |dw| \leq \int_{\mathcal{C}_z} \frac{C}{|w|^\alpha} \frac{1}{|w - z|^2} |dw|.$$

Since there exist constants C_1 and C_2 depending only on η and η' such that $C_1|z| \leq |w| \leq C_2|z|$ on $\mathcal{C}_z$,

$$\int_{\mathcal{C}_z} \left\| \frac{u_w}{(w - z)^2} \right\|_E |dw| \leq \frac{C}{C_1^\alpha |z|^\alpha} \int_{\mathcal{C}_z} \frac{1}{|w - z|^2} |dw| = \frac{C}{C_1^\alpha |z|^\alpha} \frac{C'}{|z|}.$$

Therefore $\int_{\mathcal{C}_z} \frac{u_w}{(w - z)^2} dw$ converges in E and $\left\| \frac{du_z}{dz} \right\|_E \leq \frac{c(\eta, \eta', \alpha)}{|z|^{\alpha+1}}$. $\square$

Now, let us begin the iteration of the inhomogeneous operator L using the equation $Lu_z = -\frac{du_z}{dz}$ where $u_z = e^{-zL}f$ for $f \in L^2$. For this purpose let us fix $(\Gamma_l)_{l \geq 0}$ a family of strictly decreasing embedded sectors with $\Gamma_0 = \Gamma_\eta$.

* Let $z \in \Gamma_\eta$, one has $u_z \in L^2$ with $\|u_z\|_2 \leq \|f\|_2$ for all $f \in L^2$. Since $\|\nabla^m u_z\|_2 \leq |z|^{-1/2} \|u_{z/2}\|_2$ (see [8], Proposition 2) then $\nabla^m u_z \in L^2$ with $\|\nabla^m u_z\|_2 \leq |z|^{-1/2} \|f\|_2$. It follows then by Lemma 3.2 that $\nabla^m \left(\frac{du_z}{dz} \right) \in L^2 \subset L^{2,0}$ with

$$\left\| \nabla^m \left(\frac{du_z}{dz} \right) \right\|_2 \leq C |z|^{-3/2} \|f\|_2, \quad z \in \Gamma_1.$$

Hence, Corollary 3.1 yields

$$\nabla^m u_z = -\nabla^m L^{-1} \left(\frac{du_z}{dz} \right) \in L^2 \cap L^{2,\tau_1} \text{ where } \tau_1 = \inf(2, n - 2m)$$

and then $\nabla^m u_z \in L^2 \cap L_{\psi}^{2,\tau_1-\varepsilon}$ since $L^{2,\tau_1} \subset L_{\psi}^{2,\tau_1-\varepsilon}$ by condition (iii) of Theorem 2.1.

* We start again with $\tau_1 - \varepsilon$. The same argument yields $\nabla^m \left(\frac{du_z}{dz} \right) \in L^2 \cap L_{\psi}^{2,\tau_1-\varepsilon}$ and then $\nabla^m u_z \in L^2 \cap L_{\psi}^{2,\tau_2}$, where $\tau_2 = \inf(\tau_1 - \varepsilon + 2, n - 2m)$, and so forth until the step k_0 characterized by $\tau_{k_0-1} = \tau_{k_0-2} + 2$ (i.e. $\tau_{k_0-2} + 2 < n - 2m$) and $\tau_{k_0} = n - 2m$ (i.e. $\tau_{k_0-1} + 2 > n - 2m$) and then

$$\nabla^m u_z \in L^2 \cap L_{\psi}^{2,n-2m}. \quad (15)$$

3.4 Regularity of the semigroup e^{-zL}

Here, we intend to show that u_z belongs to the generalized Hölder space $\mathcal{C}_{\tilde{\psi}}$. For this purpose, we need the following result.

Proposition 3.1. *Let $\phi : (0, R_0] \longrightarrow \mathbb{R}^{+*}$ be a function such that*

1. ϕ is almost increasing,
2. ϕ verifies the doubling property,
3. $\lim_{t \rightarrow 0} \phi(t) = 0$,
4. $\int_0^{R_0} \frac{\phi(t)}{t} dt < \infty$.

Under these assumptions, if $\nabla u \in L_{\phi}^{2,n-2}$ then $u \in \mathcal{C}_{\tilde{\phi}}$ with

$$|u(x) - u(y)| \leq C \|\nabla u\|_{L_{\phi}^{2,n-2}} \tilde{\phi}(|x - y|)$$

where

$$\tilde{\phi}(r) = \phi(r) + \int_0^r \frac{\phi(t)}{t} dt.$$

Proof. Let $R \in (0, R_0]$ and $u_{x,s} = |B(x, s)|^{-1} \int_{B(x,s)} u(y) dy$. Writing

$$|u(x) - u(y)| \leq |u(x) - u_{x,2R}| + |u_{x,2R} - u_{y,2R}| + |u(y) - u_{y,2R}|.$$

* *Step 1 (Estimation of $|u(\cdot) - u_{\cdot,2R}|$):* Let $0 < r < R$. One has

$$|u_{x_0,R} - u_{x_0,r}|^2 \leq 2(|u(x) - u_{x_0,R}|^2 + |u(x) - u_{x_0,r}|^2)$$

and

$$r^n |u_{x_0,R} - u_{x_0,r}|^2 \leq 2 \left(\int_{B(x_0,R)} |u(x) - u_{x_0,R}|^2 dx + \int_{B(x_0,r)} |u(x) - u_{x_0,r}|^2 dx \right).$$

Through Poincaré inequality and the definition of $L_\phi^{2,n-2}$ it follows that

$$\begin{aligned} |u_{x_0,R} - u_{x_0,r}|^2 &\leq 2r^{-n} \left(R^2 \int_{B(x_0,R)} |\nabla u(x)|^2 dx + r^2 \int_{B(x_0,r)} |\nabla u(x)|^2 dx \right) \\ &\leq 4r^{-n} R^2 \int_{B(x_0,R)} |\nabla u(x)|^2 dx \\ &\leq 4r^{-n} R^2 R^{n-2} \phi^2(R) \|\nabla u\|_{L_\phi^{2,n-2}}^2 \end{aligned}$$

and then

$$|u_{x_0,R} - u_{x_0,r}| \leq 2 \left(\frac{R}{r} \right)^{n/2} \phi(R) \|\nabla u\|_{L_\phi^{2,n-2}}. \quad (16)$$

Now, let $R_i = 2^{-i}R$; it follows by (16) that

$$|u_{x_0,R_i} - u_{x_0,R_{i+1}}| \leq 2^{1+n/2} \phi(2^{-i}R) \|\nabla u\|_{L_\phi^{2,n-2}}.$$

By the doubling property, one gets for $k < h$ that

$$|u_{x_0,R_k} - u_{x_0,R_h}| \leq C(n) \phi(R_k) \|\nabla u\|_{L_\phi^{2,n-2}} \leq C(n) \phi(R_h) \|\nabla u\|_{L_\phi^{2,n-2}}.$$

Since $\lim_{t \rightarrow 0} \phi(t) = 0$, then $(u_{x_0,R_h})_h$ is a Cauchy sequence. By $\tilde{u}(x_0)$ we denote its limit (i.e. $\lim_{h \rightarrow +\infty} u_{x_0,R_h} = \tilde{u}(x_0)$).

On the other hand,

$$|u_{x_0,R_k} - u_{x_0,R_h}| \leq C(n) \|\nabla u\|_{L_\phi^{2,n-2}} \sum_{i=0}^{\infty} \phi(2^{-i}R) \approx C(n) \|\nabla u\|_{L_\phi^{2,n-2}} \int_0^1 \frac{\phi(Rt)}{t} dt.$$

Thus

$$|u_{x_0,R_k} - u_{x_0,R_h}| \leq C(n) \|\nabla u\|_{L_\phi^{2,n-2}} \int_0^R \frac{\phi(t)}{t} dt. \quad (17)$$

Since $\lim_{r \rightarrow 0^+} u_{x,r} = u$ in L^1 -norm, then $\tilde{u} = u$.

Eventually, by taking the limit in (17) ($h \rightarrow \infty$) and choosing $k = 0$, we get

$$|u(x) - u_{x,R}| \leq C(n) \|\nabla u\|_{L_\phi^{2,n-2}} \int_0^R \frac{\phi(t)}{t} dt. \quad (18)$$

* *Step 2 (Estimation of $|u_{x,2R} - u_{y,2R}|$):* Writing $|u_{x,2R} - u_{y,2R}| \leq |u_{x,2R} - u(z)| + |u(z) - u_{y,2R}|$ and integrating, on $B(x, 2R) \cap B(y, 2R) := B_1 \cap B_2$, with respect to z , we get

$$|B_1 \cap B_2| (|u_{x,2R} - u_{y,2R}|) \leq \int_{B_1 \cap B_2} |u(z) - u_{x,2R}| dz + \int_{B_1 \cap B_2} |u(z) - u_{y,2R}| dz$$

and then

$$|u_{x,2R} - u_{y,2R}| \leq |B_1 \cap B_2|^{-1} \left(\int_{B_1} |u(z) - u_{x,2R}| dz + \int_{B_2} |u(z) - u_{y,2R}| dz \right).$$

Using successively Cauchy-Schwarz and Poincaré inequalities we obtain

$$\begin{aligned} \int_{B_1} |u(z) - u_{x,2R}| dz &\leq CR^{n/2} \left(\int_{B_1} |u(z) - u_{x,2R}|^2 dz \right)^{1/2} \leq CR^{n/2+1} \left(\int_{B_1} |\nabla u|^2 dz \right)^{1/2} \\ &\leq CR^{n/2+1} (2R)^{\frac{n-2}{2}} \phi(2R) \|\nabla u\|_{L_\phi^{2,n-2}} = C(n) R^n \phi(2R) \|\nabla u\|_{L_\phi^{2,n-2}}. \end{aligned}$$

Similarly,

$$\int_{B_2} |u(z) - u_{y,2R}| dz \leq C(n) R^n \phi(2R) \|\nabla u\|_{L_\phi^{2,n-2}}.$$

Since $B(x, R) \subset B_1 \cap B_2$ then $|B_1 \cap B_2|^{-1} \leq |B(x, R)|^{-1} = R^{-n}$ and

$$|u_{x,2R} - u_{y,2R}| \leq C(n) \phi(2R) \|\nabla u\|_{L_\phi^{2,n-2}}.$$

It follows by the doubling property that

$$|u_{x,2R} - u_{y,2R}| \leq C'(n) \|\nabla u\|_{L_\phi^{2,n-2}} \phi(R).$$

Eventually, combining the previous steps and choosing $R = |x - y|$ yield

$$|u(x) - u(y)| \leq C \|\nabla u\|_{L_\phi^{2,n-2}} \tilde{\phi}(|x - y|)$$

and Proposition 3.1 is proved. $\square$

Now, we are able to state our result. Indeed, by applying Poincaré inequality (8), it follows from (15) that $\nabla u_z \in L^2 \cap \mathcal{L}_{m-1,\psi}^{2,n-2} \simeq L^2 \cap L_\psi^{2,n-2}$. Therefore, $u_z \in \mathcal{C}_{\tilde{\psi}}$ thanks to Proposition 3.1 and then $e^{-zL} : L^2 \rightarrow \mathcal{C}_{\tilde{\psi}}$.

On the other hand, we have $e^{-zL} : L^2 \rightarrow L^\infty$. To see that, consider the equation (11) with $g_0 \in L^{2,\eta}$, $g_\alpha \in L^{2,\lambda}$, $h_\alpha \in L^{2,\mu}$ and follow the argument in [7] since in the proof of Theorem 3.1, (13) yields

$$\int_{B_r} |\nabla^m u|^2 \leq 2c_1 c \left(\frac{r}{R}\right)^{n-2m+\epsilon} \int_{B_R} |\nabla^m u|^2 + 2 \int_{B_R} |\nabla^m z|^2,$$

since ψ is almost increasing.

This is also true for L^* (i.e. $e^{-zL^*} : L^2 \rightarrow L^\infty$). It follows by duality that $e^{-zL} : L^1 \rightarrow L^2$ and then

$$e^{-zL} : L^1 \rightarrow L^\infty \quad (19)$$

and

$$e^{-zL} : L^1 \rightarrow \mathcal{C}_{\tilde{\psi}}. \quad (20)$$

3.5 Perturbation techniques

Let $\mathcal{F}_m = \{\phi \in \mathcal{C}_0^\infty(\mathbb{R}^n, \mathbb{R}) / \|D^\alpha \phi\|_\infty \leq 1, \forall \alpha \in \mathbb{N}^n, 1 \leq |\alpha| \leq m\}$ and consider the operator $L_0^{s_0} = s_0^{2m} T^{-1} L_0 T$. It is obvious to see that if $L_0 \in \mathcal{E}_{2m}(\delta_0, 0, M_0)$ then $L_0^{s_0} \in \mathcal{E}_{2m}(\delta_0, 0, M_0)$ for all $s_0 > 0$.

We apply the result of the previous step to the operator $L_{0,\xi\phi}^{s_0} = e^{\xi\phi} L_0^{s_0} e^{-\xi\phi}$ for $\xi \in \mathbb{R}$, $s_0 > 0$ and $\phi \in \mathcal{F}_m$. Note that $L_{0,\xi\phi}^{s_0} \in \mathcal{L}_{2m}(\delta_0, 0, M_0, c\xi^{2m})$ (indeed, in general if $L \in \mathcal{L}_{2m}(\delta_0, 0, M_0, M)$ and $\phi \in \mathcal{F}_m$ then $L_{\xi\phi} = e^{\xi\phi} L e^{-\xi\phi} \in \mathcal{L}_{2m}(\delta_0, 0, M_0, M + c\xi^{2m})$ for a constant c depending on n and m). In addition, by virtue of a result due to Davies ([10], Lemma 6),

Lemma 3.3. *There exists a constant $c_1 \geq 0$ such that*

$$\|e^{-tL_{0,\xi\phi}^{s_0}}\|_{L^2 \rightarrow L^2} \leq e^{c_1 \xi^{2m} t},$$

for all $t > 0$, $\xi \in \mathbb{R}$, $s_0 > 0$ and $\phi \in \mathcal{F}_m$.

Combining Lemma 3.3 to the result of the previous step, we obtain

$$\|e^{-\frac{1}{2}L_{0,\xi\phi}^{s_0}}\|_{L^2 \rightarrow L^\infty} \leq C(\xi, m) e^{\frac{c_1}{2} \xi^{2m}},$$

and

$$\|e^{-\frac{1}{2}L_{0,\xi\phi}^{s_0}}\|_{L^1 \rightarrow L^2} \leq C(\xi, m) e^{\frac{c_1}{2} \xi^{2m}}.$$

Eventually, since the property (19) holds for $L_{0,\xi\phi}^{s_0}$ (i.e. $e^{-zL_{0,\xi\phi}^{s_0}} : L^1 \rightarrow L^\infty$) then we get by the semigroup property

$$\begin{aligned} \|e^{-L_{0,\xi\phi}^{s_0}}\|_{L^1 \rightarrow L^\infty} &\leq \|e^{-\frac{1}{2}L_{0,\xi\phi}^{s_0}}\|_{L^1 \rightarrow L^2} \|e^{-\frac{1}{2}L_{0,\xi\phi}^{s_0}}\|_{L^2 \rightarrow L^\infty} \\ &\leq C^2(\xi, m) e^{c_1 \xi^{2m}}, \end{aligned}$$

which is equivalent to

$$|K_1^{L_0^{s_0}, \xi\phi}(x, y)| \leq C^2(\xi, m) e^{c_1 \xi^{2m}},$$

where $K_1^{L_0^{s_0}, \xi\phi}(x, y)$ stands for the kernel of $e^{-L_0^{s_0}, \xi\phi}$.

On the other hand, since

$$K_1^{L_0^{s_0}, \xi\phi}(x, y) = e^{\xi\phi(x)} K_1^{L_0^{s_0}}(x, y) e^{-\xi\phi(y)},$$

it follows then that

$$|K_1^{L_0^{s_0}}(x, y)| \leq C^2(\xi, m) \exp\left(c_1 \xi^{2m} + \xi (\phi(y) - \phi(x))\right),$$

for all $x, y \in \mathbb{R}^n$.

Fixing x, y , choosing $\phi \in \mathcal{F}_m$ such that $\phi(y) - \phi(x) = -\frac{1}{2}|x - y|$ and minimizing with respect to ξ yield

$$|K_1^{L_0^{s_0}}(x, y)| \leq C \exp\left(-a |x - y|^{\frac{2m}{2m-1}}\right)$$

uniformly for all $s_0 > 0$.

3.6 Scaling argument

Formally

$$e^{-tL_0^{s_0}} = e^{-ts_0^{2m}T^{-1}L_0T} = T^{-1}e^{-ts_0^{2m}L_0}T.$$

It then follows

$$K_t^{L_0^{s_0}}(x, y) = s_0^n K_{s_0^{2m}t}^{L_0}(s_0x, s_0y),$$

or

$$K_t^{L_0}(x, y) = s_0^{-n} K_{t/s_0^{2m}}^{L_0^{s_0}}\left(\frac{x}{s_0}, \frac{y}{s_0}\right)$$

where $K_t^{L_0}(x, y)$ and $K_t^{L_0^{s_0}}(x, y)$ are respectively the kernels of e^{-tL_0} and $e^{-tL_0^{s_0}}$. Therefore

$$K_t^{L_0}(x, y) = t^{-\frac{n}{2m}} K_1^{L_0^{s_0}}\left(\frac{x}{t^{1/2m}}, \frac{y}{t^{1/2m}}\right) \quad (21)$$

and afterwards

$$|K_t^{L_0}(x, y)| \leq \frac{C}{t^{\frac{n}{2m}}} \exp\left\{-a \left(\frac{|x - y|}{t^{\frac{1}{2m}}}\right)^{\frac{2m}{2m-1}}\right\}.$$

3.7 Hölder regularity

According to (20), we get in particular for $L_0^{t^{1/2m}}$ that if $h \in \mathbb{R}^n$ then

$$|K_1^{L_0^{t^{1/2m}}}(x+h, y) - K_1^{L_0^{t^{1/2m}}}(x, y)| \leq C \tilde{\psi}(|h|).$$

It then follows that

$$|K_t^{L_0}(x+h, y) - K_t^{L_0}(x, y)| \leq \frac{C}{t^{\frac{n}{2m}}} \tilde{\psi}\left(\frac{|h|}{t^{1/2m}}\right)$$

thanks to (21). The regularity with respect to y is similar and Theorem 2.1 is proved.

4 Remarks

1. We can relax the ellipticity condition by replacing (4) by the one we frequently meet in practice

$$\Re \int_{\mathbb{R}^n} \sum_{|\alpha|=|\beta|=m} a_{\alpha\beta} D^\beta u \overline{D^\alpha v} dx \geq \delta_0 \int_{\mathbb{R}^n} |\nabla^m u|^2 dx - \lambda_0 \int_{\mathbb{R}^n} |u|^2 dx.$$

In this case, the inequality of Lemma 3.3 becomes

$$\|e^{-tL_{0,\xi\phi}^{s_0}}\|_{L^2 \rightarrow L^2} \leq e^{(\lambda_0 + c_1 \xi^{2m})t}$$

and an adaptation of the argument used above gives us

$$|K_t^{L_0}(x, y)| \leq \frac{C}{t^{\frac{n}{2m}}} \exp\left\{-\theta\left(\frac{|x-y|}{t^{\frac{1}{2m}}}\right)^{\frac{2m}{2m-1}} + c_2 t\right\}.$$

2. It seems possible to derive from (D_ψ) upper gaussian bounds for higher derivatives $D^\gamma K_t(x, y)$ ($\gamma \in \mathbb{N}^n$, $|\gamma| < m$) of the heat kernel. For the techniques see [7] (resp. [29]) for derivatives taken with respect to x or y (resp. x and y simultaneously).

3. According to the remark used in section 3.4 to derive $e^{-zL} : L^2 \rightarrow L^\infty$, we can obtain from (D_ψ) the same estimates and regularity on higher derivatives of the kernel as in [7] provided ψ is almost increasing.

References

- [1] S. Agmon, *Lectures on Elliptic Boundary Problems*, Van Nostrand, New York, 1965.
- [2] S. Agmon, A. Douglis, and L. Nirenberg, *Estimates near the boundary for solutions of elliptic partial differential equations satisfying general boundary conditions*, I, II, Comm. Pure Appl. Math. **12** (1959), 623–727.
- [3] D. Aronson, *Bounds for fundamental solutions of a parabolic equation*, Bull. Amer. Math. Soc. **73** (1967), 890–896.
- [4] P. Auscher, *Regularity theorems and heat kernels for elliptic operators*, J. London. Math. Soc., **54** (1996), 284–296.
- [5] P. Auscher, T. Coulhon and Ph. Tchamitchian, *Absence de principe du maximum pour certaines équations paraboliques complexes*, Coll. Math., **71** (1996), 87–95.
- [6] P. Auscher, A. McIntosh and Ph. Tchamitchian, *Noyau de la chaleur d'opérateurs elliptiques complexes*, Math. Research Letters, **1** (1994), 37–45.
- [7] P. Auscher and M. Qafsaoui, *Equivalence between regularity theorems and heat kernel estimates for higher order elliptic operators and systems under divergence form*, J. Funct. Anal. **177** (2000), 310–364.
- [8] P. Auscher and Ph. Tchamitchian, *Square root problem for divergence operators and related topics*, Astérisque **249** (1998).
- [9] S. Campanato, *Sistemi ellittici in forma divergenza. Regolarità all'interno*, Ann. Scuola Normale Sup. Pisa (1980).
- [10] E.B. Davies, *Uniformly elliptic operators with measurable coefficients*, J. Funct. Anal. **132** (1995), 141–169.
- [11] E. De Giorgi, *Sulla differenziabilità e l'analiticità delle estremali degli integrali multipli regolari*, Mem. Accad. Sci. Torino Cl. Sci. Fis. Mat. Natur. (3) **3** (1957), 25–43.
- [12] E. Fabes and D. Stroock, *A new proof of Moser's parabolic inequality via the old ideas of Nash*, Arch. Rational Mech. Anal. **96** (1986), 327–338.

- [13] A. Friedman, *Partial Differential Equations*, Krieger, Melbourne, FL, 1983.
- [14] M. Giaquinta, *Multiple integrals in the calculus of variations and nonlinear elliptic systems*, Lectures in Mathematics, Annals of Math. Studies, Vol. **105**, Princeton Univ. Press, Princeton, 1983.
- [15] M. Giaquinta, *Introduction to regularity theory for nonlinear elliptic systems*, Lectures in Mathematics, Birkhäuser, ETH Zürich, 1993.
- [16] M. Giaquinta and G. Modica, *Regularity results for some classes of higher order nonlinear elliptic systems*, J. Reine Angew. Math. **311/312** (1979), 145–169.
- [17] D. Gilbarg and N.S. Trudinger, *Elliptic Partial Differential Equations of Second Order*, Springer-Verlag, Berlin/New York, 1983.
- [18] A. Grigor'yan, *Upper bounds of derivatives of the heat kernel on a arbitrary complete manifold*, J. Funct. Anal. **127** (1995), 363–389.
- [19] Q. Huang, *Estimates on generalized Morrey spaces $L_{\phi}^{2,\lambda}$ and BMO_{ψ} for linear elliptic systems*, Indiana Univ. Math. J. **45** (1996), 397–439.
- [20] T. Kato, *Perturbation Theory for Linear Operators*, Springer-Verlag, New York, 1966.
- [21] V.G. Maz'ya, S.A. Nazarov and B.A. Plamenevskii, *Absence of De Giorgi-type theorems for strongly elliptic equations with complex coefficients*, J. Math. Sov. **28** (1985), 726–739.
- [22] A. McIntosh and A. Nahmod, *Heat kernel estimates and functional calculus of $-b\Delta$* , Math. Scand. **87**, N. 2 (2000), 287–319.
- [23] N.G. Meyers, *An L^p -estimate for the gradient of solutions of second order elliptic divergence equations*, Ann. Scuola Norm. Sup. Pisa **17** (1963), 189–206.
- [24] C.B. Morrey, *Multiple integrals in the calculus of variations*, Springer Verlag, 1966.
- [25] C.B. Morrey, *Partial regularity results for nonlinear systems*, J. Math. and Mech. (1968), 649–670.
- [26] J. Moser, *On Harnack's theorem for elliptic partial differential equations*, Comm. Pure Appl. Math. **14** (1961), 577–691.

- [27] J. Moser, *A Harnack inequality for parabolic differential equations*, Comm. Pure and Appl. Math. **17** (1964), 101–134; correction to *A Harnack inequality for parabolic differential equations*, Comm. Pure and Appl. Math. **17** (1964), 232–236.
- [28] J. Nash, *Continuity of solutions of parabolic and elliptic equations*, Amer. J. Math. **80** (1958), 931–954.
- [29] M. Qafsaoui, *On a class of higher order operators with complex coefficients, elliptic in the sense of Gårding inequality*, to appear in Differential and Integral Equations.
- [30] G. Stampacchia, *Equations elliptiques du second ordre à coefficients discontinus*, Séminaire sur les équations à dérivées partielles, Collège de France, 1963.

Exponential operators and solution of pseudo-classical evolution problems

Caterina Cassisa, Paolo E. Ricci,

Università di Roma “La Sapienza”, Dipartimento di Matematica, P.le A. Moro, 2 00185
Roma, Italia - e-mail: cassisa@uniroma1.it, riccip@uniroma1.it

Ilia Tavkhelidze

“I.N. Vekua” Institute of Applied Mathematics, Tbilisi State University
2, University Street – 380043 – Tbilisi (Georgia) - e-mail: iliko@viam.hepi.edu.ge

Abstract

We use the multi-dimensional polynomials considered by Hermite, and subsequently studied by P. Appell and J. Kampé de Fériet, in order to obtain explicit solutions of pseudo-classical PDE problems in the half-plane $y > 0$. We consider systems of PDE, including some problems with degeneration on the x -axis.

2000 Mathematics Subject Classification. 33C45, 44A45, 35G15.

Key words and phrases. Hermite-Kampé de Fériet polynomials, Operational calculus, Pseudo-hyperbolic and Pseudo-circular functions, Boundary value problems.

1 Introduction

In preceding articles [3], [4], [5], by using an operatorial approach, we showed that the solution of many classical or pseudo-classical boundary value problems in the half plane for PDE, (with constant coefficients and analytic boundary, or initial, data) can be expressed in terms of the higher order Hermite-Kampé de Fériet (shortly H-KdF) polynomials. The relevant results were extended to the multi-dimensional case in [13], [14].

In this article we extend the results of our article [5] to the pseudo-classical case too.

We start recalling properties of the H-KdF polynomials and extending technical tools considered in [3], [4] to the case under examination.

2 Hermite-Kampé de Fériet polynomials

2.1 Definitions

We recall the definitions of the H-KdF polynomials [1], [10], [16], in the two-dimensional case.

Put $D := \frac{d}{dx}$, and consider the shift operator

$$e^{yD} f(x) = f(x + y) = \sum_{n=0}^{\infty} \frac{y^n}{n!} f^{(n)}(x), \quad (2.1)$$

(see e.g. [17], p. 171), the second equation being meaningful for analytic functions. Note that,

- if $f(x) = x^m$, then $e^{yD} x^m = (x + y)^m$;
- if $f(x) = \sum_{m=0}^{\infty} a_m x^m$, then $e^{yD} f(x) = \sum_{m=0}^{\infty} a_m (x + y)^m$.

Definition 2.1 *The Hermite polynomials in two variables $H_m^{(1)}(x, y)$ are then defined by*

$$H_m^{(1)}(x, y) := (x + y)^m. \quad (2.2)$$

Consequently,

- if $f(x) = \sum_{m=0}^{\infty} a_m x^m$, then

$$e^{yD} f(x) = \sum_{m=0}^{\infty} a_m H_m^{(1)}(x, y). \quad (2.3)$$

Consider now the exponential containing the second derivative, defined for an analytic function f , as follows:

$$e^{yD^2} f(x) = \sum_{n=0}^{\infty} \frac{y^n}{n!} f^{(2n)}(x). \quad (2.4)$$

Note that

- if $f(x) = x^m$, then for $n = 0, 1, \dots, \left\lfloor \frac{m}{2} \right\rfloor$ we can write:
 $D^{2n} x^m = m(m-1) \cdots (m-2n+1) x^{m-2n} = \frac{m!}{(m-2n)!} x^{m-2n}$ and therefore

$$e^{yD^2} x^m = \sum_{n=0}^{\left\lfloor \frac{m}{2} \right\rfloor} \frac{y^n}{n!} \frac{m!}{(m-2n)!} x^{m-2n} \quad (2.5)$$

- if $f(x) = \sum_{m=0}^{\infty} a_m x^m$, then $e^{yD^2} f(x) = \sum_{m=0}^{\infty} a_m H_m^{(2)}(x, y)$.

Definition 2.2 The H-KdF polynomials in two variables $H_m^{(2)}(x, y)$ are then defined by

$$H_m^{(2)}(x, y) := m! \sum_{n=0}^{\left[\frac{m}{2}\right]} \frac{y^n x^{m-2n}}{n!(m-2n)!} \quad (2.6)$$

Considering, in general, the exponential raised to the j -th derivative we have:

$$e^{yD^j} f(x) = \sum_{n=0}^{\infty} \frac{y^n}{n!} f^{(jn)}(x) \quad (2.7)$$

and therefore:

- if $f(x) = x^m$, then for $n = 0, 1, \dots, \left[\frac{m}{j}\right]$ it follows:
 $D^{jn} x^m = m(m-1) \cdots (m-jn+1) x^{m-jn} = \frac{m!}{(m-jn)!} x^{m-jn}$ so that

$$e^{yD^j} x^m = \sum_{n=0}^{\left[\frac{m}{j}\right]} \frac{y^n}{n!} \frac{m!}{(m-jn)!} x^{m-jn} \quad (2.8)$$

- if $f(x) = \sum_{m=0}^{\infty} a_m x^m$, then $e^{yD^j} f(x) = \sum_{m=0}^{\infty} a_m H_m^{(j)}(x, y)$,

Definition 2.3 The H-KdF polynomials in two variables $H_m^{(j)}(x, y)$ are then defined by

$$H_m^{(j)}(x, y) := m! \sum_{n=0}^{\left[\frac{m}{j}\right]} \frac{y^n x^{m-jn}}{n!(m-jn)!} \quad (2.9)$$

2.2 Properties

In a number of articles by G. Dattoli et al., (see e.g. [6], [7], [8]), by using the so called *monomiality principle*, the following properties for the two-variable H-KdF polynomials $H_m^{(j)}(x, y)$, $j \geq 2$ have been recovered (the case when $j = 1$ reducing to results about simple powers).

- Operational definition

$$H_m^{(j)}(x, y) = e^{y \frac{\partial^j}{\partial x^j}} x^m = \left(x + jy \frac{\partial^{j-1}}{\partial x^{j-1}} \right)^m (1). \quad (2.10)$$

- Generating function

$$\sum_{n=0}^{\infty} H_n^{(j)}(x, y) \frac{t^n}{n!} = e^{xt+yt^j}. \quad (2.11)$$

In the case when $j = 2$, (see [18]), the H-KdF polynomials $H_m^{(2)}(x, y)$ admit the following

- Integral representation

$$H_m^{(2)}(x, y) = \frac{1}{2\sqrt{\pi y}} \int_{-\infty}^{+\infty} \xi^m e^{-\frac{(x-\xi)^2}{4y}} d\xi, \quad (2.12)$$

which is a particular case of the so called Gauss-Weierstrass (or Poisson) transform.

For the case when $j > 2$ see [11], [12].

3 A decomposition theorem

In [15], in a general form, the following theorem was shown:

Theorem 3.1 *Fix the integer $r \geq 2$, and denote by*

$$\omega_{h,r} := \exp\left(\frac{2\pi i h}{r}\right), \quad (h = 0, 1, \dots, r-1), \quad (3.1)$$

the complex roots of unity.

Consider an analytic function of the complex variable z , defined in a circular neighborhood of the origin, by means of the series expansion

$$\varphi(z) = \sum_{n=0}^{\infty} a_n z^n, \quad (3.2)$$

and the series

$$\Pi_{[h,r]} \varphi(z) = \varphi_h(z; r) := \sum_{n=0}^{\infty} a_{rn+h} z^{rn+h}, \quad (3.3)$$

called the components of φ with respect to the cyclic group of order r , then the representation formula

$$\varphi_h(z; r) = \frac{1}{r} \sum_{j=0}^{r-1} \frac{\varphi(z \omega_{j,r})}{\omega_{jh,r}} \quad (3.4)$$

holds true.

For $r = 2$ and $h = 0$ or $h = 1$ the functions (3.3) reduce to the even or odd components, respectively, of the function $\varphi(z)$. In general, the $\varphi_h(z; r)$ functions are characterized by the symmetry property

$$\varphi_h(z \omega_{1,r}; r) = \omega_{h,r} \varphi_h(z; r) \quad (h = 0, 1, \dots, r-1) \quad (3.5)$$

with respect to the roots of unity (see also [2]).

In particular, assuming $\varphi(z) := e^z$, the well known decomposition of the exponential function in terms of the so called *pseudo-hyperbolic functions* appear (see [9]). For shortness, omitting in the sequel the label r , i.e. assuming $\omega_h := \omega_{h,r}$, $\varphi_h(z) := \varphi_h(z; r)$, and so on, we can write:

$$f_h(z) := \Pi_{[h,r]} e^z = \sum_{n=0}^{\infty} \frac{z^{rn+h}}{(rn+h)!}, \quad (3.6)$$

so that the exponential is decomposed as the sum:

$$e^z = \sum_{h=0}^{r-1} f_h(z), \quad (3.7)$$

and the following properties hold true:

$$f_0(0) = 1; \quad f_h(0) = 0 \quad (\text{if } h \neq 0), \quad (3.8)$$

$$f_h(\omega_1 z) = \omega_h f_h(z), \quad (\text{symmetry property}), \quad (3.9)$$

$$D_z f_h(z) = f_{h-1}(z), \quad (\text{differentiation rule}), \quad (3.10)$$

where the indices are assumed to be congruent $(\text{mod } r)$, so that, $\forall h$, the pseudo-hyperbolic functions are solutions of the differential equation:

$$w^{(r)}(z) - w(z) = 0. \quad (3.11)$$

The *pseudo-circular functions* are obtained by introducing any complex r -th root σ_0 of the number -1 , and putting:

$$g_h(z) := \sigma_0^{-h} f_h(\sigma_0 z). \quad (3.12)$$

The pseudo-circular functions are given by the series expansions:

$$g_h(z) = \sum_{n=0}^{\infty} (-1)^n \frac{z^{rn+h}}{(rn+h)!}, \quad (3.13)$$

and satisfy the following properties:

$$g_0(0) = 1; \quad g_h(0) = 0 \quad (\text{if } h \neq 0), \quad (3.14)$$

$$g_h(\omega_1 z) = \omega_h g_h(z), \quad (\text{symmetry property}), \quad (3.15)$$

$$D_z g_0(z) = -g_{r-1}(z), \quad D_z g_h(z) = g_{h-1}(z), \quad (\text{if } h \neq 0), \quad (3.16)$$

(differentiation rule)

where the indices are assumed to be congruent $(\text{mod } r)$, so that $\forall h$, the pseudo-circular functions are solutions of the differential equation:

$$w^{(r)}(z) + w(z) = 0. \quad (3.17)$$

For further properties, generalizing the ordinary trigonometrical rules, see [15].

4 Pseudo-hyperbolic or pseudo-circular functions of the derivative operator

The properties of the pseudo-hyperbolic or pseudo-circular functions of the derivative operator (see [15]), can be easily extended to the operational case. We give in the following a list of the relevant results, which can be deduced in the same way as in the functional case, considering the commutative property of the powers of D .

The pseudo-hyperbolic functions of the derivative operator are defined by the series expansions:

$$f_h(D^j) := \Pi_{[h,r]} e^{D^j} = \sum_{n=0}^{\infty} \frac{D^{j(rn+h)}}{(rn+h)!}, \quad (h = 0, 1, \dots, r-1) \quad (4.1)$$

so that the exponential is decomposed as the sum:

$$e^{D^j} = \sum_{h=0}^{r-1} f_h(D^j), \quad (4.2)$$

and the following properties hold true:

$$f_h(\omega_1 D^j) = \omega_h f_h(D^j), \quad (\text{symmetry property}), \quad (4.3)$$

Considering two independent variables x and y and denoting by D_x and D_y the relevant differentiations, it follows:

$$D_y f_h(y D_x^j) = D_x^j f_{h-1}(y D_x^j), \quad (\text{differentiation rule}), \quad (4.4)$$

where the indices are assumed to be congruent $(\text{mod } r)$, so that, $\forall h$, the pseudo-hyperbolic functions are solutions of the abstract differential equation:

$$D_y^r f_h(y D_x^j) - D_x^{rj} f_h(y D_x^j) = 0. \quad (4.5)$$

The pseudo-circular functions of the derivative operator are given by the series expansions:

$$g_h(D^j) = \sum_{n=0}^{\infty} (-1)^n \frac{D^{j(rn+h)}}{(rn+h)!}, \quad (4.6)$$

and satisfy the following properties:

$$g_h(\omega_1 D^j) = \omega_h g_h(D^j), \quad (\text{symmetry property}), \quad (4.7)$$

$$D_y g_0(y D_x^j) = -D_x^j g_{r-1}(y D_x^j), \quad D_y g_h(y D_x^j) = D_x^j g_{h-1}(y D_x^j), \quad (\text{if } h \neq 0), \quad (4.8)$$

(differentiation rule),

where the indices are assumed to be congruent (mod r), so that $\forall h$, the pseudo-circular functions are solutions of the abstract differential equation:

$$D_y^r g_h(y D_x^j) + D_x^j g_h(y D_x^j) = 0. \quad (4.9)$$

4.1 Connections with the H-KdF polynomials

Fixing the integer r , we consider now the action of the pseudo-hyperbolic and pseudo-circular functions on analytic functions, showing relations with the H-KdF. polynomials.

Theorem 4.1 *For any real number ℓ , for any $h=0, 1, \dots, r-1$, and for any positive integer j , denoting by x and y ($y > 0$) independent variables, and by $D := D_x$, the action of the pseudo-hyperbolic function $f_h(y^\ell D^j)$ on the power x^m ($m \in \mathbf{N}$), is given by*

$$f_h(y^\ell D^j) x^m = m! \sum_{n=0}^{\left[\frac{m-jh}{jr}\right]} \frac{y^{\ell(rn+h)} x^{m-j(rn+h)}}{(rn+h)!(m-j(rn+h))!} = \Pi_{[h,r]_{y^\ell}} [H_m^{(j)}(x, y^\ell)], \quad (4.10)$$

where $\Pi_{[h,r]_{y^\ell}} [K(x, y^\ell)]$, ($h = 0, 1, \dots, r-1$), denote the components of the function K with respect to the cyclic group of order r , and relevant to the variable y^ℓ , (assuming x as a parameter).

Then, if $q(x) = \sum_{m=0}^{\infty} a_m x^m$, we can write

$$\begin{aligned} f_h(y^\ell D^j) q(x) &= \sum_{m=0}^{\infty} a_m \Pi_{[h,r]_{y^\ell}} [H_m^{(j)}(x, y^\ell)] = \\ &= \sum_{m=0}^{\infty} m! a_m \sum_{n=0}^{\left[\frac{m-jh}{jr}\right]} \frac{y^{\ell(rn+h)} x^{m-j(rn+h)}}{(rn+h)!(m-j(rn+h))!}. \end{aligned} \quad (4.11)$$

For the pseudo-circular functions, using the above notations, we have:

Theorem 4.2 *For any real number ℓ , for any $h=0,1,\dots,r-1$, and for any positive integer j , denoting by x and y ($y > 0$) independent variables, and by $D := D_x$, the action of the pseudo-circular function $g_h(y^\ell D^j)$ on the power x^m ($m \in \mathbf{N}$), is given by*

$$g_h(y^\ell D^j)x^m = m! \sum_{n=0}^{\left[\frac{m-jh}{jr}\right]} (-1)^n \frac{y^{\ell(rn+h)} x^{m-j(rn+h)}}{(rn+h)!(m-j(rn+h))!} = \frac{1}{\sigma_0^h} \Pi_{[h,r]_{y^\ell}} \left[H_m^{(j)}(x, \sigma_0 y^\ell) \right], \quad (4.12)$$

and furthermore, if $q(x) = \sum_{m=0}^{\infty} a_m x^m$, then

$$\begin{aligned} g_h(y^\ell D^j)q(x) &= \frac{1}{\sigma_0^h} \sum_{m=0}^{\infty} a_m \Pi_{[h,r]_{y^\ell}} \left[H_m^{(j)}(x, \sigma_0 y^\ell) \right] = \\ &= \sum_{m=0}^{\infty} m! a_m \sum_{n=0}^{\left[\frac{m-jh}{jr}\right]} (-1)^h \frac{y^{\ell(rn+h)} x^{m-j(rn+h)}}{(rn+h)!(m-j(rn+h))!}, \end{aligned} \quad (4.13)$$

so that, if the considered coefficients a_m and variables x, y are real, the resulting expansions are also real.

5 Convergence results

In this section we recall an uniform estimate, with respect to j , for the convergence of series involving the H-KdF polynomials $H_n^{(j)}(x, y^\ell)$, for every real $\ell \geq 1$:

$$\sum_{n=0}^{\infty} a_n H_n^{(j)}(x, y^\ell). \quad (5.1)$$

Theorem 5.1 *For every real $\ell \geq 1$, $j \geq 2$, $-\infty < x < +\infty$, $-\infty < y < +\infty$, $n = 0, 1, 2, \dots$, the following estimate holds true:*

$$|H_n^{(j)}(x, y^\ell)| \leq n! \exp \left\{ |x| + |y^\ell| \right\}, \quad (5.2)$$

The proof is derived by the same method used in the book of Widder [18], p. 166, for the case $\ell = 1$ $j = 2$ (see [5]).

In the present article, we limit ourselves to consider, as boundary (or initial) data, analytic functions $f(x) = \sum_{n=0}^{\infty} a_n x^n$ for which the coefficients a_n tend to zero sufficiently fast, in order to guarantee the convergence of the series expansion (5.1). To this aim, we only use the following theorem (see [5])

Theorem 5.2 *Suppose there exists a number $\alpha > 1$, such that the coefficients a_n satisfy the following estimate:*

$$|a_n| = O\left(\frac{1}{n^\alpha n!}\right), \quad (5.3)$$

then, for every real $\ell \geq 1$ and integer j , the series expansion (5.1) is absolutely and uniformly convergent in every bounded region of the (x, y) plane.

Remark 5.1 The condition (5.3) include analytic functions with polynomial growth at infinity, but not the exponential function e^x , whereas, in the case $\ell = 1$ $j = 2$, the book of Widder [18] include all functions belonging to the so called Huygens class H^0 however, by the point of view of the Applied Analysis, the considered conditions are sufficient to cover all realistic situations.

When the function $f(x) = \sum_{n=0}^{\infty} a_n x^n$ is decomposed with respect to the cyclic group of order r , the same estimate of Theorem 5.2 is sufficient to guarantee the convergence of the relevant expansions, according to the property:

$$\left| \Pi_{[h,r]} \left(\sum_{n=0}^{\infty} |a_n| x^n \right) \right| \leq \left| \sum_{n=0}^{\infty} |a_n| x^n \right|, \quad (h = 0, 1, \dots, r-1).$$

6 Systems solved in terms of pseudo-hyperbolic operators

For introducing our subject in a more friendly way, we start considering the particular case when $r = 3$.

Consider the system:

$$\left\{ \begin{array}{l} \frac{\partial S_h}{\partial y} = \frac{\partial S_{h-1}}{\partial x} \quad (h = 0, 1, 2) \quad \text{in the half plane } y > 0, \\ S_0(x, 0) = q_0(x) \\ \frac{\partial S_1}{\partial y}(x, 0) = q'_0(x) \\ \frac{\partial^2 S_2}{\partial y^2}(x, 0) = q''_0(x), \end{array} \right. \quad (6.1)$$

where $q_0(x)$ is a given analytic function. We assume that the considered indices are congruent *mod.* 3, so that $S_3 \equiv S_0$.

By using the same technique developed in our article [5], we find

Theorem 6.1 *The operational solution of system (6.1) is given by the following functions:*

$$\begin{cases} S_0(x, y) = f_0(yD) q_0(x) , \\ S_1(x, y) = f_1(yD) q_0(x) , \\ S_2(x, y) = f_2(yD) q_0(x) . \end{cases} \quad (6.2)$$

Proof. A straightforward computation gives:

$$\begin{aligned} \frac{\partial S_0}{\partial y} &= f_2(yD) Dq_0(x) = D_x f_2(yD) q_0(x) = \frac{\partial S_2}{\partial x} , \\ \frac{\partial S_1}{\partial y} &= f_0(yD) Dq_0(x) = D_x f_0(yD) q_0(x) = \frac{\partial S_0}{\partial x} , \\ \frac{\partial S_2}{\partial y} &= f_1(yD) Dq_0(x) = D_x f_1(yD) q_0(x) = \frac{\partial S_1}{\partial x} , \end{aligned}$$

and the boundary conditions are trivially satisfied.

Remark 6.1 Note that the given boundary conditions (6.1) can be essentially derived from the first one: $S_0(x, 0) = q_0(x)$, by extending the considered equations on the boundary $y = 0$.

Remark 6.2 Note that the system (6.1) has connections with the similar one considered in [4] (eqs. (7.1)-(7.2)), but there the boundary conditions were given by three independent functions, whereas in (6.1) only the function q_0 appears.

A more general problem which can be solved is as follows:

$$\begin{cases} \frac{\partial S_h}{\partial y} = \frac{\partial S_{h-1}}{\partial x} & (h = 0, 1, 2) & \text{in the half plane } y > 0, \\ S_0(x, 0) = q_0(x) \\ S_1(x, 0) = q_1(x) \\ S_2(x, 0) = q_2(x) , \end{cases} \quad (6.3)$$

where $q_h(x)$ ($h = 0, 1, 2$) are given analytic functions, assuming again that the considered indices are congruent *mod.* 3.

In this case, we find

Theorem 6.2 *The operational solution of system (6.3) is given by the following functions:*

$$\begin{cases} S_0(x, y) = f_0(yD) q_0(x) + f_1(yD) q_2(x) + f_2(yD) q_1(x) , \\ S_1(x, y) = f_0(yD) q_1(x) + f_1(yD) q_0(x) + f_2(yD) q_2(x) , \\ S_2(x, y) = f_0(yD) q_2(x) + f_1(yD) q_1(x) + f_2(yD) q_0(x) . \end{cases} \quad (6.4)$$

6.1 The case when r is an arbitrary integer

For any fixed integral r , consider the system

$$\begin{cases} \frac{\partial S_k}{\partial y} = \frac{\partial S_{k-1}}{\partial x}, & (k = 0, 1, \dots, r-1) \text{ in the half plane } y > 0, \\ S_k(x, 0) = q_k(x), \end{cases} \quad (6.5)$$

where $q_k(x)$, $(k = 0, 1, \dots, r-1)$, are given analytic functions, assuming that the considered indices are congruent *mod. r*, so that $S_r \equiv S_0$.

By using the same technique, we find the general result:

Theorem 6.3 *The operational solution of the system (6.5) is given by the following r -tuples of functions:*

$$S_k(x, y) = \sum_{h=0}^{r-1} f_h(yD) q_{r-h+k}(x), \quad (k = 0, \dots, r-1). \quad (6.6)$$

7 Higher order pseudo-hyperbolic systems

For any real number $\ell \geq 1$ and integer $j \geq 1$, we consider now the system:

$$\begin{cases} \frac{\partial S_k}{\partial y} = \ell y^{\ell-1} \frac{\partial^j S_{k-1}}{\partial x^j}, & (k = 0, 1, \dots, r-1) \text{ in the half plane } y > 0, \\ S_k(x, 0) = q_k(x), \end{cases} \quad (7.1)$$

where the $q_k(x)$ are given analytic functions. We assume again that $k = 0, 1, \dots, r-1$ and that the considered indices are congruent *mod. r*, so that $S_r \equiv S_0$.

Note that equations in system (7.1), assuming $\ell > 1$, degenerate on the boundary when $y \rightarrow 0$.

Theorem 7.1 *The operational solution of system (7.1) is given by the following r -tuples of functions:*

$$S_k(x, y) = \sum_{h=0}^{r-1} f_h(y^\ell D^j) q_{r+k-h}(x), \quad (k = 0, \dots, r-1). \quad (7.2)$$

Remark 7.1 Theorem 7.1 generalizes some results already obtained in [4]. More precisely, if $\ell = 1, j = 1, r = 3$ we find eqs. (7.1)-(7.2), if $\ell = 1, j = 1$ we find eqs. (8.1)-(8.2) in ref. [4]. Furthermore, if $\ell = 1, j = s$, we recover (9.1)-(9.2) of the same reference.

8 Pseudo-circular systems

For any fixed integral r , consider, in the half-plane $y > 0$, the system

$$\begin{cases} \frac{\partial S_0}{\partial y} + \ell y^{\ell-1} \frac{\partial^j S_{r-1}}{\partial x^j} = 0 \\ \frac{\partial S_k}{\partial y} - \ell y^{\ell-1} \frac{\partial^j S_{k-1}}{\partial x^j} = 0 & (k = 1, 2, \dots, r-1), \\ S_k(x, 0) = q_k(x) & (k = 0, 1, \dots, r-1), \end{cases} \quad (8.1)$$

where the $q_k(x)$ are given analytic functions. We assume again that $k = 0, 1, \dots, r-1$ and that the considered indices are congruent *mod*. r .

By using same methods as in the preceding section, but relevant to pseudo-circular operators, we find the result:

Theorem 8.1 *The operational solution of system (8.1) is given by the following r -tuples of functions:*

$$S_k(x, y) = \sum_{h=0}^k g_h \left(y^\ell D^j \right) q_{r+k-h}(x) - \sum_{h=k+1}^{r-1} g_h \left(y^\ell D^j \right) q_{r+k-h}(x), \quad (8.2)$$

$(k = 0, \dots, r-1)$

Remark 8.1 Theorem 8.1 generalizes some results already obtained in [4]. More precisely, assuming $\ell = 1, j = 1, r = 3$ we find eqs. (10.1)-(10.2) in ref. [4]. Furthermore, if $\ell = 1, j = s$, we recover (11.1)-(11.2) of the same reference.

Acknowledgements

This article was concluded during the visit of Dr. Ilia Tavkhelidze, partially supported by the visiting professors program of the Ateneo Roma “La Sapienza”.

References

- [1] P. Appell, J. Kampé de Fériet, *Fonctions Hypergéométriques et Hypersphériques. Polynômes d’Hermite*, Gauthier-Villars, Paris, 1926.
- [2] Y. Ben Cheikh, Decomposition of some complex functions with respect to the cyclic group of order n , *Appl. Math. Inform.*, **4**, 30–53, (1999).
- [3] C. Cassisa, P.E. Ricci, I. Tavkhelidze, Operational identities for circular and hyperbolic functions and their generalizations, *Georgian Math. J.*, **10**, 45–56, (2003).

- [4] C. Cassisa, P.E. Ricci, I. Tavkhelidze, An operatorial approach to solutions of BVP in the half-plane, *J. Concr. Appl. Math.*, **1**, 37–62, (2003).
- [5] C. Cassisa, P.E. Ricci, I. Tavkhelidze, Exponential operators for solving evolution problems with degeneration, *J. Appl. Funct. Anal.*, (to appear).
- [6] G. Dattoli, A. Torre, P.L. Ottaviani, L. Vázquez, Evolution operator equations: integration with algebraic and finite difference methods. Application to physical problems in classical and quantum mechanics, *La Rivista del Nuovo Cimento*, **2**, (1997).
- [7] G. Dattoli, A. Torre, C. Carpanese, Operational rules and arbitrary order Hermite generating functions, *J. Math. Anal. Appl.*, **227**, 98–111, (1998).
- [8] G. Dattoli, A. Torre, Operational methods and two variable Laguerre polynomials, *Atti Acc. Sc. Torino*, **132**, 1–7, (1998).
- [9] A. Erdélyi, W. Magnus, F. Oberhettinger, F.G. Tricomi, *Higher Transcendental Functions*, McGraw-Hill, New York, 1953.
- [10] H.W. Gould, A.T. Hopper, Operational formulas connected with two generalizations of Hermite Polynomials, *Duke Math. J.*, **29**, 51–62, (1962).
- [11] D.T. Haimo and C. Markett, A representation theory for solutions of a higher order heat equation, I, *J. Math. Anal. Appl.* **168**, 89–107, (1992).
- [12] D.T. Haimo and C. Markett, A representation theory for solutions of a higher order heat equation, II, *J. Math. Anal. Appl.* **168**, 289–305, (1992).
- [13] G. Maroscia, P.E. Ricci, Hermite-Kampé de Fériet polynomials and solutions of Boundary Value Problems in the half-space, *J. Concr. Appl. Math.*, (to appear).
- [14] G. Maroscia, P.E. Ricci, Explicit solutions of multidimensional pseudo-classical BVP in the half-space, *Math. Comput. Modelling*, **40** (2004), 667–698.
- [15] P.E. Ricci, Le Funzioni Pseudo-Iperboliche e Pseudo-trigonometriche, *Pubbl. Ist. Mat. Appl. Fac. Ingegneria Univ. Stud. Roma*, **12**, 37–49, (1978).
- [16] H.M. Srivastava, H.L. Manocha, *A treatise on generating functions*, Wiley, New York, 1984.
- [17] D.V. Widder, *An Introduction to Transform Theory*, Academic Press, New York, 1971.
- [18] D.V. Widder, *The Heat Equation*, Academic Press, New York, 1975.

Newton sum rules of the zeros of some associated orthogonal q-polynomials

C.G.Kokologiannaki¹, P.D.Siafarikas² and I.D.Stabolas^{3,4}

Department of Mathematics, University of Patras, 26500 Patras, Greece
e-mails:¹chrykok@math.upatras.gr, ²panos@math.upatras.gr, ³stabol@math.upatras.gr

Abstract

By using a functional analytic method we evaluate the Newton sum rules of the zeros of some classes of associated orthogonal q-polynomials, in terms of the coefficients of the three-term recurrence relation, which they satisfy. As particular cases we obtain the Newton sum rules of some associated orthogonal polynomials found recently.

Keywords: Newton sum rules, associated orthogonal q-polynomials.

MSC2000: 33C47, 33D45.

1 Introduction

In the mathematical physics literature the moments of the distribution of zeros of polynomials have received a great deal of attention [8]. In order to know the distribution of zeros of polynomials, it is desirable to learn as much as possible about the density of their zeros. This can be done by studying the moments around the origin of the density of zeros of polynomials from which valuable information can be derived [5].

The evaluation of the moments of zero distribution or equivalently the Newton sum rules of the zeros of classical orthogonal polynomials have been studied in [2, 7, 8, 9, 11, 12, 18, 21, 23, 28], of semiclassical orthogonal polynomials in [22, 32], of associated and co-recursive associated orthogonal polynomials in [18, 24, 29], of the scaled co-recursive associated orthogonal polynomials in [16, 25], of quasiorthogonal polynomials of the classical class in [34] and of relativistic polynomials in [26].

Lately there has been an increasing interest in the orthogonal q-polynomials due to their applications in various areas of mathematics and physics (see [1] and the references there in). The classical orthogonal q-polynomials were introduced by Hahn

⁴Partially supported by the Greek State Scholarships Foundation (I.K.Y.)

in 1949 [10] when in analogy to the classical orthogonal polynomials he was interested in finding all the orthogonal polynomial sequences, the q -derivatives of which, defined by $D_q f(x) = \frac{f(x) - f(qx)}{(1-q)x}$, are also orthogonal. Hahn obtained the first results as q -hypergeometric series, after solving a Sturm-Liouville type equation in q -differences. In 1985, Andrews and Askey [4] continued the work of Hahn from the hypergeometric point of view. They showed that all the classical orthogonal polynomials can be obtained as limit cases of the q -Racah or the Askey-Wilson polynomials which are defined in terms of basic hypergeometric series ${}_4\phi_3$. For more information regarding the q -polynomials see [3, 19]. Recently various results were derived in [20, 27, 30] concerning the monotonicity and convexity properties, and various inequalities regarding the zeros of some classes of associated orthogonal q -polynomials. We recall that the associated orthogonal q -polynomials are obtained from the orthogonal q -polynomials after replacing n by $n+c$ in the coefficients of the recurrence relation that they satisfy. In [1] the authors study the discrete density of zeros (i.e the number of zeros per unit of zero interval) of the q -polynomials and its asymptotic limit.

As far as the classical orthogonal polynomials are concerned, it is known that they satisfy a second order differential equation while the corresponding associated and the co-recursive associated ones satisfy a fourth order differential equation [29, 33]. This property has been used in [6, 7, 8, 13, 22, 23, 28, 29, 33, 34] to evaluate the moments of zero distribution of these polynomials, in terms of the coefficients of the differential equation that they satisfy. However, the orthogonal q -polynomials do not satisfy a differential equation but a q -difference equation of the following form

$$\sigma(x)D_q D_{1/q} y(x) + \tau(x)D_q y(x) + \lambda_{q,n} y(x) = 0,$$

where $\sigma(x)$ and $\tau(x)$ are polynomials of at most second and first degree respectively.

In the case of the associated orthogonal q -polynomials, a functional analytic method based on the three term recurrence relation that they satisfy, can be used to determine the Newton sum rules $\sum_{n=0}^{N-1} \lambda_n^k(c|q)$, $k = 1, 2, \dots$, of the zeros $\lambda_n(c|q)$ of the polynomials under consideration. This method has been introduced in [16] to obtain the Newton sum rules of the scaled co-recursive associated polynomials. In the present paper using this method we give the explicit expressions for the Newton sum rules $\sum_{n=0}^{N-1} \lambda_n^k(c|q)$, $k = 1, 2, 3, 4$ of the zeros of some families of associated orthogonal q -polynomials. We mention that by use of this method it is possible, although tedious, to calculate the Newton sum rules for any desired k .

In section 2 we describe briefly the method we use and in section 3 we determine the Newton sum rules for the zeros of the associated q-Laguerre, associated q-Charlier, associated continuous q-Ultraspherical, q-Lommel, associated q-Meixner, and associated Al-Salam-Carlitz I, II polynomials. From these, if we let $q \rightarrow 1^-$, we obtain as particular cases the Newton sum rules for the zeros of the associated Laguerre, associated Charlier, associated Ultraspherical and associated Meixner polynomials found recently [2, 12, 16, 18, 24, 25]. Also, from the Newton sum rules for the zeros of the q-Lommel polynomials, we obtain the corresponding results for the zeros of q-Bessel functions and from them the known Rayleigh sums [31] for the zeros of Bessel functions.

2 Preliminaries

A sequence $\{Q_n(x|q)\}_{n=0}^{\infty}$ of orthogonal q-polynomials satisfies a three term recurrence relation of the form

$$\alpha_n(q)Q_{n+1}(x|q) + \beta_n(q)Q_{n-1}(x|q) + b_n(q)Q_n(x|q) = xQ_n(x|q), \quad n = 0, 1, 2, \dots$$

$$Q_{-1}(x|q) = 0, \quad Q_0(x|q) = 1,$$

with $\alpha_n(q)$, $\beta_n(q)$, $b_n(q)$ real sequences and $\alpha_{n-1}(q)\beta_n(q) > 0$. We will always assume that $0 < q < 1$.

If we set $P_n(x|q) = U_n Q_n(x|q)$, where $U_n = \sqrt{\frac{\alpha_{n-1}(q)}{\beta_n(q)}} U_{n-1}$, $U_{-1} = 0$, $U_0 = 1$, we obtain the corresponding q-orthonormal polynomials $P_n(x|q)$ which satisfy the recurrence relation

$$a_n(q)P_{n+1}(x|q) + a_{n-1}(q)P_{n-1}(x|q) + b_n(q)P_n(x|q) = xP_n(x|q), \quad (2.1)$$

$$P_{-1}(x|q) = 0, \quad P_0(x|q) = 1,$$

where $a_n(q) = \sqrt{\alpha_n(q)\beta_{n+1}(q)}$. Obviously the polynomials $P_n(x|q)$ and $Q_n(x|q)$ have the same zeros.

The q-associated orthonormal polynomials $P_n^{(c)}(x|q)$ are obtained after replacing n by $n + c$, for arbitrary real $c \geq 0$ or $c > -1$ in the coefficients $a_n(q)$ and $b_n(q)$ of (2.1), i.e. they satisfy the following recurrence relation:

$$a_{n+c}(q)P_{n+1}^{(c)}(x|q) + a_{n+c-1}(q)P_{n-1}^{(c)}(x|q) + b_{n+c}(q)P_n^{(c)}(x|q) = xP_n^{(c)}(x|q), \quad (2.2)$$

$$P_{-1}^{(c)}(x|q) = 0, \quad P_0^{(c)}(x|q) = 1.$$

The functional analytic method that we use is, briefly, the following:

Let e_n , $n = 0, 1, \dots, N-1$ be an orthonormal base in a finite dimensional Hilbert space H_N , V the truncated shift operator $Ve_n = e_{n+1}$, $n = 0, 1, \dots, N-2$ and $Ve_{N-1} = 0$ and V^* its adjoint $V^*e_n = e_{n-1}$, $n = 1, 2, \dots, N-1$, $V^*e_0 = 0$. Also let $A^{(c)}(q)$, $B^{(c)}(q)$ be the diagonal operators $A^{(c)}(q)e_n = a_{n+c}(q)e_n$ and $B^{(c)}(q)e_n = b_{n+c}(q)e_n$, $n = 0, 1, \dots, N-1$.

According to [14, 15] the zeros $\lambda_n(c|q)$ of the polynomials $P_n^{(c)}(x|q)$ defined by (2.2) are the eigenvalues of the operator

$$T^{(c)}(q) = A^{(c)}(q)V^* + VA^{(c)}(q) + B^{(c)}(q), \quad (2.3)$$

i.e.

$$(A^{(c)}(q)V^* + VA^{(c)}(q) + B^{(c)}(q))x_n(c|q) = \lambda_n(c|q)x_n(c|q), \quad (2.4)$$

$$\|x_n(c|q)\| = 1, \quad n = 0, 1, \dots, N-1,$$

and vice versa.

In the following, we use an important result in the operator theory, for a symmetric operator M in a finite dimensional real Hilbert space, e.g. the space H_N : “The sum $\sum_{n=0}^{N-1} (Me_n, e_n)$ is independent of the orthonormal base e_n , $n = 0, 1, \dots, N-1$.” As in [16] it can be proved that, if $x_n(c|q)$, $n = 0, 1, \dots, N-1$, is the complete orthonormal system of eigenvectors of the operator $T^{(c)}(q)$ in H_N , then

$$\sum_{n=0}^{N-1} ((T^{(c)})^k(q)e_n, e_n) = \sum_{n=0}^{N-1} ((T^{(c)})^k(q)x_n(c|q), x_n(c|q)) = \sum_{n=0}^{N-1} \lambda_n^k(c|q). \quad (2.5)$$

From (2.5), we obtain for $k = 1, 2, 3, 4$, correspondingly

$$\sum_{n=0}^{N-1} \lambda_n(c|q) = \sum_{n=0}^{N-1} b_{n+c}(q), \quad (2.6)$$

$$\sum_{n=0}^{N-1} \lambda_n^2(c|q) = 2 \sum_{n=0}^{N-2} a_{n+c}^2(q) + \sum_{n=0}^{N-1} b_{n+c}^2(q), \quad (2.7)$$

$$\sum_{n=0}^{N-1} \lambda_n^3(c|q) = 3 \sum_{n=0}^{N-2} a_{n+c}^2(q)b_{n+c}(q) + 3 \sum_{n=0}^{N-2} a_{n+c}^2(q)b_{n+c+1}(q) + \sum_{n=0}^{N-1} b_{n+c}^3(q), \quad (2.8)$$

and

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^4(c|q) = & \sum_{n=0}^{N-1} b_{n+c}^4(q) + 2 \sum_{n=0}^{N-2} a_{n+c}^4(q) + 4 \sum_{n=0}^{N-3} a_{n+c}^2(q)a_{n+c+1}^2(q) + \\ & + 4 \sum_{n=0}^{N-2} a_{n+c}^2(q)b_{n+c}^2(q) + 4 \sum_{n=0}^{N-2} a_{n+c}^2(q)b_{n+c+1}^2(q) + 4 \sum_{n=0}^{N-2} a_{n+c}^2(q)b_{n+c}(q)b_{n+c+1}(q). \end{aligned} \quad (2.9)$$

3 Main results

3.1 Associated q-Laguerre polynomials

The associated q-Laguerre polynomials $L_n^{(c)}(x; a|q)$ satisfy the recurrence relation (2.2) with

$$a_{n+c}(q) = \sqrt{\frac{(1 - q^{n+c+1})(1 - q^{n+c+a+1})}{q^{4n+4c+2a+3}}}, \quad b_{n+c}(q) = \frac{1 - q^{n+c+1} + q - q^{n+c+a+1}}{q^{2n+2c+a+1}},$$

for $0 < q < 1$. From (2.6)-(2.9) we obtain, after some manipulations, the relations:

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n(c|q) &= (1+q)q^{-2c-a-1} \sum_{n=0}^{N-1} q^{-2n} - (1+q^a)q^{-c-a} \sum_{n=0}^{N-1} q^{-n} = \\ &= \frac{1 - q^{2N} - (1+q^a)q^{c+N}(1 - q^N)}{q^{2c+2N+a-1}(1 - q)}, \end{aligned} \quad (3.1)$$

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^2(c|q) &= 2q^{-4c-2a-3} \sum_{n=0}^{N-2} q^{-4n} - 2(1+q^a)q^{-3c-2a-2} \sum_{n=0}^{N-2} q^{-3n} + \\ &+ 2q^{-2c-a-1} \sum_{n=0}^{N-2} q^{-2n} + (1+q)^2 q^{-4c-2a-2} \sum_{n=0}^{N-1} q^{-4n} + \\ &+ (1+q^a)^2 q^{-2c-2a} \sum_{n=0}^{N-1} q^{-2n} - 2(1+q)(1+q^a)q^{-3c-2a-1} \sum_{n=0}^{N-1} q^{-3n} = \quad (3.2) \\ &= (-q - 2q^2 + 2q^{4N} + 2q^{1+c+N} + 2q^{2+c+N} + 2q^{1+a+c+N} + 2q^{2+a+c+N} - \\ &- q^{1+2c+2N} - 2q^{1+a+2c+2N} - 2q^{2+a+2c+2N} - q^{1+2a+2c+2N} + q^{1+4N} - \\ &- 2q^{c+4N} - 2q^{1+c+4N} - 2q^{a+c+4N} - 2q^{1+a+c+4N} + q^{1+2c+4N} + \\ &+ 2q^{a+2c+4N} + 2q^{1+a+2c+4N} + q^{1+2a+2c+4N}) / ((q^2 - 1)q^{2a+4c+4N-1}), \end{aligned}$$

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^3(c|q) &= -\frac{(3 + 3q + 3q^2 + q^3)}{q^{6c+3a}(1 - q^3)} + \frac{(1 + 3q + 3q^2 + 3q^3)}{q^{6N+6c+3a-3}(1 - q^3)} + \\ &+ \frac{3(1+q)(1+q^a)}{q^{5c+3a}(1 - q)} - \frac{3(1+q)(1+q^a)}{q^{5N+5c+3a-3}(1 - q)} + \\ &+ \frac{3(1+2q^a+q^{2a}+q^{1+a})}{q^{4N+4c+3a-3}(1 - q)} - \frac{3(q+q^a+2q^{1+a}+q^{1+2a})}{q^{4c+3a}(1 - q)} - \\ &- \frac{(1+3q^a+3q^{2a}+q^{3a}+3q^{1+a}+3q^{2+a}+3q^{1+2a}+3q^{2+2a})}{q^{3N+3c+3a-3}(1 - q^3)} + \\ &+ \frac{(q^2+3q^a+3q^{2a}+3q^{1+a}+3q^{2+a}+3q^{1+2a}+3q^{2+2a}+q^{2+3a})}{q^{3c+3a-1}(1 - q^3)}, \end{aligned} \quad (3.3)$$

and

$$\begin{aligned}
\sum_{n=0}^{N-1} \lambda_n^4(c|q) = & - \frac{(4 + 4q + 8q^2 + 8q^3 + 6q^4 + 4q^5 + q^6)}{q^{8c+4a+2}(1-q^4)} + \\
& + \frac{(1 + 4q + 6q^2 + 8q^3 + 8q^4 + 4q^5 + 4q^6)}{q^{8N+8c+4a-4}(1-q^4)} - \\
& \frac{4(1 + 2q + q^2 + q^3)(1 + q^a)}{q^{7N+7c+4a-4}(1-q)} + \frac{4(1 + q + 2q^2 + q^3)(1 + q^a)}{q^{7c+4a+2}(1-q)} + \\
& \frac{4(1 + q^a)(q^2 + q^a + 2q^{1+a} + 2q^{2+a} + q^{2+2a})}{q^{5c+4a+1}(1-q)} - \\
& \frac{4(1 + 3q^a + 3q^{2a} + q^{3a} + 2q^{1+a} + q^{2+a} + 2q^{1+2a} + q^{2+2a})}{q^{5N+5c+4a-4}(1-q)} + \\
& \frac{2(3 + 6q + 4q^2 + 2q^3 + 6q^a + 3q^{2a} + 14q^{1+a} + 12q^{2+a})}{q^{6N+6c+4a-4}(1-q^2)} + \\
& \frac{2(6q^{3+a} + 2q^{4+a} + 6q^{1+2a} + 4q^{2+2a} + 2q^{3+2a})}{q^{6N+6c+4a-4}(1-q^2)} - \\
& \frac{2(2q + 4q^2 + 6q^3 + 3q^4 + 2q^a + 6q^{1+a} + 12q^{2+a})}{q^{6c+4a+2}(1-q^2)} - \\
& \frac{2(14q^{3+a} + 6q^{4+a} + 2q^{1+2a} + 4q^{2+2a} + 6q^{3+2a} + 3q^{4+2a})}{q^{6c+4a+2}(1-q^2)} + \\
& \frac{(1 + 4q^a + 6q^{2a} + 4q^{3a} + q^{4a} + 4q^{1+a} + 4q^{2+a} + 4q^{3+a} + 8q^{1+2a})}{q^{4N+4c+4a-4}(1-q^4)} + \\
& \frac{(10q^{2+2a} + 8q^{3+2a} + 4q^{4+2a} + 4q^{1+3a} + 4q^{2+3a} + 4q^{3+3a})}{q^{4N+4c+4a-4}(1-q^4)} - \\
& \frac{(q^4 + 4q^{2a} + 4q^{1+a} + 4q^{2+a} + 4q^{3+a} + 4q^{4+a} + 8q^{1+2a} + 10q^{2+2a})}{q^{4c+4a}(1-q^4)} - \\
& \frac{(8q^{3+2a} + 6q^{4+2a} + 4q^{1+3a} + 4q^{2+3a} + 4q^{3+3a} + 4q^{4+3a} + q^{4+4a})}{q^{4c+4a}(1-q^4)}
\end{aligned} \tag{3.4}$$

respectively.

Remark 3.1.1. If we divide (3.1) by $1 - q$, (3.2) by $(1 - q)^2$, (3.3) by $(1 - q)^3$, (3.4) by $(1 - q)^4$ and let $q \rightarrow 1^-$, we obtain the Newton sum rules for the zeros of the associated Laguerre polynomials, which have been found in [16, 24]. Also, for $c = 0$ we obtain the Newton sum rules for the Laguerre polynomials which were given in [2].

3.2 Associated q-Charlier polynomials

The associated q-Charlier polynomials $C_n^{(c)}(q^{-x}; a|q)$ and the associated q-Laguerre polynomials $L_n^{(c)}(x; a|q)$ are related in the following way

$$L_n^{(c)}(x; a|q) = \frac{C_n^{(c)}(-x; -q^{-a}|q)}{(q; q)_n}.$$

Hence, if $\lambda_n^L(a, c|q)$ and $\lambda_n^C(a, c|q)$ are the zeros of the associated q-Laguerre and associated q-Charlier polynomials respectively, it holds

$$\sum_{n=0}^{N-1} \left(\lambda_n^L \left(-\frac{\ln(-a)}{\ln q}, c|q \right) \right)^k = \sum_{n=0}^{N-1} (-q^{-\lambda_n^C(a, c|q)})^k, \quad k = 1, 2, \dots \quad (3.5)$$

By using relation (3.5) we can find the following sums $\sum_{n=0}^{N-1} y_n^k(c; a|q)$, where $y_n(c; a|q) = q^{-\lambda_n^C(c; a|q)} - 1$, for the associated q-Charlier polynomials. From these, by setting $\alpha(1 - q)$ instead of a , dividing by $(1 - q)^k$ and taking the limit $q \rightarrow 1^-$ we get the Newton sum rules for the associated Charlier polynomials, which for $k = 1, 2, 3, 4$ are the following:

$$\sum_{n=0}^{N-1} \lambda_n(\alpha, c) = \frac{1}{2} N (N + 2c + 2\alpha - 1), \quad (3.6)$$

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^2(\alpha, c) = & \alpha^2 N + \frac{N(1 + 6(c - 1)c - 3N + 6cN + 2N^2)}{6} + \\ & + 2\alpha((N - 1)N + c(2N - 1)), \end{aligned} \quad (3.7)$$

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^3(\alpha, c) = & \alpha^3 N + \frac{N(-1 + 2c + N)(2c^2 + 2c(-1 + N) + (-1 + N)N)}{4} + \\ & + 3\alpha((-1 + N)^2 N + c^2(-2 + 3N) + c(-1 + N)(-1 + 3N)) + \\ & + \frac{3\alpha^2(3(-1 + N)N + c(-4 + 6N))}{2}, \end{aligned} \quad (3.8)$$

and

$$\begin{aligned}
\sum_{n=0}^{N-1} \lambda_n^4(\alpha, c) &= \alpha^4 N + \left((c-1)^2 c^2 - \frac{1}{30} \right) N + (c-1) c (2c-1) N^2 + \\
&+ \frac{(1+6(c-1)c) N^3}{3} + \left(c - \frac{1}{2} \right) N^4 + \frac{N^5}{5} + 4\alpha^3 (2(N-1)N + c(4N-3)) + \\
&+ 4\alpha ((N-1)^3 N + c^3 (4N-3) + c(N-1)^2 (4N-1) + c^2 (3-9N+6N^2)) + \\
&+ 2\alpha^2 ((-1+N)N(-7+6N) + c^2 (-17+18N) + 2c(3+N(-13+9N)))
\end{aligned} \quad (3.9)$$

Remark 3.2.1. For $c = 0$, from relations (3.6)-(3.9) we obtain the Newton sum rules for the zeros of Charlier polynomials, which were also given in [2].

3.3 Associated continuous q-Ultraspherical polynomials

The associated continuous q-Ultraspherical polynomials satisfy the recurrence relation (2.2) with

$$a_{n+c}(q) = \frac{1}{2} \sqrt{\frac{(1-q^{n+c+1})(1-\lambda^2 q^{n+c})}{(1-\lambda q^{n+c})(1-\lambda q^{n+c+1})}} \quad b_{n+c}(q) = 0.$$

From (2.6) and (2.8) we have

$$\sum_{n=0}^{N-1} \lambda_n(c|q) = \sum_{n=0}^{N-1} \lambda_n^3(c|q) = 0.$$

From (2.7) and (2.9), using the relation

$$\begin{aligned}
\frac{(1-q^{n+c+1})(1-\lambda^2 q^{n+c})}{(1-\lambda q^{n+c})(1-\lambda q^{n+c+1})} &= 1 + \frac{q^{n+c}(1-\lambda)(\lambda-q)}{(1-\lambda q^{n+c})(1-\lambda q^{n+c+1})} = \\
&= 1 + \frac{(1-\lambda)(\lambda-q)}{\lambda(1-q)} \left[\frac{1}{1-\lambda q^{n+c}} - \frac{1}{1-\lambda q^{n+c+1}} \right],
\end{aligned} \quad (3.10)$$

after some manipulations we obtain:

$$\begin{aligned}
\sum_{n=0}^{N-1} \lambda_n^2(c|q) &= \frac{1}{2} \sum_{n=0}^{N-1} \frac{(1-q^{n+c+1})(1-\lambda^2 q^{n+c})}{(1-\lambda q^{n+c})(1-\lambda q^{n+c+1})} = \\
&= \frac{N-1}{2} + \frac{(1-\lambda)(\lambda-q)}{2\lambda(1-q)} \left[\frac{1}{1-\lambda q^c} - \frac{1}{1-\lambda q^{c+N-1}} \right]
\end{aligned} \quad (3.11)$$

and

$$\begin{aligned}
 \sum_{n=0}^{N-1} \lambda_n^4(c|q) &= 2 \frac{1}{4^2} \sum_{n=0}^{N-2} \left\{ 1 + \frac{(1-\lambda)(\lambda-q)}{\lambda(1-q)} \left[\frac{1}{1-\lambda q^{n+c}} - \frac{1}{1-\lambda q^{n+c+1}} \right] \right\}^2 + \\
 &+ 4 \frac{1}{4^2} \sum_{n=0}^{N-3} \left\{ 1 + \frac{(1-\lambda)(\lambda-q)}{\lambda(1-q)} \left[\frac{1}{1-\lambda q^{n+c}} - \frac{1}{1-\lambda q^{n+c+1}} \right] \right\} \\
 &\quad \left\{ 1 + \frac{(1-\lambda)(\lambda-q)}{\lambda(1-q)} \left[\frac{1}{1-\lambda q^{n+c+1}} - \frac{1}{1-\lambda q^{n+c+2}} \right] \right\} = \\
 &= \frac{3N-5}{8} + \frac{(1-\lambda)^2(\lambda-q)^2}{8\lambda^2(1-q)^2} \times \\
 &\quad \left[\left(\frac{1}{1-\lambda q^{c+N-2}} - \frac{1}{1-\lambda q^{c+N-1}} \right)^2 + \frac{1}{(1-\lambda q^c)^2} - \frac{1}{(1-\lambda q^{c+N-2})^2} \right] \\
 &+ \frac{(1-\lambda)(\lambda-q)}{4\lambda(1-q)} \left[\frac{1}{1-\lambda q^{c+1}} - \frac{2}{1-\lambda q^{c+N-1}} + \frac{2}{1-\lambda q^c} - \frac{1}{1-\lambda q^{c+N-2}} \right] \\
 &+ \frac{(1-\lambda)^2(\lambda-q)^2}{4\lambda^2(1-q)^2(1-q^2)} \left[\frac{q}{1-\lambda q^{c+1}} - \frac{1}{1-\lambda q^c} + \frac{1}{1-\lambda q^{c+N-2}} - \frac{q}{1-\lambda q^{c+N-1}} \right].
 \end{aligned} \tag{3.12}$$

Remark 3.3.1. Setting in (3.11) and (3.12) q^λ instead of λ and taking the limit $q \rightarrow 1^-$, we find the Newton sum rules for the associated Ultraspherical polynomials, which can also be obtained, as special cases ($a = b = \lambda - 1/2$) from the corresponding formulas for the associated Jacobi polynomials which are given in [16, 24, 25].

Remark 3.3.2. The associated continuous q-Legendre polynomials $P_n^{(c)}(x|q)$ are related to the associated continuous q-Ultraspherical polynomials $C_n^{(c)}(x; \lambda|q)$ in the following way

$$C_n^{(c)}(x; q^{\frac{1}{2}}|q) = q^{-\frac{1}{4}n} P_n^{(c)}(x|q).$$

Therefore, the Newton sum rules for the zeros of the associated continuous q-Legendre polynomials can be obtained from the corresponding Newton sum rules for the zeros of the associated continuous q-Ultraspherical polynomials by setting $\lambda = q^{\frac{1}{2}}$. Moreover, taking the limit $q \rightarrow 1^-$ we get the Newton sum rules for the zeros of the associated Legendre polynomials, which for $c = 0$ (Legendre polynomials) were also given in [2].

Remark 3.3.3. The continuous associated q-Hermite polynomials can be obtained from the continuous associated q-Ultraspherical polynomials, by setting $\lambda = 0$. Thus the Newton sum rules for the zeros of the continuous associated q-Hermite polynomials are obtained from (3.11) and (3.12) for $\lambda = 0$. From these, by setting $x\sqrt{\frac{1-q}{2}}$ instead of x and letting $q \rightarrow 1^-$ we obtain the Newton sum rules for the associated Hermite polynomials

found in [12, 16, 18, 24]. Moreover, for $c = 0$ we find the Newton sum rules of the zeros of Hermite polynomials which were also given in [2].

3.4 q-Lommel polynomials

The q-Lommel polynomials satisfy the recurrence relation (2.2) with

$$a_n(\nu|q) = \frac{1}{2} \sqrt{\frac{q^{\nu+n}}{(1-q^{\nu+n+1})(1-q^{\nu+n})}}, \quad b_n(\nu|q) = 0.$$

From (2.6) and (2.8) we obtain

$$\sum_{n=0}^{N-1} \lambda_n(\nu|q) = \sum_{n=0}^{N-1} \lambda_n^3(\nu|q) = 0.$$

Also, from (2.7), using the relation

$$\frac{q^{n+\nu}}{(1-q^{n+\nu+1})(1-q^{\nu+n})} = \frac{1}{1-q} \left[\frac{1}{1-q^{\nu+n}} - \frac{1}{1-q^{\nu+n+1}} \right], \quad (3.13)$$

we obtain

$$\sum_{n=0}^{N-1} \lambda_n^2(\nu|q) = \frac{1}{2(1-q)} \left[\frac{1}{1-q^\nu} - \frac{1}{1-q^{\nu+N-1}} \right] = \frac{q^\nu(1-q^{N-1})}{2(1-q)(1-q^\nu)(1-q^{\nu+N-1})}. \quad (3.14)$$

Also, from (2.9) using (3.13) and the relations

$$\begin{aligned} \frac{1}{(1-q^{\nu+n+1})(1-q^{\nu+n+2})} &= \frac{1}{1-q} \left[\frac{1}{1-q^{\nu+n+1}} - \frac{q}{1-q^{\nu+n+2}} \right], \\ \frac{1}{(1-q^{\nu+n})(1-q^{\nu+n+2})} &= \frac{1}{1-q^2} \left[\frac{1}{1-q^{\nu+n}} - \frac{q^2}{1-q^{\nu+n+2}} \right], \end{aligned}$$

after some manipulations we obtain

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^4(\nu|q) &= \frac{1}{2^3} \frac{q^{2\nu+2(N-2)}}{(1-q^{\nu+N-1})^2(1-q^{\nu+N-2})^2} + \\ &+ \frac{1}{2^3} \sum_{n=0}^{N-3} \left[\frac{q^{2(\nu+n)}}{(1-q^{\nu+n+1})^2(1-q^{\nu+n})^2} + \frac{2q^{2(\nu+n)+1}}{(1-q^{\nu+n})(1-q^{\nu+n+1})^2(1-q^{\nu+n+2})} \right] = \\ &= \frac{1}{2^3} \frac{q^{2\nu+2(N-2)}}{(1-q^{\nu+N-1})^2(1-q^{\nu+N-2})^2} + \\ &+ \frac{1}{2^3} \left[\frac{1}{(1-q^\nu)^2} - \frac{1}{(1-q^{\nu+N-2})^2} + \frac{2}{1-q^2} \left(\frac{1}{1-q^{\nu+N-2}} - \frac{1}{1-q^\nu} \right) + \right. \\ &\quad \left. + \frac{2q}{1-q^2} \left(\frac{1}{1-q^{\nu+1}} - \frac{1}{1-q^{\nu+N-1}} \right) \right] \frac{1}{(1-q)^2} \end{aligned} \quad (3.15)$$

Remark 3.4.1. It is known [17] that when $N \rightarrow \infty$ then $\lambda_n(\nu|q) \rightarrow \pm \frac{1}{j_{\nu-1,n}(q)}$, where $j_{\nu,n}(q)$ are the zeros of the q-Bessel function. Thus, from (3.14) and (3.15) for $N \rightarrow \infty$ we obtain

$$\sum_{n=0}^{\infty} \frac{1}{j_{\nu,n}^2(q)} = \frac{q^{\nu+1}}{4(1-q)(1-q^{\nu+1})}, \quad (3.16)$$

$$\sum_{n=0}^{\infty} \frac{1}{j_{\nu,n}^4(q)} = \frac{1}{2^3} \left[\frac{1}{(1-q^{\nu+1})^2} - \frac{2}{(1+q)(1-q^{\nu+1})(1-q^{\nu+2})} + \frac{1-q}{1+q} \right] \frac{1}{(1-q)^2}. \quad (3.17)$$

If we multiply (3.16) by $(1-q)^2$, (3.17) by $(1-q)^4$ and let $q \rightarrow 1^-$, we obtain the known [31] Rayleigh sums for the zeros $j_{\nu,n}$ of Bessel functions.

3.5 Associated q-Meixner polynomials

The associated q-Meixner polynomials satisfy the recurrence relation (2.2) for $x = q^{-x} - 1$ and

$$a_n(r, b, c|q) = \sqrt{\frac{r(1-q^{n+c+1})(r+q^{n+c+1})(1-bq^{n+c+1})}{q^{4n+4c+3}}},$$

$$b_n(r, b, c|q) = q^{-n-c} + r[1+q-(1+b)q^{n+c+1}]q^{-2n-2c-1} - 1,$$

where $c > -1$, $r > 0$ and $b < q^{-c-1}$. Setting $y_n(r, b, c|q) = q^{-\lambda_n(r, b, c|q)} - 1$ and after some manipulations, from (2.6) - (2.8) we obtain

$$\begin{aligned} \sum_{n=0}^{N-1} y_n(r, b, c|q) &= \sum_{n=0}^{N-1} b_n(r, b, c|q) = \sum_{n=0}^{N-1} (-1) + q^{-c} \sum_{n=0}^{N-1} q^{-n} + \\ &+ rq^{-2c-1} \sum_{n=0}^{N-1} q^{-2n} + rq^{-2c} \sum_{n=0}^{N-1} q^{-2n} - r(1+b)q^{-c} \sum_{n=0}^{N-1} q^{-n} = \quad (3.18) \\ &= -N + \frac{q^{-c-N+1}}{1-q} [(1-q^N)(1-r-rb) + rq^{-c-N}(1-q^{2N})], \end{aligned}$$

$$\begin{aligned} \sum_{n=0}^{N-1} y_n^2(r, b, c|q) &= 2 \sum_{n=0}^{N-2} a_n^2(r, b, c|q) + \sum_{n=0}^{N-1} b_n^2(r, b, c|q) = 2r^2 q^{-4c-3} \sum_{n=0}^{N-2} q^{-4n} + \\ &+ 2r(1-r-rb)q^{-3c-2} \sum_{n=0}^{N-2} q^{-3n} - 2r(1+b-rb)q^{-2c-1} \sum_{n=0}^{N-2} q^{-2n} + 2rbq^{-c} \sum_{n=0}^{N-2} q^{-n} + \end{aligned}$$

$$\begin{aligned}
& + \sum_{n=0}^{N-1} 1 + r^2(1+q)^2 q^{-4c-2} \sum_{n=0}^{N-1} q^{-4n} + 2r(1+q)(1-r-rb)q^{-3c-1} \sum_{n=0}^{N-1} q^{-3n} \\
& + (1-r-rb)^2 q^{-2c} \sum_{n=0}^{N-1} q^{-2n} - 2r(1+q)q^{-2c-1} \sum_{n=0}^{N-1} q^{-2n} + 2(1-r-rb)q^{-c} \sum_{n=0}^{N-1} q^{-n} = \\
& = N - \frac{q^{-4c+1}(2+q)r^2}{1-q^2} + \frac{q^{-4N-4c+2}(1+2q)r^2}{1-q^2} - \frac{2q^{-3c+1}r(1-r-rb)}{1-q} + \\
& + \frac{2q^{-3N-3c+2}r(1-r-rb)}{1-q} + \frac{2q^{-c+1}(1-r-2rb)}{1-q} - \\
& - \frac{2q^{-N-c+1}(1-r-rb-rbq)}{1-q} - \\
& - \frac{q^{-2c+1}(q-4r-2br-4qr-2bqr+2br^2+qr^2+2bqr^2+b^2r^2q)}{1-q^2} + \\
& + \frac{q^{-2N-2c+1}(q-2(1+q)(1+q+bq)r+q((1+b)^2+2bq)r^2)}{1-q^2},
\end{aligned} \tag{3.19}$$

and

$$\sum_{n=0}^{N-1} y_n^3(r, b, c|q) = A + Br + \Gamma r^2 + \Delta r^3, \tag{3.20}$$

where

$$\begin{aligned}
A & = -N - \frac{3q^{1-c}(1-q^{-N})}{1-q} + \frac{3q^{2-2c}(1-q^{-2N})}{(1-q^2)} - \frac{q^{3-3c}(1-q^{-3N})}{(1-q^3)}, \\
B & = \frac{3q^{1-4c-4N}}{1-q} \{ q^2 - q^{4N} + q^{c+N} [(3+b)q^{3N} - q(2+q+bq)] + q^{3(c+N)} [-1 + q^N + \\
& + b(-1-2q+3q^N)] + q^{2(c+N)} [1 - 3(1+b)q^{2N} + q(2+b(2+q))] \}, \\
\Gamma & = \frac{3(1+q)}{(-1+q)q^{5c}} - \frac{3q^{3-5c-5N}(1+q)}{-1+q} + \frac{3(b+3q+5bq+b^2q)}{(-1+q)q^{3c}} - \\
& - \frac{3q^{1-2c}(3b+b^2+q+3bq+2b^2q)}{(-1+q)(1+q)} - \frac{3q^{2-3c-3N}(2+2b+q+3bq+b^2q+bq^2)}{-1+q} - \\
& - \frac{3(1+b+5q+3bq+3q^2+2bq^2)}{(-1+q)q^{4c}(1+q)} + \\
& + \frac{3q^{2-2c-2N}(1+2b+b^2+3bq+b^2q+bq^2+b^2q^2)}{(-1+q)(1+q)} + \\
& + \frac{3q^{2-4c-4N}(1+4q+2bq+3q^2+3bq^2+q^3+bq^3)}{(-1+q)(1+q)},
\end{aligned}$$

and

$$\begin{aligned} \Delta = & \frac{-3(1+b)(1+q)}{(-1+q)q^{5c}} + \frac{3(1+b)q^{3-5c-5N}(1+q)}{-1+q} - \frac{3q^{3-4c-4N}(1+2b+b^2+bq)}{-1+q} + \\ & + \frac{3(b+q+2bq+b^2q)}{(-1+q)q^{4c}} - \frac{(1+b)q^{1-3c}(3b+3bq+q^2+2bq^2+b^2q^2)}{(-1+q)(1+q+q^2)} + \\ & + \frac{q^{3-3c-3N}(1+3b+3b^2+b^3+3bq+3b^2q+3bq^2+3b^2q^2)}{(-1+q)(1+q+q^2)} + \\ & + \frac{3+3q+3q^2+q^3}{(-1+q)q^{6c}(1+q+q^2)} - \frac{q^{3-6c-6N}(1+3q+3q^2+3q^3)}{(-1+q)(1+q+q^2)}. \end{aligned}$$

Remark 3.5.1. Setting $q^{\beta-1}$, and $\frac{\mu}{1-\mu}$ instead of b and r respectively and taking the limit $q \rightarrow 1^-$ in (3.18) - (3.20), we obtain the following Newton sum rules for the zeros of the associated Meixner polynomials.

$$\sum_{n=0}^{N-1} \lambda_n(\beta, \mu, c) = \frac{N(1-N-(N-1+2\beta)\mu-2c(1+\mu))}{2(\mu-1)}, \quad (3.21)$$

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^2(\beta, \mu, c) = & \frac{1}{6(-1+\mu)^2} [N(1+6(-1+c)c-3N+6cN+2N^2) \\ & + 2(6c^2(2N-1)+6c(2N-1)(N+\beta-1)+(N-1)N(4N+6\beta-5))\mu \\ & + N(1+6c^2+2N^2+6(-1+\beta)\beta+6c(-1+N+2\beta)+N(-3+6\beta))\mu^2], \end{aligned} \quad (3.22)$$

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^3(\beta, \mu, c) = & \frac{1}{4(1-\mu)^3} [(N(2c+N-1)(2c^2+2c(N-1)+(N-1)N)+ \\ & 3(4c^3(3N-2)+2c^2(3N-2)(3N+2\beta-3)+(N-1)^2N(3N+4\beta-4)+ \\ & c(N-1)(2-2\beta+3N(2N+2\beta-3)))\mu+3(4c^3(-2+3N)+2c^2(-2+3N) \\ & \times (3N+4\beta-3)+(N-1)N(4+N(3N-7)-10\beta+8N\beta+6\beta^2)+ \\ & 2c(6N^3+6\beta-2-4\beta^2+3N^2(4\beta-5)+N(11+6(\beta-3)\beta)))\mu^2+ \\ & N(2c+N+2\beta-1)(2c^2+N^2+2(\beta-1)\beta+ \\ & N(2\beta-1)+2c(N+2\beta-1))\mu^3]. \end{aligned} \quad (3.23)$$

Also, from the sum $\sum_{n=0}^{N-1} (q^{-\lambda_n(r,b,c|q)} - 1)^4$, which can be obtained in a similar way, but which we omit due to its complexity and lack of space, we get the following Newton sum rule of the fourth power of the zeros of the associated Meixner polynomials

$$\sum_{n=0}^{N-1} \lambda_n^4(\beta, \mu, c) = \frac{1}{30(-1+\mu)^4} \{A + B\mu + \Gamma\mu^2 + \Delta\mu^3 + E\mu^4\}, \quad (3.24)$$

where

$$A = N(-1 + 30c^4 + 60c^3(-1 + N) + 30c(-1 + N)^2N + 30c^2(-1 + N)(-1 + 2N) + N^2(10 + 3N(-5 + 2N))),$$

$$B = 4[30c^4(-3 + 4N) + 30c^3(-3 + 4N)(-2 + 2N + \beta) + (-1 + N)N(-31 + 89N - 81N^2 + 24N^3 + 30(-1 + N)^2\beta) + 30c(-1 + N)^2(1 - \beta + 2N(-3 + 2N + 2\beta)) + 30c^2(-1 + N)(4 - 3\beta + N(-13 + 8N + 6\beta))],$$

$$\Gamma = 6\{10c^4(-17 + 18N) + 20c^3(-17 + 18N)(-1 + N + \beta) + (-1 + N)N[-79 + N(181 + N(-139 + 36N)) + 150\beta + 10N(-23 + 9N)\beta + 10(-7 + 6N)\beta^2] + 10c^2[-25 + 36N^3 + (40 - 17\beta)\beta + 6N(-4 + \beta)(-4 + 3\beta) + 3N^2(-35 + 18\beta)] + 20c(-1 + N + \beta)(-4 + 3\beta + N(22 - 13\beta + N(-26 + 9N + 9\beta)))\},$$

$$\Delta = 4\{30c^4(-3 + 4N) + 30c^3(-3 + 4N)(-2 + 2N + 3\beta) + 30c[(-1 + N)^2(1 - 6N + 4N^2) + 2(-1 + N)(2 - 9N + 6N^2)\beta + 6(-1 + N)(-1 + 2N)\beta^2 + (-3 + 4N)\beta^3] + 30c^2[8N^3 - (2 - 3\beta)^2 + 3N^2(-7 + 6\beta) + N(17 + 6\beta(-5 + 2\beta))] + (-1 + N)N \times [-31 + 24N^3 + 9N^2(-9 + 10\beta) + 30\beta(4 + \beta(-5 + 2\beta)) + N(89 + 30\beta(-7 + 4\beta))]\},$$

and

$$E = N[-1 + 10N^2 - 60cN^2 + 60c^2N^2 - 15N^3 + 30cN^3 + 6N^4 - 60N^2\beta + 120cN^2\beta + 30N^3\beta + 60N^2\beta^2 + 30(-1 + c + \beta)^2(c + \beta)^2 + 30N(-1 + c + \beta)(c + \beta)(-1 + 2c + 2\beta)].$$

Moreover for $c = 0$ from (3.21)-(3.24), we obtain the Newton sum rules for the Meixner polynomials. The first three of them were also given in [2].

Remark 3.5.2. The associated q -Charlier polynomials can be obtained from the associated q -Meixner polynomials for $b = 0$ and $r = a$. Thus from (3.18)-(3.20) we immediately find the corresponding formulas for the associated q -Charlier polynomials. Moreover, setting $a = \alpha(1 - q)$ and taking the limit $q \rightarrow 1^-$ we obtain the Newton sum rules for the associated Charlier polynomials (see (3.6)-(3.9)).

Remark 3.5.3. The associated Al-Salam-Carlitz II polynomials can be obtained from the associated q -Meixner polynomials by setting $b = -t^{-1}a$ and taking the limit $t \rightarrow 0^+$. Then from (3.18)-(3.20), after some manipulations, we find

$$\sum_{n=0}^{N-1} \lambda_n(a, c|q) = \frac{a+1}{q^c} \sum_{n=0}^{N-1} q^{-n} = \frac{(a+1)(1-q^N)}{q^{c+N-1}(1-q)}, \quad (3.25)$$

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^2(a, c|q) &= \frac{2a}{q^{2c+1}} \sum_{n=0}^{N-2} \frac{1-q^{n+c+1}}{q^{2n}} + \frac{(a+1)^2}{q^{2c}} \sum_{n=0}^{N-1} q^{-2n} = \\ &= \frac{2a(1-q^{2N-2}) - 2a(1+q)(1-q^{N-1})q^{c+N+1} + (a+1)^2(1-q^{2N})q}{q^{2c+2N-3}(1-q^2)}, \end{aligned} \quad (3.26)$$

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^3(a, c|q) &= -\frac{3a(1+a)}{(1-q)q^{2N+2c-3}} + \frac{(1+a)((1+a)^2 + 3aq(1+q))}{q^{3N+3c-3}(1-q^3)} - \\ &\quad - \frac{(1+a)(3a(1+q+q^{c+1}) + q^2(1+a)^2 - 3aq^c(1+q)^2)}{q^{3c-1}(1-q^3)}. \end{aligned} \quad (3.27)$$

Also, from the formula of $\sum_{n=0}^{N-1} (q^{-\lambda_n(r,b,c|q)} - 1)^4$, we find

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^4(a, c|q) &= \frac{(1+a)^4 q^{4-4c-4N} (1-q^{4N})}{1-q^4} - \frac{2aq^{-4c-4N}}{1-q^4} \{-aq^6 - 2aq^8 + \\ &\quad + q^{4N} (2a + aq^2) + 2(1+a)^2 (1+q+q^2) (-q^5 + q^{1+4N}) + \\ &\quad + aq^{2(c+N)} (q+q^3) (q^{2N} (2+q) - q^3 (1+2q)) + 2q^{c+N} (1+q) \times \\ &\quad \times (1+q^2) (-q^{3N} (a + (1+a)^2 q) + q^4 (1+a(2+a+q)))\}. \end{aligned} \quad (3.28)$$

3.6 Associated Al-Salam Carlitz I polynomials

The associated Al-Salam Carlitz I polynomials satisfy the recurrence relation (2.2) with

$$a_{n+c}(a|q) = \sqrt{-a(1 - q^{n+c+1})q^{n+c}}, \quad b_{n+c}(a|q) = (1 + a)q^{n+c}, \quad a < 0.$$

From equations (2.6)-(2.9) we find

$$\sum_{n=0}^{N-1} \lambda_n(a, c|q) = \frac{q^c(1 + a)(1 - q^N)}{1 - q},$$

$$\sum_{n=0}^{N-1} \lambda_n^2(a, c|q) = \frac{q^{2c}(1 + a)^2(1 - q^{2N})}{1 - q^2} + \frac{2aq^{2c+1}(1 - q^{2N-2})}{1 - q^2} - \frac{2aq^c(1 - q^{N-1})}{1 - q},$$

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^3(a, c|q) = & -\frac{3a(1 + a)q^{2c-2}(q^2 - q^{2N} + q^{3N+c})}{1 - q} - \\ & - \frac{q^{3N+3c}(a^2 - a + 1)(a + 1)}{1 - q^3} + \frac{(1 + a)q^{3c}[(1 + a)^2 + 3aq(1 + q)]}{1 - q^3}, \end{aligned}$$

and

$$\begin{aligned} \sum_{n=0}^{N-1} \lambda_n^4(a, c|q) = & \frac{4aq^{3N+3c-4}(a + q + 2aq + a^2q)}{1 - q} - \frac{2a^2q^{2N+2c-3}(2 + q)}{1 - q^2} - \\ & - \frac{q^{4(N+c-1)}[4a^2 + 4a(1 + a)^2q + 2a(2 + a)(1 + 2a)q^2 + 4a(1 + a)^2q^3 + (1 + a)^4q^4]}{1 - q^4} + \\ & + \frac{q^{2c}[q^{2c} + a^4q^{2c} - 4aq^c(1 + q)(1 + q^2)(1 - q^c) - 4a^3q^c(1 + q)(1 + q^2)(1 - q^c)]}{1 - q^4} - \\ & - \frac{q^{2c}\{2a^2(1 + q^2)(1 - q^c)[1 + 2q - q^c(3 + 2q(2 + q))]\}}{1 - q^4}. \end{aligned}$$

References

- [1] R. Álvarez-Nodarse, E. Buendia and J.S. Dehesa, The distribution of zeros of general q-polynomials, *J. Phys. A*, 30, 6743-6768 (1997).
- [2] R. Álvarez-Nodarse, J. Dehesa, Distributions of zeros of discrete and continuous polynomials from their recurrence relation, *Appl. Math. Comput.*, 128, 167-190 (2002).

- [3] R. Álvarez-Nodarse, *Polinomios hipergeometricos y q-polinomios*, Monografias del Seminario Matematico “Garcia de Galdeano” Numero 26. Prensas Universitarias de Zaragoza, Spain, 2003. ISBN 84-7733-637-7 (in spanish).
- [4] G.E.Andrews and R.Askey, Classical orthogonal polynomials, in *Polynomes Orthogonaux et Applications* (C.Brezinski et al., eds.), Lecture Notes in Mathematics 1171, Springer-Verlag, Berlin, 1985, pp. 33-62.
- [5] M.C. Bosca and J.S. Dehesa, Rational Jacobi matrices and certain quantum mechanical problems, *J. Phys. A: Math. Gen.*, 17, 3487-3491 (1984).
- [6] E.Buendia, J.S.Dehesa and F.J. Gálver, The distribution of zeros of the polynomial eigenfunctions of ordinary differential operators of arbitrary order, in *Orthogonal Polynomials and their Applications* (M. Alfaro et al., eds.), Lecture Notes in Mathematics 1329, Springer, Berlin, 1986, pp. 222-235.
- [7] E.Buendia, J.S.Dehesa and A.Sanchez-Buendia, On the zeros of eigenfunctions of polynomial differential operators, *J. Math. Phys.*, 26 (11), 2729-2736 (1985).
- [8] K.M.Case, Sum rules for the zeros of polynomials I-II, *J.Math.Phys.*, 21, 702-714 (1980).
- [9] J.H.Dehesa, E.Buendia and M.A.Sanchez-Buendia, On the polynomial solutions of ordinary differential equations of the fourth order, *J.Math. Phys.*, 26 (7), 1547-1552 (1985).
- [10] W.Hahn, Uber uthogonalpolynome, die q-Differenzgleichungen genugen, *Math. Nachr.*, 2, 4-34 (1949).
- [11] B.Germano and P.E.Ricci, Representation formulas for the moments of the density of zeros of orthogonal polynomials, *Mathematiche (Catania)*, 48, 77-86 (1993).
- [12] B.Germano, P.Natalini and P.E.Ricci, Computing the moments of the density of zeros for orthogonal polynomials, *Comput. Math. Appl.* (issue on Concrete Analysis), 30, 69-81 (1995).
- [13] E. Godoy, J. S. Dehesa and A. Zarzo, Density of zeros of orthogonal polynomials. A study with MATHEMATICA, in *Orthogonal Polynomials and their Applications* (L. Arias et al., eds.), Univ. of Oviedo, 1989, pp. 136-154.

- [14] E.K.Ifantis and P.D.Siafarikas, Differential inequalities for the largest zero of Laguerre and Ultraspherical polynomials, in *Actas del VI Symposium on Polinomios Orthogonales y Applicationes*, Gijon, Spain, 1989, pp. 187-197.
- [15] E.K.Ifantis and P.D.Siafarikas, Differential inequalities and monotonicity properties of the zeros of associated Laguerre and Hermite polynomials, *Ann. Numer. Math.*, 2, 79-91 (1995).
- [16] E.K.Ifantis, C.G.Kokologiannaki and P.D.Siafarikas, Newton sum rules and monotonicity properties of scaled co-recursive associated polynomials, *Methods Anal.*, 3 (4), 486-497 (1996).
- [17] M.E.H.Ismail and M.E.Muldoon, On the variation with respect to a parameter of zeros of Bessel and q-Bessel functions, *J. Math. Anal. Appl.*, 135, 187-207 (1988).
- [18] T.Isoni, P.Natalini and P.E.Ricci, Symbolic computation of Newton sum rules for the zeros of polynomial eigenfunctions of linear differential operators, *Numerical Algorithms*, 28, 215-227 (2001).
- [19] R.Koekoek and R.F.Swarttouw, *The Askey-scheme of hypergeometric orthogonal polynomials and its q-analogue*, Delft University of Technology, Faculty of Information Technology and Systems, Dep. of Technical Mathematics and Informatics, Report no. 98-17 (1998).
- [20] C.G.Kokologiannaki, P.D.Siafarikas and J.D.Stabolas, Monotonicity properties and inequalities of the zeros of q-associated polynomials, in *Nonlinear Analysis and Applications* (R.P.Agarwal and D.O'Regan, eds.), Vol. II, 2004 pp. 671-726
- [21] W.Lang, On sums of powers of zeros of polynomials, *J. Comput. Appl. Math.*, 89, 237-256 (1998).
- [22] P.Natalini, Sul calcolo dei momenti della densita degli zeri dei polinomi ortogonali classici e semiclassici, *Calcolo*, 31, 127-144 (1994).
- [23] P.Natalini and P.E.Ricci, Computation of Newton sum rules for polynomial solutions of ODE with polynomial coefficients, *Riv. Mat. Univ. Parma*, 6 (3), 69-76 (2000).

- [24] P.Natalini and P.E.Ricci, Computation of Newton sum rules for the associated and co-recursive classical orthogonal polynomials, *J.Comput. Appl. Math.*, 133, 495-505 (2001).
- [25] P.Natalini and P.E.Ricci, Newton sum rules of polynomials defined by a three-term recurrence relation, *Comput. Math. Appl.*, 42, 767-771 (2001).
- [26] P.Natalini and P.E.Ricci, The moments of the density of zeros of relativistic Hermite and Laguerre polynomials, *Comput. Math. Appl.*, 31 (6), 87-96 (1996).
- [27] E.N.Petropoulou, P.D.Siafarikas and J.D.Stabolas, Convexity results for the largest zero and functions involving the largest zero of q-associated polynomials, *Integral Transform Spec. Funct.*, (to appear).
- [28] P.E.Ricci, Sum rules for zeros of polynomials and generalized Lucas polynomials, *J.Math. Phys.*, 34 (10), 4884-4891 (1993).
- [29] A.Ronveaux, A.Zarzo and E.Godoy, Fourth order differential equations satisfied by the generalized co-recursive of all classical orthogonal polynomials. A study of their distribution of zeros, *J. Comput. Appl. Math.*, 59, 295-328 (1995).
- [30] P.D.Siafarikas, J.D.Stabolas and L.Velasquez, Differential inequalities of functions involving the lowest zero of some associated orthogonal q-polynomials, *Integral Transform Spec. Funct.*, (to appear).
- [31] G.N.Watson, *A treatise on the theory of Bessel functions*, Cambridge Univ. Press, London/New York, 1965.
- [32] A.Zarzo, J.S.Dehesa, A.Ronveaux, Newton sum rules of zeros of semiclassical orthogonal polynomials, *J. Comput. Appl. Math.* 33, 85-96 (1990).
- [33] A. Zarzo, A. Ronveaux and E. Godoy, Fourth order differential equation satisfied by the associated of any order of all classical orthogonal polynomials. A study of their distribution of zeros, *J. Comput. Appl. Math.*, 49, 349-359 (1993).
- [34] A. Zarzo, R. J. Yáñez, A. Ronveaux, J. S. Dehesa, Algebraic and spectral properties of some quasiorthogonal polynomials encountered in quantum radiation, *J. Math. Phys.*, 36, 5179-5197 (1995).

Kondurar Theorem and Itô formula in Riesz spaces

A. Boccuto - D. Candeloro - E. Kubińska

Department of Mathematics and Computer Sciences,

via Vanvitelli 1, I-06123 Perugia (Italy)

e-mail: boccuto@dipmat.unipg.it, boccuto@yahoo.it;

Department of Mathematics and Computer Sciences,

via Vanvitelli 1, I-06123 Perugia (Italy)

e-mail: candelor@dipmat.unipg.it;

Nowy Sącz Business School- National-Louis University,

ul. Zielona 27, PL-33300 Nowy Sącz (Poland), and

Department of Mathematics, Silesian University,

ul. Bankowa 14, PL-40007 Katowice (Poland)

e-mail: kubinska@wsb-nlu.edu.pl

Abstract

In this paper we formulate a version of the Kondurar theorem and a generalization of the Itô formula for functions taking values in Riesz spaces with respect to a convergence, satisfying suitable axioms. When the involved space is the space of measurable functions, both convergence

almost everywhere and convergence in probability are included. Finally we present some comments and possible applications of the Itô formula.

2000 AMS Subject Classification: 28B15, 28B05, 28B10, 46G10.

Key words: Riesz spaces, convergence, Riemann-Stieltjes integration, Kondurar theorem, Itô formula.

1 Introduction

In this paper we establish the Kondurar theorem and the Itô formula in Riesz spaces.

The classical (scalar) version of the Kondurar theorem ensures the existence of the Riemann-Stieltjes integral $\int_a^b f dg$, as soon as f and g are Hölder-continuous, of orders α and β respectively, with $\alpha + \beta > 1$.

Here we chose $[a, b] = [0, 1]$, and fixed the Riesz space setting as a *product triple* (R_1, R_2, R) such that f is R_1 -valued, g is R_2 -valued, and the integral takes values in R .

More generally Riemann-Stieltjes integration is investigated also with respect to an interval function q in the place of g : we obtain meaningful extensions of the Kondurar theorem, and deduce a version of the Itô formula for integrands in a Riesz space. The classical stochastic integrals (Itô, Stratonovich, Backward) and the classical Itô formula are included.

Finally, we present some examples in order to illustrate the possible applications.

Our warmest thanks to the referees for their helpful suggestions.

2 Preliminaries

Definition 2.1 Throughout the paper, R will denote a *Dedekind complete* Riesz space. An axiomatic notion of convergence in R is defined by choosing a suitable *order ideal* $\mathcal{N}$ (in the sense of [6], p. 93) on the Riesz product space $\mathbb{R}^N$. More precisely, this ideal $\mathcal{N}$ is assumed to satisfy the following additional conditions:

- a) Every element of $\mathcal{N}$ is a *bounded* sequence.
- b) Every sequence in R whose terms eventually vanish belongs to $\mathcal{N}$.
- c) Every sequence of the form $(\frac{u}{n})_n$, with $u \in R$ and $u \geq 0$, belongs to $\mathcal{N}$.

By means of the ideal $\mathcal{N}$, it is possible to define the concept of convergence for nets in the following manner. Let $(T, \geq)$ be any directed set. Given any bounded net $\Psi : T \rightarrow R$, we say that it $\mathcal{N}$ -converges to an element y of R (or that it admits y as $\mathcal{N}$ -limit) if there exists a sequence $(s_n)_n$ of elements of T such that the sequence $(p_n)_n$ defined by $p_n = \sup_{t \in T, t \geq s_n} |\Psi(t) - y|$ belongs to $\mathcal{N}$. The ideal $\mathcal{N}$ then coincides with the space of all bounded sequences in R which converge to zero. For $\mathcal{N}$ -convergence all usual properties of limits hold; in particular, *uniqueness* of the limit and Cauchy criterion for convergence: a bounded net $\Psi : T \rightarrow R$ converges if and only if there exists a sequence $(s_n)_n$ of elements of T such that the sequence $(q_n)_n$ defined by $q_n = \sup_{s, t \in T, s \wedge t \geq s_n} |\Psi(s) - \Psi(t)|$ belongs to $\mathcal{N}$. Moreover, for a fixed directed set T , the space of all bounded nets of the form $\Psi : T \rightarrow R$ for which the limit exists is a Riesz subspace of R^T , and the operator which assigns to each

convergent net the corresponding limit is linear and order-preserving.

Now, let us denote by $\mathcal{J}$ the family of all (nontrivial) closed subintervals of the interval $[0, 1]$, and by $\mathcal{D}$ the family of all finite decompositions of $[0, 1]$, $D = \{0 = t_0 < t_1 < \dots < t_{n-1} < t_n = 1\}$, or also $D = \{I_1, \dots, I_n\}$, where $I_i = [t_{i-1}, t_i]$ for each $i = 1, \dots, n$.

Given any decomposition $D = \{I_1, \dots, I_n\}$, the *mesh* of D is the number $\delta(D) := \max\{|I_i| : i = 1, \dots, n\}$, where $|I|$ as usual denotes the *length* of the interval I . Then the set $\mathcal{D}$ can be endowed with the *mesh* ordering: $D_1 \geq D_2$ iff $\delta(D_1) \leq \delta(D_2)$; this makes $\mathcal{D}$ a directed set, and, unless otherwise specified, convergence of any net $\Psi : \mathcal{D} \rightarrow R$ will be always related to this order relation. However, besides this order relation, a weaker ordering will be helpful: for any two decompositions, D_1, D_2 , we say that D_2 is *finer* than D_1 , and write $D_2 \succ D_1$, if every interval from D_2 is entirely contained in some interval from D_1 . Moreover, we shall say that a decomposition D is *rational* if all its endpoints are rational, and that a decomposition D is *equidistributed* if all intervals in D have the same length. If D is equidistributed, and consists of 2^n intervals, we say that D is *dyadic* of order n . Equidistributed and dyadic decompositions of *any* interval $[a, b]$ are similarly defined. In general, the family of all decompositions of any interval $[a, b] \subset [0, 1]$ will be denoted by $\mathcal{D}_{[a,b]}$.

From now on, we shall assume that an ideal $\mathcal{N}$ has been fixed, according with Definition 2.1, and consequently all limits we shall introduce (unless otherwise specified) are related to $\mathcal{N}$ -convergence, without a particular notation.

Definition 2.2 Let $q : \mathcal{J} \rightarrow R$ be any interval function. We say that q is *integrable*, if the net $S(D) = \sum_{I \in D} q(I)$ is convergent to some element $Y := \int q$.

A very useful tool to prove integrability is the Cauchy property. We first introduce a notation: for every interval function $q : \mathcal{J} \rightarrow R$, we set

$$OB(q)(I) = \sup \left\{ \left| \sum_{J \in D} q(J) - \sum_{H \in D'} q(H) \right| : D, D' \in \mathcal{D}_I \right\}.$$

Now we have the following:

Theorem 2.3 Assume that $q : \mathcal{J} \rightarrow R$ is any interval function. The following three conditions are equivalent:

(i) q is integrable;

(ii) $(r_n)_n \in \mathcal{N}$, where

$$r_n := \sup \left\{ |S(D) - S(D_0)| : D, D_0 \in \mathcal{D}, \delta(D_0) \leq \frac{1}{n}, D \succ D_0 \right\};$$

(iii) $\int OB(q) = 0$.

Proof. It is easy to prove that (i) implies (ii). To prove the implication (ii) $\Rightarrow$ (iii), one can proceed in a similar fashion as in the proof of Theorem 1.4 in [4].

So we only prove that (iii) implies (i). By supposition, we have $(\kappa_n)_n \in \mathcal{N}$, where $\kappa_n = \sup \left\{ \sum_{I \in D} OB(q)(I) : \delta(D) \leq \frac{1}{n} \right\}$. Let us fix n , and choose two decompositions, D_0 and D , such that $\delta(D_0) \leq \frac{1}{n}$ and $D \succ D_0$. From

$$S(D_0) - S(D) = \sum_{I \in D_0} \left[q(I) - \sum_{J \in D, J \subset I} q(J) \right]$$

we obtain $|S(D_0) - S(D)| \leq \kappa_n$. From this one can easily deduce that

$$|S(D_1) - S(D_2)| \leq 2\kappa_n, \tag{1}$$

as soon as $\delta(D_1)$ and $\delta(D_2)$ are less than $\frac{1}{n}$. Now define

$$M_n = \sup \left\{ S(D) : \delta(D) \leq \frac{1}{n} \right\}, \quad m_n = \inf \left\{ S(D) : \delta(D) \leq \frac{1}{n} \right\}.$$

Of course, $m_n \leq M_n$. Moreover, thanks to (1), it follows easily that $M_n - m_n \leq 2\kappa_n$. Then $\inf_{k \in \mathbb{N}} M_k = \sup_{k \in \mathbb{N}} m_k$, and now we prove that the common value L is the integral: $L = \int q$. For every integer n , and each decomposition D_0 with $\delta(D_0) \leq \frac{1}{n}$, we have:

$$s(D_0) - L \leq s(D_0) - m_n \leq M_n - m_n \leq 2\kappa_n;$$

$$L - s(D_0) \leq M_n - s(D_0) \leq M_n - m_n \leq 2\kappa_n.$$

Therefore $|L - s(D_0)| \leq 2\kappa_n$, and this completes the proof, by arbitrariness of D_0 . $\square$

Our next goal is to prove that, if an interval function is integrable in $[a, b]$, then it is integrable in any subinterval, and the integral is an additive function.

Theorem 2.4 *Let $q : \mathcal{J} \rightarrow R$ be an integrable function. Then, for every subinterval $J \subset [0, 1]$, the function q_J is integrable, where q_J is defined as $q_J(I) = q(I \cap J)$, as soon as $I \cap J$ is nondegenerate, and 0 otherwise. Moreover, if $\{J_1, J_2\}$ is any decomposition of some interval $J \subset [0, 1]$, we have*

$$\int_J q = \int_{J_1} q + \int_{J_2} q. \text{ (Here, as usual, } \int_J q \text{ means the integral of } q_J \text{)}$$

Proof. Fix any interval $J \subset [0, 1]$, and, for every positive integer n , define

$$\sigma_n(J) := \sup \left\{ \left| \sum_{I \in D'_J} q(I) - \sum_{H \in D''_J} q(H) \right| : D'_J, D''_J \in \mathcal{D}_J, \delta(D'_J) \leq \frac{1}{n}, \delta(D''_J) \leq \frac{1}{n} \right\}.$$

It is easy to deduce that $\sigma_n(J) \leq \sigma_n([0, 1])$. As $(\sigma_n([0, 1]))_n \in \mathcal{N}$, the same holds for $(\sigma_n(J))_n$. Therefore, by Theorem 2.3, it follows that q_J is integrable.

To prove additivity, fix an interval $J = [a, b]$, an integer $n > 0$, and a point $c \in]a, b[$. Now, choose any decomposition D_0 of $[a, b]$, such that $\delta(D_0) \leq \frac{1}{n}$, and such that c is one of its intermediate points; next, denote by D' the decomposition consisting of the intervals from D_0 that are contained in $[a, c]$ and by D'' the decomposition consisting of the remaining intervals from D_0 . Clearly we have

$$\left| \int_J q - \int_{[a,c]} q - \int_{[c,b]} q \right| =$$

$$\left| \int_J q - \sum_{I \in D_0} q(I) + \sum_{I \in D'} q(I) - \int_{[a,c]} q + \sum_{I \in D''} q(I) - \int_{[c,b]} q \right| \leq \gamma_n,$$

where $(\gamma_n)_n$ is a suitable element from $\mathcal{N}$. By arbitrariness of n , we can infer that

$$\left| \int_J q - \int_{[a,c]} q - \int_{[c,b]} q \right| = 0, \text{ i.e. additivity of the integral. } \quad \square$$

A further consequence of the previous result is the following.

Theorem 2.5 *Let $q : \mathcal{J} \rightarrow R$ be any integrable function, and let us denote by ψ its integral function, i.e. $\psi(I) = \int_I q$, $I \in \mathcal{J}$. Then the function $|q - \psi|$ has null integral.*

Proof. Let us denote by Z the function $Z = |q - \psi|$. We must show that the sequence $(\beta_n)_n$ belongs to $\mathcal{N}$, where $\beta_n = \sup \left\{ \sum_{I \in D} Z(I) : \delta(D) \leq \frac{1}{n} \right\}$. For every $n > 0$, let us set $\sigma_n := \sup \left\{ \sum_{I \in D} OB(q)(I) : \delta(D) \leq \frac{1}{n} \right\}$. Fix now an integer $n > 0$, and a decomposition D_0 , with $\delta(D_0) \leq \frac{1}{n}$. For every element $I \in D_0$, we have $q(I) - \psi(I) = \lim_{D \in \mathcal{D}_I} (q(I) - S(D))$ (intended as $\mathcal{N}$ -limit of

a net, depending on I), and therefore

$$\sum_{I \in D_0} Z(I) = \sum_{I \in D_0} \left| \lim_{D \in \mathcal{D}_I} (q(I) - S(D)) \right| \leq \sum_{I \in D_0} OB(q)(I) \leq \sigma_n$$

according with all properties of convergence. But $(\sigma_n)_n \in \mathcal{N}$ by Theorem 2.3, hence also $(\beta_n)_n \in \mathcal{N}$, and the proof is finished. $\square$

We now introduce some structural assumptions, which will be needed later.

Assumptions 2.6 Let R_1, R_2, R be three Dedekind complete Riesz spaces.

We say that (R_1, R_2, R) is a *product triple* if there exists a map $\cdot : R_1 \times R_2 \rightarrow R$, which we will call *product*, such that

$$2.6.1) \quad (r_1 + s_1) \cdot r_2 = r_1 \cdot r_2 + s_1 \cdot r_2,$$

$$2.6.2) \quad r_1 \cdot (r_2 + s_2) = r_1 \cdot r_2 + r_1 \cdot s_2,$$

$$2.6.3) \quad [r_1 \geq s_1, r_2 \geq 0] \Rightarrow [r_1 \cdot r_2 \geq s_1 \cdot r_2],$$

$$2.6.4) \quad [r_1 \geq 0, r_2 \geq s_2] \Rightarrow [r_1 \cdot r_2 \geq r_1 \cdot s_2] \quad \text{for all } r_j, s_j \in R_j, j = 1, 2;$$

$$2.6.5) \quad \text{if } (a_\lambda)_{\lambda \in \Lambda} \text{ is any net in } R_2 \text{ and } b \in R_1, \text{ then } [a_\lambda \downarrow 0, b \geq 0] \Rightarrow [b \cdot a_\lambda \downarrow 0];$$

$$2.6.6) \quad \text{if } (a_\lambda)_{\lambda \in \Lambda} \text{ is any net in } R_1 \text{ and } b \in R_2, \text{ then } [a_\lambda \downarrow 0, b \geq 0] \Rightarrow [a_\lambda \cdot b \downarrow 0].$$

A Dedekind complete Riesz space R is called an *algebra* if (R, R, R) is a product triple.

Let now (R_1, R_2, R) be a *product triple* of Riesz spaces. Given a bounded function $g : [0, 1] \rightarrow R_2$, we can associate with g its *jump* function $\Delta(g)$:

$\Delta(g)([a, b]) := g(b) - g(a)$. Moreover, the interval function

$$\omega(g)(I) = \sup\{|\Delta(g)([u, v])| : [u, v] \subset I\}$$

is called the *oscillation* of g in the interval I .

Definition 2.7 Let $f : [0, 1] \rightarrow R_1$ and $q : \mathcal{J} \rightarrow R_2$ be two functions. In order to properly define the Riemann-Stieltjes integral of f with respect to q , we need a slight complication of the directed set $\mathcal{D}$. In fact, we shall denote by $\tilde{\mathcal{D}}$ the set of all pairs (D, τ) , where D runs in $\mathcal{D}$ and τ is a mapping which associates to every element $I \in D$ an arbitrary point $\tau_I \in I$. The set $\tilde{\mathcal{D}}$ will be ordered according with the usual *mesh* ordering on $\mathcal{D}$, i.e. $(D, \tau) \leq (D_1, \tau_1)$ iff $\delta(D_1) \leq \delta(D)$. Given an element $(D, \tau) \in \tilde{\mathcal{D}}$, we call *Riemann-Stieltjes sum* of f with respect to q (and denote it by $S(D, \tau)$) the following quantity:

$$S(D, \tau) = \sum_{I \in D} f(\tau_I) q(I).$$

Definition 2.8 Assume that $f : [0, 1] \rightarrow R_1$ and $q : \mathcal{J} \rightarrow R_2$ are two functions. We say that f is *Riemann-Stieltjes integrable* with respect to q , if the net $(S(D, \tau))_{(D, \tau) \in \tilde{\mathcal{D}}}$ is $\mathcal{N}$ -convergent, i.e. if there exists an element $Y \in R$, such that the sequence $(\pi_n)_n$ belongs to $\mathcal{N}$, where

$$\pi_n = \sup \left\{ |Y - S(D, \tau)| : (D, \tau) \in \tilde{\mathcal{D}}, \delta(D) \leq \frac{1}{n} \right\}, \quad n \in \mathbb{N}.$$

If this is the case, the element $Y \in R$ is called the *Riemann-Stieltjes integral* and denoted by $Y := (RS) \int_0^1 f dq$. When $q = \Delta(g)$, for some suitable function $g : [0, 1] \rightarrow R_2$, integrability of f with respect to q will be expressed by saying that f is Riemann-Stieltjes integrable w.r.t. g , and its integral will be denoted by $(RS) \int_0^1 f dg$.

3 The Kondurar theorem in Riesz spaces

Definition 3.1 Let $f : [0, 1] \rightarrow R_1$ and $q : \mathcal{J} \rightarrow R_2$ be two fixed functions.

We say that f and q satisfy *assumption (C)* if the interval function $\phi : \mathcal{J} \rightarrow R$, defined as $\phi(I) = \omega(f)(I)|q(I)|$, is integrable, and its integral is 0.

The next result is easy.

Proposition 3.2 Let $f : [0, 1] \rightarrow R_1$ and $q : \mathcal{J} \rightarrow R_2$ be two bounded functions, satisfying *assumption (C)*. Then the Riemann-Stieltjes integral

$(RS) \int_0^1 f \, dq$ exists in R if and only if the interval function $Q(I) = f(a_I)q(I)$ is integrable, where a_I denotes the left endpoint of I : if this is the case, the two integrals coincide.

In the next theorems, the concept of Hölder-continuous function will be crucial:

Definition 3.3 Let R be an arbitrary Riesz space. We say that an interval function q is *Hölder-continuous* (of order γ) if there exist a unit $u \in R$, that is an element $R \ni u \geq 0$, $u \neq 0$, and a real constant $\gamma > 0$, such that $|q(I)| \leq |I|^\gamma u$ for all $I \subset [0, 1]$. As usual, a function $g : [0, 1] \rightarrow R$ is said to be *Hölder-continuous* of order γ if the function $\Delta(g)$ is.

We shall assume in addition that the interval function q is *additive*, i.e. $q([a, b]) = q([a, c]) + q([c, b])$ as soon as $0 \leq a < c < b \leq 1$. This means also that $q = \Delta(g)$ for some suitable bounded function $g : [0, 1] \rightarrow R_2$, but we shall maintain our notation, without mentioning g .

The next step is the following result.

Proposition 3.4 *Assume that $f : [0, 1] \rightarrow R_1$ and $q : \mathcal{J} \rightarrow R_2$ are two Hölder-continuous functions, for which assumption (C) is satisfied. Suppose also that q is additive, and that $(\rho_n)_n \in \mathcal{N}$, where*

$$\rho_n = \sup \left\{ \left| \sum_{J \in D} f(a_J)q(J) - \sum_{I \in D_0} f(a_I)q(I) \right| : \right. \\ \left. \mathcal{D} \ni D_0, D \text{ rational, } \delta(D_0) \leq \frac{1}{n}, D \succ D_0 \right\}$$

for every $n \in \mathbb{N}$. Then, there exists in R the (RS)-integral of f w.r.t. q .

Proof. Since assumption (C) is satisfied, from Proposition 3.2 it's enough to prove that the R -valued interval function $Q(I) := f(a_I)q(I)$ is integrable. As usual, for every decomposition D of $[0, 1]$, we write $S(D) = \sum_{I \in D} Q(I) = \sum_{I \in D} f(a_I)q(I)$. By using Hölder-continuity and boundedness of f and q (which is a simple consequence), one can obtain the following *key tool*, which is rather technical, but not difficult.

Key Tool: There exists a unit $w \in R$ such that, for every $n \in \mathbb{N}$ and every $D \in \mathcal{D}$ with $\delta(D) < \frac{1}{n}$, one can find a rational decomposition D_0 such that $\delta(D_0) < \frac{1}{n}$ and

$$|S(D) - S(D_0)| \leq \frac{w}{n}. \quad (2)$$

Let us now turn to the proof that Q is integrable. We first observe that

$$\sup \left\{ |S(D_1) - S(D_2)| : D_1 \text{ and } D_2 \text{ rational, } \delta(D_1) < \frac{1}{n}, \delta(D_2) < \frac{1}{n} \right\} \leq 2\rho_n$$

for every n ; hence, using the key tool above, we see also that

$$\sup \left\{ |S(D) - S(D')| : D \in \mathcal{D}, D' \in \mathcal{D}, \delta(D) < \frac{1}{n}, \delta(D') < \frac{1}{n} \right\} \leq 2\rho_n + 2\frac{w}{n}$$

for every n . Now, the assertion follows from Theorem 2.3. $\square$

Unless otherwise specified, we always assume that q is additive, that f and q are Hölder-continuous of order α and β respectively, and related units u_1 and u_2 respectively. Moreover, we shall require that $\gamma := \alpha + \beta > 1$.

Lemma 3.5 *Let $[a, b]$ be any sub-interval of $[0, 1]$, and let us denote by c its midpoint. Then*

$$|Q([a, b]) - (Q([a, c]) + Q([c, b]))| \leq u_1 u_2 \left(\frac{b-a}{2} \right)^\gamma.$$

Proof. We have:

$$\begin{aligned} & |Q([a, b]) - (Q([a, c]) + Q([c, b]))| = |f(a)q([a, b]) - f(a)q([a, c]) \\ & - f(c)q([c, b])| = |f(a)q([c, b]) - f(c)q([c, b])| \\ & \leq |f(c) - f(a)|q([c, b])| \leq u_1 u_2 \left(\frac{b-a}{2} \right)^{\alpha+\beta}. \quad \square \end{aligned}$$

Proposition 3.6 *Let $[a, b]$ be any subinterval of $[0, 1]$, and assume that D is a dyadic decomposition of $[a, b]$, of order n . Then there exists a positive element $v \in R$, independent of a, b and n , such that $|Q([a, b]) - \sum_{I \in D} Q(I)| \leq v |b-a|^\gamma$.*

Proof. Let us consider the dyadic decompositions of $[a, b]$ of order $1, 2, \dots, n$ and denote them by $D_1, D_2, \dots, D_n = D$ respectively. Then we have

$$|Q([a, b]) - \sum_{I \in D} Q(I)| \leq |Q([a, b]) - \sum_{I \in D_1} Q(I)| + \sum_{i=2}^n \left| \sum_{J \in D_i} Q(J) - \sum_{I \in D_{i-1}} Q(I) \right|.$$

Thanks to Lemma 3.5, the first summand in the right-hand side is less than $u_1 u_2 \left(\frac{b-a}{2} \right)^\gamma$. By the same Lemma, for every index i from 2 to n , we have

$$\begin{aligned}
|\sum_{J \in D_i} Q(J) - \sum_{I \in D_{i-1}} Q(I)| &\leq u_1 u_2 2^{i-1} \left(\frac{b-a}{2^i} \right)^\gamma, \text{ hence} \\
|Q([a, b]) - \sum_{I \in D} Q(I)| &\leq u_1 u_2 \left(\frac{b-a}{2} \right)^\gamma \\
+ \sum_{i=2}^n u_1 u_2 2^{i-1} \left(\frac{b-a}{2^i} \right)^\gamma &= \frac{1}{2} u_1 u_2 (b-a)^\gamma \sum_{i=1}^n \left(\frac{1}{2^{\gamma-1}} \right)^i \\
&\leq \frac{1}{2} u_1 u_2 (b-a)^\gamma \sum_{i=1}^{\infty} \left(\frac{1}{2^{\gamma-1}} \right)^i = \frac{1}{2(2^{\gamma-1} - 1)} (b-a)^\gamma u_1 u_2.
\end{aligned}$$

The assertion follows, by choosing $v = \frac{1}{2(2^{\gamma-1} - 1)} u_1 u_2$. $\square$

Proposition 3.7 *Let $[a, b]$ be any subinterval of $[0, 1]$, and assume that D is an equidistributed decomposition of $[a, b]$, consisting of N elements. Then*

$$|Q([a, b]) - \sum_{I \in D} Q(I)| \leq r(b-a)^\gamma$$

for some unit $r \in R$, which does not depend on N , D and $[a, b]$.

Proof. Of course, if $N = 2^h$ for some integer h , the result follows from Proposition 3.6. So, let us assume that N is not a power of 2, and write N in dyadic expansion, $N = e_0 2^0 + e_1 2^1 + \dots + e_h 2^h$, where $(e_0, e_1, \dots, e_h)$ is a suitable element of $\{0, 1\}^{h+1}$, and $h = \lfloor \log_2 N \rfloor$. Clearly then, $e_h = 1$. Now let us define, for $j = 1, 2, \dots, h$:

$$t_0 := a, \quad t_{h+1} := b, \quad \text{and} \quad t_j := a + \sum_{i=1}^j e_{i-1} 2^{i-1} \frac{b-a}{N}.$$

Let us denote by D_0 the decomposition of $[a, b]$ whose intermediate points are $t_0, t_1, \dots, t_{h+1}$. From Proposition 3.6, it follows:

$$|S(D_0) - S(D)| \leq \sum_{I \in D_0} |Q(I) - \sum_{J \in D, J \subset I} Q(J)| \leq \sum_{I \in D_0} v |I|^\gamma \leq$$

$$\leq v \sum_{i=0}^h \left(2^i \frac{b-a}{N} \right)^\gamma \leq v(b-a)^\gamma \frac{2^{(h+1)\gamma}}{N^\gamma} \leq v(b-a)^\gamma 2^\gamma.$$

Thus, setting $v_1 := 2^\gamma v$, we get $|S(D_0) - S(D)| \leq v_1(b-a)^\gamma$.

Let us now evaluate $|Q([a, b]) - S(D_0)|$. We find

$$\begin{aligned} |Q([a, b]) - S(D_0)| &\leq |(f(a) - f(t_1))q([t_1, t_2]) + \\ &\quad + (f(a) - f(t_2))q([t_2, t_3]) + \dots + (f(a) - f(t_h))q([t_h, b])| \leq \\ &\leq \sum_{j=1}^h |f(a) - f(t_j)| |q([t_j, t_{j+1}])| \leq \sum_{j=1}^h u_1 u_2 (t_j - a)^\alpha (t_{j+1} - t_j)^\beta. \end{aligned}$$

For all $j = 1, 2, \dots, h$ we have easily $t_j - a \leq 2^j \frac{b-a}{N}$, hence

$$|Q([a, b]) - S(D_0)| \leq \sum_{j=1}^h u_1 u_2 \left(2^j \frac{b-a}{N} \right)^\alpha \left(2^j \frac{b-a}{N} \right)^\beta \leq u_1 u_2 2^\gamma (b-a)^\gamma.$$

Thus, setting $v_2 := u_1 u_2 2^\gamma$, we get $|Q([a, b]) - S(D_0)| \leq v_2(b-a)^\gamma$, and finally

we obtain the assertion by setting $r = v_1 + v_2$. $\square$

We are now ready for the main result.

Theorem 3.8 (The Kondurar Theorem) *Let $f : [0, 1] \rightarrow R_1$ and $q : \mathcal{J} \rightarrow R_2$ be Hölder-continuous of order α and β respectively, with $\gamma = \alpha + \beta > 1$, and suppose that q is additive. Then f is (RS) -integrable w.r.t. q .*

Proof. We first observe that the real valued interval function $W(I) := |I|^\gamma$ is integrable, and its integral is null. This immediately implies that f and q satisfy assumption (C). Let r be the unit in R given by Proposition 3.7, and set $b_n := \frac{2r}{n}$. Fix n and choose *any* rational decomposition D_0 such that $\delta(D_0) \leq \frac{1}{n}$ and $\sum_{J \in D} |J|^\gamma \leq \frac{1}{n}$ for every decomposition D , finer than D_0 . If D is any rational decomposition, finer than D_0 , then there exists an integer N

such that the equidistributed decomposition D^* , consisting of N subintervals, is finer than D . Of course, D^* is also finer than D_0 , so we have, by Proposition 3.7:

$$|S(D^*) - S(D_0)| \leq \sum_{I \in D_0} r|I|^\gamma \leq \frac{1}{n}r \text{ and } |S(D^*) - S(D)| \leq \sum_{J \in D} r|J|^\gamma \leq \frac{1}{n}r.$$

Therefore, by axioms of 2.1, we obtain:

$$(\sup \{|S(D) - S(D_0)| : D, D_0 \text{ are rational, } \delta(D_0) \leq 1/n, D \succ D_0\})_n \in \mathcal{N},$$

and the theorem is proved, thanks to Proposition 3.4. $\square$

More concretely, we have the following version of the previous theorem.

Theorem 3.9 *Let $f : [0, 1] \rightarrow R_1$ and $g : [0, 1] \rightarrow R_2$ be two Hölder-continuous functions, of order α and β respectively, and assume that $\alpha + \beta > 1$. Then f is (RS) -integrable with respect to g .*

However, we can obtain a more general result.

Theorem 3.10 *Assume that $f : [0, 1] \rightarrow R_1$ is any Hölder-continuous function of order α , and $q : \mathcal{J} \rightarrow R_2$ is any integrable interval function (not necessarily additive). If the integral function $\psi(I) = \int_I q$ is Hölder-continuous of order β , and $\alpha + \beta > 1$, then f is (RS) -integrable with respect to q and to ψ , and the two integrals coincide.*

Proof. In 2.4 we proved that ψ is an additive function; hence, if it is Hölder-continuous of order β , and $\alpha + \beta > 1$, we can deduce that f is (RS) -integrable with respect to ψ , thanks to 3.8. Moreover, we also know from 2.5 that $|q - \psi|$ has null integral. This fact, together with boundedness of f , implies

that $\int_J f d(q - \psi) = 0$, for every interval $J \subset [0, 1]$. Now the conclusion is obvious. $\square$

Remark 3.11 Assume that W is the standard Brownian Motion, on a probability space $(X, \mathcal{B}, P)$, and $g : [0, 1] \rightarrow L^0(X, \mathcal{B}, P)$ is the function which associates to each $t \in [0, 1]$ the random variable $W_t \in L^0$. We can endow the Riesz space L^0 with the convergence in probability (which satisfies all assumptions in Definition 2.1); thus the interval function $q(I) = (\Delta(g)(I))^2$ is well-known (from the literature) to be integrable, and its integral function is $\psi(I) = |I|$. So, Theorem 3.10 implies integrability with respect to $q = \Delta(g)^2$ of every function f which is Hölder continuous (of any order), and also of more general continuous functions (but we shall not discuss this here).

Remark 3.12 The previous existence theorems remain true, if the definition 2.8 is modified, replacing the value $f(\tau_I)$ with any element θ_I such that $\inf\{f(x) : x \in I\} \leq \theta_I \leq \sup\{f(x) : x \in I\}$. This mainly rests on Proposition 3.2 and on the fact that $|\theta_I - f(a_I)| \leq \omega(f)(I)$.

4 The Itô formula in Riesz spaces

In [4] we can find some results about the Itô formula for Riesz space-valued functions of one variable. Here we deal with functions of two variables; from now on we shall assume that our Riesz space R is an algebra, hence the involved functions and their integrals will take values in R .

Definition 4.1 Let R be an algebra, and $f : [0, 1] \times R \rightarrow R$ a fixed function.

We say that f satisfies *Taylor's formula* of order 1 if there exist two R -valued functions, defined on $[0, 1] \times R$ and denoted by $\frac{\partial f}{\partial t}, \frac{\partial f}{\partial x}$, such that for every $(t, x) \in [0, 1] \times R$, $h \in \mathbb{R}$ and $k \in R$, such that $t + h \in [0, 1]$, we have:

$$f(t + h, x + k) = f(t, x) + \left(h \frac{\partial f}{\partial t} + k \frac{\partial f}{\partial x} \right) (t, x) + (|h| + |k|) B(t, x, h, k),$$

where B is a suitable R -valued function, defined on all points of the type (t, x, h, k) with $t + h \in [0, 1]$, and bounded on bounded sets.

Moreover, we say that f satisfies *Taylor's formula* of order 2 if f satisfies Taylor's formula of the first order, and moreover there exist three more functions, denoted by $\frac{\partial^2 f}{\partial t^2}, \frac{\partial^2 f}{\partial t \partial x}, \frac{\partial^2 f}{\partial x^2}$, such that for every (t, x) in $[0, 1] \times R$, $h \in \mathbb{R}$ and $k \in R$, $(t + h \in [0, 1])$, we have

$$\begin{aligned} f(t + h, x + k) - f(t, x) &= h \frac{\partial f}{\partial t}(t, x) + k \frac{\partial f}{\partial x}(t, x) \\ &+ \frac{1}{2} \left[h^2 \frac{\partial^2 f}{\partial t^2}(t, x) + 2hk \frac{\partial^2 f}{\partial t \partial x}(t, x) + k^2 \frac{\partial^2 f}{\partial x^2}(t, x) \right] \\ &+ [h^2 + |hk| + |k|^2] B(t, x, h, k), \end{aligned} \quad (3)$$

where B is a suitable function, bounded on bounded sets.

We now introduce a "weaker" concept of integrability for Riesz space-valued functions. As above, $\mathcal{N}$ will always denote *any* fixed ideal as in 2.1.

Definitions 4.2 Let $F : [0, 1] \times R \rightarrow R$ and $g : [0, 1] \rightarrow R$ be two functions.

Following the Stratonovich approach ([8]), for every $\lambda \in [0, 1]$ we say that $F(\cdot, g(\cdot))$ is $\langle \lambda \rangle$ -integrable with respect to g (or to $q(I) := \Delta(g)(I)$) if there

exists an element $Y_\lambda \in R$ such that

$$\left(\sup \left\{ \left| Y_\lambda - \sum_{I \in D} F(\lambda u_i + (1-\lambda)v_i, \lambda g(u_i) + (1-\lambda)g(v_i)) \Delta(g)(I) \right| : \delta(D) \leq \frac{1}{n} \right\} \right)_n \in \mathcal{N},$$

where $D = \{([u_i, v_i]) : i = 1, \dots, N\}$ is the involved decomposition of $[0, 1]$.

We denote the $\langle \lambda \rangle$ -integral of a function F with respect to g by the symbol

$$\langle \lambda \rangle \int_0^1 F dg.$$

Similarly, we say that the function $F(\cdot, g(\cdot))$ is $\langle \lambda \rangle$ -integrable with respect to

$q(I) := (\Delta(g)(I))^2$ if there exists an element $J_\lambda \in R$ such that

$$\left(\sup \left\{ \left| J_\lambda - \sum_{I \in D} F(\lambda u_i + (1-\lambda)v_i, \lambda g(u_i) + (1-\lambda)g(v_i)) (\Delta(g)(I))^2 \right| : \delta(D) \leq \frac{1}{n} \right\} \right)_n \in \mathcal{N},$$

where $D = \{([u_i, v_i]) : i = 1, \dots, N\}$ is the involved decomposition of $[0, 1]$.

In this case, we write $\langle \lambda \rangle \int_0^1 F(s, g(s)) (dg)^2(s) := J_\lambda(F)$. We now prove the

following:

Proposition 4.3 *Let R be an algebra, and $f : [0, 1] \times R \rightarrow R$, $(t, x) \mapsto f(t, x)$, satisfy Taylor's formula of order 1; suppose also that $g : [0, 1] \rightarrow R$ is Hölder-continuous of order β , where $\beta > \frac{1}{3}$. If f is $\langle 1 \rangle$ -integrable with respect to $q(I) := (\Delta(g)(I))^2$, i.e. there exists the integral $J_1(f)$, then for every $\lambda \in [0, 1]$ the integral $J_\lambda(f)$ exists in R , and $J_\lambda(f) = J_1(f)$.*

Proof. Let $D = \{[u_i, v_i], i = 1, \dots, n\}$ be a decomposition of the interval $[0, 1]$

and $\lambda \in [0, 1]$. By hypotheses, we have:

$$\sum_{i=1}^n f(\lambda u_i + (1-\lambda)v_i, \lambda g(u_i) + (1-\lambda)g(v_i)) \cdot q([u_i, v_i]) = \sum_{i=1}^n f(u_i, g(u_i)) \cdot$$

$$q([u_i, v_i]) +$$

$$+ \sum_{i=1}^n [f(\lambda u_i + (1-\lambda)v_i, \lambda g(u_i) + (1-\lambda)g(v_i)) - f(u_i, g(u_i))] \cdot q([u_i, v_i]) =$$

$$\sum_{i=1}^n \left[f(u_i, g(u_i)) + (1-\lambda)(v_i - u_i) \frac{\partial f}{\partial t}(u_i, g(u_i)) + \right.$$

$$\left. (1-\lambda)(\Delta(g)([u_i, v_i])) \frac{\partial f}{\partial x}(u_i, g(u_i)) \right] \cdot q([u_i, v_i]) + \sum_{i=1}^n B([u_i, v_i]) |v_i - u_i|^\tau =$$

$$V_1 + V_2 + V_3 + V_4, \text{ where } B : \{I\} \rightarrow R \text{ is a suitable interval function, bounded}$$

on bounded sets (with bound independent on λ), and $\tau > 1$ by virtue of Hölder-continuity of g (we denote the consecutive terms of the above sum by V_1, V_2, V_3, V_4). The first part V_1 tends to $J_1(f)$ as $\delta(D)$ tends to 0. Now we can conclude, observing that the expressions V_2, V_3, V_4 are negligible because we can present them in the form $\sum_{i=1}^n B_0([u_i, v_i]) |v_i - u_i|^\zeta$, where $B_0 : \mathcal{J} \rightarrow R$ is bounded on bounded sets and $\zeta > 1$. $\square$

We now prove our generalization of the Itô formula in the context of Riesz spaces:

Theorem 4.4 *Let R be an algebra, $F : [0, 1] \times R \rightarrow R$ satisfy Taylor's formula of order 2 and $\frac{\partial F}{\partial t} : [0, 1] \times R \rightarrow R$, $\frac{\partial^2 F}{\partial x^2} : [0, 1] \times R \rightarrow R$, satisfy Taylor's formula of order 1. Assume that $g : [0, 1] \rightarrow R$ is Hölder-continuous of order $\beta > \frac{1}{3}$. Assume also that $q(I) = (\Delta(g)(I))^2$ is integrable and that its integral function is Hölder-continuous of order $\gamma > 1 - \beta$. Then*

- 1) *for every $\lambda \in [0, 1]$ the integral $J_\lambda := \langle \lambda \rangle \int_0^1 \frac{\partial^2 F}{\partial x^2}(s, g(s)) (dg)^2(s)$ exists in R , and is independent on λ ;*
- 2) *for every $\lambda \in [0, 1]$ the integral $\langle \lambda \rangle \int_0^1 \frac{\partial F}{\partial x}(s, g(s)) dg(s)$ exists in R ;*

3) the following formula holds:

$$\begin{aligned} & F(1, g(1)) - F(0, g(0)) = \\ &= \int_0^1 \frac{\partial F}{\partial t}(s, g(s)) ds + \langle \lambda \rangle \int_0^1 \frac{\partial F}{\partial x}(s, g(s)) dg(s) + \frac{1}{2} (2\lambda - 1) J, \end{aligned}$$

where J denotes any of the integrals J_λ above.

Proof. The function $t \mapsto \frac{\partial^2 F}{\partial x^2}(t, g(t))$ is Hölder-continuous of order β , because of Taylor's formula and of Hölder-continuity of g . Thus, using Theorem 3.10, we deduce that the function $t \mapsto \frac{\partial^2 F}{\partial x^2}(t, g(t))$ is (RS) -integrable with respect to g , and thus by virtue of Proposition 4.3 (where $f = \frac{\partial^2 F}{\partial x^2}$), for every $\lambda \in [0, 1]$ the integral $J := \langle \lambda \rangle \int_0^1 \frac{\partial^2 F}{\partial x^2}(s, g(s)) (dg)^2(s)$ exists in R and is independent on λ . This shows the assertion 1).

Let now $D = \{[u_i, v_i], i = 1, \dots, n\}$ be a decomposition of $[0, 1]$ and fix $\lambda \in [0, 1]$. We have:

$$\begin{aligned} & F(1, g(1)) - F(0, g(0)) = \\ & \sum_{i=1}^n [F(v_i, g(v_i)) - F(\lambda u_i + (1 - \lambda) v_i, \lambda g(u_i) + (1 - \lambda) g(v_i))] - \\ & \sum_{i=1}^n [F(u_i, g(u_i)) - F(\lambda u_i + (1 - \lambda) v_i, \lambda g(u_i) + (1 - \lambda) g(v_i))]. \end{aligned}$$

Now we shall apply Taylor's formula of order 2:

$$\begin{aligned} & F(1, g(1)) - F(0, g(0)) = \\ & \left\{ \sum_{i=1}^n \frac{\partial F}{\partial t}(\lambda u_i + (1 - \lambda) v_i, \lambda g(u_i) + (1 - \lambda) g(v_i)) \cdot [\lambda (v_i - u_i)] + \right. \\ & \sum_{i=1}^n \frac{\partial F}{\partial x}(\lambda u_i + (1 - \lambda) v_i, \lambda g(u_i) + (1 - \lambda) g(v_i)) \cdot [\lambda (\Delta(g)([u_i, v_i]))] + \\ & \sum_{i=1}^n \frac{\partial^2 F}{\partial t \partial x}(\lambda u_i + (1 - \lambda) v_i, \lambda g(u_i) + (1 - \lambda) g(v_i)) \\ & \cdot [\lambda^2 (v_i - u_i) (\Delta(g)([u_i, v_i]))] + \\ & \left. \sum_{i=1}^n \frac{\partial^2 F}{\partial t^2}(\lambda u_i + (1 - \lambda) v_i, \lambda g(u_i) + (1 - \lambda) g(v_i)) \cdot \left[\frac{1}{2} \lambda^2 (v_i - u_i)^2 \right] + \right\} \end{aligned}$$

$$\begin{aligned}
& \sum_{i=1}^n \frac{\partial^2 F}{\partial x^2} (\lambda u_i + (1-\lambda) v_i, \lambda g(u_i) + (1-\lambda) g(v_i)) \cdot \left[\frac{1}{2} \lambda^2 q([u_i, v_i]) \right] + \\
& \sum_{i=1}^n B_1([u_i, v_i]) |v_i - u_i|^\gamma \} - \\
& \{ \sum_{i=1}^n \frac{\partial F}{\partial t} (\lambda u_i + (1-\lambda) v_i, \lambda g(u_i) + (1-\lambda) g(v_i)) \cdot [(\lambda - 1)(v_i - u_i)] + \\
& \sum_{i=1}^n \frac{\partial F}{\partial x} (\lambda u_i + (1-\lambda) v_i, \lambda g(u_i) + (1-\lambda) g(v_i)) \cdot [(\lambda - 1)(\Delta(g)([u_i, v_i]))] + \\
& \sum_{i=1}^n \frac{\partial^2 F}{\partial t \partial x} (\lambda u_i + (1-\lambda) v_i, \lambda g(u_i) + (1-\lambda) g(v_i)) \\
& \cdot [(1-\lambda)^2 (v_i - u_i)(\Delta(g)([u_i, v_i]))] + \\
& \sum_{i=1}^n \frac{\partial^2 F}{\partial t^2} (\lambda u_i + (1-\lambda) v_i, \lambda g(u_i) + (1-\lambda) g(v_i)) \cdot \left[\frac{1}{2} (1-\lambda)^2 (v_i - u_i)^2 \right] + \\
& \sum_{i=1}^n \frac{\partial^2 F}{\partial x^2} (\lambda u_i + (1-\lambda) v_i, \lambda g(u_i) + (1-\lambda) g(v_i)) \cdot \left[\frac{1}{2} (1-\lambda)^2 q([u_i, v_i]) \right] + \\
& \sum_{i=1}^n B_2([u_i, v_i]) |v_i - u_i|^\gamma \},
\end{aligned}$$

where $B_1, B_2 : \{I\} \rightarrow R$ are suitable interval functions, bounded on bounded sets (with bound independent on λ), and $\gamma > 1$. By collecting the similar terms we obtain:

$$\begin{aligned}
& F(1, g(1)) - F(0, g(0)) = \\
& \sum_{i=1}^n \frac{\partial F}{\partial t} (\lambda u_i + (1-\lambda) v_i, \lambda g(u_i) + (1-\lambda) g(v_i)) \cdot (v_i - u_i) + \\
& \sum_{i=1}^n \frac{\partial F}{\partial x} (\lambda u_i + (1-\lambda) v_i, \lambda g(u_i) + (1-\lambda) g(v_i)) \cdot (\Delta(g)([u_i, v_i])) + \\
& \sum_{i=1}^n \frac{\partial^2 F}{\partial t \partial x} (\lambda u_i + (1-\lambda) v_i, \lambda g(u_i) + (1-\lambda) g(v_i)) \\
& \cdot [(2\lambda - 1)(v_i - u_i)(\Delta(g)([u_i, v_i]))] + \\
& \sum_{i=1}^n \frac{\partial^2 F}{\partial t^2} (\lambda u_i + (1-\lambda) v_i, \lambda g(u_i) + (1-\lambda) g(v_i)) \cdot \left[\frac{1}{2} (2\lambda - 1)(v_i - u_i)^2 \right] + \\
& \sum_{i=1}^n \frac{\partial^2 F}{\partial x^2} (\lambda u_i + (1-\lambda) v_i, \lambda g(u_i) + (1-\lambda) g(v_i)) \cdot \left[\frac{1}{2} (2\lambda - 1) q([u_i, v_i]) \right] + \\
& \sum_{i=1}^n (B_1([u_i, v_i]) - B_2([u_i, v_i])) |v_i - u_i|^\gamma =
\end{aligned}$$

$S_1 + S_2 + S_3 + S_4 + S_5 + S_6$ (where $S_1, S_2, \dots, S_6$ denote the consecutive terms of the above sum, and the expressions S_3, S_4, S_6 are negligible, as previously observed).

Moreover, since $\frac{\partial F}{\partial t}$ satisfies Taylor's formula of order 1, we have:

$$S_1 = \sum_{i=1}^n \frac{\partial F}{\partial t}(u_i, g(u_i)) \cdot (v_i - u_i) + W, \text{ where } W \text{ as usual is negligible.}$$

The first summand, and hence S_1 , tends to $\int_0^1 \frac{\partial F}{\partial t}(s, g(s)) ds$ as n tends to $+\infty$ (the function $\frac{\partial F}{\partial t}(\cdot, g(\cdot))$ is Hölder-continuous, so $\frac{\partial F}{\partial t}(\cdot, g(\cdot))$ is (RS) -integrable w.r.t. dt). Now,

$$S_5 = \frac{1}{2} (2\lambda - 1) \sum_{i=1}^n \frac{\partial^2 F}{\partial x^2}(\lambda u_i + (1 - \lambda) v_i, \lambda g(u_i) + (1 - \lambda) g(v_i)) q([u_i, v_i])$$

tends to $\frac{1}{2} (2\lambda - 1) J$, as already observed. Hence it follows that

$$S_2 = \sum_{i=1}^n \frac{\partial F}{\partial x}(\lambda u_i + (1 - \lambda) v_i, \lambda g(u_i) + (1 - \lambda) g(v_i)) \cdot (\Delta(g)([u_i, v_i])) \quad (4)$$

converges, and its limit is the $\langle \lambda \rangle$ -integral of $\frac{\partial F}{\partial x}(\cdot, g(\cdot))$ with respect to g , that is the assertion 2). The formula in assertion 3) follows now easily. $\square$

Remark 4.5

I) Sometimes, the $\langle \lambda \rangle$ -integral exists, without assumptions on the function g .

For example, let $F(t, x) = x^2$, and $\lambda = \frac{1}{2}$: so we are looking for the integral $\langle \frac{1}{2} \rangle \int_0^1 2g(s)dg(s)$. By definition, this integral is the limit of

$$2 \sum_{I \in D} [g(u_{i+1}) - g(u_i)] \left[\frac{g(u_{i+1}) + g(u_i)}{2} \right] \quad (5)$$

as $\delta(D) \rightarrow 0$, where $D = \{[u_i, u_{i+1}] : i = 0, \dots, n\}$. But the quantity (5) always coincides with $g^2(1) - g^2(0)$, so we obtain $\langle \frac{1}{2} \rangle \int_0^1 2g(s)dg(s) = g^2(1) - g^2(0)$.

II) On the other hand, once we have a function g satisfying the hypotheses of Theorem 4.4, at least in particular spaces R it is possible to find many other functions $\tilde{g}$ with the same properties. For example, let $R = L^0(X, \mathcal{B}, P)$ be as in the Remark 3.11 and assume that $g : [0, 1] \rightarrow R$ is the standard Brownian

Motion. Now, choose any C^2 -function $\phi : \mathbb{R} \rightarrow \mathbb{R}$ and define, for every $t \in [0, 1]$: $\tilde{g}(t)(x) = \phi(g(t)(x))$ for almost all $x \in X$. We shall see that $\Delta(\tilde{g})^2$ is integrable, and its integral function is Lipschitz. Indeed, if $0 \leq u < v \leq 1$, we get (a.e.)

$$(\tilde{g}(v) - \tilde{g}(u))(x) = \phi(g(v)(x)) - \phi(g(u)(x)) = (g(v)(x) - g(u)(x))\phi'(g(\tau)(x)),$$

where τ is a suitable point in $[u, v]$, depending also on x . So we have

$$[\tilde{g}(v) - \tilde{g}(u)]^2(x) = [g(v)(x) - g(u)(x)]^2 (\phi'(g(\tau)(x)))^2. \quad (6)$$

(We remark here that $(\phi'(g(\tau)(x)))^2$ is an element of R , between the extrema of the function $t \mapsto f(t) = \phi'(g(t))^2$ in the interval $[u, v]$). Now consider the function $t \mapsto f(t) = \phi'(g(t))^2$, for all $t \in [0, 1]$. This is a Hölder-continuous function, hence Riemann-Stieltjes integrable with respect to $q = \Delta(g)^2$ by Theorem 3.10. From this, (6) and the Remark 3.12, we deduce that $(\Delta(\tilde{g})(I))^2$ is integrable, and its integral function coincides with the Riemann integral $\int_I \phi'(g(t))^2 dt$ (which can be computed pathwise), and this is clearly a Lipschitz function of interval.

References

- [1] A. BOCCUTO, *Abstract Integration in Riesz Spaces*, Tatra Mountains Math. Publ. **5** (1995), 107-124.
- [2] A. BOCCUTO, *Differential and Integral Calculus in Riesz Spaces*, Tatra Mountains Math. Publ. **14** (1998), 293-323.
- [3] A. BOCCUTO, *A Perron-type integral of order 2 for Riesz spaces*, Math. Slovaca **51** (2001), 185-204.

- [4] D. CANDELOORO, *Riemann-Stieltjes Integration in Riesz Spaces*, Rend. Mat. (Roma) **16** (1996), 563-585.
- [5] V. KONDURAR, *Sur l'intégrale de Stieltjes*, Mat. Sbornik **2** (1937), 361-366.
- [6] W. A. J. LUXEMBURG & A. C. ZAAANEN, *Riesz Spaces, I*, (1971), North-Holland Publ. Co.
- [7] Z. SCHUSS, *Theory and Applications of Stochastic Differential Equations*, (1980), John Wiley & Sons, New York.
- [8] R. L. STRATONOVICH, *A new representation for stochastic integrals and equations*, Siam J. Control **4** (1966), 362-371.
- [9] B. Z. VULIKH, *Introduction to the theory of partially ordered spaces*, (1967), Wolters - Noordhoff Sci. Publ., Groningen.
- [10] A. C. ZAAANEN, *Riesz spaces, II*, (1983), North-Holland Publishing Co.

Bounds On Triangular Discrimination, Harmonic Mean and Symmetric Chi-square Divergences

Inder Jeet Taneja

Departamento de Matemática
Universidade Federal de Santa Catarina
88.040-900 Florianópolis, SC, Brazil.
e-mail: taneja@mtm.ufsc.br
http: <http://www.mtm.ufsc.br/~taneja>

Abstract

There are many information and divergence measures exist in the literature on information theory and statistics. The most famous among them are Kullback-Leiber [15] *relative information* and Jeffreys [14] *J-divergence*. The measures like *Bhattacharya distance*, *Hellinger discrimination*, *Chi-square divergence*, *triangular discrimination* and *harmonic mean divergence* are also famous in the literature on statistics. In this paper we have obtained bounds on *triangular discrimination* and *symmetric chi-square divergence* in terms of *relative information of type s* using Csiszár's *f-divergence*. A relationship among *triangular discrimination* and *harmonic mean divergence* is also given.

Key words: *Relative information of type s; Harmonic mean divergence; Triangular discrimination; Symmetric Chi-square divergence; Csiszár's f-divergence; Information inequalities.*

1 Introduction

Let

$$\Gamma_n = \left\{ P = (p_1, p_2, \dots, p_n) \left| p_i > 0, \sum_{i=1}^n p_i = 1 \right. \right\}, \quad n \geq 2,$$

be the set of all complete finite discrete probability distributions. For all $P, Q \in \Gamma_n$, the following measures are well known in the literature on information theory and statistics:

- **Bhattacharya Distance** (Bhattacharya [2])

$$B(P||Q) = \sum_{i=1}^n \sqrt{p_i q_i}. \quad (1)$$

- **Hellinger discrimination** (Hellinger [13])

$$h(P||Q) = 1 - B(P||Q) = \frac{1}{2} \sum_{i=1}^n (\sqrt{p_i} - \sqrt{q_i})^2. \quad (2)$$

- χ^2 -**Divergence** (Pearson [18])

$$\chi^2(P||Q) = \sum_{i=1}^n \frac{(p_i - q_i)^2}{q_i} = \sum_{i=1}^n \frac{p_i^2}{q_i} - 1. \quad (3)$$

- **Relative Information** (Kullback and Leibler [15])

$$K(P||Q) = \sum_{i=1}^n p_i \ln\left(\frac{p_i}{q_i}\right). \quad (4)$$

The above four measures can be obtained as particular or limiting case of the *relative information of type s*. This measure is given by

- **Relative Information of Type s**

$$\Phi_s(P||Q) = \begin{cases} K_s(P||Q) = [s(s-1)]^{-1} \left[\sum_{i=1}^n p_i^s q_i^{1-s} - 1 \right], & s \neq 0, 1 \\ K(Q||P) = \sum_{i=1}^n q_i \ln\left(\frac{q_i}{p_i}\right), & s = 0 \\ K(P||Q) = \sum_{i=1}^n p_i \ln\left(\frac{p_i}{q_i}\right), & s = 1 \end{cases}. \quad (5)$$

The measure (5) admits the following interesting particular cases:

- (i) $\Phi_{-1}(P||Q) = \frac{1}{2} \chi^2(Q||P)$.
- (ii) $\Phi_0(P||Q) = K(Q||P)$.
- (iii) $\Phi_{1/2}(P||Q) = 4 [1 - B(P||Q)] = 4h(P||Q)$.
- (iv) $\Phi_1(P||Q) = K(P||Q)$.

$$(v) \quad \Phi_2(P||Q) = \frac{1}{2}\chi^2(P||Q).$$

Thus we observe that $\Phi_2(P||Q) = \Phi_{-1}(Q||P)$ and $\Phi_1(P||Q) = \Phi_0(Q||P)$.

For more studies on the measure (5) refer to Liese and Vajda [16], Vajda [27], Taneja [19], [20], [22] and Cerone et al. [4].

Recently Taneja [23] and Taneja and Kumar [25] studied the (5) and obtained bounds in terms of the measures (1)-(4). Here we shall extend the our study for the other measures known in the literature as *triangular discrimination*, *harmonic mean divergence* and *symmetric chi-square divergence*.

The *triangular discrimination* is given by

$$\Delta(P||Q) = \sum_{i=1}^n \frac{(p_i - q_i)^2}{p_i + q_i}. \quad (6)$$

After simplification, we can write

$$\Delta(P||Q) = 2 [1 - W(P||Q)], \quad (7)$$

where

$$W(P||Q) = \sum_{i=1}^n \frac{2p_i q_i}{p_i + q_i}, \quad (8)$$

is the well known *harmonic mean divergence*.

We observe that the measures (3) and (4) are not symmetric with respect to probability distributions. The symmetric version of the measure (4) famous as Jeffreys-Kullback-Leiber *J-divergence* is given by

$$J(P||Q) = K(P||Q) + K(Q||P). \quad (9)$$

Let us consider the *symmetric chi-square divergence* given by

$$\Psi(P||Q) = \chi^2(P||Q) + \chi^2(Q||P) = \sum_{i=1}^n \frac{(p_i - q_i)^2 (p_i + q_i)}{p_i q_i}. \quad (10)$$

Dragomir [12] studied the measure (10) and obtained interesting result relating it to *triangular discrimination* and *J-divergence*.

Some studies on the measures (6) and (8) can be seen in Dragomir [10], [11] and Topsøe [26]. Recently, Taneja [23] and Taneja and Kumar [25] studied the measure (5) and obtained bounds in terms of the measures (1)-(4). Similar kind of bounds on the measure (9) are recently obtained by Taneja [24]. In this paper, we shall extend the our study for *triangular discrimination* and *symmetric chi-square divergence*. In order to obtain bounds on these measures we make use of Csiszár's [5] *f-divergence*.

2 Csiszár's f -Divergence and Its Particular Cases

Given a convex function $f : [0, \infty) \rightarrow \mathbb{R}$, the f -divergence measure introduced by Csiszár [5] is given by

$$C_f(P||Q) = \sum_{i=1}^n q_i f\left(\frac{p_i}{q_i}\right), \quad (11)$$

where $P, Q \in \Gamma_n$.

It is well known in the literature [5] that *if f is convex and normalized, i.e., $f(1) = 0$, then the Csiszár's function $C_f(P||Q)$ is nonnegative and convex in the pair of probability distribution $(P, Q) \in \Gamma_n \times \Gamma_n$.*

Here below we shall give the measures (6) and (10) being examples of the measure (11).

Example 2.1. (*Triangular discrimination*). Let us consider

$$f_{\Delta}(x) = \frac{(x-1)^2}{x+1}, \quad x \in (0, \infty) \quad (12)$$

in (11), we have

$$C_f(P||Q) = \Delta(P||Q) = \sum_{i=1}^n \frac{(p_i - q_i)^2}{p_i + q_i},$$

where $\Delta(P||Q)$ is as given by (6).

Moreover,

$$f'_{\Delta}(x) = \frac{(x-1)(x+3)}{(x+1)^2} \quad (13)$$

and

$$f''_{\Delta}(x) = \frac{8}{(x+1)^3}. \quad (14)$$

Thus we have $f''_{\Delta}(x) > 0$ for all $x > 0$, and hence, $f_{\Delta}(x)$ is *strictly convex* for all $x > 0$. Also, we have $f_{\Delta}(1) = 0$. In view of this we can say that the *triangular discrimination* is *nonnegative* and *convex* in the pair of probability distributions $(P, Q) \in \Gamma_n \times \Gamma_n$.

Example 2.2. (*Symmetric chi-square divergence*). Let us consider

$$f_{\Psi}(x) = \frac{(x-1)^2(x+1)}{x}, \quad x \in (0, \infty) \quad (15)$$

in (11), we have

$$C_f(P||Q) = \Psi(P||Q) = \sum_{i=1}^n \frac{(p_i - q_i)^2(p_i + q_i)}{p_i q_i},$$

where $\Psi(P||Q)$ is as given by (10).

Moreover,

$$f'_{\Psi}(x) = \frac{(x-1)(2x^2+x+1)}{x^2} \quad (16)$$

and

$$f''_{\Psi}(x) = \frac{2(x^3+1)}{x^3}. \quad (17)$$

Thus we have $f''_{\Psi}(x) > 0$ for all $x > 0$, and hence, $f_{\Psi}(x)$ is *strictly convex* for all $x > 0$. Also, we have $f_{\Psi}(1) = 0$. In view of this we can say that the *symmetric chi-square divergence* is *nonnegative* and *convex* in the pair of probability distributions $(P, Q) \in \Gamma_n \times \Gamma_n$.

3 Csiszár's f -Divergence and Relative Information of Type s

During past years Dragomir done a lot of work giving bounds on Csiszár's f -divergence. Here below we shall summarize the some his results [6], [7], [9].

Theorem 3.1. *Let $f : \mathbb{R}_+ \rightarrow [0, \infty)$ be differentiable convex and normalized i.e., $f(1) = 0$. If $P, Q \in \Gamma_n$, then we have*

$$0 \leq C_f(P||Q) \leq \rho_{C_f}(P||Q), \quad (18)$$

where $\rho_{C_f}(P||Q)$ is given by

$$\rho_{C_f}(P||Q) = C_{f'}\left(\frac{P^2}{Q}||P\right) - C_{f'}(P||Q) = \sum_{i=1}^n (p_i - q_i) f'\left(\frac{p_i}{q_i}\right). \quad (19)$$

If $P, Q \in \Gamma_n$ are such that

$$0 < r \leq \frac{p_i}{q_i} \leq R < \infty, \quad \forall i \in \{1, 2, \dots, n\},$$

for some r and R with $0 < r \leq 1 \leq R < \infty$, then we have the following inequalities:

$$0 \leq C_f(P||Q) \leq \alpha_{C_f}(r, R), \quad (20)$$

$$0 \leq C_f(P||Q) \leq \beta_{C_f}(r, R) \quad (21)$$

and

$$\begin{aligned} 0 &\leq \beta_{C_f}(r, R) - C_f(P||Q) \\ &\leq \gamma_{C_f}(r, R) [(R-1)(1-r) - \chi^2(P||Q)] \leq \alpha_{C_f}(r, R), \end{aligned} \quad (22)$$

where

$$\alpha_{C_f}(r, R) = \frac{1}{4}(R-r)^2 \gamma_{C_f}(r, R), \quad (23)$$

$$\beta_{C_f}(r, R) = \frac{(R-1)f(r) + (1-r)f(R)}{R-r} \quad (24)$$

and

$$\gamma_{C_f}(r, R) = \frac{f'(R) - f'(r)}{R-r}. \quad (25)$$

The following proposition is due to Taneja [23] and Taneja and Kumar [25] and is a consequence of the above theorem.

Proposition 3.1. *Let $P, Q \in \Gamma_n$ and $s \in \mathbb{R}$, then we have*

$$0 \leq \Phi_s(P||Q) \leq \rho_{\Phi_s}(P||Q), \quad (26)$$

where

$$\begin{aligned} \rho_{\Phi_s}(P||Q) &= C_{\phi'_s} \left(\frac{P^2}{Q} || P \right) - C_{\phi'_s}(P||Q) \\ &= \begin{cases} (s-1)^{-1} \sum_{i=1}^n (p_i - q_i) \left(\frac{p_i}{q_i} \right)^{s-1}, & s \neq 1 \\ \sum_{i=1}^n (p_i - q_i) \ln \left(\frac{p_i}{q_i} \right), & s = 1 \end{cases}. \end{aligned} \quad (27)$$

If there exists r, R ($0 < r \leq 1 \leq R < \infty$) such that

$$0 < r \leq \frac{p_i}{q_i} \leq R < \infty, \quad \forall i \in \{1, 2, \dots, n\},$$

then we have the following inequalities

$$0 \leq \Phi_s(P||Q) \leq \alpha_{\Phi_s}(r, R), \quad (28)$$

$$0 \leq \Phi_s(P||Q) \leq \beta_{\Phi_s}(r, R) \quad (29)$$

and

$$0 \leq \beta_{\Phi_s}(r, R) - \Phi_s(P||Q) \quad (30)$$

$$\leq \gamma_{\Phi_s}(r, R) [(R-1)(1-r) - \chi^2(P||Q)] \leq \alpha_{\Phi_s}(r, R),$$

where

$$\alpha_{\Phi_s}(r, R) = \frac{1}{4}(R-r)^2 \gamma_{\Phi_s}(r, R), \quad (31)$$

$$\beta_{\Phi_s}(r, R) = \begin{cases} \frac{(R-1)(r^s-1)+(1-r)(R^s-1)}{(R-r)s(s-1)}, & s \neq 0, 1 \\ \frac{(R-1) \ln \frac{1}{r} + (1-r) \ln \frac{1}{R}}{(R-r)}, & s = 0 \\ \frac{(R-1)r \ln r + (1-r)R \ln R}{(R-r)}, & s = 1 \end{cases} \quad (32)$$

and

$$\gamma_{\Phi_s}(r, R) = \begin{cases} \frac{R^{s-1} - r^{s-1}}{(R-r)(s-1)}, & s \neq 1 \\ \frac{\ln R - \ln r}{R-r}, & s = 1 \end{cases}, \quad (33)$$

We can also write $\gamma_{\Phi_s}(r, R)$ as follows

$$\gamma_{\Phi_s}(r, R) = \begin{cases} L_{s-2}^{s-2}(r, R), & s \neq 1 \\ L_{-1}^{-1}(r, R) & s = 1 \end{cases}, \quad (34)$$

where $L_p(a, b)$ is the famous (Bullen, Mitrinović and Vasić [3]) *p-logarithmic power mean* given by

$$L_p(a, b) = \begin{cases} \left[\frac{b^{p+1} - a^{p+1}}{(p+1)(b-a)} \right]^{\frac{1}{p}}, & p \neq -1, 0 \\ \frac{b-a}{\ln b - \ln a}, & p = -1 \\ \frac{1}{e} \left[\frac{b^b}{a^a} \right]^{\frac{1}{b-a}}, & p = 0 \end{cases}, \quad (35)$$

for all $p \in \mathbb{R}$, $a \neq b$.

The expression (27) admits the following particular cases:

- (i) $\rho_{\Phi_{-1}}(P||Q) = 3\Phi_3(Q||P) - \frac{1}{2}\chi^2(Q||P)$.
- (ii) $\rho_{\Phi_0}(P||Q) = \chi^2(Q||P)$.
- (iii) $\rho_{\Phi_1}(P||Q) = J(P||Q)$.
- (iv) $\rho_{\Phi_{1/2}}(P||Q) = 2 \sum_{i=1}^n (q_i - p_i) \sqrt{\frac{q_i}{p_i}}$
- (v) $\rho_{\Phi_2}(P||Q) = \chi^2(P||Q)$

The expression (32) admits the following particular cases:

- (i) $\beta_{\Phi_{-1}}(P||Q) = \frac{(R-1)(1-r)}{2rR}$.
- (ii) $\beta_{\Phi_0}(r, R) = \frac{(R-1) \ln \frac{1}{r} + (1-r) \ln \frac{1}{R}}{R-r}$.
- (iii) $\beta_{\Phi_1}(r, R) = \frac{(R-1)r \ln r + (1-r)R \ln R}{R-r}$.

$$(iv) \beta_{\Phi_{1/2}}(P||Q) = \frac{4(\sqrt{R}-1)(1-\sqrt{r})}{\sqrt{R}+\sqrt{r}}.$$

$$(v) \beta_{\Phi_2}(P||Q) = \frac{(R-1)(1-r)}{2}.$$

The following theorem is due to Taneja [23] and Taneja and Kumar [25].

Theorem 3.2. *Let $f : I \subset \mathbb{R}_+ \rightarrow [0, \infty)$ the generating mapping is normalized, i.e., $f(1) = 0$ and satisfy the assumptions:*

- (i) *f is twice differentiable on (r, R) , where $0 \leq r \leq 1 \leq R \leq \infty$;*
- (ii) *there exists real constants m, M such that $m < M$ and*

$$m \leq x^{2-s} f''(x) \leq M, \quad \forall x \in (r, R), \quad s \in \mathbb{R}. \quad (36)$$

If $P, Q \in \Gamma_n$ are discrete probability distributions satisfying the assumption

$$0 < r \leq \frac{p_i}{q_i} \leq R < \infty,$$

then we have the inequalities:

$$m\Phi_s(P||Q) \leq C_f(P||Q) \leq M\Phi_s(P||Q), \quad (37)$$

$$\begin{aligned} m [\rho_{\Phi_s}(P||Q) - \Phi_s(P||Q)] \\ \leq \rho_{C_f}(P||Q) - C_f(P||Q) \\ \leq M [\rho_{\Phi_s}(P||Q) - \Phi_s(P||Q)] \end{aligned} \quad (38)$$

and

$$\begin{aligned} m [\beta_{\Phi_s}(r, R) - \Phi_s(P||Q)] \\ \leq \beta_{C_f}(r, R) - C_f(P||Q) \\ \leq M [\beta_{\Phi_s}(r, R) - \Phi_s(P||Q)], \end{aligned} \quad (39)$$

where $C_f(P||Q)$, $\Phi_s(P||Q)$, $\rho_{C_f}(P||Q)$, $\rho_{\Phi_s}(P||Q)$, $\beta_{C_f}(r, R)$ and $\beta_{\Phi_s}(r, R)$ are as given by (11), (5), (19), (27), (24) and (32) respectively.

The above theorem unifies some of the results studied by Dragomir [8], [10], [11].

In the papers Taneja [23] and Taneja and Kumar [25] considered the particular values of s and Φ_s by taking $s = -1$, $s = 0$, $s = \frac{1}{2}$, $s = 1$ and $s = 2$. The aim here is to obtain results by taking different values of f given by examples 2.1-2.2, and then obtain particular cases for different values of s .

Remark 3.1. *If it is not specified, from now onwards, it is understood that, if there are r, R then $0 < r \leq \frac{p_i}{q_i} \leq R < \infty$, $\forall i \in \{1, 2, \dots, n\}$, with $0 < r \leq 1 \leq R < \infty$ where $P = (p_1, p_2, \dots, p_n) \in \Gamma_n$ and $Q = (q_1, q_2, \dots, q_n) \in \Gamma_n$.*

4 Triangular Discrimination and Inequalities

In this section, we shall give bounds on *triangular discrimination* based on the Theorems 3.1 and 3.2.

Theorem 4.1. *For all $P, Q \in \Gamma_n$, we have the following inequalities*

$$0 \leq \Delta(P||Q) \leq \rho_\Delta(P||Q), \quad (40)$$

where

$$\rho_\Delta(P||Q) = \sum_{i=1}^n \left(\frac{p_i - q_i}{p_i + q_i} \right)^2 (p_i + 3q_i). \quad (41)$$

If there exists r, R ($0 < r \leq 1 \leq R < \infty$) such that

$$0 < r \leq \frac{p_i}{q_i} \leq R < \infty, \quad \forall i \in \{1, 2, \dots, n\},$$

then we have the following inequalities:

$$0 \leq \Delta(P||Q) \leq \alpha_\Delta(r, R), \quad (42)$$

$$0 \leq \Delta(P||Q) \leq \beta_\Delta(r, R) \quad (43)$$

and

$$\begin{aligned} 0 &\leq \beta_\Delta(r, R) - \Delta(P||Q) \\ &\leq \gamma_\Delta(r, R) [(R-1)(1-r) - \chi^2(P||Q)] \leq \alpha_\Delta(r, R), \end{aligned} \quad (44)$$

where

$$\alpha_\Delta(r, R) = \frac{1}{4}(R-r) \left[\frac{(R-1)(R+3)}{(R+1)^2} + \frac{(1-r)(r+3)}{(r+1)^2} \right], \quad (45)$$

$$\beta_\Delta(r, R) = \frac{2(R-1)(1-r)}{(R+1)(1+r)} \quad (46)$$

and

$$\gamma_\Delta(r, R) = (R-r)^{-1} \left[\frac{(R-1)(R+3)}{(R+1)^2} + \frac{(1-r)(r+3)}{(r+1)^2} \right]. \quad (47)$$

Proof. Follows from the Theorem 3.1 by considering f by f_Δ and making necessary calculations. $\square$

Theorem 4.2. *Let $P, Q \in \Gamma_n$ and $s \in \mathbb{R}$. Let there exists r, R ($0 < r \leq 1 \leq R < \infty$) such that $0 < r \leq \frac{p_i}{q_i} \leq R < \infty, \forall i \in \{1, 2, \dots, n\}$.*

(a) For $s \leq -1$, we have the following inequalities:

$$\frac{8r^{2-s}}{(r+1)^3} \Phi_s(P||Q) \leq \Delta(P||Q) \leq \frac{8R^{2-s}}{(R+1)^3} \Phi_s(P||Q), \quad (48)$$

$$\begin{aligned} & \frac{8r^{2-s}}{(r+1)^3} [\rho_{\Phi_s}(P||Q) - \Phi_s(P||Q)] \\ & \leq \Delta^*(P||Q) \leq \frac{8R^{2-s}}{(R+1)^3} [\rho_{\Phi_s}(P||Q) - \Phi_s(P||Q)] \end{aligned} \quad (49)$$

and

$$\begin{aligned} & \frac{8r^{2-s}}{(r+1)^3} [\beta_{\Phi_s}(r, R) - \Phi_s(P||Q)] \\ & \leq \beta_{\Delta}(r, R) - \Delta(P||Q) \\ & \leq \frac{8R^{2-s}}{(R+1)^3} [\beta_{\Phi_s}(r, R) - \Phi_s(P||Q)]. \end{aligned} \quad (50)$$

(b) For $s \geq 2$, we have the following inequalities:

$$\frac{8R^{2-s}}{(R+1)^3} \Phi_s(P||Q) \leq \Delta(P||Q) \leq \frac{8r^{2-s}}{(r+1)^3} \Phi_s(P||Q), \quad (51)$$

$$\begin{aligned} & \frac{8R^{2-s}}{(R+1)^3} [\rho_{\Phi_s}(P||Q) - \Phi_s(P||Q)] \\ & \leq \Delta^*(P||Q) \leq \frac{8r^{2-s}}{(r+1)^3} [\rho_{\Phi_s}(P||Q) - \Phi_s(P||Q)] \end{aligned} \quad (52)$$

and

$$\begin{aligned} & \frac{R^{1-s}}{(R+1)^2} [\beta_{\Phi_s}(r, R) - \Phi_s(P||Q)] \\ & \leq \beta_{\Delta}(r, R) - \Delta(P||Q) \\ & \leq \frac{r^{1-s}}{(r+1)^2} [\beta_{\Phi_s}(r, R) - \Phi_s(P||Q)], \end{aligned} \quad (53)$$

where

$$\Delta^*(P||Q) = \rho_{\Delta}(P||Q) - \Delta(P||Q) = 2 \sum_{i=1}^n q_i \left(\frac{p_i - q_i}{p_i + q_i} \right)^2. \quad (54)$$

Proof. Let us consider

$$g_{\Delta}(x) = x^{2-s} f''_{\Delta}(x) = \frac{8x^{2-s}}{(x+1)^3}, \quad x \in (0, \infty), \quad (55)$$

where $f''_{\Delta}(x)$ is as given by (14).

We have

$$g'_{\Delta}(x) = -\frac{8x^{1-s}[(s+1)x + (s-2)]}{(x+1)^4} \begin{cases} \geq 0, & s \leq -1 \\ \leq 0, & s \geq 2 \end{cases}. \quad (56)$$

In view of (56), we conclude the followings:

$$m = \inf_{x \in [r, R]} g(x) = \min_{x \in [r, R]} g(x) = \begin{cases} \frac{8r^{2-s}}{(r+1)^3}, & s \leq -1 \\ \frac{8R^{2-s}}{(R+1)^3}, & s \geq 2 \end{cases} \quad (57)$$

and

$$M = \sup_{x \in [r, R]} g(x) = \max_{x \in [r, R]} g(x) = \begin{cases} \frac{8R^{2-s}}{(R+1)^3}, & s \leq -1 \\ \frac{8r^{2-s}}{(r+1)^3}, & s \geq 2 \end{cases}. \quad (58)$$

From (57) and (58) and Theorem 3.2, we have the required proof. $\square$

The following propositions are the particular cases of the above theorem.

Proposition 4.1. *We have the following bounds in terms of χ^2 -divergence:*

$$\frac{4r^3}{(r+1)^3} \chi^2(Q||P) \leq \Delta(P||Q) \leq \frac{4R^3}{(R+1)^3} \chi^2(Q||P), \quad (59)$$

$$\begin{aligned} & \frac{8r^3}{(r+1)^3} [3\Phi_3(Q||P) - \chi^2(Q||P)] \\ & \leq \Delta^*(P||Q) \leq \frac{8R^3}{(R+1)^3} [3\Phi_3(Q||P) - \chi^2(Q||P)] \end{aligned} \quad (60)$$

and

$$\begin{aligned} & \frac{4r^3}{(r+1)^3} \left[\frac{(R-1)(1-r)}{rR} - \chi^2(Q||P) \right] \\ & \leq \frac{2(R-1)(1-r)}{(R+1)(1+r)} - \Delta(P||Q) \\ & \leq \frac{4R^3}{(R+1)^3} \left[\frac{(R-1)(1-r)}{rR} - \chi^2(Q||P) \right]. \end{aligned} \quad (61)$$

Proof. Take $s = -1$ in (48), (49) and (50) we get respectively (59), (60) and (61). $\square$

Proposition 4.2. *We have the following bounds in terms of χ^2 -divergence:*

$$\frac{4}{(R+1)^3} \chi^2(P||Q) \leq \Delta(P||Q) \leq \frac{4}{(r+1)^3} \chi^2(P||Q), \quad (62)$$

$$\frac{4}{(R+1)^3} \chi^2(P||Q) \leq \Delta^*(P||Q) \leq \frac{4}{(r+1)^3} \chi^2(P||Q) \quad (63)$$

and

$$\begin{aligned}
& \frac{4}{(R+1)^3} [(R-1)(1-r) - \chi^2(P||Q)] \\
& \leq \frac{2(R-1)(1-r)}{(R+1)(1+r)} - \Delta(P||Q) \\
& \leq \frac{4}{(r+1)^3} [(R-1)(1-r) - \chi^2(P||Q)].
\end{aligned} \tag{64}$$

Proof. Take $s = 2$ in (51), (52) and (53) we get respectively (62), (63) and (64). $\square$

We observe that the Theorem 4.2 is not valid for $s = 0, \frac{1}{2}$ and 1. These particular values of s we shall do separately. In these cases, we don't have inequalities on both sides as in the case of Propositions 4.1 and 4.2.

Proposition 4.3. *The following inequalities hold:*

$$0 \leq \Delta(P||Q) \leq \frac{32}{27} K(Q||P), \tag{65}$$

$$0 \leq \Delta^*(P||Q) \leq \frac{32}{27} [\chi^2(Q||P) - K(Q||P)] \tag{66}$$

and

$$\begin{aligned}
0 & \leq \frac{32}{27} K(Q||P) - \Delta(P||Q) \\
& \leq \frac{32}{27} \frac{(R-1) \ln \frac{1}{r} + (1-r) \ln \frac{1}{R}}{R-r} - \frac{2(R-1)(1-r)}{(R+1)(1+r)}
\end{aligned} \tag{67}$$

Proof. For $s = 0$ in (55), we have

$$g_{\Delta}(x) = \frac{8x^2}{(x+1)^3}. \tag{68}$$

This gives

$$g'_{\Delta}(x) = -\frac{8x(x-2)}{(x+1)^4} \begin{cases} \geq 0, & x \leq 2 \\ \leq 0, & x \geq 2 \end{cases}. \tag{69}$$

Thus we conclude that the function $g_W(x)$ given by (68) is increasing in $x \in (0, 2)$ and decreasing in $x \in (2, \infty)$, and hence

$$M = \sup_{x \in (0, \infty)} g_{\Delta}(x) = \max_{x \in (0, \infty)} g_{\Delta}(x) = g_{\Delta}(2) = \frac{32}{27}. \tag{70}$$

Now (70) together with (37), (38) and (39) give respectively (65), (66) and (67). $\square$

Proposition 4.4. *The following inequalities hold:*

$$0 \leq \Delta(P||Q) \leq 4 h(P||Q), \quad (71)$$

$$0 \leq \Delta^*(P||Q) \leq 2 \sum_{i=1}^n (q_i - p_i) \sqrt{\frac{q_i}{p_i}} - 4 h(P||Q) \quad (72)$$

and

$$\begin{aligned} 0 &\leq 4 h(P||Q) - \Delta(P||Q) \\ &\leq \frac{4(\sqrt{R}-1)(1-\sqrt{r})}{\sqrt{R}+\sqrt{r}} - \frac{2(R-1)(1-r)}{(R+1)(1+r)}. \end{aligned} \quad (73)$$

Proof. For $s = \frac{1}{2}$ in (55), we have

$$g_{\Delta}(x) = \frac{8x^{3/2}}{(x+1)^3}. \quad (74)$$

This gives

$$g'_{\Delta}(x) = -\frac{12\sqrt{x}(x-1)}{(x+1)^4} \begin{cases} \geq 0, & x \leq 1 \\ \leq 0, & x \geq 1 \end{cases}. \quad (75)$$

Thus we conclude that the function $g_{\Delta}(x)$ given by (74) is increasing in $x \in (0, 1)$ and decreasing in $x \in (1, \infty)$, and hence

$$M = \sup_{x \in (0, \infty)} g_{\Delta}(x) = \max_{x \in (0, \infty)} g_{\Delta}(x) = g_{\Delta}(1) = 1. \quad (76)$$

Now (76) together with (37), (38) and (39) give respectively (71), (72) and (73). $\square$

Proposition 4.5. *We have following inequalities:*

$$0 \leq \Delta(P||Q) \leq \frac{32}{27} K(P||Q), \quad (77)$$

$$0 \leq \Delta^*(P||Q) \leq \frac{32}{27} K(Q||P) \quad (78)$$

and

$$\begin{aligned} 0 &\leq \frac{32}{27} K(P||Q) - \Delta(P||Q) \\ &\leq \frac{32}{27} \frac{(R-1)r \ln r + (1-r)R \ln R}{R-r} - \frac{2(R-1)(1-r)}{(R+1)(1+r)}. \end{aligned} \quad (79)$$

Proof. For $s = 1$ in (55), we have

$$g_{\Delta}(x) = \frac{8x}{(x+1)^3}. \quad (80)$$

This gives

$$g'_W(x) = -\frac{8(2x-1)}{(x+1)^4} = \begin{cases} \geq 0, & x \leq \frac{1}{2} \\ \leq 0, & x \geq \frac{1}{2} \end{cases}. \quad (81)$$

Thus we conclude that the function $g_{\Delta}(x)$ given by (80) is increasing in $x \in (0, \frac{1}{2})$ and decreasing in $x \in (\frac{1}{2}, \infty)$, and hence

$$M = \sup_{x \in (0, \infty)} g_{\Delta}(x) = \max_{x \in (0, \infty)} g_{\Delta}(x) = g(\frac{1}{2}) = \frac{32}{27}. \quad (82)$$

Now (82) together with (37), (38) and (39) give respectively (77), (78) and (79). $\square$

Remark 4.1. In view of relation (7) and Propositions 4.1-4.5, we have the following main bounds on harmonic mean divergence:

$$\frac{2r^3}{(r+1)^3} \chi^2(Q||P) \leq 1 - W(P||Q) \leq \frac{2R^3}{(R+1)^3} \chi^2(Q||P), \quad (83)$$

$$\frac{2}{(R+1)^3} \chi^2(P||Q) \leq 1 - W(P||Q) \leq \frac{2}{(r+1)^3} \chi^2(P||Q), \quad (84)$$

$$0 \leq 1 - W(P||Q) \leq \frac{16}{27} K(Q||P), \quad (85)$$

$$0 \leq 1 - W(P||Q) \leq 2 h(P||Q) \quad (86)$$

and

$$0 \leq 1 - W(P||Q) \leq \frac{16}{27} K(P||Q). \quad (87)$$

The inequalities (84) were also studied by Dragomir [8]. The inequalities (87) can be seen in Dragomir [10]. The inequalities (71) can be seen in Dragomir [11] and Topsøe [26]. The inequalities (77) can be in and Dragomir [10].

5 Symmetric Chi-square Divergence and Inequalities

In this section, we shall give bounds on *symmetric chi-square divergence* based on the Theorems 3.1 and 3.2.

Theorem 5.1. For all $P, Q \in \Gamma_n$, we have the following inequalities:

$$0 \leq \Psi(P||Q) \leq \rho_{\Psi}(P||Q), \quad (88)$$

where

$$\rho_{\Psi}(P||Q) = \Psi(P||Q) + \sum_{i=1}^n \frac{(p_i - q_i)^2(p_i^2 + q_i^2)}{p_i^2 q_i}. \quad (89)$$

If there exists r, R ($0 < r \leq 1 \leq R < \infty$) such that

$$0 < r \leq \frac{p_i}{q_i} \leq R < \infty, \quad \forall i \in \{1, 2, \dots, n\},$$

then we have the following inequalities:

$$0 \leq \Psi(P||Q) \leq \alpha_{\Psi}(r, R), \quad (90)$$

$$0 \leq \Psi(P||Q) \leq \beta_{\Psi}(r, R) \quad (91)$$

and

$$\begin{aligned} 0 &\leq \beta_{\Psi}(r, R) - \Psi(P||Q) \\ &\leq \gamma_{\Psi}(r, R) [(R-1)(1-r) - \chi^2(P||Q)] \leq \alpha_{\Psi}(r, R), \end{aligned} \quad (92)$$

where

$$\alpha_{\Psi}(r, R) = \frac{1}{4}(R-r)^2 [2L_2^{-1}(r, R) - L_1^{-1}(r, R)], \quad (93)$$

$$\beta_{\Psi}(r, R) = (R-1)(1-r)(R+r) \quad (94)$$

and

$$\gamma_{\Psi}(r, R) = 2L_2^{-1}(r, R) - L_1^{-1}(r, R). \quad (95)$$

Proof. Follows from Theorem 3.1 by considering f by f_{Ψ} and making necessary calculations. $\square$

Theorem 5.2. Let $P, Q \in \Gamma_n$ and $s \in \mathbb{R}$. Let there exists r, R ($0 \leq r \leq 1 \leq R \leq \infty$) such that $0 < r \leq \frac{p_i}{q_i} \leq R < \infty, \forall i \in \{1, 2, \dots, n\}$.

(a) For $s \leq -1$, we have the following inequalities:

$$\frac{2(r^3 + 1)}{r^{1+s}} \Phi_s(P||Q) \leq \Psi(P||Q) \leq \frac{2(R^3 + 1)}{R^{1+s}} \Phi_s(P||Q), \quad (96)$$

$$\begin{aligned} &\frac{2(r^3 + 1)}{r^{1+s}} [\rho_{\Phi_s}(P||Q) - \Phi_s(P||Q)] \\ &\leq \Psi^*(P||Q) \leq \frac{2(R^3 + 1)}{R^{1+s}} [\rho_{\Phi_s}(P||Q) - \Phi_s(P||Q)] \end{aligned} \quad (97)$$

and

$$\begin{aligned} \frac{2(r^3 + 1)}{r^{1+s}} [\beta_{\Phi_s}(r, R) - \Phi_s(P||Q)] \\ \leq \beta_{\Psi}(r, R) - \Psi(P||Q) \\ \leq \frac{2(R^3 + 1)}{R^{1+s}} [\beta_{\Phi_s}(r, R) - \Phi_s(P||Q)]. \end{aligned} \quad (98)$$

(b) For $s \geq 2$, we have the following inequalities:

$$\frac{2(R^3 + 1)}{R^{1+s}} \Phi_s(P||Q) \leq \Psi(P||Q) \leq \frac{2(r^3 + 1)}{r^{1+s}} \Phi_s(P||Q), \quad (99)$$

$$\begin{aligned} \frac{2(r^3 + 1)}{r^{1+s}} [\rho_{\Phi_s}(P||Q) - \Phi_s(P||Q)] \\ \leq \Psi^*(P||Q) \leq \frac{2(R^3 + 1)}{R^{1+s}} [\rho_{\Phi_s}(P||Q) - \Phi_s(P||Q)] \end{aligned} \quad (100)$$

and

$$\begin{aligned} \frac{2(R^3 + 1)}{R^{1+s}} [\beta_{\Phi_s}(r, R) - \Phi_s(P||Q)] \\ \leq \beta_{\Psi}(r, R) - \Psi(P||Q) \\ \leq \frac{2(r^3 + 1)}{r^{1+s}} [\beta_{\Phi_s}(r, R) - \Phi_s(P||Q)], \end{aligned} \quad (101)$$

where

$$\Psi^*(P||Q) = \rho_{\Psi}(P||Q) - \Psi(P||Q) = \sum_{i=1}^n \frac{(p_i - q_i)^2 (p_i^2 + q_i^2)}{p_i^2 q_i}. \quad (102)$$

Proof. Let us consider

$$g_{\Psi}(x) = x^{2-s} f_{\Psi}''(x) = 2x^{-1-s}(x^3 + 1), \quad x \in (0, \infty), \quad (103)$$

where $f_{\Psi}''(x)$ is as given by (17).

We have

$$g'_{\Psi}(x) = -2x^{-2-s} [(s-2)x^3 + (s+1)] \begin{cases} \geq 0, & s \leq -1 \\ \leq 0, & s \geq 2 \end{cases}. \quad (104)$$

From (104), we conclude the followings:

$$m = \inf_{x \in [r, R]} g_{\Psi}(x) = \min_{x \in [r, R]} g_{\Psi}(x) = \begin{cases} \frac{2(r^3+1)}{r^{1+s}}, & s \leq -1 \\ \frac{2(R^3+1)}{R^{1+s}}, & s \geq 2 \end{cases} \quad (105)$$

and

$$M = \sup_{x \in [r, R]} g_{\Psi}(x) = \max_{x \in [r, R]} g_{\Psi}(x) = \begin{cases} \frac{2(R^3+1)}{R^{1+s}}, & s \leq -1 \\ \frac{2(r^3+1)}{r^{1+s}}, & s \geq 2 \end{cases} \quad (106)$$

In view of (105) and (106) and Theorem 3.2, we have the required proof. $\square$

Proposition 5.1. *We have the following bounds in terms of χ^2 -divergence:*

$$(r^3 + 1)\chi^2(Q||P) \leq \Psi(P||Q) \leq (R^3 + 1)\chi^2(Q||P), \quad (107)$$

$$\begin{aligned} 2(r^3 + 1) [3\Phi_3(Q||P) - \chi^2(Q||P)] \\ \leq \Psi^*(P||Q) \leq 2(R^3 + 1) [3\Phi_3(P||Q) - \chi^2(Q||P)] \end{aligned} \quad (108)$$

and

$$\begin{aligned} (r^3 + 1) \left[\frac{(R-1)(1-r)}{rR} - \chi^2(Q||P) \right] \\ \leq (R-1)(1-r)(R+r) - \Psi(P||Q) \\ \leq (R^3 + 1) \left[\frac{(R-1)(1-r)}{rR} - \chi^2(Q||P) \right]. \end{aligned} \quad (109)$$

Proof. Take $s = -1$ in (96), (97) and (98) we get respectively (107), (108) and (109). $\square$

Proposition 5.2. *The following bounds on in terms of χ^2 -divergence hold:*

$$\frac{R^3 + 1}{R^3} \chi^2(P||Q) \leq \Psi(P||Q) \leq \frac{r^3 + 1}{r^3} \chi^2(P||Q), \quad (110)$$

$$\frac{R^3 + 1}{R^3} \chi^2(P||Q) \leq \Psi^*(P||Q) \leq \frac{r^3 + 1}{r^3} \chi^2(P||Q) \quad (111)$$

and

$$\begin{aligned} \frac{R^3 + 1}{R^3} [(R-1)(1-r) - \chi^2(P||Q)] \\ \leq \beta_\Psi(r, R) - \Psi(P||Q) \\ \leq \frac{r^3 + 1}{r^3} [(R-1)(1-r) - \chi^2(P||Q)]. \end{aligned} \quad (112)$$

Proof. Take $s = 2$ in (99), (100) and (101) we get respectively (110), (111) and (112). $\square$

We observe that the above two propositions follows from Theorem 5.2 immediately by taking $s = -1$ and $s = 2$ respectively. But still there are another values of s such as $s = 0$, $s = 1$ and $s = \frac{1}{2}$ for which we can obtain bounds. These values are studied below.

Proposition 5.3. *We have following bounds in terms of relative information:*

$$0 \leq 3\sqrt[3]{2} K(Q||P) \leq \Psi(P||Q), \quad (113)$$

$$0 \leq 3\sqrt[3]{2} [\chi^2(Q||P) - K(Q||P)] \leq \Psi^*(P||Q) \quad (114)$$

and

$$\begin{aligned} 0 \leq \Psi(P||Q) - 3\sqrt[3]{2} K(Q||P) \\ \leq (R-1)(1-r)(R+r) - 3\sqrt[3]{2} \frac{(R-1) \ln \frac{1}{r} + (1-r) \ln \frac{1}{R}}{R-r}. \end{aligned} \quad (115)$$

Proof. For $s = 0$ in (103), we have

$$g_{\Psi}(x) = \frac{2(x^3 + 1)}{x}, \quad (116)$$

This gives

$$\begin{aligned} g'_{\Psi}(x) &= \frac{2(2x^3 - 1)}{x^2} \\ &= \frac{2(\sqrt[3]{2} x - 1)(\sqrt[3]{4} x^2 + \sqrt[3]{2} x + 1)}{x^2} \begin{cases} \geq 0, & x \geq \frac{1}{\sqrt[3]{2}} \\ \leq 0, & x \leq \frac{1}{\sqrt[3]{2}} \end{cases}. \end{aligned} \quad (117)$$

Thus we conclude that the function $g_{\Psi}(x)$ given by (116) is decreasing in $x \in (0, \frac{1}{\sqrt[3]{2}})$ and increasing in $x \in (\frac{1}{\sqrt[3]{2}}, \infty)$, and hence

$$m = \inf_{x \in (0, \infty)} g_{\Psi}(x) = \min_{x \in (0, \infty)} g_{\Psi}(x) = g_{\Psi}\left(\frac{1}{\sqrt[3]{2}}\right) = 3\sqrt[3]{2}. \quad (118)$$

Now (118) together with (37), (38) and (39) give respectively (113), (114) and (115). $\square$

Proposition 5.4. *We have following bounds in terms of Hellinger's discrimination:*

$$0 \leq 16 h(P||Q) \leq \Psi(P||Q), \quad (119)$$

$$0 \leq 16 \left[\frac{1}{2} \sum_{i=1}^n (q_i - p_i) \sqrt{\frac{q_i}{p_i}} - h(Q||P) \right] \leq \Psi^*(P||Q) \quad (120)$$

and

$$\begin{aligned} 0 &\leq \Psi(P||Q) - 16 h(P||Q) \\ &\leq (R - 1)(1 - r)(R + r) - \frac{16(\sqrt{R} - 1)(1 - \sqrt{r})}{\sqrt{R} + \sqrt{r}}. \end{aligned} \quad (121)$$

Proof. For $s = \frac{1}{2}$ in (103), we have

$$g_{\Psi}(x) = \frac{2(x^3 + 1)}{x^{3/2}}. \quad (122)$$

This gives

$$g'_{\Psi}(x) = \frac{3(x^3 - 1)}{x^{5/2}} = \frac{3(x - 1)(x^2 + x + 1)}{x^{5/2}} \begin{cases} \geq 0, & x \geq 1 \\ \leq 0, & x \leq 1 \end{cases}. \quad (123)$$

Thus we conclude that the function $g_{\Psi}(x)$ given by (122) is decreasing in $x \in (0, 1)$ and increasing in $x \in (1, \infty)$, and hence

$$m = \inf_{x \in (0, \infty)} g_{\Psi}(x) = \min_{x \in (0, \infty)} g_{\Psi}(x) = g_{\Psi}(1) = 4. \quad (124)$$

Now (124) together with (37), (38) and (39) give respectively (119), (120) and (121). $\square$

Proposition 5.5. *We have the following bounds in terms of relative information:*

$$0 \leq 3\sqrt[3]{2} K(P||Q) \leq \Psi(P||Q), \quad (125)$$

$$0 \leq 3\sqrt[3]{2} K(Q||P) \leq \Psi^*(P||Q) \quad (126)$$

and

$$\begin{aligned} 0 &\leq \Psi(P||Q) - 3\sqrt[3]{2} K(P||Q) \\ &\leq (R-1)(1-r)(R+r) - 3\sqrt[3]{2} \frac{(R-1)r \ln r + (1-r)R \ln R}{R-r}. \end{aligned} \quad (127)$$

Proof. For $s = 1$ in (103), we have

$$g_{\Psi}(x) = \frac{2(x^3 + 1)}{x^2}. \quad (128)$$

This gives

$$\begin{aligned} g'_{\Psi}(x) &= \frac{2(x^3 - 2)}{x^3} \\ &= \frac{2(x - \sqrt[3]{2})(x^2 + \sqrt[3]{2}x + \sqrt[3]{4})}{x^3} \begin{cases} \geq 0, & x \geq \sqrt[3]{2} \\ \leq 0, & x \leq \sqrt[3]{2} \end{cases}. \end{aligned} \quad (129)$$

Thus we conclude that the function $g_{\Psi}(x)$ given by (128) is decreasing in $x \in (0, \sqrt[3]{2})$ and increasing in $x \in (\sqrt[3]{2}, \infty)$, and hence

$$m = \inf_{x \in (0, \infty)} g_{\Psi}(x) = \min_{x \in (0, \infty)} g_{\Psi}(x) = g_{\Psi}(\sqrt[3]{2}) = 3\sqrt[3]{2}. \quad (130)$$

Now (130) together with (37), (38) and (39) give respectively (125), (126) and (127). $\square$

References

- [1] N.S. BARNETT, P. CERENE and S. S. DRAGOMIR, Some New Inequalities for Hermite-Hadamard Divergence in Information Theory, <http://rgmia.vu.edu.au>, *RGMA Research Report Collection*, **(5)**(4)(2002), Article 8.
- [2] A. BHATTACHARYYA, Some Analogues to the Amount of Information and Their uses in Statistical Estimation, *Sankhya*, **8**(1946), 1-14.
- [3] BULLEN, P.S., D.S. MITRINOVIĆ and P.M. VASIĆ, *Means and Their Inequalities*, Kluwer Academic Publishers, 1988.
- [4] P. CERONE, S.S. DRAGOMIR and F. ÖSTERREICHER, Bounds on Extended f-Divergence for a Variety of Classes, <http://rgmia.vu.edu.au>, *RGMA Research Report Collection*, **6**(1)(2003), Article 7.

- [5] I. CSISZÁR, Information Type Measures of Differences of Probability Distribution and Indirect Observations, *Studia Math. Hungarica*, **2**(1967), 299-318.
- [6] S. S. DRAGOMIR, Some Inequalities for the Csiszár Φ -Divergence - *Inequalities for Csiszár f-Divergence in Information Theory* - Monograph - Chapter I - Article 1 - <http://rgmia.vu.edu.au/monographs/csiszar.htm>.
- [7] S. S. DRAGOMIR, A Converse Inequality for the Csiszár Φ -Divergence- *Inequalities for Csiszár f-Divergence in Information Theory* - Monograph - Chapter I - Article 2 - <http://rgmia.vu.edu.au/monographs/csiszar.htm>.
- [8] S. S. DRAGOMIR, Some Inequalities for (m, M)-Convex Mappings and Applications for the Csiszár Φ -Divergence in Information Theory- *Inequalities for Csiszár f-Divergence in Information Theory* - Monograph - Chapter I - Article 3 - <http://rgmia.vu.edu.au/monographs/csiszar.htm>.
- [9] S. S. DRAGOMIR, Other Inequalities for Csiszár Divergence and Applications - *Inequalities for Csiszár f-Divergence in Information Theory* - Monograph - Chapter I - Article 4 - <http://rgmia.vu.edu.au/monographs/csiszar.htm>.
- [10] S. S. DRAGOMIR, Upper and Lower Bounds for Csiszár's f-divergence in terms of the Kullback-Leibler Distance and Applications - *Inequalities for Csiszár f-Divergence in Information Theory* - Monograph - Chapter II - Article 1 - <http://rgmia.vu.edu.au/monographs/csiszar.htm>.
- [11] S. S. DRAGOMIR, Upper and Lower Bounds for Csiszár's f-Divergence in terms of Hellinger Discrimination and Applications - *Inequalities for Csiszár f-Divergence in Information Theory* - Monograph - Chapter II - Article 2 - <http://rgmia.vu.edu.au/monographs/csiszar.htm>.
- [12] S. S. DRAGOMIR, J. SUNDE and C. BUSE, New Inequalities for Jeffreys Divergence Measure, *Tamsui Oxford Journal of Mathematical Sciences*, **16**(2)(2000), 295-309.
- [13] E. HELLINGER, Neue Begründung der Theorie der quadratischen Formen von unendlichen vielen Veränderlichen, *J. Reine Aug. Math.*, **136**(1909), 210-271.
- [14] H. JEFFREYS, An Invariant Form for the Prior Probability in Estimation Problems, *Proc. Roy. Soc. Lon., Ser. A*, **186**(1946), 453-461.
- [15] S. KULLBACK and R.A. LEIBLER, On Information and Sufficiency, *Ann. Math. Statist.*, **22**(1951), 79-86.
- [16] F. LIESE and I. VAJDA, *Convex Statistical Decision Rule*, Teubner-Texte zur Mathematik, Band 95, Leipzig, 1987.
- [17] J. LIN, Divergence Measures Based on the Shannon Entropy, *IEEE Trans. on Inform. Theory*, **IT-37**(1991), 145-151.

- [18] K. PEARSON, On the Criterion that a given system of deviations from the probable in the case of correlated system of variables is such that it can be reasonable supposed to have arisen from random sampling, *Phil. Mag.*, **50**(1900), 157-172.
- [19] I.J. TANEJA, Some Contributions to Information Theory – I (A Survey: On Measures of Information), *Journal of Combinatorics, Information and System Sciences*, **4**(1979), 253-274.
- [20] I.J. TANEJA, On Generalized Information Measures and Their Applications, Chapter in: *Advances in Electronics and Electron Physics*, Ed. P.W. Hawkes, Academic Press, **76**(1989), 327-413.
- [21] I.J. TANEJA, New Developments in Generalized Information Measures, Chapter in: *Advances in Imaging and Electron Physics*, Ed. P.W. Hawkes, **91**(1995), 37-135.
- [22] I.J. TANEJA, *Generalized Information Measures and their Applications*, on line book: <http://www.mtm.ufsc.br/~taneja/book/book.html>, 2001.
- [23] I.J. TANEJA, Generalized Relative Information and Information Inequalities, *Journal of Inequalities in Pure and Applied Mathematics*. **5**(1)(2004), Article 21, 1-19. Also in: *RGMIA Research Report Collection*, <http://rgmia.vu.edu.au>, **6**(3)(2003), Article 10.
- [24] I.J. TANEJA, Bounds on Non-Symmetric Divergence Measures in terms of Symmetric Divergence Measures - To appear in: *Journal of Combinatorics, Information and System Sciences*.
- [25] I.J. TANEJA and P. KUMAR, Relative Information of Type s , Csiszár f -Divergence, and Information Inequalities, *Information Sciences*, **166**(1-4)(2004), 105-125. Also in: <http://rgmia.vu.edu.au>, *RGMIA Research Report Collection*, **6**(3)(2003), Article 12.
- [26] F. TOPSØE, Some Inequalities for Information Divergence and Related Measures of Discrimination, *IEEE Trans. on Inform. Theory*, **IT-46**(2000), 1602-1609. Also in *RGMIA Research Report Collection*, **2**(1)(1999), Article 9 – <http://sci.vu.edu.au/~rgmia>.
- [27] I. VAJDA, *Theory of Statistical Inference and Information*, Kluwer Academic Press, London, 1989.

Instructions to Contributors
Journal of Concrete and Applicable Mathematics

A quarterly international publication of Eudoxus Press, LLC, of TN.

Editor in Chief: George Anastassiou

Department of Mathematical Sciences
 University of Memphis
 Memphis, TN 38152-3240, U.S.A.

1. Manuscripts hard copies in triplicate, and in English, should be submitted to the Editor-in-Chief:

Prof. George A. Anastassiou
 Department of Mathematical Sciences
 The University of Memphis
 Memphis, TN 38152, USA.
 Tel. 001. 901.678.3144
 e-mail: ganastss@memphis.edu.

Authors may want to recommend an associate editor the most related to the submission to possibly handle it.

Also authors may want to submit a list of six possible referees, to be used in case we cannot find related referees by ourselves.

2. Manuscripts should be typed using any of TEX, LaTeX, AMS-TEX, or AMS-LaTeX and according to EUDOXUS PRESS, LLC. LATEX STYLE FILE. (VISIT www.msci.memphis.edu/~ganastss/jcaam/ to save a copy of the style file.) They should be carefully prepared in all respects. Submitted copies should be brightly printed (not dot-matrix), double spaced, in ten point type size, on one side high quality paper 8(1/2)x11 inch. Manuscripts should have generous margins on all sides and should not exceed 24 pages.

3. Submission is a representation that the manuscript has not been published previously in this or any other similar form and is not currently under consideration for publication elsewhere. A statement transferring from the authors (or their employers, if they hold the copyright) to Eudoxus Press, LLC, will be required before the manuscript can be accepted for publication. The Editor-in-Chief will supply the necessary forms for this transfer. Such a written transfer of copyright, which previously was assumed to be implicit in the act of submitting a manuscript, is necessary under the U.S. Copyright Law in order for the publisher to carry through the dissemination of research results and reviews as widely and effectively as possible.

4. The paper starts with the title of the article, author's name(s) (no titles or degrees), author's affiliation(s) and e-mail addresses. The affiliation should comprise the department, institution (usually university or company), city, state (and/or nation) and mail code.

The following items, 5 and 6, should be on page no. 1 of the paper.

5. An abstract is to be provided, preferably no longer than 150 words.
 6. A list of 5 key words is to be provided directly below the abstract. Key words should express the precise content of the manuscript, as they are used for indexing purposes. The main body of the paper should begin on page no. 1, if possible.

7. All sections should be numbered with Arabic numerals (such as: 1. INTRODUCTION).

Subsections should be identified with section and subsection numbers (such as 6.1. Second-Value Subheading).

If applicable, an independent single-number system (one for each category) should be used to label all theorems, lemmas, propositions, corollaries, definitions, remarks, examples, etc. The label (such as Lemma 7) should be typed with paragraph indentation, followed by a period and the lemma itself.

8. Mathematical notation must be typeset. Equations should be numbered consecutively with Arabic numerals in parentheses placed flush right, and should be thusly referred to in the text [such as Eqs.(2) and (5)]. The running title must be placed at the top of even numbered pages and the first author's name, et al., must be placed at the top of the odd numbered pages.

9. Illustrations (photographs, drawings, diagrams, and charts) are to be numbered in one consecutive series of Arabic numerals. The captions for illustrations should be typed double space. All illustrations, charts, tables, etc., must be embedded in the body of the manuscript in proper, final, print position. In particular, manuscript, source, and PDF file version must be at camera ready stage for publication or they cannot be considered.

Tables are to be numbered (with Roman numerals) and referred to by number in the text. Center the title above the table, and type explanatory footnotes (indicated by superscript lowercase letters) below the table.

10. List references alphabetically at the end of the paper and number them consecutively. Each must be cited in the text by the appropriate Arabic numeral in square brackets on the baseline.

References should include (in the following order):

initials of first and middle name, last name of author(s)

title of article,

name of publication, volume number, inclusive pages, and year of publication.

Authors should follow these examples:

Journal Article

1. H.H.Gonska, Degree of simultaneous approximation of bivariate functions by Gordon operators, (journal name in italics) *J. Approx. Theory*, 62,170-191(1990).

Book

2. G.G.Lorentz, (title of book in italics) *Bernstein Polynomials* (2nd ed.), Chelsea, New York, 1986.

Contribution to a Book

3. M.K.Khan, Approximation properties of beta operators, in (title of book in italics) *Progress in Approximation Theory* (P.Nevai and A.Pinkus, eds.), Academic Press, New York, 1991, pp.483-495.

11. All acknowledgements (including those for a grant and financial support) should occur in one paragraph that directly precedes the References section.

12. Footnotes should be avoided. When their use is absolutely necessary, footnotes should be numbered consecutively using Arabic numerals and should be typed at the bottom of the page to which they refer. Place a line above the footnote, so that it is set

off from the text. Use the appropriate superscript numeral for citation in the text.

13. After each revision is made please again submit three hard copies of the revised manuscript, including in the final one. And after a manuscript has been accepted for publication and with all revisions incorporated, manuscripts, including the TEX/LaTeX source file and the PDF file, are to be submitted to the Editor's Office on a personal-computer disk, 3.5 inch size. Label the disk with clearly written identifying information and properly ship, such as:

Your name, title of article, kind of computer used, kind of software and version number, disk format and files names of article, as well as abbreviated journal name.

Package the disk in a disk mailer or protective cardboard. Make sure contents of disks are identical with the ones of final hard copies submitted!

Note: The Editor's Office cannot accept the disk without the accompanying matching hard copies of manuscript. No e-mail final submissions are allowed! The disk submission must be used.

14. Effective 1 Nov. 2005 the journal's page charges are \$8.00 per PDF file page. Upon acceptance of the paper an invoice will be sent to the contact author. The fee payment will be due one month from the invoice date. The article will proceed to publication only after the fee is paid. The charges are to be sent, by money order or certified check, in US dollars, payable to Eudoxus Press, LLC, to the address shown in the Scope and Prices section.

No galleys will be sent and the contact author will receive one(1) electronic copy of the journal issue in which the article appears.

15. This journal will consider for publication only papers that contain proofs for their listed results.

TABLE OF CONTENTS, JOURNAL OF CONCRETE AND APPLICABLE
MATHEMATICS, VOL.4, NO.1, 2006

FROM GENERALIZED DE GIORGI ESTIMATES TO UPPER GAUSSIAN BOUNDS FOR THE HEAT KERNEL ASSOCIATED TO HIGHER ORDER ELLIPTIC OPERATORS, M.QAFSAOUI,.....	9
EXPONENTIAL OPERATORS AND SOLUTION OF PSEUDO-CLASSICAL EVOLUTION PROBLEMS, C.CASSISA, P.RICCI,.....	33
NEWTON SUM RULES OF THE ZEROS OF SOME ASSOCIATED ORTHOGONAL q-POLYNOMIALS, C.KOKOLOGIANNAKI, P.SIAFARICAS, I.STABOLAS,.....	47
KONDURAR THEOREM AND ITO FORMULA IN RIESZ SPACES, A.BOCCUTO, D.CANDELORO, E.KUBINSKA,.....	67
BOUNDS ON TRIANGULAR DISCRIMINATION, HARMONIC MEAN AND SYMMETRIC CHI-SQUARE DIVERGENCES, I.TANEJA,.....	91

VOLUME 4,NUMBER 2 APRIL 2006

ISSN:1548-5390 PRINT,1559-176X ONLINE



**JOURNAL
OF CONCRETE
AND APPLICABLE
MATHEMATICS**

EUDOXUS PRESS,LLC

SCOPE AND PRICES OF THE JOURNAL
Journal of Concrete and Applicable Mathematics

A quartely international publication of **Eudoxus Press, LLC**

Editor in Chief: George Anastassiou
Department of Mathematical Sciences,
University of Memphis
Memphis, TN 38152, U.S.A.
ganastss@memphis.edu

The main purpose of the "Journal of Concrete and Applicable Mathematics" is to publish high quality original research articles from all subareas of Non-Pure and/or Applicable Mathematics and its many real life applications, as well connections to other areas of Mathematical Sciences, as long as they are presented in a Concrete way. It welcomes also related research survey articles and book reviews. A sample list of connected mathematical areas with this publication includes and is not restricted to: Applied Analysis, Applied Functional Analysis, Probability theory, Stochastic Processes, Approximation Theory, O.D.E, P.D.E, Wavelet, Neural Networks, Difference Equations, Summability, Fractals, Special Functions, Splines, Asymptotic Analysis, Fractional Analysis, Inequalities, Moment Theory, Numerical Functional Analysis, Tomography, Asymptotic Expansions, Fourier Analysis, Applied Harmonic Analysis, Integral Equations, Signal Analysis, Numerical Analysis, Optimization, Operations Research, Linear Programming, Fuzzyness, Mathematical Finance, Stochastic Analysis, Game Theory, Math. Physics aspects, Applied Real and Complex Analysis, Computational Number Theory, Graph Theory, Combinatorics, Computer Science Math. related topics, combinations of the above, etc. In general any kind of Concretely presented Mathematics which is Applicable fits to the scope of this journal. Working Concretely and in Applicable Mathematics has become a main trend in many recent years, so we can understand better and deeper and solve the important problems of our real and scientific world. "Journal of Concrete and Applicable Mathematics" is a peer-reviewed International Quarterly Journal. We are calling for papers for possible publication. The contributor should send three copies of the contribution to the editor in-Chief typed in TEX, LATEX double spaced. [See: Instructions to Contributors]

Journal of Concrete and Applicable Mathematics(JCAAM)
ISSN:1548-5390 PRINT, 1559-176X ONLINE.

is published in January, April, July and October of each year by
EUDOXUS PRESS, LLC,
1424 Beaver Trail Drive, Cordova, TN 38016, USA,
Tel. 001-901-751-3553
anastassioug@yahoo.com
<http://www.EudoxusPress.com>.

Annual Subscription Current Prices: For USA and Canada, Institutional: Print \$250, Electronic \$220, Print and Electronic \$310. Individual: Print \$77, Electronic \$60, Print & Electronic \$110. For any other part of the world add \$25 more to the above prices for Print.

Single article PDF file for individual \$8. Single issue in PDF form for individual \$25.
The journal carries page charges \$8 per page of the pdf file of an article, payable upon acceptance of the article within one month and before publication.
No credit card payments. Only certified check, money order or international check in US dollars are acceptable.
Combination orders of any two from JoCAAA, JCAAM, JAFa receive 25% discount, all three receive 30% discount.

Copyright©2004 by Eudoxus Press, LLC all rights reserved. JCAAM is printed in USA.

JCAAM is reviewed and abstracted by AMS Mathematical Reviews, MATHSCI, and Zentralblatt MATH.

It is strictly prohibited the reproduction and transmission of any part of JCAAM and in any form and by any means without the written permission of the publisher. It is only allowed to educators to Xerox articles for educational purposes. The publisher assumes no responsibility for the content of published papers.

JCAAM IS A JOURNAL OF RAPID PUBLICATION

Editorial Board Associate Editors

Editor in -Chief:
George Anastassiou
Department of Mathematical Sciences
The University Of Memphis
Memphis, TN 38152, USA
tel. 901-678-3144, fax 901-678-2480
e-mail ganastss@memphis.edu
www.msci.memphis.edu/~ganastss/jcaam
Areas: Approximation Theory,
Probability, Moments, Wavelet,
Neural Networks, Inequalities, Fuzzyness.

Associate Editors:

1) Ravi Agarwal
Florida Institute of Technology
Applied Mathematics Program
150 W. University Blvd.
Melbourne, FL 32901, USA
agarwal@fit.edu
Differential Equations, Difference
Equations, inequalities

2) Shair Ahmad
University of Texas at San Antonio
Division of Math. & Stat.
San Antonio, TX 78249-0664, USA
SAHMAD@utsa.edu
Differential Equations, Mathematical
Biology

3) Drumi D. Bainov
Medical University of Sofia
P.O. Box 45, 1504 Sofia, Bulgaria
drumibainov@yahoo.com
Differential Equations, Optimal Control,
Numerical Analysis, Approximation Theory

4) Carlo Bardaro
Dipartimento di Matematica & Informatica
Universita' di Perugia
Via Vanvitelli 1
06123 Perugia, ITALY
tel. +390755855034, +390755853822,
fax +390755855024
bardaro@unipg.it ,
bardaro@dipmat.unipg.it
Functional Analysis and Approximation Th.,

19) Rupert Lasser
Institut fur Biomathematik & Biomertie, GSF
-National Research Center for environment and
health
Ingolstaedter landstr.1
D-85764 Neuherberg, Germany
lasser@gsf.de
Orthogonal Polynomials, Fourier Analysis,
Mathematical Biology

20) Alexandru Lupas
University of Sibiu
Faculty of Sciences
Department of Mathematics
Str. I. Ratiu nr. 7
2400-Sibiu, Romania
lupas@ulbsibiu.ro
Classical Analysis, Inequalities,
Special Functions, Umbral Calculus,
Approximation Th., Numerical Analysis
and Methods

21) Ram N. Mohapatra
Department of Mathematics
University of Central Florida
Orlando, FL 32816-1364
tel. 407-823-5080
ramm@mail.ucf.edu
Real and Complex analysis, Approximation Th.,
Fourier Analysis, Fuzzy Sets and Systems

22) Rainer Nagel
Arbeitsbereich Funktionalanalysis
Mathematisches Institut
Auf der Morgenstelle 10
D-72076 Tuebingen
Germany
tel. 49-7071-2973242
fax 49-7071-294322
rana@fa.uni-tuebingen.de
evolution equations, semigroups, spectral th.,
positivity

23) Panos M. Pardalos
Center for Appl. Optimization
University of Florida
303 Weil Hall
P.O. Box 116595
Gainesville, FL 32611-6595

Summability, Signal Analysis, Integral
Equations,
Measure Th., Real Analysis

5) Francoise Bastin
Institute of Mathematics
University of Liege
4000 Liege
BELGIUM
f.bastin@ulg.ac.be
Functional Analysis, Wavelets

6) Paul L. Butzer
RWTH Aachen
Lehrstuhl A für Mathematik
D-52056 Aachen
Germany
tel. 0049/241/80-94627 office,
0049/241/72833 home,
fax 0049/241/80-92212
Butzer@rwth-aachen.de
Approximation Th., Sampling Th., Signals,
Semigroups of Operators, Fourier Analysis

7) Yeol Je Cho
Department of Mathematics Education
College of Education
Gyeongsang National University
Chinju 660-701
KOREA
tel. 055-751-5673 Office,
055-755-3644 home,
fax 055-751-6117
yjcho@nongae.gsnu.ac.kr
Nonlinear operator Th., Inequalities,
Geometry of Banach Spaces

8) Sever S. Dragomir
School of Communications and Informatics
Victoria University of Technology
PO Box 14428
Melbourne City M.C
Victoria 8001, Australia
tel 61 3 9688 4437, fax 61 3 9688 4050
sever.dragomir@vu.edu.au,
sever@sci.vu.edu.au
Math. Analysis, Inequalities, Approximation
Th.,
Numerical Analysis, Geometry of Banach
Spaces,
Information Th. and Coding

9) A.M. Fink
Department of Mathematics
Iowa State University
Ames, IA 50011-0001, USA

tel. 352-392-9011
pardalos@ufl.edu
Optimization, Operations Research

24) Svetlozar T. Rachev
Dept. of Statistics and Applied Probability
Program
University of California, Santa Barbara
CA 93106-3110, USA
tel. 805-893-4869
rachev@pstat.ucsb.edu
AND
Chair of Econometrics and Statistics
School of Economics and Business Engineering
University of Karlsruhe
Kollegium am Schloss, Bau II, 20.12, R210
Postfach 6980, D-76128, Karlsruhe, Germany
tel. 011-49-721-608-7535
rachev@lsoe.uni-karlsruhe.de
Mathematical and Empirical Finance,
Applied Probability, Statistics and Econometrics

25) Paolo Emilio Ricci
Universita' degli Studi di Roma "La Sapienza"
Dipartimento di Matematica-Istituto
"G. Castelnuovo"
P.le A. Moro, 2-00185 Roma, ITALY
tel. ++39 0649913201, fax ++39 0644701007
riccip@uniroma1.it, Paoloemilio.Ricci@uniroma1.it
Orthogonal Polynomials and Special functions,
Numerical Analysis, Transforms, Operational
Calculus,
Differential and Difference equations

26) Cecil C. Rousseau
Department of Mathematical Sciences
The University of Memphis
Memphis, TN 38152, USA
tel. 901-678-2490, fax 901-678-2480
ccrousse@memphis.edu
Combinatorics, Graph Th.,
Asymptotic Approximations,
Applications to Physics

27) Tomasz Rychlik
Institute of Mathematics
Polish Academy of Sciences
Chopina 12, 87100 Torun, Poland
T.Rychlik@impan.gov.pl
Mathematical Statistics, Probabilistic
Inequalities

28) Bl. Sendov
Institute of Mathematics and Informatics
Bulgarian Academy of Sciences
Sofia 1090, Bulgaria

tel.515-294-8150
fink@math.iastate.edu
Inequalities, Ordinary Differential
Equations

10) Sorin Gal
Department of Mathematics
University of Oradea
Str. Armatei Romane 5
3700 Oradea, Romania
galso@uoradea.ro
Approximation Th., Fuzzyness, Complex
Analysis

11) Jerome A. Goldstein
Department of Mathematical Sciences
The University of Memphis,
Memphis, TN 38152, USA
tel. 901-678-2484
jgoldste@memphis.edu
Partial Differential Equations,
Semigroups of Operators

12) Heiner H. Gonska
Department of Mathematics
University of Duisburg
Duisburg, D-47048
Germany
tel. 0049-203-379-3542 office
gonska@informatik.uni-duisburg.de
Approximation Th., Computer Aided
Geometric Design

13) Dmitry Khavinson
Department of Mathematical Sciences
University of Arkansas
Fayetteville, AR 72701, USA
tel. (479) 575-6331, fax (479) 575-8630
dmitry@uark.edu
Potential Th., Complex Analysis, Holomorphic
PDE, Approximation Th., Function Th.

14) Virginia S. Kiryakova
Institute of Mathematics and Informatics
Bulgarian Academy of Sciences
Sofia 1090, Bulgaria
virginia@diogenes.bg
Special Functions, Integral Transforms,
Fractional Calculus

15) Hans-Bernd Knoop
Institute of Mathematics
Gerhard Mercator University
D-47048 Duisburg
Germany
tel. 0049-203-379-2676

bSENDov@bas.bg
Approximation Th., Geometry of Polynomials,
Image Compression

29) Igor Shevchuk
Faculty of Mathematics and Mechanics
National Taras Shevchenko
University of Kyiv
252017 Kyiv
UKRAINE
shevchuk@univ.kiev.ua
Approximation Theory

30) H.M. Srivastava
Department of Mathematics and Statistics
University of Victoria
Victoria, British Columbia V8W 3P4
Canada
tel. 250-721-7455 office, 250-477-6960 home,
fax 250-721-8962
harimsri@math.uvic.ca
Real and Complex Analysis, Fractional Calculus
and Appl.,
Integral Equations and Transforms, Higher
Transcendental
Functions and Appl., q-Series and q-Polynomials,
Analytic Number Th.

31) Ferenc Szidarovszky
Dept. Systems and Industrial Engineering
The University of Arizona
Engineering Building, 111
PO Box 210020
Tucson, AZ 85721-0020, USA
szidar@sie.arizona.edu
Numerical Methods, Game Th., Dynamic Systems,
Multicriteria Decision making,
Conflict Resolution, Applications
in Economics and Natural Resources
Management

32) Gancho Tachev
Dept. of Mathematics
Univ. of Architecture, Civil Eng. and Geodesy
1 Hr. Smirnenski blvd
BG-1421 Sofia, Bulgaria
Approximation Theory

33) Manfred Tasche
Department of Mathematics
University of Rostock
D-18051 Rostock
Germany
manfred.tasche@mathematik.uni-rostock.de
Approximation Th., Wavelet, Fourier Analysis,
Numerical Methods, Signal Processing,

knoop@math.uni-duisburg.de
Approximation Theory, Interpolation

16) Jerry Koliha
Dept. of Mathematics & Statistics
University of Melbourne
VIC 3010, Melbourne
Australia
koliha@unimelb.edu.au
Inequalities, Operator Theory,
Matrix Analysis, Generalized Inverses

17) Mustafa Kulenovic
Department of Mathematics
University of Rhode Island
Kingston, RI 02881, USA
kulenm@math.uri.edu
Differential and Difference Equations

18) Gerassimos Ladas
Department of Mathematics
University of Rhode Island
Kingston, RI 02881, USA
gladas@math.uri.edu
Differential and Difference Equations

Image Processing, Harmonic Analysis

34) Chris P. Tsokos
Department of Mathematics
University of South Florida
4202 E. Fowler Ave., PHY 114
Tampa, FL 33620-5700, USA
profcpt@math.usf.edu, profcpt@chuma1.cas.usf.edu
Stochastic Systems, Biomathematics,
Environmental Systems, Reliability Th.

35) Lutz Volkmann
Lehrstuhl II fuer Mathematik
RWTH-Aachen
Templergraben 55
D-52062 Aachen
Germany
volkm@math2.rwth-aachen.de
Complex Analysis, Combinatorics, Graph Theory

Power series of fuzzy numbers with cross product and applications to fuzzy differential equations

Adrian Ban* and Barnabás Bede**

*Department of Mathematics,
University of Oradea,
Str. Armatei Romane 5,
410087 Oradea, Romania
E-mail: aiban@uoradea.ro,

**Department of Mechanical and System Engineering,
Budapest Tech, Népszínház u. 8,
H-1081 Budapest, Hungary
E-mail: bede.barna@bgk.bmf.hu

Abstract

In a recent paper ([2]) the authors introduced and studied the main properties of the cross product of fuzzy numbers. The aim of this paper is to study power series with the cross product. The properties of convergence and the expected value of fuzzy power series are mainly approached. As an important application, a solution of the fuzzy differential equation $y'(x) = a \odot y(x) \oplus b(x)$, where a is a strict positive fuzzy number and b is a continuous fuzzy number valued mapping, is given.

2000 AMS Subject Classification: 26E50, 03E72, 34G10

Keywords and phrases: fuzzy number, cross product, power series, fuzzy differential equation

1 Introduction

Fuzzy numbers allow us to model the linguistic expressions appeared in different scientific areas mainly because of dependence on human judgement, finite resolution of measuring instruments, finite representation of numbers in computers. This explains the increasing interest on theoretical aspects of fuzzy arithmetic in the last years, especially directed to: operations over fuzzy numbers and properties ([4], [16], [19], [20], [22], [23], [24], [26], [27], [30], [31]), ranking of fuzzy numbers ([6], [11], [13], [28], [36]), canonical representation of fuzzy numbers ([8], [9], [12], [14], [17]).

Generally, the addition (or triangular norm-based addition) of fuzzy numbers and the scalar multiplication are considered. The operations of product type of fuzzy numbers are few studied because the known definitions lead to difficult handling.

The main disadvantage is that the shape of $L - R$ type fuzzy numbers is not preserved. In many papers (see e.g. [5], [33], [35]) approximative methods are given to estimate the result of the product of two fuzzy numbers, but these may lead to large computational errors after successive applications.

That is why in the paper [2] we introduced and studied a new product, called cross product, over fuzzy numbers and explained the advantages of the use of this one.

In this paper we prove that good results can be obtained on power series of fuzzy numbers with cross product.

The paper is divided as follows. In Section 2 we present some necessary basic concepts including the definition of a fuzzy number, addition and scalar multiplication of fuzzy numbers, distance between fuzzy numbers, definition and properties of the cross product of fuzzy numbers (proved in [2]). The third section contains the definition of a fuzzy power series with cross product as a particular fuzzy-number-valued function series, important examples and results on convergence and expected value of a fuzzy power series. As an application, in Section 4, the fuzzy differential equation $y'(x) = a \odot y(x) \oplus b(x)$, where a is a strict positive fuzzy number and b is a continuous fuzzy-number valued mapping with all the values strict positive or all strict negative is studied. Finally, a consequence of the main result of this section and an example are given.

2 Basic concepts

Firstly, let us recall the following parametric definition of a fuzzy number (see [10], [29]).

Definition 2.1. A fuzzy number is a pair $(\underline{u}, \bar{u})$ of functions $\underline{u}^r, \bar{u}^r, 0 \leq r \leq 1$ which satisfy the following requirements:

- (i) $\underline{u}^r$ is a bounded monotonic increasing left continuous function;
- (ii) $\bar{u}^r$ is a bounded monotonic decreasing left continuous function;
- (iii) $\underline{u}^r \leq \bar{u}^r$, for every $r \in [0, 1]$.

We denote by $\mathbb{R}_{\mathcal{F}}$ the set of all fuzzy numbers. Obviously $\mathbb{R} \subset \mathbb{R}_{\mathcal{F}}$ (here $\mathbb{R}$ is understood as $\mathbb{R} = \{\tilde{x} = \chi_{\{x\}}; x \text{ is usual real number}\}$). For arbitrary $u = (\underline{u}, \bar{u}) \in \mathbb{R}_{\mathcal{F}}, v = (\underline{v}, \bar{v}) \in \mathbb{R}_{\mathcal{F}}$ and $k \in \mathbf{R}$ we define the addition of u and v , $u \oplus v$, and the scalar multiplication of u by k , $k \cdot u$ or $u \cdot k$, as

$$\begin{aligned} (u \oplus v)^r &= \underline{u}^r + \underline{v}^r, \\ (\overline{u \oplus v})^r &= \bar{u}^r + \bar{v}^r \end{aligned}$$

$$\begin{aligned} (k \cdot u)^r &= \begin{cases} k\underline{u}^r, & \text{if } k \geq 0, \\ k\bar{u}^r, & \text{if } k < 0, \end{cases} \\ (\overline{k \cdot u})^r &= \begin{cases} k\bar{u}^r, & \text{if } k \geq 0, \\ k\underline{u}^r, & \text{if } k < 0. \end{cases} \end{aligned}$$

We denote by $\ominus u = (-1) \cdot u$ the symmetric of $u \in \mathbb{R}_{\mathcal{F}}$ and $\frac{1}{\mu} \cdot u$ by $\frac{u}{\mu}$ if $\mu \neq 0$.

A fuzzy number $u \in \mathbb{R}_{\mathcal{F}}$ is said to be positive if $\underline{u}^1 \geq 0$, strict positive if $\underline{u}^1 > 0$, negative if $\bar{u}^1 \leq 0$ and strict negative if $\bar{u}^1 < 0$. We say that u and v have the same sign if they are both strict positive or both strict negative. If u is positive (negative) then $\ominus u$ is negative (positive).

The following definitions and results are useful in the paper.

Definition 2.2. For arbitrary fuzzy numbers $u = (\underline{u}, \bar{u})$ and $v = (\underline{v}, \bar{v})$ the quantity

$$D(u, v) = \sup_{0 \leq r \leq 1} \{\max\{|\underline{u}^r - \underline{v}^r|, |\bar{u}^r - \bar{v}^r|\}\}$$

is called the distance between u and v .

It is well-known (see e.g. [7]) that $(\mathbb{R}_{\mathcal{F}}, D)$ is a complete metric space and D has the following properties:

- (i) $D(u \oplus w, v \oplus w) = D(u, v)$, for all $u, v, w \in \mathbb{R}_{\mathcal{F}}$;
- (ii) $D(k \cdot u, k \cdot v) = |k| D(u, v)$, for all $u, v \in \mathbb{R}_{\mathcal{F}}$ and $k \in \mathbb{R}$;
- (iii) $D(u \oplus v, w \oplus t) \leq D(u, w) + D(v, t)$, for all $u, v, w, t \in \mathbb{R}_{\mathcal{F}}$.

If we denote $\|u\|_{\mathcal{F}} = D(u, \tilde{0})$, $\forall u \in \mathbb{R}_{\mathcal{F}}$, where $\tilde{0} = \chi_{\{0\}}$ is the neutral element with respect to $\oplus$, then $\|\cdot\|_{\mathcal{F}}$ has the properties of an usual norm on $\mathbb{R}_{\mathcal{F}}$, i.e., $\|u\|_{\mathcal{F}} = \tilde{0}$ if and only if $u = \tilde{0}$, $\|k \cdot u\|_{\mathcal{F}} = |k| \|u\|_{\mathcal{F}}$, $\|u \oplus v\|_{\mathcal{F}} \leq \|u\|_{\mathcal{F}} + \|v\|_{\mathcal{F}}$, $|\|u\|_{\mathcal{F}} - \|v\|_{\mathcal{F}}| \leq D(u, v)$.

Definition 2.3. (see [1]) A fuzzy-number-valued function $f : [a, b] \rightarrow \mathbb{R}_{\mathcal{F}}$ is called Riemann integrable to $I \in \mathbb{R}_{\mathcal{F}}$, if for any $\varepsilon > 0$, there exists $\delta > 0$ such that for any division $P = \{[u, v], \xi\}$ of $[a, b]$ with the norm $\Delta(P) < \delta$, we have

$$D\left(\sum_P^{\oplus} (v - u) \odot f(\xi), I\right) < \varepsilon,$$

where $\sum^{\oplus}$ is the addition with respect to $\oplus$ in $\mathbb{R}_{\mathcal{F}}$.

We write $I = (FR) \int_a^b f(x) dx$.

Let $x, y \in \mathbb{R}_{\mathcal{F}}$. If there exists $z \in \mathbb{R}_{\mathcal{F}}$ such that $x = y \oplus z$, then we call z the H-difference of x and y , denoted by $x - y$.

Definition 2.4. ([18], [32]) A function $f : A \rightarrow \mathbb{R}_{\mathcal{F}}$ is said to be differentiable at $x \in A$ if there exists $f'(x) \in \mathbb{R}_{\mathcal{F}}$ such that the limits $\lim_{h \searrow 0} \frac{f(x+h) - f(x)}{h}$ and $\lim_{h \searrow 0} \frac{f(x) - f(x-h)}{h}$ exist and are equal to $f'(x)$. We call $f'(x)$ the derivative of f at x . If f is differentiable at any $x \in A$ we call f differentiable.

Lemma 2.5. (i) Let $f : [a, b] \rightarrow \mathbb{R}_{\mathcal{F}}$ be differentiable at $c \in [a, b]$. Then f is continuous at c (see [1, Lemma 5, Section 1]).

(ii) Let $f, g : (a, b) \rightarrow \mathbb{R}_{\mathcal{F}}$ be continuous functions. If the H-difference function $f - g$ exists on (a, b) then $f - g$ is continuous on (a, b) (see [1, Lemma 1, Section 2]).

(iii) Let I be an open interval of $\mathbb{R}$, and let $f, g : I \rightarrow \mathbb{R}_{\mathcal{F}}$ be differentiable functions with the derivatives f', g' . Then $(f \oplus g)'$ exists and $(f \oplus g)' = f' \oplus g'$ (see [1, Proposition 5, Section 2]).

The concept of expected value of a fuzzy number was introduced and studied in [12]-[14] as a good representation of a fuzzy number as a crisp

number. In our notations, if u is a fuzzy number then the expected value $EV(u)$ is given by (see Definition 2 and Lemma 3 in [12])

$$EV(u) = \frac{1}{2} \int_0^1 (\underline{u}^r + \bar{u}^r) dr.$$

It is obvious that $EV(\ominus u) = -EV(u)$, for every fuzzy number u .

Let us denote $\mathbb{R}_{\mathcal{F}}^* = \{u \in \mathbb{R}_{\mathcal{F}}; u \text{ is strict positive or strict negative}\}$ and let $\mathbb{R}_{\mathcal{F}}^{\odot} = \mathbb{R}_{\mathcal{F}}^* \cup \{\tilde{0}\}$.

Definition 2.6. ([2]) The binary operation $\odot$ on $\mathbb{R}_{\mathcal{F}}^{\odot}$ defined by $w = u \odot v$, $[w]^r = [\underline{w}^r, \bar{w}^r]$, for every $r \in [0, 1]$, where

$$\underline{w}^r = \begin{cases} \underline{u}^r \underline{v}^1 + \underline{u}^1 \underline{v}^r - \underline{u}^1 \underline{v}^1, & \text{if } \underline{u}^1 > 0 \text{ and } \underline{v}^1 > 0 \\ \bar{u}^r \underline{v}^1 + \bar{u}^1 \underline{v}^r - \bar{u}^1 \underline{v}^1, & \text{if } \underline{u}^1 < 0 \text{ and } \bar{v}^1 < 0 \\ \underline{u}^r \bar{v}^1 + \underline{u}^1 \bar{v}^r - \underline{u}^1 \bar{v}^1, & \text{if } \bar{u}^1 < 0 \text{ and } \underline{v}^1 > 0 \\ \bar{u}^r \bar{v}^1 + \bar{u}^1 \bar{v}^r - \bar{u}^1 \bar{v}^1, & \text{if } \bar{u}^1 < 0 \text{ and } \bar{v}^1 < 0 \end{cases}$$

and

$$\bar{w}^r = \begin{cases} \bar{u}^r \bar{v}^1 + \bar{u}^1 \bar{v}^r - \bar{u}^1 \bar{v}^1, & \text{if } \underline{u}^1 > 0 \text{ and } \underline{v}^1 > 0 \\ \underline{u}^r \bar{v}^1 + \underline{u}^1 \bar{v}^r - \underline{u}^1 \bar{v}^1, & \text{if } \underline{u}^1 > 0 \text{ and } \bar{v}^1 < 0 \\ \bar{u}^r \underline{v}^1 + \bar{u}^1 \underline{v}^r - \bar{u}^1 \underline{v}^1, & \text{if } \bar{u}^1 < 0 \text{ and } \underline{v}^1 > 0 \\ \underline{u}^r \underline{v}^1 + \underline{u}^1 \underline{v}^r - \underline{u}^1 \underline{v}^1, & \text{if } \bar{u}^1 < 0 \text{ and } \bar{v}^1 < 0 \end{cases}$$

for any $u, v \in \mathbb{R}_{\mathcal{F}}^*$ and $u \odot \tilde{0} = \tilde{0} \odot u = \tilde{0}$ for any $u \in \mathbb{R}_{\mathcal{F}}^{\odot}$ is called the cross product of fuzzy numbers.

Remark 2.7. 1) Starting from the formulas of calculus in the case u and v strict positive we can deduce: $u \odot v = \ominus(u \odot (\ominus v))$ if u is strict positive and v is strict negative, $u \odot v = \ominus((\ominus u) \odot v)$ if u is strict negative and v is strict positive and $u \odot v = (\ominus u) \odot (\ominus v)$ if u and v are strict negative. The cross product $u \odot v$ is strict positive if u and v have the same sign and strict negative contrariwise.

2) The cross product extends the scalar multiplication of fuzzy numbers because $\mathbb{R} \subset \mathbb{R}_{\mathcal{F}}^{\odot}$.

3) Similar operation can be defined on the set $\mathbb{R}_{\mathcal{F}}^{\#} = \{u \in \mathbb{R}_{\mathcal{F}}; \text{there exist a unique } x_0 \in \mathbb{R} \text{ such that } u(x_0) = 1\}$ as follows $w = u \odot v$, $[w]^r = [\underline{w}^r, \bar{w}^r]$, for every $r \in [0, 1]$, where $\underline{w}^r = \underline{u}^r \underline{v}^1 + \underline{u}^1 \underline{v}^r - \underline{u}^1 \underline{v}^1$ and $\bar{w}^r = \bar{u}^r \bar{v}^1 + \bar{u}^1 \bar{v}^r - \bar{u}^1 \bar{v}^1$ for every u, v positive fuzzy numbers. Let also, $u \odot v = \ominus(u \odot (\ominus v))$ if u

is positive and v is negative, $u \odot v = \ominus((\ominus u) \odot v)$ if u is negative and v is positive and $u \odot v = (\ominus u) \odot (\ominus v)$ if u and v are negative. We observe that the results of this paper hold with respect to this operation too. The main algebraic properties of the cross product are given by the following.

Theorem 2.8. ([2]) If $u, v, w \in \mathbb{R}_{\mathcal{F}}^{\odot}$ then

- (i) $(\ominus u) \odot v = u \odot (\ominus v) = \ominus(u \odot v)$;
- (ii) $u \odot v = v \odot u$;
- (iii) $(u \odot v) \odot w = u \odot (v \odot w)$;
- (iv) If u and v are strict positive or u and v are strict negative then $(u \oplus v) \odot w = (u \odot w) \oplus (v \odot w)$.

We present other important properties of the cross product.

Theorem 2.9. ([2]) (i) If u, v have the same sign and $w \in \mathbb{R}_{\mathcal{F}}^{\odot}$ then

$$D(w \odot u, w \odot v) \leq K_w D(u, v),$$

where $K_w = \max\{|\bar{w}^1|, |\underline{w}^1|\} + \bar{w}^0 - \underline{w}^0$.

- (ii) If $u, v \in \mathbb{R}_{\mathcal{F}}^{\#}$ and we denote $\underline{u}^1 = \bar{u}^1 = u^1$, $\underline{v}^1 = \bar{v}^1 = v^1$ then

$$EV(u \odot v) = u^1 EV(v) + v^1 EV(u) - u^1 v^1.$$

- (iii) If $u \in \mathbb{R}_{\mathcal{F}}^{\#}$ then

$$EV(u^{\odot n}) = n(u^1)^{n-1} EV(u) - (n-1)(u^1)^n$$

for every $n \in \mathbb{N}^*$, where $u^{\odot n} = \underbrace{u \odot \dots \odot u}_{n \text{ times}}$.

- (iv) If (u_n) is a sequence of strict positive or strict negative fuzzy numbers such that $u_n \rightarrow u \in \mathbb{R}_{\mathcal{F}}^{\odot}$ then $c \odot u_n \rightarrow c \odot u$, for every $c \in \mathbb{R}_{\mathcal{F}}^{\odot}$ (the convergence in the metric D is considered).

- (v) If $f, g : [a, b] \rightarrow \mathbb{R}_{\mathcal{F}}$ are continuous at $x_0 \in [a, b]$ and $f(x_0), g(x_0)$ are strict positive or strict negative then $f \odot g$ is continuous at x_0 .

- (vi) If $f, g : [a, b] \rightarrow \mathbb{R}_{\mathcal{F}}$ are differentiable at $x_0 \in [a, b]$, $f(x_0), g(x_0)$ are strict positive or strict negative and for sufficiently small $h > 0$ the H -differences $f(x_0 + h) - f(x_0)$ and $f(x_0) - f(x_0 - h)$ exist and have the same sign as $f(x_0)$, the H -differences $g(x_0 + h) - g(x_0)$ and $g(x_0) - g(x_0 - h)$ exist and have the same sign as $g(x_0)$, then $f \odot g$ is differentiable at x_0 and

$$(f \odot g)'(x_0) = f(x_0) \odot g'(x_0) \oplus f'(x_0) \odot g(x_0).$$

(vii) If $f : [a, b] \rightarrow \mathbb{R}_{\mathcal{F}}$ is differentiable at $x_0 \in [a, b]$, $f(x_0)$ is strict positive or strict negative and for sufficiently small $h > 0$ the H -differences $f(x_0 + h) - f(x_0)$ and $f(x_0) - f(x_0 - h)$ exist and have the same sign as $f(x_0)$ then $c \odot f$ is differentiable at x_0 for any $c \in \mathbb{R}_{\mathcal{F}}^{\odot}$ and

$$(c \odot f)'(x_0) = c \odot f'(x_0).$$

3 Power series with cross product

A series of fuzzy-number-valued functions is a function series $\sum_n^{\oplus} f_n$, where $f_n : A(\subseteq \mathbb{R} \text{ or } \mathbb{R}_{\mathcal{F}}) \rightarrow \mathbb{R}_{\mathcal{F}}$, and $\sum^{\oplus}$ means the countable sum with respect to the addition of fuzzy numbers. The concepts of convergence and uniform convergence are considered in the metric space $(\mathbb{R}_{\mathcal{F}}, D)$. We say that the series $\sum_n^{\oplus} f_n$ is absolutely (uniformly) convergent if $\sum_n \|f_n\|_{\mathcal{F}}$ is (uniformly) convergent. Because $(\mathbb{R}_{\mathcal{F}}, \oplus, \cdot)$ is not a linear space over $\mathbb{R}$, $(\mathbb{R}_{\mathcal{F}}, \|\cdot\|_{\mathcal{F}})$ is not a normed space. We need to prove some lemmas which are analogous to some classical results in normed spaces.

Lemma 3.1. (i) If the series $\sum_n^{\oplus} f_n$ is absolutely convergent on a set A then it is convergent on A .

(ii) If the series $\sum_n^{\oplus} f_n$ is absolutely uniformly convergent on a set A then it is uniformly convergent on A .

Proof. (i) Let $x \in A$ and $\sum_n \|f_n(x)\|_{\mathcal{F}}$ be convergent. By [34, Theorem 1] we have

$$\left[\sum_n^{\oplus} f_n(x) \right]^r = \sum_n [f_n(x)]^r = \left[\sum_n \underline{f_n(x)}^r, \sum_n \overline{f_n(x)}^r \right].$$

Then

$$\sum_n \|f_n(x)\|_{\mathcal{F}} = \sum_n D(f_n(x), \tilde{0}) = \sum_n \sup_{r \in [0,1]} \max\{|\underline{f_n(x)}^r|, |\overline{f_n(x)}^r|\}.$$

For any $r \in [0, 1]$ we obtain

$$\left| \sum_n \underline{f_n(x)}^r \right| \leq \sum_n |\underline{f_n(x)}^r| \leq \sum_n \|f_n(x)\|_{\mathcal{F}}.$$

Since the last series converges, it follows that $\sum_n \overline{f_n(x)}^r$ converges. The same reasoning implies that $\sum_n \overline{f_n(x)}^r$ converges.

(ii) For any $\varepsilon > 0$ there exists $N(\varepsilon) \in \mathbb{N}$ such that

$$\|f_{n+1}(x) \oplus \dots \oplus f_{n+p}(x)\|_{\mathcal{F}} \leq \|f_{n+1}(x)\|_{\mathcal{F}} + \dots + \|f_{n+p}(x)\|_{\mathcal{F}} < \varepsilon,$$

$\forall n \geq N(\varepsilon), p \geq 1, \forall x \in A$. The rest is obvious by Cauchy criterion and the invariance to translation of the Hausdorff metric D . $\square$

Lemma 3.2. *If $f_n : A(\subset \mathbb{R} \text{ or } \mathbb{R}_{\mathcal{F}}) \rightarrow \mathbb{R}_{\mathcal{F}}$ are continuous functions and if the series of fuzzy-number-valued functions $\sum_n^{\oplus} f_n$ is uniformly convergent to $f : A \rightarrow \mathbb{R}_{\mathcal{F}}$ then f is continuous.*

Proof. By the properties of the distance D we have for $n \in \mathbb{N}$

$$\begin{aligned} D(f(x), f(y)) &\leq D\left(f(x), \sum_{k=0}^n^{\oplus} f_k(x)\right) \\ &+ D\left(\sum_{k=0}^n^{\oplus} f_k(x), \sum_{k=0}^n^{\oplus} f_k(y)\right) + D\left(\sum_{k=0}^n^{\oplus} f_k(y), f(y)\right) \\ &\leq D\left(f(x), \sum_{k=0}^n^{\oplus} f_k(x)\right) + \sum_{k=0}^n D(f_k(x), f_k(y)) \\ &+ D\left(\sum_{k=0}^n^{\oplus} f_k(y), f(y)\right) \end{aligned}$$

for every $x, y \in A$ and $n \in \mathbb{N}$.

The property of uniform convergence applied to the first and the last term and the continuity of each function f_k , $k = 0, \dots, n$ complete the proof. $\square$

In what follows we study fuzzy power series as particular series of fuzzy-number-valued functions by using the cross product of fuzzy numbers.

Definition 3.3. The function series $\sum_{n=0}^{\infty} a_n \odot x^{\odot n}$, where $a_n, x \in \mathbb{R}_{\mathcal{F}}^{\odot}$, $n \in \mathbb{N}$ and $x^{\odot n}$ denotes $\underbrace{x \odot \dots \odot x}_{n \text{ times}}$ if $n \geq 1$ and by convention $x^{\odot 0} = \tilde{1} = \chi_{\{1\}}$, is called $\odot$ -power series.

The following theorem is an analogous of Abel's first theorem.

Theorem 3.4. Let $\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n}$ be a $\odot$ -power series. Then there exist $R_1, R_2 \in \mathbb{R}$, $R_1 \leq R_2$ such that

(i) $\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n}$ converges on $B(\tilde{0}, R_1) \cap (\mathbb{R}_{\mathcal{F}}^{\odot})$, where $B(\tilde{0}, R_1) = \{x \in \mathbb{R}_{\mathcal{F}}; \|x\|_{\mathcal{F}} < R_1\}$;

(ii) $\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n}$ converges uniformly on $\overline{B(\tilde{0}, R)} \cap (\mathbb{R}_{\mathcal{F}}^{\odot})$, $\forall R < R_1$ ($\overline{A}$ means the closure of the set A);

(iii) If a_n and x have the same sign then $\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n}$ diverges for every x with $|\bar{x}^1| > R_2$ or $|\underline{x}^1| > R_2$.

Proof. (i) For $x \in \mathbb{R}_{\mathcal{F}}^*$, by Theorem 2.9, (i) we have:

$$\|a_n \odot x^{\odot n}\|_{\mathcal{F}} = D(a_n \odot x^{\odot n}, \tilde{0} \odot x^{\odot n}) \leq K_x^n D(a_n, 0) = K_x^n \|a_n\|_{\mathcal{F}} \quad (1)$$

where $K_x = \max\{|\bar{x}^1|, |\underline{x}^1|\} + \bar{x}^0 - \underline{x}^0$. The case $x = \tilde{0}$ is obvious. Let us denote by r_1 the ray of convergence of the classical power series $\sum_{n=0}^{\infty} \|a_n\|_{\mathcal{F}} u^n$. Then the series $\sum_{n=0}^{\infty} \|a_n\|_{\mathcal{F}} u^n$ converges for $u \in (-r_1, r_1)$ and converges uniformly on $[-r, r]$ for all $r < r_1$. Let $R_1 = \frac{1}{3}r_1$ and $x \in B(\tilde{0}, R_1) \subseteq \mathbb{R}_{\mathcal{F}}$ i.e. $\|x\|_{\mathcal{F}} < R_1$. Then

$$\sup_{r \in [0,1]} \max\{|\underline{x}^r|, |\bar{x}^r|\} < R_1$$

therefore

$$K_x \leq \max\{|\bar{x}^1|, |\underline{x}^1|\} + |\bar{x}^0| + |\underline{x}^0| < 3R_1 = r_1.$$

By (1) and Lemma 3.1, (i) it follows that, for $x \in B(\tilde{0}, R_1)$, the series $\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n}$ converges.

(ii) It follows from Lemma 3.1, (ii) and the above proof.

(iii) Let $\bar{R}_2$ be the ray of convergence of the series $\sum_n \bar{a}_n^1 u^n$ and $\underline{R}_2$ be the ray of convergence of the series $\sum_n \underline{a}_n^1 u^n$. We denote $R_2 = \max\{\bar{R}_2, \underline{R}_2\}$. Let $|\bar{x}^1| > R_2$ or $|\underline{x}^1| > R_2$ and let us suppose that the $\odot$ -power series $\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n}$ converges. If a_n and x are positive, then by [34, Theorem 1]

we have

$$\left[\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n} \right]^1 = \left[\sum_{n=0}^{\infty} \underline{a}_n^1 (\underline{x}^1)^n, \sum_{n=0}^{\infty} \bar{a}_n^1 (\bar{x}^1)^n \right].$$

But at least one of the series $\sum_{n=0}^{\infty} \underline{a}_n^1 (\underline{x}^1)^n$ and $\sum_{n=0}^{\infty} \bar{a}_n^1 (\bar{x}^1)^n$ diverges, which is a contradiction. As a conclusion $\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n}$ is divergent. $\square$

Remark 3.5. On $\mathbb{R}_{\mathcal{F}}^{\#}$ a similar result can be obtained.

Corollary 3.6. *If $a_n \in \mathbb{R}_{\mathcal{F}}^{\odot}$, $n \in \mathbb{N}$ and $x \in \mathbb{R}$ then there exists R_1 such that the series $\sum_{n=0}^{\infty} \oplus a_n \cdot x^n$ converges on $(-R_1, R_1)$, converges uniformly on $[-R, R]$, $\forall R < R_1$ and diverges on $\mathbb{R} \setminus [-R_1, R_1]$.*

Proof. Since $\|a_n \cdot x^n\|_{\mathcal{F}} = |x|^n \cdot \|a_n\|_{\mathcal{F}}$, choosing R_1 the ray of convergence of the crisp power series $\sum_{n=0}^{\infty} \|a_n\|_{\mathcal{F}} u^n$, we obtain the statement of the corollary. $\square$

Remark 3.7. If $\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n}$ converges uniformly on a set $A \subset \mathbb{R}_{\mathcal{F}}^{\odot}$ then it is continuous on A . Indeed, by Theorem 2.9, (i) and (v) we obtain that $a_n \odot x^{\odot n}$ is continuous as a function of x and by Lemma 3.2 it follows that $\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n}$ is continuous.

Remark 3.8. If $\sum_{n=0}^{\infty} \|a_n\|_{\mathcal{F}} u^n$ converges for any $u \in \mathbb{R}$ then $\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n}$ converges for all $x \in \mathbb{R}_{\mathcal{F}}^{\odot}$ and the convergence is uniform on any closed ball in this set.

Theorem 3.9. *If $\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n}$ is a $\odot$ -power series with $a_n \neq \tilde{0}$, for every $n \in \mathbb{N}$, then the number R_1 in Theorem 3.4 is $R_1 = \frac{1}{3} \lim_{n \rightarrow \infty} (\|a_n\|_{\mathcal{F}})^{-\frac{1}{n}}$ or $R_1 = \frac{1}{3} \lim_{n \rightarrow \infty} \frac{\|a_n\|_{\mathcal{F}}}{\|a_{n+1}\|_{\mathcal{F}}}$.*

Proof. Directly from Theorem 3.4

and Abel's second Theorem for crisp series. $\square$

The following examples of $\odot$ -power series define some fuzzy analogous of exponential, logarithmic, sine and cosine functions.

Example 3.10. (i) *The series $\sum_{n=0}^{\infty} \oplus \frac{1}{n!} \cdot x^{\odot n}$ converges for any $x \in \mathbb{R}_{\mathcal{F}}^{\odot}$ because the number R_1 in Theorem 3.4 is $R_1 = \frac{1}{3} \lim_{n \rightarrow \infty} \frac{\|a_n\|_{\mathcal{F}}}{\|a_{n+1}\|_{\mathcal{F}}} = +\infty$. We denote its sum by $e_{\odot}^x$.*

(ii) *The series $\sum_{n=1}^{\infty} \oplus \frac{(-1)^{n+1}}{n} \cdot x^{\odot n}$ converges for any strict positive or strict negative fuzzy number $x \in B(\tilde{0}, \frac{1}{3})$ because $\lim_{n \rightarrow \infty} \frac{\|a_n\|_{\mathcal{F}}}{\|a_{n+1}\|_{\mathcal{F}}} = 1$. We denote its sum by $\ln_{\odot}(1+x)$.*

(iii) *The series $\sum_{n=0}^{\infty} \oplus \frac{(-1)^n}{(2n+1)!} \cdot x^{\odot(2n+1)}$ converges for any $x \in \mathbb{R}_{\mathcal{F}}^{\odot}$ and we denote its sum by $\sin_{\odot} x$.*

(iv) *The series $\sum_{n=0}^{\infty} \oplus \frac{(-1)^n}{(2n)!} \cdot x^{\odot(2n)}$ converges for any $x \in \mathbb{R}_{\mathcal{F}}^{\odot}$ and we denote its sum by $\cos_{\odot} x$.*

Similar results can be obtained for the operation $\odot$ defined on $\mathbb{R}_{\mathcal{F}}^{\#}$.

In the following theorem we compute the expected value of a $\odot$ -power series.

Theorem 3.11. (i) If the power series $\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n}$ converges, where $a_n, x \in \mathbb{R}_{\mathcal{F}}^{\#}$, then

$$EV \left(\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n} \right) = EV(a_0) + \sum_{n=1}^{\infty} (a_n^1 n (x^1)^{n-1} (EV(x) - x^1) + EV(a_n) (x^1)^n);$$

(ii) If, in addition, a_n are symmetric fuzzy numbers (i.e. $a_n^1 - \underline{a_n}^r = \overline{a_n}^r - a_n^1$ for all $r \in [0, 1]$) then

$$EV \left(\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n} \right) = EV(a_0) + \sum_{n=1}^{\infty} EV(a_n) [n (x^1)^{n-1} EV(x) - (n-1) (x^1)^n];$$

(iii) If, in addition, x is a symmetric fuzzy number then

$$EV \left(\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n} \right) = \sum_{n=0}^{\infty} EV(a_n) (EV(x))^n.$$

Proof. (i) We observe that if (u_n) is a sequence of fuzzy numbers converging to u then $EV(u_n) \rightarrow EV(u)$. We get

$$EV \left(\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n} \right) = \lim_{n \rightarrow \infty} EV \left(\sum_{k=0}^n \oplus a_k \odot x^{\odot k} \right) = \lim_{n \rightarrow \infty} \sum_{k=0}^n EV(a_k \odot x^{\odot k}).$$

Then by Theorem 2.9, (ii) we have

$$EV \left(\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n} \right) = \lim_{n \rightarrow \infty} \sum_{k=0}^n (a_k^1 EV(x^{\odot k}) + (EV(a_k) - a_k^1) (x^1)^k).$$

By Theorem 2.9, (iii) we obtain

$$EV \left(\sum_{n=0}^{\infty} \oplus a_n \odot x^{\odot n} \right)$$

$$\begin{aligned}
&= EV(a_0) + \lim_{n \rightarrow \infty} \sum_{k=1}^n (a_k^1 [k(x^1)^{k-1} EV(x) - (k-1)(x^1)^k] + (EV(a_k) - a_k^1)(x^1)^k) \\
&= EV(a_0) + \lim_{n \rightarrow \infty} \sum_{k=1}^n (a_k^1 k(x^1)^{k-1} (EV(x) - x^1) + EV(a_k)(x^1)^k),
\end{aligned}$$

which proves (i). For (ii) we observe that $EV(a_n) = a_n^1$, $\forall n \in \mathbb{N}$, and for (iii) we have $EV(x) = x^1$. $\square$

Example 3.12. We compute the expected values of the $\odot$ -power series in Example 3.10. We have

$$\begin{aligned}
EV(e_{\odot}^x) &= EV\left(\sum_{n=0}^{\infty} \frac{1}{n!} \cdot x^{\odot n}\right) = 1 + \sum_{n=1}^{\infty} \frac{1}{n!} [n(x^1)^{n-1} EV(x) - (n-1)(x^1)^n] \\
&= 1 + EV(x) \sum_{n=1}^{\infty} \frac{(x^1)^{n-1}}{(n-1)!} - \sum_{n=1}^{\infty} \frac{(n-1)(x^1)^n}{n!} \\
&= EV(x)e^{x^1} - x^1 e^{x^1} + e^{x^1} = e^{x^1} (1 + EV(x) - x^1),
\end{aligned}$$

for any $x \in \mathbb{R}_{\mathcal{F}}^{\#}$.

We observe that if x is symmetric then $EV(e_{\odot}^x) = e^{EV(x)}$.

Also,

$$\begin{aligned}
EV(\ln_{\odot}(1+x)) &= \sum_{n=1}^{\infty} \frac{(-1)^{n+1}}{n} [n(x^1)^{n-1} EV(x) - (n-1)(x^1)^n] \\
&= \sum_{n=1}^{\infty} (-1)^{n+1} (EV(x) - x^1)(x^1)^{n-1} + \ln(1+x^1) \\
&= \frac{EV(x) - x^1}{1 - x^1} + \ln(1+x^1),
\end{aligned}$$

for any $x \in \mathbb{R}_{\mathcal{F}}^{\#} \cap B(\tilde{0}, \frac{1}{3})$. In addition, if x is symmetric then $EV(\ln_{\odot}(1+x)) = \ln(1+EV(x))$.

It is easy to check that

$$EV(\sin_{\odot} x) = (EV(x) - x^1) \cos x^1 + \sin x^1$$

and

$$EV(\cos_{\odot} x) = (x^1 - EV(x)) \sin x^1 + \cos x^1,$$

for every $x \in \mathbb{R}_{\mathcal{F}}^{\#}$.

For almost symmetric fuzzy numbers, or fuzzy numbers with small support, the expected values of the power series presented in Example 3.10 are near to $e^{EV(x)}$, $\ln(1 + EV(x))$, $\sin(EV(x))$, $\cos(EV(x))$, respectively. This suggest us the possible use of these series as definitions for fuzzy exponential, logarithmic, sine and cosine functions.

4 Application to fuzzy differential equations

In this section we study the fuzzy differential equation

$$y'(x) = a \odot y(x) \oplus b(x), \quad (2)$$

where a is a strict positive fuzzy number and $b : \mathbb{R} \rightarrow \mathbb{R}_{\mathcal{F}}$ is continuous and all its values are strict positive or all are strict negative.

First we prove the following lemma on differentials of function sequences.

Lemma 4.1. *Let $(f_n)_{n \in \mathbb{N}}$, $f_n : [a, b] \rightarrow \mathbb{R}_{\mathcal{F}}$ be a sequence of differentiable functions converging uniformly to $f : [a, b] \rightarrow \mathbb{R}_{\mathcal{F}}$. If there exists $\beta > 0$ such that the H-differences $f_n(x + h) - f_n(x)$, $f_n(x) - f_n(x - h)$ exist for all $0 < h < \beta$ and $n \in \mathbb{N}$ and if f'_n is uniformly convergent to g then f is differentiable and $f' = g$.*

Proof. Let $x \in [a, b]$. Since f_n is differentiable, the H-differences $f_n(x + h) - f_n(x)$, $n \in \mathbb{N}$ exist for sufficiently small h , i.e., there exists $\beta_n \in \mathbb{R}$ such that for $0 < h < \beta_n$ the H-differences $f_n(x + h) - f_n(x)$, $n \in \mathbb{N}$ exist. We can choose β_n such that $\inf_{n \in \mathbb{N}} \beta_n \geq \beta$. Then for $0 < h < \beta$ there exists $u_n(h)$ such that $u_n(h) \oplus f_n(x) = f_n(x + h)$, for every $n \in \mathbb{N}$. We prove that the function sequence (u_n) is uniformly convergent on $(0, \beta)$. Indeed, by the invariance to translation of the Hausdorff distance we obtain

$$\begin{aligned} D(u_n(h), u_m(h)) &= D(u_n(h) \oplus f_n(x), u_m(h) \oplus f_n(x)) \\ &= D(f_n(x + h), u_m(h) \oplus f_n(x)) \\ &= D(f_n(x + h) \oplus f_m(x), u_m(h) \oplus f_m(x) \oplus f_n(x)) \\ &= D(f_n(x + h) \oplus f_m(x), f_m(x + h) \oplus f_n(x)), \end{aligned}$$

for every $m, n \in \mathbb{N}$ and $h \in (0, \beta)$. Then by the property (iii) of D (see Section 2) we obtain

$$D(u_n(h), u_m(h)) \leq D(f_m(x), f_n(x)) + D(f_n(x + h), f_m(x + h)).$$

The uniform convergence of the sequence (f_n) implies (u_n) uniformly fundamental. Since $(\mathbb{R}_{\mathcal{F}}, D)$ is complete it follows that (u_n) is uniformly convergent. Let us denote u the limit of the function sequence (u_n) , that is $\lim_{n \rightarrow \infty} u_n(h) = u(h)$ for every $h \in (0, \beta)$. Since $f_n(x+h) = f_n(x) \oplus u_n(h)$ for every $n \in \mathbb{N}$ and $h \in (0, \beta)$, we have

$$\begin{aligned} D(f(x+h), f(x) \oplus u(h)) &\leq \\ D(f(x+h), f_n(x+h)) + D(f_n(x) \oplus u_n(h), f(x) \oplus u(h)) \\ &\leq D(f(x+h), f_n(x+h)) + D(f_n(x), f(x)) + D(u_n(h), u(h)), \end{aligned}$$

for every $n \in \mathbb{N}$ and $h \in (0, \beta)$. Passing to limit with $n \rightarrow \infty$ we obtain that $f(x) \oplus u(h) = f(x+h)$ for $0 < h < \beta$, i.e., the H-difference $f(x+h) - f(x)$ exists for $0 < h < \beta$. Also, we get

$$\begin{aligned} D\left(g(x), \frac{f(x+h) - f(x)}{h}\right) &\leq D(g(x), f'_n(x)) \\ &+ D\left(f'_n(x), \frac{f_n(x+h) - f_n(x)}{h}\right) \\ &+ D\left(\frac{f_n(x+h) - f_n(x)}{h}, \frac{f(x+h) - f(x)}{h}\right), \end{aligned} \quad (3)$$

for every $n \in \mathbb{N}$ and $h \in (0, \beta)$. The properties of the Hausdorff distance D in Section 2 imply

$$\begin{aligned} D\left(\frac{f_n(x+h) - f_n(x)}{h}, \frac{f(x+h) - f(x)}{h}\right) &= \\ &= \frac{1}{h} D(f_n(x+h) - f_n(x), f(x+h) - f(x)) = \\ &= \frac{1}{h} D(f_n(x+h) \oplus f(x), f(x+h) \oplus f_n(x)) \leq \\ &\leq \frac{1}{h} [D(f_n(x+h), f(x+h)) + D(f_n(x), f(x))], \end{aligned} \quad (4)$$

for every $n \in \mathbb{N}$ and $h \in (0, \beta)$. Passing to limit with $n \rightarrow \infty$ in (3) and taking into account (4) we obtain

$$D\left(g(x), \frac{f(x+h) - f(x)}{h}\right) \leq \lim_{n \rightarrow \infty} D\left(f'_n(x), \frac{f_n(x+h) - f_n(x)}{h}\right)$$

for every $h \in (0, \beta)$, and finally passing to limit with $h \searrow 0$ we obtain

$$D \left(g(x), \lim_{h \searrow 0} \frac{f(x+h) - f(x)}{h} \right) \leq \lim_{n \rightarrow \infty} D \left(f'_n(x), \lim_{h \searrow 0} \frac{f_n(x+h) - f_n(x)}{h} \right),$$

i.e., the limit $\lim_{h \searrow 0} \frac{f(x+h) - f(x)}{h}$ exists and it is equal to $g(x)$. Analogously can be proved that $\lim_{h \searrow 0} \frac{f(x) - f(x-h)}{h} = g(x)$ and the proof is complete. $\square$

Lemma 4.2. *If a is a strict positive fuzzy number and $e_{\odot}^{a \cdot (x-x_0)}$ is defined by the power series in Example 3.10 then $e_{\odot}^{a \cdot (x-x_0)}$ is differentiable and*

$$(e^{a \cdot (x-x_0)})' = a \odot e_{\odot}^{a \cdot (x-x_0)},$$

for every $x, x_0 \in \mathbb{R}$, $x > x_0$.

Proof. Because $(\lambda + \mu) \cdot u = \lambda \cdot u \oplus \mu \cdot u$ for every $\lambda, \mu \in \mathbb{R}$, $\lambda, \mu > 0$ and u a fuzzy number, we get

$$\frac{1}{k!} \cdot a^{\odot k} \cdot (x - x_0 + h)^k = \frac{1}{k!} \cdot a^{\odot k} \cdot (x - x_0)^k \quad (5)$$

$$\oplus \frac{1}{k!} \cdot a^{\odot k} \cdot [(x - x_0 + h)^k - (x - x_0)^k],$$

for all $h > 0$, $k \in \mathbb{N}$ and $x, x_0 \in \mathbb{R}$, $x > x_0$. Let $n \in \mathbb{N}^*$. In the above relation we obtain

$$\sum_{k=0}^n \oplus \frac{1}{k!} \cdot a^{\odot k} \cdot (x - x_0 + h)^k = \sum_{k=0}^n \oplus \frac{1}{k!} \cdot a^{\odot k} \cdot (x - x_0)^k \quad (6)$$

$$\oplus \sum_{k=0}^n \oplus \frac{1}{k!} \cdot a^{\odot k} \cdot [(x - x_0 + h)^k - (x - x_0)^k],$$

for all $h > 0$, $x, x_0 \in \mathbb{R}$, $x > x_0$. Let us denote $S_n(x) = \sum_{k=0}^n \oplus \frac{1}{k!} \cdot a^{\odot k} \cdot (x - x_0)^k$. Then the H-differences

$$S_n(x+h) - S_n(x) = \sum_{k=0}^n \oplus \frac{1}{k!} \cdot a^{\odot k} \cdot [(x - x_0 + h)^k - (x - x_0)^k]$$

exist for all $h > 0$ (that is the first condition in Lemma 4.1

is verified), $x, x_0 \in \mathbb{R}$, $x > x_0$. Because

$$(x - x_0)^k = (x - x_0 - h + h)^k = (x - x_0 - h)^k + v,$$

with $v > 0$ for all $h > 0$, $h < x - x_0$, we similarly obtain the existence of the H-differences

$$\begin{aligned} \frac{1}{k!} \cdot a^{\odot k} \cdot (x - x_0)^k - \frac{1}{k!} \cdot a^{\odot k} \cdot (x - x_0 - h)^k &= \\ &= \frac{1}{k!} \cdot a^{\odot k} \cdot [(x - x_0)^k - (x - x_0 - h)^k], \end{aligned} \quad (7)$$

for all $0 < h < x - x_0$, $x, x_0 \in \mathbb{R}$, $x > x_0$. The same reasoning as above implies that

$$S_n(x) - S_n(x - h) = \sum_{k=0}^n \frac{1}{k!} \cdot a^{\odot k} \cdot [(x - x_0)^k - (x - x_0 - h)^k],$$

i.e., the H-differences $S_n(x) - S_n(x - h)$ exist for all $0 < h < x - x_0$. The condition in Lemma 4.1

is satisfied with $\beta = x - x_0 > 0$. Also, by (5) and the properties of the metric D we get

$$\begin{aligned} &\lim_{h \searrow 0} \frac{1}{h} \cdot \left[\frac{1}{k!} \cdot a^{\odot k} \cdot (x - x_0 + h)^k - \frac{1}{k!} \cdot a^{\odot k} \cdot (x - x_0)^k \right] \\ &= \lim_{h \searrow 0} \frac{1}{h} \frac{1}{k!} \cdot a^{\odot k} \cdot [(x - x_0 + h)^k - (x - x_0)^k] = \frac{1}{k!} \cdot a^{\odot k} \cdot k(x - x_0)^{k-1} \end{aligned}$$

and by (7),

$$\begin{aligned} &\lim_{h \searrow 0} \frac{1}{h} \cdot \left[\frac{1}{k!} \cdot a^{\odot k} \cdot (x - x_0)^k - \frac{1}{k!} \cdot a^{\odot k} \cdot (x - x_0 - h)^k \right] \\ &= \lim_{h \searrow 0} \frac{1}{h} \frac{1}{k!} \cdot a^{\odot k} \cdot [(x - x_0)^k - (x - x_0 - h)^k] = \frac{1}{k!} \cdot a^{\odot k} \cdot k(x - x_0)^{k-1}, \end{aligned}$$

for every $x, x_0 \in \mathbb{R}$, $x > x_0$ and $k \in \mathbb{N}^*$. We obtain

$$\left(\frac{1}{k!} \cdot a^{\odot k} \cdot (x - x_0)^k \right)' = \frac{1}{(k-1)!} \cdot a^{\odot k} \cdot (x - x_0)^{k-1}, \quad \forall k \in \mathbb{N}^*.$$

Because the terms $a^{\odot(k-1)} \cdot \frac{(x-x_0)^{k-1}}{(k-1)!}$ are strict positive for every $k \in \mathbb{N}^*$, Lemma 2.5, (iii) implies

$$S'_n(x) = \sum_{k=1}^n \frac{1}{(k-1)!} \cdot a^{\odot k} \cdot (x-x_0)^{k-1} = a \odot S_{n-1}(x).$$

Because $(S_n)_{n \in \mathbb{N}}$ converges uniformly to $e_{\odot}^{a \cdot (x-x_0)}$ and S_n is a strict positive fuzzy number for every $n \in \mathbb{N}$, it follows, by Theorem 2.9, (iv), that $(S'_n)_{n \in \mathbb{N}}$ converges uniformly to $a \odot e_{\odot}^{a \cdot (x-x_0)}$. Finally, Lemma 4.1 implies $e_{\odot}^{a \cdot (x-x_0)}$ differentiable and $\left(e_{\odot}^{a \cdot (x-x_0)}\right)' = a \odot e_{\odot}^{a \cdot (x-x_0)}$, for every $x, x_0 \in \mathbb{R}$, $x > x_0$. $\square$

Theorem 4.3. *The fuzzy-number-valued function $y : \mathbb{R} \rightarrow \mathbb{R}_{\mathcal{F}}$ given by*

$$y(x) = c \odot e_{\odot}^{a \cdot (x-x_0)} \oplus (FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt, \quad (8)$$

where a is a strict positive fuzzy number, is a solution of the equation (2) for all $x > x_0$ and $c \in \mathbb{R}_{\mathcal{F}}^{\odot}$ having the same sign as $b(t)$, $t > x_0$ (if $c \neq \tilde{0}$).

Proof. First we observe that, by Remark 3.7, $e_{\odot}^{a \cdot (x-t)}$ is continuous for $t \in [x_0, x]$. Then in Theorem 2.9, (v) it follows that $b(t) \odot e_{\odot}^{a \cdot (x-t)}$ is continuous as a fuzzy-number-valued function of t , and by [7, Theorem 3.7]

it is integrable. So, we have proved that the function y is well defined. If the last term in (8) is differentiable, then by Theorem 2.9, (vii), Lemma 2.5, (iii) and Lemma 4.2

we get

$$y'(x) = c \odot a \odot e_{\odot}^{a \cdot (x-x_0)} \oplus \left((FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt \right)' \quad (9)$$

Let us prove that $(FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt$ is differentiable and to compute its derivative. We have

$$\begin{aligned} (FR) \int_{x_0}^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt &= (FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt \\ &\quad \oplus (FR) \int_x^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt. \end{aligned}$$

Since the H-difference $u(t) = e_{\odot}^{a \cdot (x+h-t)} - e_{\odot}^{a \cdot (x-t)}$ exists (by Lemma 4.2) and it is continuous (by Lemma 2.5, (ii)) for any $h > 0$ and $t \in [x_0, x]$, it follows that $u(t) \oplus e_{\odot}^{a \cdot (x-t)} = e_{\odot}^{a \cdot (x+h-t)}$. We observe that $u(t)$ and $e_{\odot}^{a \cdot (x-t)}$ are strict positive and by Theorem 2.8, (iv) we have $b(t) \odot u(t) \oplus b(t) \odot e_{\odot}^{a \cdot (x-t)} = b(t) \odot e_{\odot}^{a \cdot (x+h-t)}$ with all the terms continuous and so integrable (see [7, Theorem 3.7]). Then we get [7, Theorem 2.5 and Theorem 2.6]

$$(FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt = (FR) \int_{x_0}^x b(t) \odot u(t) dt \\ \oplus (FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt$$

and it follows

$$(FR) \int_{x_0}^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt = (FR) \int_{x_0}^x b(t) \odot u(t) dt \\ \oplus (FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt \oplus (FR) \int_x^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt.$$

Then the H-difference $(FR) \int_{x_0}^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt - (FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt$ exists and we have

$$\lim_{h \searrow 0} \frac{1}{h} \cdot \left((FR) \int_{x_0}^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt - (FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt \right) \\ = \lim_{h \searrow 0} \frac{1}{h} \cdot \left((FR) \int_{x_0}^x b(t) \odot u(t) dt \oplus (FR) \int_x^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt \right) \\ = \lim_{h \searrow 0} \frac{1}{h} \cdot \left((FR) \int_{x_0}^x b(t) \odot \left(e_{\odot}^{a \cdot (x+h-t)} - e_{\odot}^{a \cdot (x-t)} \right) dt \right. \\ \left. \oplus (FR) \int_x^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt \right) \quad (10)$$

By [7, Theorem 2.5 and Theorem 2.6] we obtain

$$\lim_{h \searrow 0} \frac{1}{h} \cdot \left((FR) \int_{x_0}^x b(t) \odot \left(e_{\odot}^{a \cdot (x+h-t)} - e_{\odot}^{a \cdot (x-t)} \right) dt \right) =$$

$$= \lim_{h \searrow 0} (FR) \int_{x_0}^x b(t) \odot \frac{e_{\odot}^{a \cdot (x+h-t)} - e_{\odot}^{a \cdot (x-t)}}{h} dt$$

Also, [7, Remark 3.1 and Remark 3.2]
and Theorem 2.9, (i) imply

$$\begin{aligned} D \left((FR) \int_{x_0}^x b(t) \odot \frac{e_{\odot}^{a \cdot (x+h-t)} - e_{\odot}^{a \cdot (x-t)}}{h} dt, (FR) \int_{x_0}^x b(t) \odot a \odot e_{\odot}^{a \cdot (x-t)} dt \right) \leq \\ \int_{x_0}^x D \left(b(t) \odot \frac{e_{\odot}^{a \cdot (x+h-t)} - e_{\odot}^{a \cdot (x-t)}}{h}, b(t) \odot a \odot e_{\odot}^{a \cdot (x-t)} \right) dt \leq \\ \int_{x_0}^x K_{b(t)} D \left(\frac{e_{\odot}^{a \cdot (x+h-t)} - e_{\odot}^{a \cdot (x-t)}}{h}, a \odot e_{\odot}^{a \cdot (x-t)} \right) dt, \end{aligned}$$

where $K_{b(t)} = \sup_{t \in [x_0, x]} \{ \max\{ |\underline{b(t)}|^1, |\overline{b(t)}|^1 \} + \overline{b(t)}^0 - \underline{b(t)}^0 \}$. Lemma 4.2 implies

$$\lim_{h \searrow 0} D \left(\frac{e_{\odot}^{a \cdot (x+h-t)} - e_{\odot}^{a \cdot (x-t)}}{h}, a \odot e_{\odot}^{a \cdot (x-t)} \right) = 0,$$

therefore

$$\begin{aligned} \lim_{h \searrow 0} \frac{1}{h} \cdot \left((FR) \int_{x_0}^x b(t) \odot \left(e_{\odot}^{a \cdot (x+h-t)} - e_{\odot}^{a \cdot (x-t)} \right) dt \right) = \\ = (FR) \int_{x_0}^x b(t) \odot a \odot e_{\odot}^{a \cdot (x-t)} dt. \end{aligned}$$

In relation (10) we get

$$\begin{aligned} \lim_{h \searrow 0} \frac{1}{h} \cdot \left((FR) \int_{x_0}^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt - (FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt \right) = \\ = (FR) \int_{x_0}^x b(t) \odot a \odot e_{\odot}^{a \cdot (x-t)} dt \oplus \lim_{h \searrow 0} \frac{1}{h} \cdot (FR) \int_x^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt. \end{aligned}$$

Theorem 2.6, Remark 3.2 in [7]

and the properties of the distance D imply

$$D \left(b(x), \frac{1}{h} \cdot (FR) \int_x^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt \right) =$$

$$\begin{aligned}
&= D \left((FR) \int_x^{x+h} \frac{1}{h} \cdot b(x) dt, \frac{1}{h} \cdot (FR) \int_x^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt \right) \leq \\
&\leq \int_x^{x+h} D \left(\frac{1}{h} \cdot b(x), b(t) \odot \frac{1}{h} \cdot e_{\odot}^{a \cdot (x+h-t)} \right) dt \\
&= \frac{1}{h} \int_x^{x+h} D \left(b(x), b(t) \odot e_{\odot}^{a \cdot (x+h-t)} \right) dt.
\end{aligned}$$

In addition,

$$\begin{aligned}
D \left(b(x), b(t) \odot e_{\odot}^{a \cdot (x+h-t)} \right) &\leq D(b(x), b(t)) + D \left(b(t), b(t) \odot e_{\odot}^{a \cdot (x+h-t)} \right) \leq \\
&(\text{Theorem 2.9, (i)}) \leq \omega(b, h) + K_{b(t)} D(\tilde{1}, e_{\odot}^{a \cdot (x+h-t)}) \quad (11)
\end{aligned}$$

for every $t \in [x, x+h]$, with ω and $K_{b(t)}$ as above. We get

$$\begin{aligned}
&D \left(b(x), \frac{1}{h} \cdot (FR) \int_x^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt \right) \\
&\leq \frac{1}{h} \int_x^{x+h} \omega(b, h) dt + K_{b(t)} \frac{1}{h} \int_x^{x+h} D(\tilde{1}, e_{\odot}^{a \cdot (x+h-t)}) dt.
\end{aligned}$$

We also observe that

$$\begin{aligned}
D(\tilde{1}, e_{\odot}^{a \cdot (x+h-t)}) &= \lim_{n \rightarrow \infty} D \left(\tilde{1}, \tilde{1} \oplus \frac{1}{1!} a \cdot (x+h-t) \oplus \dots \oplus \frac{1}{n!} a^{\odot n} \cdot (x+h-t)^n \right) \\
&\quad (12) \\
&= \lim_{n \rightarrow \infty} D \left(\tilde{0}, \frac{1}{1!} a \cdot (x+h-t) \oplus \dots \oplus \frac{1}{n!} a^{\odot n} \cdot (x+h-t)^n \right) \\
&\leq \lim_{n \rightarrow \infty} \left(\left\| \frac{1}{1!} a \cdot (x+h-t) \right\|_{\mathcal{F}} + \dots + \left\| \frac{1}{n!} a^{\odot n} \cdot (x+h-t)^n \right\|_{\mathcal{F}} \right) \\
&\leq \lim_{n \rightarrow \infty} \left(|x+h-t| \left\| \frac{1}{1!} \cdot a \right\|_{\mathcal{F}} + \dots + |x+h-t|^n \left\| \frac{1}{n!} \cdot a^{\odot n} \right\|_{\mathcal{F}} \right),
\end{aligned}$$

therefore

$$\frac{1}{h} \int_x^{x+h} D(\tilde{1}, e_{\odot}^{a \cdot (x+h-t)}) dt$$

$$\begin{aligned}
&\leq \frac{1}{h} \lim_{n \rightarrow \infty} \left(\left\| \frac{1}{1!} \cdot a \right\|_{\mathcal{F}} \int_x^{x+h} |x+h-t| dt + \dots \right. \\
&\quad \left. + \left\| \frac{1}{n!} \cdot a^{\odot n} \right\|_{\mathcal{F}} \int_x^{x+h} |x+h-t|^n dt \right) \\
&= \frac{1}{h} \lim_{n \rightarrow \infty} \left(\frac{h^2}{2} \left\| \frac{1}{1!} \cdot a \right\|_{\mathcal{F}} + \dots + \frac{h^{n+1}}{n+1} \left\| \frac{1}{n!} \cdot a^{\odot n} \right\|_{\mathcal{F}} \right).
\end{aligned}$$

Passing to limit we obtain

$$\begin{aligned}
&\lim_{h \searrow 0} \frac{1}{h} \int_x^{x+h} D(\tilde{1}, e_{\odot}^{a \cdot (x+h-t)}) dt \leq \\
&\leq \lim_{n \rightarrow \infty} \lim_{h \searrow 0} \frac{1}{h} \left(\frac{h^2}{2} \left\| \frac{1}{1!} \cdot a \right\|_{\mathcal{F}} + \dots + \frac{h^{n+1}}{n+1} \left\| \frac{1}{n!} \cdot a^{\odot n} \right\|_{\mathcal{F}} \right) = 0.
\end{aligned}$$

Since b is continuous, $\omega(b, h) \rightarrow 0$ for $h \searrow 0$ (see e.g. [21]). We obtain

$$D \left(b(x), \lim_{h \searrow 0} \frac{1}{h} \cdot (FR) \int_x^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt \right) = 0 \quad (13)$$

for every $x > x_0$, which means $\lim_{h \searrow 0} \frac{1}{h} \cdot (FR) \int_x^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt = b(x)$, for every $x > x_0$. Then in (10) we have

$$\begin{aligned}
&\lim_{h \searrow 0} \frac{1}{h} \cdot \left((FR) \int_{x_0}^{x+h} b(t) \odot e_{\odot}^{a \cdot (x+h-t)} dt - (FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt \right) \\
&= (FR) \int_{x_0}^x b(t) \odot a \odot e_{\odot}^{a \cdot (x-t)} dt \oplus b(x), \quad (14)
\end{aligned}$$

for every $x > x_0$.

For the symmetric case we observe that

$$\begin{aligned}
&(FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt = (FR) \int_{x_0}^{x-h} b(t) \odot e_{\odot}^{a \cdot (x-t)} dt \\
&\quad \oplus (FR) \int_{x-h}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt,
\end{aligned}$$

the H-difference $v(t) = e_{\odot}^{a \cdot (x-t)} - e_{\odot}^{a \cdot (x-h-t)}$ exists and it is continuous for any $t \in [x_0, x-h]$ and $0 < h < x - x_0$ (see Lemma 2.5, (ii)). Also, by Lemma

4.2, $v(t)$ and $e_{\odot}^{a \cdot (x-h-t)}$ are strict positive and we get

$$\begin{aligned} (FR) \int_{x_0}^{x-h} b(t) \odot e_{\odot}^{a \cdot (x-t)} dt &= (FR) \int_{x_0}^{x-h} b(t) \odot v(t) dt \\ &\oplus (FR) \int_{x_0}^{x-h} b(t) \odot e_{\odot}^{a \cdot (x-h-t)} dt, \end{aligned}$$

for every $0 < h < x - x_0$, therefore

$$\begin{aligned} (FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt &= (FR) \int_{x_0}^{x-h} b(t) \odot v(t) dt \\ &\oplus (FR) \int_{x_0}^{x-h} b(t) \odot e_{\odot}^{a \cdot (x-h-t)} dt \oplus (FR) \int_{x-h}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt. \end{aligned}$$

Then the H-difference $(FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt - (FR) \int_{x_0}^{x-h} b(t) \odot e_{\odot}^{a \cdot (x-h-t)} dt$ exists and similar to (10) we have

$$\begin{aligned} &\lim_{h \searrow 0} \frac{1}{h} \cdot \left((FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt - (FR) \int_{x_0}^{x-h} b(t) \odot e_{\odot}^{a \cdot (x-h-t)} dt \right) \\ &= \lim_{h \searrow 0} \frac{1}{h} \left((FR) \int_{x_0}^{x-h} b(t) \odot \left(e_{\odot}^{a \cdot (x-t)} - e_{\odot}^{a \cdot (x-h-t)} \right) dt \right) \quad (15) \\ &\quad \oplus \lim_{h \searrow 0} \frac{1}{h} \left((FR) \int_{x-h}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt \right) \\ &= (FR) \int_{x_0}^x b(t) \odot a \odot e_{\odot}^{a \cdot (x-t)} dt \oplus \lim_{h \searrow 0} \frac{1}{h} \cdot (FR) \int_{x-h}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt. \end{aligned}$$

Similar to the symmetric case we get

$$\begin{aligned} &D \left(b(x), \frac{1}{h} \cdot (FR) \int_{x-h}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt \right) \\ &\leq \frac{1}{h} \int_{x-h}^x D(b(x), b(t) \odot e_{\odot}^{a \cdot (x-t)}) dt, \end{aligned}$$

for every $0 < h < x - x_0$. As in (11) we have

$$D \left(b(x), b(t) \odot e_{\odot}^{a \cdot (x-t)} \right) \leq \omega(b, h) + K_{b(t)} D(\tilde{1}, e_{\odot}^{a \cdot (x-t)})$$

for every $t \in [x - h, x]$, $0 < h < x - x_0$. Similar to (12) we obtain

$$D(\tilde{1}, e_{\odot}^{a \cdot (x-t)}) \leq \lim_{n \rightarrow \infty} \left(|x - t| \left\| \frac{1}{1!} \cdot a \right\|_{\mathcal{F}} + \dots + |x - t|^n \left\| \frac{1}{n!} \cdot a^{\odot n} \right\|_{\mathcal{F}} \right),$$

for every $t \in [x - h, x]$, $0 < h < x - x_0$. Similar to (13) we get

$$\lim_{h \searrow 0} \frac{1}{h} \cdot (FR) \int_{x-h}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt = b(x),$$

for every $x > x_0$. Then in (15) we have

$$\begin{aligned} \lim_{h \searrow 0} \frac{1}{h} \cdot \left((FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt - (FR) \int_{x_0}^{x-h} b(t) \odot e_{\odot}^{a \cdot (x-h-t)} dt \right) \\ = (FR) \int_{x_0}^x b(t) \odot a \odot e_{\odot}^{a \cdot (x-t)} dt \oplus b(x), \end{aligned} \quad (16)$$

for every $x > x_0$, and by (14) and (16) we obtain that

$$\left((FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt \right)' = (FR) \int_{x_0}^x b(t) \odot a \odot e_{\odot}^{a \cdot (x-t)} dt \oplus b(x),$$

for every $x > x_0$. In (9) we get

$$y'(x) = c \odot a \odot e_{\odot}^{a \cdot (x-x_0)} \oplus (FR) \int_{x_0}^x b(t) \odot a \odot e_{\odot}^{a \cdot (x-t)} dt \oplus b(x)$$

Because c and $b(t)$ have the same sign it is easy to prove that the fuzzy-number-valued function given by (9) is a solution of the fuzzy differential equation (2) and the proof is complete. $\square$

It is immediate the following result.

Corollary 4.4. *The fuzzy-number-valued function $y : \mathbb{R} \rightarrow \mathbb{R}_{\mathcal{F}}$ given by*

$$y(x) = y_0 \odot e^{a \cdot (x-x_0)} \oplus (FR) \int_{x_0}^x b(t) \odot e_{\odot}^{a \cdot (x-t)} dt,$$

where a is a strict positive fuzzy number, is a solution of the fuzzy initial value problem

$$\begin{cases} y'(x) = a \odot y(x) \oplus b(x) \\ y(x_0) = y_0 \end{cases}, \quad (17)$$

for all $x > x_0$, where $y_0 \in \mathbb{R}_{\mathcal{F}}^{\odot}$, $b : \mathbb{R} \rightarrow \mathbb{R}_{\mathcal{F}}$ is continuous, y_0 and $b(x)$ have the same sign for all $x > x_0$ (if $y_0 \neq \tilde{0}$).

The following example gives the solution of a fuzzy initial value problem with triangular fuzzy numbers as data. We recall that for $a < b < c$, $a, b, c \in \mathbb{R}$ the triangular fuzzy number $u = (a, b, c)$ determined by a, b, c is given such that $\underline{u}^r = a + (b - a)r$ and $\bar{u}^r = c - (c - b)r$ are the endpoints of the r -level sets, for every $r \in [0, 1]$. If $u_i = (a_i, b_i, c_i)$, $i \in \{1, 2\}$ are triangular fuzzy numbers then

$$u_1 \odot u_2 = (a_1 b_2 + a_2 b_1 - b_1 b_2, b_1 b_2, c_1 b_2 + c_2 b_1 - b_1 b_2) \text{ (see [3])}$$

and

$$u_1 \oplus u_2 = (a_1 + a_2, b_1 + b_2, c_1 + c_2).$$

If $u = (a, b, c)$ is a triangular fuzzy number and $k \in \mathbb{R}$ then $k \cdot u = (ka, kb, kc)$ if k is positive and $k \cdot u = (kc, kb, ka)$ if k is negative.

Example 4.5. Let $a = (\alpha(1 - \varepsilon), \alpha, \alpha(1 + \varepsilon))$, $\alpha, \varepsilon > 0$ and $b = (\beta(1 - \varepsilon'), \beta, \beta(1 + \varepsilon'))$, $\beta, \varepsilon' > 0$, strict positive triangular fuzzy numbers and let us consider the fuzzy initial value problem

$$\begin{cases} y'(x) = a \odot y(x) \oplus b \\ y(0) = \tilde{1} \end{cases} \quad (18)$$

Firstly, let us compute $e_{\odot}^{a \cdot x}$ for $x > 0, x \in \mathbb{R}$. Because (see [3])

$$(a \cdot x)^{\odot n} = (\alpha^n x^n (1 - n\varepsilon), \alpha^n x^n, \alpha^n x^n (1 + n\varepsilon)),$$

by direct computation we get

$$\begin{aligned} e_{\odot}^{a \cdot x} &= \sum_{n=0}^{\infty} \frac{1}{n!} \cdot (a \cdot x)^{\odot n} \\ &= \sum_{n=0}^{\infty} \left(\frac{\alpha^n x^n}{n!} (1 - n\varepsilon), \frac{\alpha^n x^n}{n!}, \frac{\alpha^n x^n}{n!} (1 + n\varepsilon) \right) \\ &= \left(\sum_{n=0}^{\infty} \frac{\alpha^n x^n}{n!} (1 - n\varepsilon), \sum_{n=0}^{\infty} \frac{\alpha^n x^n}{n!}, \sum_{n=0}^{\infty} \frac{\alpha^n x^n}{n!} (1 + n\varepsilon) \right) \\ &= \left(\sum_{n=0}^{\infty} \frac{\alpha^n x^n}{n!} - \varepsilon \sum_{n=1}^{\infty} \frac{\alpha^n x^n}{(n-1)!}, \sum_{n=0}^{\infty} \frac{\alpha^n x^n}{n!}, \sum_{n=0}^{\infty} \frac{\alpha^n x^n}{n!} + \varepsilon \sum_{n=1}^{\infty} \frac{\alpha^n x^n}{(n-1)!} \right) \end{aligned}$$

$$= (e^{\alpha x}(1 - \varepsilon \alpha x), e^{\alpha x}, e^{\alpha x}(1 + \varepsilon \alpha x)).$$

It follows that

$$\begin{aligned} (FR) \int_0^x b \odot e_{\odot}^{a \cdot (x-t)} dt &= [2, \text{Theorem 4.9, (iv)}] = b \odot (FR) \int_0^x e_{\odot}^{a \cdot (x-t)} dt \\ &= b \odot (FR) \int_0^x (e^{\alpha(x-t)}(1 - \varepsilon \alpha(x-t)), e^{\alpha(x-t)}, e^{\alpha(x-t)}(1 + \varepsilon \alpha(x-t))) dt \\ &= [7, \text{Theorem 3.2}] \\ &= b \odot \left(\int_0^x e^{\alpha(x-t)}(1 - \varepsilon \alpha(x-t)) dt, \int_0^x e^{\alpha(x-t)} dt, \int_0^x e^{\alpha(x-t)}(1 + \varepsilon \alpha(x-t)) dt \right) \\ &= (\beta(1 - \varepsilon'), \beta, \beta(1 + \varepsilon')) \odot \\ &\quad \odot \left(\frac{e^{\alpha x} - 1}{\alpha}(1 + \varepsilon) - \varepsilon x e^{\alpha x}, \frac{e^{\alpha x} - 1}{\alpha}, \frac{e^{\alpha x} - 1}{\alpha}(1 - \varepsilon) + \varepsilon x e^{\alpha x} \right) \\ &= \beta \cdot \left(\frac{e^{\alpha x} - 1}{\alpha}(1 + \varepsilon - \varepsilon') - \varepsilon x e^{\alpha x}, \frac{e^{\alpha x} - 1}{\alpha}, \frac{e^{\alpha x} - 1}{\alpha}(1 + \varepsilon' - \varepsilon) + \varepsilon x e^{\alpha x} \right) \end{aligned}$$

Then

$$\begin{aligned} y(x) &= e^{\alpha x} \cdot (1 - \varepsilon \alpha x, 1, 1 + \varepsilon \alpha x) \oplus \\ &\quad \oplus \beta \cdot \left(\frac{e^{\alpha x} - 1}{\alpha}(1 + \varepsilon - \varepsilon') - \varepsilon x e^{\alpha x}, \frac{e^{\alpha x} - 1}{\alpha}, \frac{e^{\alpha x} - 1}{\alpha}(1 + \varepsilon' - \varepsilon) + \varepsilon x e^{\alpha x} \right) \end{aligned}$$

is the solution of the problem (18).

Acknowledgement. We express our thanks to Professor S.G. Gal, whose comments helped us to improve the paper.

References

- [1] G.A. Anastassiou, On H-Fuzzy Differentiation, *Mathematica Balkanica* (New Series), vol. 16, Fasc. 1-4, 153-193 (2002).
- [2] A.I. Ban, B. Bede, Properties of the cross product of fuzzy numbers, submitted.
- [3] A.I. Ban, B. Bede, Cross product of $L-R$ fuzzy numbers and applications, *Anal. Math. Univ. Oradea* (fasc. math.) 9, 95-108 (2002).

- [4] B. Bouchon-Meunier, O. Kosheleva, V. Kreinovich, H.T. Nguyen, Fuzzy numbers are the only fuzzy sets that keep invertible operations invertible, *Fuzzy Sets and Systems* 91, 155-163 (1997).
- [5] S.M. Chen, Fuzzy system reliability analysis using fuzzy number arithmetic operations, *Fuzzy Sets and Systems* 64, 31-38 (1994).
- [6] C.H. Cheng, A new approach for ranking fuzzy numbers by distance method, *Fuzzy Sets and Systems* 95, 307-317 (1998).
- [7] Wu Congxin, Gong Zengtai, On Henstock integral of fuzzy number valued functions (I), *Fuzzy Sets and Systems* 120, 523-532 (2001).
- [8] M. Delgado, M.A. Vila, W. Voxman, On a canonical representation of fuzzy numbers, *Fuzzy Sets and Systems* 93, 125-135 (1998).
- [9] M. Delgado, M.A. Vila, W. Voxman, A fuzziness measure for fuzzy numbers: Applications, *Fuzzy Sets and Systems* 94, 205-216 (1998).
- [10] R. Goetschel, W. Voxman, Elementary fuzzy calculus, *Fuzzy Sets and Systems* 18, 31-43 (1986).
- [11] P. Grzegorzewski, Metrics and orders in space of fuzzy numbers, *Fuzzy Sets and Systems* 97, 83-94 (1998).
- [12] S. Heilpern, The expected value of a fuzzy number, *Fuzzy Sets and Systems* 47, 81-86 (1992).
- [13] S. Heilpern, Comparison of fuzzy numbers in decision making, *Tatra Mt. Math. Publ.* 6, 47-53 (1995).
- [14] S. Heilpern, Representation and application of fuzzy numbers, *Fuzzy Sets and Systems* 91, 259-268 (1997).
- [15] D.H. Hong, H.Y. Do, Fuzzy system reliability analysis by the use of T_w (the weakest t-norm) on fuzzy number arithmetic operations, *Fuzzy Sets and Systems* 90, 307-316 (1997).
- [16] D.H. Hong, S.Y. Hwang, On the convergence of T -sum of $L - R$ fuzzy numbers, *Fuzzy Sets and Systems* 63, 175-180 (1994).

- [17] A. Irion, Fuzzy rules and fuzzy functions: A combination of logic and arithmetic operations for fuzzy numbers, *Fuzzy Sets and Systems* 99, 49-56 (1998).
- [18] O. Kaleva, Fuzzy differential equations, *Fuzzy Sets and Systems* 24, 301-317 (1987).
- [19] G.J. Klir, Fuzzy arithmetic with requisite constraints, *Fuzzy Sets and Systems* 91, 165-175 (1997).
- [20] A. Kolesárová, Triangular norm-based addition of linear fuzzy numbers, *Tatra Mt. Math. Publ.* 6, 75-81 (1995).
- [21] R. A. DeVore, G.G. Lorentz, *Constructive approximation, polynomials and splines approximation*, Springer Verlag, Berlin, Heidelberg, 1993.
- [22] M. Mareš, Multiplication of fuzzy quantities, *Kybernetika* 28, 337-356 (1992).
- [23] M. Mareš, M., Algebraic equivalences over fuzzy quantities, *Kybernetika* 29, 121-132 (1993).
- [24] M. Mareš, Brief note on distributivity of triangular fuzzy numbers, *Kybernetika* 31, 451-457 (1995).
- [25] M. Mareš, Equivalentions over fuzzy quantities, *Tatra Mt. Math. Publ.* 6, 117-121 (1995).
- [26] M. Mareš, Weak arithmetics of fuzzy numbers, *Fuzzy Sets and Systems* 91, 143-153 (1997).
- [27] A. Markova, T -sum of $L - R$ fuzzy numbers, *Fuzzy Sets and Systems* 85, 379-384 (1997).
- [28] B. Matarazzo, G. Munda, New approaches for the comparison of $L - R$ fuzzy numbers: a theoretical and operational analysis, *Fuzzy Sets and Systems* 118, 407-418 (2001).
- [29] Ma Ming, M. Friedman, A. Kandel, A new fuzzy arithmetic, *Fuzzy Sets and Systems* 108, 83-90 (1999).

- [30] R. Mesiar, A note to the T -sum of $L - R$ fuzzy numbers, *Fuzzy Sets and Systems* 79, 259-261 (1996).
- [31] R. Mesiar, Shape preserving additions of fuzzy intervals, *Fuzzy Sets and Systems* 86, 73-78 (1997).
- [32] M. Puri, D. Ralescu, Differentials of fuzzy functions, *J. Math. Anal. Appl.*, 91, 552-558 (1983).
- [33] D. Singer, A fuzzy set approach to fault tree and reliability analysis, *Fuzzy Sets and Systems* 34, 145-155 (1990).
- [34] M. Stojaković, Z. Stojaković, Addition and series of fuzzy sets, *Fuzzy Sets and Systems* 83, 341-346 (1996).
- [35] M. Wagenknecht, R. Hampel, V. Schneider, Computational aspects of fuzzy arithmetics based on Archimedean t-norms, *Fuzzy Sets and Systems* 123, 49-62 (2001).
- [36] X. Wang, E.E. Kerre, Reasonable properties for the ordering of fuzzy quantities, (I), (II), *Fuzzy Sets and Systems* 118, 387-405 (2001).

On a delay integral equation in biomathematics

Alexandru Bica

Department of Mathematics, University of Oradea,
Str. Armatei Române no 5, Oradea, Romania
smbica@yahoo.com ; abica@uoradea.ro

Crăciun Iancu

Faculty of Mathematics and Computer Science,
"Babes-Bolyai" University,
Str. N. Kogalniceanu no 1, Cluj-Napoca, Romania
ciancu@math.ubbcluj.ro

Abstract

In this paper we obtain a numerical method for the solution of the delay integral equation (1) which is known in some anterior given points. In this purpose we construct a cubic fitting spline function through the least squares method on the anterior interval which realize here a suboptimal approximation of the solution. Afterwards, we use the successive approximations method combined with a quadrature rule to approximate the solution on $[0,T]$.

2000 Mathematics Subject Classification: Primary 65R20, 65D10, Secondary 45D05.

Keywords and Phrases: fitting spline, delay integral equation, perturbed trapezoidal quadrature rule.

1 Introduction

The aim of this paper is to propose a numerical method for the approximation of the smooth positive solution of the delay integral equation,

$$x(t) = \int_{t-\tau}^t f(s, x(s)) ds. \quad (1)$$

This integral equation has the origin in biomathematics where it models the spread of certain infectious diseases with a contact rate that varies seasonally. So that $x(t)$ is the proportion of infectious in the population at the time t , τ is the length of time in which an individual remains infectious and $f(t, x(t))$ is the proportion of new infections on unit time.

In [4], [14], [9] and [10] sufficient conditions for the existence of the positive continuous solution of (1), which can be periodic, were given by D. Guo, V. Lakshmikantham, L.R. Williams, R.W. Legget and R. Precup. Using the Banach's fixed point principle, I.A. Rus obtains existence and uniqueness results of this solution in [11] and [12]. In [7] C. Iancu obtains a numerical method for the approximation of the positive bounded solution of (1) using the Picard's technique of successive approximations and the trapezoidal quadrature rule. In [2] we obtain some numerical methods for the approximation of the positive smooth solution of (1) according to the properties of the function f , (that is Lipschitzian or of C^1, C^2 smoothness class). Continuing our ideas, in this paper we propose a numerical method using fitting spline functions and a perturbed trapezoidal quadrature rule in C^3 smoothness case.

In this paper we consider $\tau > 0$, and $T > 0$ fixed, suppose that we know the solution of the equation (1) in some p discrete points in the interval $[-\tau, 0]$ and we want to obtain an approximation of this solution for $t \in [-\tau, T]$. Firstly we give sufficient conditions for the existence and uniqueness of a positive bounded smooth solution of the equation (1). Using the above discrete values and the Ichida-Yoshimoto-Kiyono's method based on the least squares principle we construct a fitting spline function which approximates the solution on $[-\tau, 0]$ ([8] and [13]). Afterwards, using the knots and the parameter values obtained with this method we construct an approximation of the solution for $t \in [0, T]$ based on the successive approximations method and on a perturbed trapezoidal quadrature rule ([1] and [3]).

2 The existence and the uniqueness of the positive smooth solution

Let $T > 0$ be a fixed real number. Considering the initial condition, $x(t) = \Phi(t)$, $t \in [-\tau, 0]$, we have the initial value problem :

$$\begin{cases} x(t) = \int_{t-\tau}^t f(s, x(s))ds, & \forall t \in [0, T] \\ x(t) = \Phi(t), & \forall t \in [-\tau, 0]. \end{cases} \quad (2)$$

We suppose that $\Phi \in C[-\tau, 0]$ satisfies the condition

$$\Phi(0) = \int_{-\tau}^0 f(s, \Phi(s))ds. \quad (3)$$

Supposing that $f \in C([-\tau, T] \times \mathbb{R})$ and using the Leibniz's formula of derivation for an integral with parameters, we obtain the equivalence between (2) and the following Cauchy problem for a delay differential equation

$$\begin{cases} x'(t) = f(t, x(t)) - f(t - \tau, x(t - \tau)), & \forall t \in [0, T] \\ x(t) = \Phi(t), & \forall t \in [-\tau, 0]. \end{cases} \quad (4)$$

To obtain the existence of an unique positive, bounded solution we consider these assumptions :

- (i) $\exists a, \beta > 0$ such that Φ is continuous on $[-\tau, 0]$ and $0 < a \leq \Phi(t) \leq \beta, \forall t \in [-\tau, 0]$;
(ii) $f \in C([-\tau, T] \times [a, \beta]), f(t, x) \geq 0, \forall t \in [-\tau, T], \forall x \geq 0$ and $\exists M > 0$ such that $f(t, y) \leq M, \forall t \in [-\tau, T], \forall y \in [a, \beta]$;
(iii) $\Phi \in C^1[-\tau, 0]$ satisfies the condition (3) and

$$\Phi'(0) = f(0, b) - f(-\tau, \Phi(-\tau))$$

where $b = \Phi(0)$;

- (iv) $M\tau \leq \beta$ and there is an integrable function $g(t)$ such that

$$f(t, x) \geq g(t), \quad \forall t \in [-\tau, T], \quad \forall x \geq a,$$

with

$$\int_{t-\tau}^t g(s)ds \geq a, \quad \forall t \in [0, T] ;$$

- (v) $\exists L > 0$ such that $|f(t, x) - f(t, y)| \leq L|x - y|$ for all $t \in [-\tau, T]$ and $x, y \in [a, \beta]$.

Theorem 1 Suppose that assumptions (i)-(v) are satisfied. Then the equation (1) has a unique continuous on $[-\tau, T]$ solution $x(t)$, with $a \leq x(t) \leq \beta$, $\forall t \in [-\tau, T]$ such that $x(t) = \Phi(t)$ for $t \in [-\tau, 0]$. Moreover,

$$\max \{|x_n(t) - x(t)| : t \in [0, T]\} \longrightarrow 0$$

as $n \rightarrow \infty$ where $x_n(t) = \Phi(t)$ for $t \in [-\tau, 0]$, $n \in \mathbb{N}$, $x_0(t) = b$ and $x_n(t) = \int_{t-\tau}^t f(s, x_{n-1}(s))ds$ for $t \in [0, T]$, $n \in \mathbb{N}^*$. The solution x belongs to $C^1[-\tau, T]$.

Proof. In the functional space $C[-\tau, T]$ we consider the Bielecki's norm

$$\|x\|_B = \max\{|x(t)| \cdot e^{-\theta(t+\tau)} : t \in [-\tau, T]\}$$

for $\theta > 0$ and define the operator $A : C[-\tau, T] \longrightarrow C[-\tau, T]$

$$A(x(t)) = \begin{cases} \int_{t-\tau}^t f(s, x(s))ds, & t \in [0, T] \\ \Phi(t), & t \in [-\tau, 0]. \end{cases}$$

We have

$$\begin{aligned} |A(x_1(t)) - A(x_2(t))| &\leq \int_{t-\tau}^t |f(s, x_1(s)) - f(s, x_2(s))| ds \\ &\leq \int_{t-\tau}^t L \|x_1 - x_2\|_B \cdot e^{\theta(s+\tau)} ds \leq \frac{L}{\theta} \|x_1 - x_2\|_B \cdot e^{\theta(t+\tau)} \end{aligned}$$

for $t \in [0, T]$ and consequently

$$\|A(x_1) - A(x_2)\| \leq \frac{L}{\theta} \cdot \|x_1 - x_2\|_B$$

for $x_1, x_2 \in C[-\tau, T]$. This operator is contraction for $\theta > L$. Using the Banach's fixed point principle we obtain the existence of an unique continuous on $[-\tau, T]$ solution for (1) such that $x(t) = \Phi(t)$ for $t \in [-\tau, 0]$ and

$$\max \{|x_n(t) - x(t)| : t \in [0, T]\} \longrightarrow 0$$

as $n \rightarrow \infty$. Using now the Chebyshev's norm,

$$\|x\|_C = \max\{|x(t)| : t \in [-\tau, T]\}$$

we obtain

$$|x_n(t) - x(t)| \leq \frac{\tau^m L^m}{1 - \tau L} \cdot \|x_0 - x_1\|_C, \forall t \in [0, T].$$

From (iv) we see that

$$\begin{aligned} x(t) &= \int_{t-\tau}^t f(s, x(s)) ds \geq \int_{t-\tau}^t g(s) ds \geq a \\ x(t) &= \int_{t-\tau}^t f(s, x(s)) ds \leq \int_{t-\tau}^t M ds \leq \beta, \forall t \in [0, T]. \end{aligned}$$

Because $a \leq \Phi(t) \leq \beta, \forall t \in [-\tau, 0]$ and $x(t) = \Phi(t)$ for $t \in [-\tau, 0]$ we infer that $a \leq x(t) \leq \beta, \forall t \in [-\tau, T]$, and the solution is bounded. Since x is solution for (1) and we have $x(t) = \int_{t-\tau}^t f(s, x(s)) ds$, for all $t \in [0, T]$, and because $f \in C([-\tau, T] \times [a, \beta])$ we infer that x is derivable on $[0, T]$, with x' is continuous on $[0, T]$. From the condition (iii) follow that x is derivable with x' continuous on $[-\tau, 0]$ (including the continuity at the point $t = 0$). Then $x \in C^1[-\tau, T]$ and the proof is complete. ■

Corollary 2 *In the conditions of Theorem 1, if $f \in C^1([-\tau, T] \times [a, \beta])$, $\Phi \in C^2[-\tau, 0]$ and*

$$\begin{aligned} \Phi''(0) &= \frac{\partial f}{\partial t}(0, b) + \frac{\partial f}{\partial x}(0, b) \cdot [f(0, b) - f(-\tau, \Phi(-\tau))] - \\ &\quad - \frac{\partial f}{\partial t}(-\tau, \Phi(-\tau)) - \frac{\partial f}{\partial x}(-\tau, \Phi(-\tau)) \cdot \Phi'(-\tau) \end{aligned}$$

then $x \in C^2[-\tau, T]$.

Proof. Follows directly from the above theorem and from hypothesis. ■

3 The construction of the fitting spline function

The initial condition in (2) means that the proportion $\Phi(t)$ of infectious in population is known for $t \in [-\tau, 0]$.

But in the real world this knowledge is incomplete, more exactly we can know the function Φ only in some discret moments when we make measurements. For this reason we can consider many sufficient values of Φ , $f_1, f_2, \dots, f_p$ experimentally found at the moments $u_i, i = \overline{1, p}$ such that

$$-\tau = u_1 < u_2 < \dots < u_p = 0.$$

Because these values are affected by errors we have to find an adequate fitting spline function which approximate the function Φ on $[-\tau, 0]$, using the least squares method combined with the algorithm of K. Ichida, F. Yoshimoto, T. Kiyono (see [8] and [13]) and the spline function of C. Iancu (see [5] and [6]).

This algorithm have the knots as variables. Starting from three knots we insert another knot successively until a certain criterion is satisfied. The estimator of the variance of errors (as well as in [8]) is used to determine if a satisfactory fit is reached. Then we obtain the new knots $u^{(i)}, i = \overline{1, n+1}$. Let $h_i = u^{(i+1)} - u^{(i)}, i = \overline{1, n}$.

3.1 The cubic spline function

The spline function, above mentioned, is the following piecewise cubic polynomial on $[-\tau, 0]$ (see [5] and [6]) for which if $t \in [u^{(i)}, u^{(i+1)}]$ we have :

$$s_i(t) = \frac{M_{i+1} - M_i}{6h_i} \cdot (t - u^{(i)})^3 + \frac{M_i}{2} \cdot (t - u^{(i)})^2 + \quad (5)$$

$$+ m_i \cdot (t - u^{(i)}) + y_i, \quad \forall i = \overline{1, n-1}.$$

From [5] and [6] we have,

$$\begin{cases} M_{i+1} + 2M_i = 6 \cdot \frac{y_{i+1} - y_i - m_i \cdot h_i}{h_i^2} \\ M_{i+1} + M_i = \frac{2(m_{i+1} - m_i)}{h_i}, \quad i = \overline{1, n-1}. \end{cases} \quad (6)$$

From relation (6) we obtain the following recurrent relations

$$\begin{cases} M_{i+1} = 6 \cdot \frac{y_{i+1} - y_i}{h_i^2} - \frac{6m_i}{h_i} - 2M_i \\ m_{i+1} = 3 \cdot \frac{y_{i+1} - y_i}{h_i} - 2m_i - \frac{1}{2}M_i h_i \end{cases}, \quad i = \overline{1, n-1}. \quad (7)$$

and we can also obtain

$$\begin{cases} m_i = \frac{y_{i+1} - y_i}{h_i} - \frac{1}{6}h_i(M_{i+1} + 2M_i) \\ m_{i+1} = \frac{y_{i+1} - y_i}{h_i} + \frac{1}{6}h_i(4M_{i+1} + M_i) \end{cases}, \quad i = \overline{1, n-1}. \quad (8)$$

Replacing this value of m_i in (5), we obtain a new expresion for $s_i(t)$.

If we denote,

$$a_i(t) = \frac{(t - u^{(i)})^3}{6h_i} - \frac{h_i \cdot (t - u^{(i)})}{6}$$

$$b_i(t) = -\frac{(t - u^{(i)})^3}{6h_i} + \frac{(t - u^{(i)})^2}{2} - \frac{h_i \cdot (t - u^{(i)})}{3}$$

$$c_i(t) = \frac{t - u^{(i)}}{6} \quad , \quad d_i(t) = 1 - \frac{t - u^{(i)}}{h_i}$$

then relation (5) becomes,

$$s_i(t) = M_{i+1}a_i(t) + M_i b_i(t) + y_{i+1}c_i(t) + y_i d_i(t) . \quad (9)$$

Here, y_i, m_i, M_i are the values of s_i, s'_i, s''_i on $u^{(i)}$, $i = \overline{1, n}$. The computing algorithm may be summarized as follows ([8] and [6]):

Step 1 : Calculate least squares fitting for initial two intervals. Construct the normal equation and go to step 4.

Step 2 : Determine the interval to divide and insert a new knot.

Step 3 : Construct the normal equation.

Step 4 : Solve the normal equation.

Step 5 : If the number of parameters (y_i together $M_i, i = \overline{1, n}$) is greater than that of data, then end. Else test the criterion. If the criterion is not satisfied, go to step 2. If it is satisfied, then end.

The algorithm run if we choose a satisfactory $\varepsilon > 0$ to test the criterion.

3.2 The algorithm

First we take three knots $u^{(1)}, u^{(2)}, u^{(3)}$ where $u^{(1)} = u_1 = -\tau, u^{(3)} = u_p = 0$ and $u^{(2)}$ is determined in order that there are equal numbers of data u_i in each interval $[u^{(1)}, u^{(2)}]$ and $[u^{(2)}, u^{(3)}]$. We have

$$y_1 = s_1(u^{(1)}), \quad y_2 = s_1(u^{(2)}) = s_2(u^{(2)}), \quad y_3 = s_2(u^{(3)})$$

$$m_1 = s'_1(u^{(1)}), \quad m_2 = s'_1(u^{(2)}) = s'_2(u^{(2)}), \quad m_3 = s'_2(u^{(3)})$$

$$M_1 = s''_1(u^{(1)}), \quad M_2 = s''_1(u^{(2)}) = s''_2(u^{(2)}), \quad M_3 = s''_2(u^{(3)}).$$

The values of parameters $y_1, y_2, y_3, M_1, M_2, M_3$ will be determined by least squares fitting. From (9) we have,

$$s_1(t) = M_2 a_1(t) + M_1 b_1(t) + y_2 c_1(t) + y_1 d_1(t) \quad (10)$$

$$s_2(t) = M_3 a_2(t) + M_2 b_2(t) + y_3 c_2(t) + y_2 d_2(t).$$

Let the least and the largest values u_i of data in the interval $[u^{(1)}, u^{(2)}]$ be respectively $u_{p_1} = u_1$ and u_{q_1} . Similarly we denote u_{p_2} and $u_{q_2} = u_p$ in the interval $[u^{(2)}, u^{(3)}]$. The sum of squares of residuals is,

$$R(y_1, M_1, y_2, M_2, y_3, M_3) = \sum_{k=p_1}^{q_1} [s_1(u_k) - f_k]^2 + \sum_{k=p_2}^{q_2} [s_2(u_k) - f_k]^2. \quad (11)$$

We want to obtain the values of y_i , M_i , $i = \overline{1, 3}$ which minimize R . Differentiating this expression with the 6 parameters y_i , M_i and setting to zero, that is

$$\frac{\partial R}{\partial y_1} = 0, \quad \frac{\partial R}{\partial M_1} = 0, \quad \frac{\partial R}{\partial y_2} = 0$$

$$\frac{\partial R}{\partial M_2} = 0, \quad \frac{\partial R}{\partial y_3} = 0, \quad \frac{\partial R}{\partial M_3} = 0$$

we get a linear system in this 6 parameters, which can be written in the following normal equation :

$$A_3 z_3 = g_3 \quad (12)$$

where the subscripts denote the number of knots and,

$$z_3^T = (y_1, M_1, y_2, M_2, y_3, M_3) \in \mathbb{R}^6 \quad (13)$$

$$g_3^T = \left(\sum_{k=p_1}^{q_1} d_1(u_k) f_k, \sum_{k=p_1}^{q_1} b_1(u_k) f_k, \sum_{k=p_1}^{q_1} c_1(u_k) f_k + \sum_{k=p_2}^{q_2} d_2(u_k) f_k, \right. \\ \left. \sum_{k=p_1}^{q_1} a_1(u_k) f_k + \sum_{k=p_2}^{q_2} b_2(u_k) f_k, \sum_{k=p_2}^{q_2} c_2(u_k) f_k, \sum_{k=p_2}^{q_2} a_2(u_k) f_k \right) \in \mathbb{R}^6 \quad (14)$$

$$A_3 = \begin{pmatrix} a_{11}^{(3)} & a_{12}^{(3)} & a_{13}^{(3)} & a_{14}^{(3)} & 0 & 0 \\ a_{21}^{(3)} & a_{22}^{(3)} & a_{23}^{(3)} & a_{24}^{(3)} & 0 & 0 \\ a_{31}^{(3)} & a_{32}^{(3)} & a_{33}^{(3)} & a_{34}^{(3)} & a_{35}^{(3)} & a_{36}^{(3)} \\ a_{41}^{(3)} & a_{42}^{(3)} & a_{43}^{(3)} & a_{44}^{(3)} & a_{45}^{(3)} & a_{46}^{(3)} \\ 0 & 0 & a_{53}^{(3)} & a_{54}^{(3)} & a_{55}^{(3)} & a_{56}^{(3)} \\ 0 & 0 & a_{63}^{(3)} & a_{64}^{(3)} & a_{65}^{(3)} & a_{66}^{(3)} \end{pmatrix} \quad (15)$$

$$a_{ij}^{(3)} = a_{ji}^{(3)}, \quad i, j = \overline{1, 6}, \quad i \neq j$$

$$a_{11}^{(3)} = \sum_k [d_1(u_k)]^2, \quad a_{22}^{(3)} = \sum_k [b_1(u_k)]^2, \quad a_{12}^{(3)} = \sum_k b_1(u_k) d_1(u_k) \\ a_{13}^{(3)} = \sum_k c_1(u_k) d_1(u_k), \quad a_{14}^{(3)} = \sum_k a_1(u_k) d_1(u_k) \quad (16)$$

$$a_{23}^{(3)} = \sum_k b_1(u_k) c_1(u_k), \quad a_{24}^{(3)} = \sum_k a_1(u_k) b_1(u_k), \quad \forall k = \overline{p_1, q_1}.$$

$$a_{33}^{(3)} = \sum_{k=p_1}^{q_1} [c_1(u_k)]^2 + \sum_{k=p_2}^{q_2} [d_2(u_k)]^2, \quad a_{34}^{(3)} = \sum_{k=p_1}^{q_1} a_1(u_k) c_1(u_k) + \\ + \sum_{k=p_2}^{q_2} b_2(u_k) d_2(u_k), \quad a_{44}^{(3)} = \sum_{k=p_1}^{q_1} [a_1(u_k)]^2 + \sum_{k=p_2}^{q_2} [b_2(u_k)]^2. \quad (17)$$

$$\begin{aligned}
a_{35}^{(3)} &= \sum_k c_2(u_k) d_2(u_k), \quad a_{36}^{(3)} = \sum_k a_2(u_k) d_2(u_k) \\
a_{45}^{(3)} &= \sum_k b_2(u_k) c_2(u_k), \quad a_{46}^{(3)} = \sum_k a_2(u_k) b_2(u_k), \quad a_{55}^{(3)} = \sum_k [c_2(u_k)]^2 \\
a_{56}^{(3)} &= \sum_k a_2(u_k) c_2(u_k), \quad a_{66}^{(3)} = \sum_k [a_2(u_k)]^2, \quad \forall k = \overline{p_2, q_2}. \quad (18)
\end{aligned}$$

Solving the equation (12) by the Gauss's elimination method we obtain the vector of parameters, z_3^T , and the minimal value of R .

We denote by $R_1(y_1, M_1, y_2, M_2)$ and $R_2(y_2, M_2, y_3, M_3)$ the first sum and the second sum of R . We see that R_1 is associated to the interval $[u^{(1)}, u^{(2)}]$ and R_2 to the interval $[u^{(2)}, u^{(3)}]$. We compute R_1 and R_2 and make comparison between $\frac{R_1}{q_1 - p_1 + 1}$ and $\frac{R_2}{q_2 - p_2 + 1}$, where $q_i - p_i + 1$ is the number of data u_i from the interval $[u^{(i)}, u^{(i+1)}]$. If $\frac{R_j}{q_j - p_j + 1}$ is greater, $j \in \{1, 2\}$, then we divide the interval $[u^{(j)}, u^{(j+1)}]$ and insert here the new knot in order that the two new subintervals have equal number of data u_i . In this way we have four new knots and three intervals, three cubic polynomial as new functions, and eight parameters $y_i, M_i, i = \overline{1, 4}$. With the least squares method we minimize $R(y_1, M_1, y_2, M_2, y_3, M_3, y_4, M_4)$, which is an expression with 3 sums as in (10) and (11), obtaining the values of these eight parameters, as above. We compute $R_1(y_1, M_1, y_2, M_2)$, $R_2(y_2, M_2, y_3, M_3)$, $R_3(y_3, M_3, y_4, M_4)$ and $\frac{R_j}{q_j - p_j + 1} = \max \left\{ \frac{R_i}{q_i - p_i + 1}, i = 1, 2, 3 \right\}$, in order to divide the interval $[u^{(j)}, u^{(j+1)}]$ and to insert a new knot. The procedure continue and comes up with $n - 1$ intervals. Computing,

$$\frac{R_j}{q_j - p_j + 1} = \max \left\{ \frac{R_i}{q_i - p_i + 1}, i = \overline{1, n-1} \right\}$$

we divide the interval $[u^{(j)}, u^{(j+1)}]$ and insert a new knot. We obtain $n + 1$ new knots, $u^{(1)}, u^{(2)}, \dots, u^{(n)}, u^{(n+1)}, i = \overline{1, n}$. Similarly, we have n cubic polynomial functions $s_i, i = \overline{1, n}$ and $2(n + 1)$ parameters $y_i, M_i, i = \overline{1, n + 1}$,

$$y_1 = s_1(u^{(1)}), \quad M_1 = s_1''(u^{(1)}), \quad y_{n+1} = s_n(u^{(n+1)})$$

$$M_{n+1} = s_n''(u^{(n+1)}), \quad y_{i+1} = s_i(u^{(i+1)}) = s_{i+1}(u^{(i+1)})$$

$$M_{i+1} = s_i''(u^{(i+1)}) = s_{i+1}''(u^{(i+1)}), \quad i = \overline{1, n-1}$$

The values of the parameters were obtained by the least squares method, minimizing the similar expression $R(y_1, M_1, \dots, y_{n+1}, M_{n+1})$ and solving the normal equation $A_{n+1} z_{n+1} = g_{n+1}$, where z_{n+1}^T and g_{n+1}^T are $(n + 1)$ -vectors analogous with (13) and (14), and A_{n+1} is a symmetric matrix of dimension $(n + 1) \times (n + 1)$ analogous with A_3 . The elements of A_{n+1} have analogous form with (16), (17) and (18). This matrix has a band structure ([8] and [13]) and therefore we can apply the Householder reductions or the Gaussian elimination method to solve the normal equation.

In the following we present the criterion of convergence (as in [8] and [13]). Let,

$$\delta(n+1) = \frac{R(y_1, M_1, \dots, y_{n+1}, M_{n+1})}{p - 2(n+1)},$$

be the estimator of the variance of errors.

This estimator decrease when we insert a new knot. When $\delta(n+1) - \delta(n) < \varepsilon$, for $\varepsilon > 0$ choosen, we consider that we have reached the adequate fitting. Then the number $n+1$ is enough and we store the new knots $u^{(1)}, \dots, u^{(n+1)}$ and the new spline parameters $y_i, M_i, i = \overline{1, n+1}$.

3.3 The fitting spline function

With the knots $u^{(i)}$ and the values $y_i, M_i, i = \overline{1, n+1}$ we construct the fitting spline function $s : [-\tau, 0] \rightarrow \mathbb{R}$, which has the restrictions s_i (as in (5)) on $[u^{(i)}, u^{(i+1)}]$ for each $i = \overline{1, n}$. With the relations (8) we obtain the values $m_i, \overline{1, n+1}$ and then

$$s(u^{(i)}) = y_i, s'(u^{(i)}) = m_i, s''(u^{(i)}) = M_i, \forall i = \overline{1, n+1}.$$

We have that $s \in C^2[-\tau, 0]$. Indeed, from (5) we obtain that

$$s_i(u^{(i+1)}) = s_{i+1}(u^{(i+1)}), \quad \forall i = \overline{1, n}.$$

The relations (8) are equivalent with

$$s'_i(u^{(i+1)}) = s'_{i+1}(u^{(i+1)}), \quad \forall i = \overline{1, n}$$

and since

$$s''_i(t) = M_i + \frac{M_{i+1} - M_i}{h_i} \cdot (t - u^{(i)})$$

we have $s''_i(u^{(i+1)}) = s''_{i+1}(u^{(i+1)}), \forall i = \overline{1, n}$.

Remark 1 Unlike the papers [5] and [8], here we use only the parameters y_i, M_i . Of course, the values m_i are directly obtained from the relations (8). In [5] the parameters are y_i, m_i, M_i and the expression of the spline function is the same (5). In [8] the parameters are y_i, m_i , but the expression of the spline function is more complicated (see expression (1) in [8]). Because we choose at the beginning of the algorithm $u^{(1)} = -\tau$ and $u^{(3)} = 0$, then after the stop of this algorithm we have $u^{(1)} = -\tau$ and $u^{(n+1)} = 0$. The knots $u^{(i)}, i = \overline{1, n+1}$ are not equidistant in general. Since the positive values $f_1, f_2, \dots, f_p$ are large enough we can have $s(t) \geq 0$ for $t \in [-\tau, 0]$.

4 The approximation of the solution

From Theorem 1 it follows that the equation (1) has an unique positive, bounded and smooth solution on $[-\tau, T]$. Let φ be this solution, which can be obtained by successive approximations method on $[0, T]$.

So, we have :

$$\left\{ \begin{array}{l} \varphi_m(t) = \Phi(t) , \text{ for } t \in [-\tau, 0] , \quad \forall m \in \mathbb{N} \\ \varphi_0(t) = \Phi(0) = b = \int_{-\tau}^0 f(s, \Phi(s))ds , \\ \varphi_1(t) = \int_{t-\tau}^t f(s, \varphi_0(s))ds , \\ \dots\dots\dots \text{for } t \in [0, T] . \\ \varphi_m(t) = \int_{t-\tau}^t f(s, \varphi_{m-1}(s))ds , \\ \dots\dots\dots \end{array} \right. \quad (19)$$

To obtain the sequence of successive approximations (19) we compute the above integrals using a quadrature rule.

4.1 The quadrature rule

We have seen that the fitting spline function approximates the solution φ on the interval $[-\tau, 0]$ if $\varphi \in C^2[-\tau, 0]$, and the knots $u^{(i)}, i = \overline{1, n+1}$ cover this interval. We choose $T > 0$ such that $\exists l \in \mathbb{N}^*$ with $T = l\tau$, and we move by translation the knots $u^{(i)}, i = \overline{1, n+1}$ over the each subinterval of $[0, T]$ having the length τ . In this way we have $q+1 = l \cdot n+1$ knots on $[0, T]$. We denote these knots by $t_0, t_1, \dots, t_q$, and obtain a partition of $[0, T]$, $0 = t_0 < t_1 < \dots < t_q = T$ with

$$t_0 - \tau = -\tau = u^{(1)}, \dots, t_n - \tau = u^{(n+1)} = 0.$$

Let $h^{(i)} = u^{(i+1)} - u^{(i)}, i = \overline{1, n}$ and $h_i = t_i - t_{i-1}, i = \overline{0, q}$.

We see that $h_i = h^{(i)}, \forall i = \overline{1, n}$, and

$$h_n = h^{(1)}, \dots, h_{n+ki} = h^{(i)}, h_q = h^{(n)},$$

$\forall i = \overline{1, n}, \forall k = \overline{1, l-1}$. If we denote $t_{i-n} = u^{(i+1)}, \forall i = \overline{0, n}$ then we have $t_k - \tau = t_{k-n} \quad \forall k = \overline{1, q}$.

We apply the following quadrature rule from [1]:

$$\int_a^b F(t)dt = P_n(F, I_n) + R_n(F, I_n) , \quad (20)$$

where I_n is a division of $[a, b]$, $I_n : a = x_0 < x_1 < x_2 < \dots < x_{n-1} < x_n = b$, with $h_i = x_{i+1} - x_i, i = \overline{0, n-1}$ and

$$P_n(F, I_n) = \frac{1}{2} \sum_{i=0}^{n-1} h_i [F(x_i) + F(x_{i+1})] - \frac{1}{12} \sum_{i=0}^{n-1} h_i^2 [F'(x_{i+1}) - F'(x_i)]$$

with the remainder's estimation, in the case $F \in C^3[a, b]$:

$$|R_n(F, I_n)| \leq \frac{1}{160} \cdot \|F'''\| \cdot \sum_{i=0}^{n-1} h_i^4 . \quad (21)$$

For the integrals (19) we apply the rule (20) with the function $F_m(t) = f(t, \varphi_{m-1}(t))$, and according to (21) we need to realize an estimation for $F_m'''(t) = [f(t, x(t))]_t'''$, where $x = \varphi_{m-1}$.

After an elementary calculus we obtain :

$$F_m'''(t) = \frac{\partial^3 f}{\partial t^3}(t, x(t)) + 3 \frac{\partial^3 f}{\partial t^2 \partial x}(t, x(t))x'(t) + 3 \frac{\partial^3 f}{\partial t \partial x^2}(t, x(t)) \cdot \quad (22)$$

$$\begin{aligned} & \cdot [x'(t)]^2 + \frac{\partial^3 f}{\partial x^3}(t, x(t)) [x'(t)]^3 + 3 \frac{\partial^2 f}{\partial t \partial x}(t, x(t))x''(t) + \\ & + 3 \frac{\partial^2 f}{\partial x^2}(t, x(t))x'(t) \cdot x''(t) + \frac{\partial f}{\partial x}(t, x(t))x'''(t). \end{aligned}$$

Remark 2 $[f(t, x(t))]_t'''$ has 7 terms and $[f(t, x(t))]_t^{(IV)}$ has (after an elementary calculus) 13 terms. For this reason the use of a remainder which contains the fourth derivative is not efficient. Then we avoid the Simpson's or Newton's quadrature formulas.

Let

$$M_0 = M = \max \{|f(t, x)| : t \in [-\tau, T], x \in [a, \beta]\}$$

$$\left\| \frac{\partial^\alpha f}{\partial t^{\alpha_1} \partial x^{\alpha_2}} \right\| = \max \left\{ \left\| \frac{\partial^\alpha f(t, x)}{\partial t^{\alpha_1} \partial x^{\alpha_2}} \right\| : (t, x) \in [-\tau, T] \times [a, \beta] \right\}$$

for $\alpha = \overline{1, 3}$, $\alpha_1, \alpha_2 = \overline{0, 3}$, $\alpha_1 + \alpha_2 = \alpha$,

$$M_1 = \max \left\{ \left\| \frac{\partial f}{\partial t} \right\|, \left\| \frac{\partial f}{\partial x} \right\| \right\}$$

$$M_2 = \max \left\{ \left\| \frac{\partial^2 f}{\partial t^2} \right\|, \left\| \frac{\partial^2 f}{\partial t \partial x} \right\|, \left\| \frac{\partial^2 f}{\partial x^2} \right\| \right\}$$

$$M_3 = \max \left\{ \left\| \frac{\partial^3 f}{\partial t^3} \right\|, \left\| \frac{\partial^3 f}{\partial t^2 \partial x} \right\|, \left\| \frac{\partial^3 f}{\partial t \partial x^2} \right\|, \left\| \frac{\partial^3 f}{\partial x^3} \right\| \right\}.$$

Then from (22) we obtain :

$$\begin{aligned} |F_m'''(t)| & \leq M_3(1 + M_0)^3 + 5M_2M_1(1 + M_0)^2 + \\ & + 2M_1^2(1 + M_0) = M''', \quad \forall t \in [-\tau, T]. \end{aligned}$$

4.2 The algorithm

For the relations (19) and adapting the quadrature rule (20) at the knots $t_i, i = \overline{0, q}$ we obtain for $m \in \mathbb{N}^*$, and $k = \overline{1, q}$:

$$\varphi_m(t_k) = \int_{t_k - \tau}^{t_k} f(s, \varphi_{m-1}(s)) ds = \frac{1}{2} \sum_{i=k-n}^k h_i [f(t_i, \varphi_{m-1}(t_i)) + \quad (23)$$

$$\begin{aligned}
& + f(t_{i+1}, \varphi_{m-1}(t_{i+1})) - \frac{1}{12} \sum_{i=k-n}^k h_i^2 \cdot \left[\frac{\partial f}{\partial t}(t_{i+1}, \varphi_{m-1}(t_{i+1})) + \right. \\
& + \frac{\partial f}{\partial x}(t_{i+1}, \varphi_{m-1}(t_{i+1})) \cdot \varphi'_{m-1}(t_{i+1}) - \frac{\partial f}{\partial t}(t_i, \varphi_{m-1}(t_i)) - \\
& \left. - \frac{\partial f}{\partial x}(t_i, \varphi_{m-1}(t_i)) \cdot \varphi'_{m-1}(t_i) \right] + r_{m,k}^{(n)}(f) .
\end{aligned}$$

Also, we have

$$\left| r_{m,k}^{(n)}(f) \right| \leq \frac{M'''}{160} \cdot \sum_{i=k-n}^k h_i^4, \quad \forall k = \overline{1, q}. \quad (24)$$

if $f \in C^3([-\tau, T] \times [a, \beta])$. In (23), $\varphi_{m-1}(t_i - \tau) = s(t_i - \tau)$ and $\varphi'_{m-1}(t_i - \tau) = s'(t_i - \tau)$, for $i = \overline{1, n}$. Here, the values of the first derivative of φ_{m-1} are for $m \in \mathbb{N}, m \geq 2$,

$$\varphi'_{m-1}(t) = f(t, \varphi_{m-2}(t)) - f(t - \tau, \varphi_{m-2}(t - \tau)), \quad \forall t \in [0, T]$$

and $\varphi'_0(t_i) = \varphi'_1(t_i) = s'(t_i)$ for $t_i \in [-\tau, 0]$, $\varphi'_0(t) = 0$, $\forall t \in [0, T]$. We see that

$$\varphi'_1(t_i) = \begin{cases} f(t_i, b) - f(t_i - \tau, s(t_i)), & \forall i = \overline{1, n} \\ f(t_i, b) - f(t_i - \tau, b), & \forall i = \overline{n+1, q} \end{cases}$$

and $\varphi_0(t_i) = s(t_i)$ for $t_i \in [-\tau, 0]$, $\varphi_0(t) = b$, $\forall t \in [0, T]$.

Using for the successive approximations (19) and the formulas (23) with the remainder estimation (24), we can obtain an algorithm for the approximate solution of (2).

So, we have for $k = \overline{1, q}$:

$$\begin{aligned}
\varphi_1(t_k) &= \frac{1}{2} \sum_{i=k-n}^k h_i \cdot [f(t_i, \varphi_0(t_i)) + f(t_{i+1}, \varphi_0(t_{i+1}))] - \\
& \frac{1}{12} \sum_{i=k-n}^k h_i^2 \left[\frac{\partial f}{\partial t}(t_{i+1}, \varphi_0(t_{i+1})) + \frac{\partial f}{\partial x}(t_{i+1}, \varphi_0(t_{i+1})) \varphi'_0(t_{i+1}) \right. \\
& \left. - \frac{\partial f}{\partial t}(t_i, \varphi_0(t_i)) - \frac{\partial f}{\partial x}(t_i, \varphi_0(t_i)) \varphi'_0(t_i) \right] + r_{1,k}^{(n)}(f)
\end{aligned}$$

that is $\varphi_1(t_k) = \widetilde{\varphi}_1(t_k) + r_{1,k}^{(n)}(f)$ and

$$\begin{aligned}
\varphi_2(t_k) &= \frac{1}{2} \sum_{i=k-n}^k h_i [f(t_i, \widetilde{\varphi}_1(t_i) + r_{1,i}^{(n)}(f)) + f(t_{i+1}, \widetilde{\varphi}_1(t_{i+1}) + \\
& + r_{1,i+1}^{(n)}(f))] - \frac{1}{12} \sum_{i=k-n}^k h_i^2 \cdot \left[\frac{\partial f}{\partial t}(t_{i+1}, \widetilde{\varphi}_1(t_{i+1}) + r_{1,i+1}^{(n)}(f)) + \right.
\end{aligned}$$

$$\begin{aligned}
& + \frac{\partial f}{\partial x}(t_{i+1}, \widetilde{\varphi}_1(t_{i+1}) + r_{1,i+1}^{(n)}(f)) \cdot \varphi_1'(t_{i+1}) - \frac{\partial f}{\partial t}(t_i, \widetilde{\varphi}_1(t_i) + \\
& + r_{1,i}^{(n)}(f)) - \frac{\partial f}{\partial x}(t_i, \widetilde{\varphi}_1(t_i) + r_{1,i}^{(n)}(f)) \cdot \varphi_1'(t_i)] + r_{2,k}^{(n)}(f)
\end{aligned}$$

With the Lipschitz's property of f we have

$$\begin{aligned}
\varphi_2(t_k) &= \frac{1}{2} \sum_{i=k-n}^k h_i \cdot [f(t_i, \widetilde{\varphi}_1(t_i)) + f(t_{i+1}, \widetilde{\varphi}_1(t_{i+1}))] - \\
& - \frac{1}{12} \sum_{i=k-n}^k h_i^2 \cdot \left[\frac{\partial f}{\partial t}(t_{i+1}, \widetilde{\varphi}_1(t_{i+1})) + \frac{\partial f}{\partial x}(t_{i+1}, \widetilde{\varphi}_1(t_{i+1})) \cdot \right. \\
& \cdot (f(t_{i+1}, \varphi_0(t_{i+1})) - f(t_{i+1-n}, \varphi_0(t_{i+1-n}))) - \frac{\partial f}{\partial t}(t_i, \widetilde{\varphi}_1(t_i)) - \\
& - \frac{\partial f}{\partial x}(t_i, \widetilde{\varphi}_1(t_i)) \cdot (f(t_i, \varphi_0(t_i)) - f(t_{i-n}, \varphi_0(t_{i-n}))) \Big] + \\
& + \widetilde{r_{2,k}^{(n)}(f)} = \widetilde{\varphi_2(t_k)} + \widetilde{r_{2,k}^{(n)}(f)}, \quad \forall k = \overline{1, q}.
\end{aligned} \tag{25}$$

Let $h = \nu(I_n) = \max \{h^{(i)}, i = \overline{1, n}\}$. Then, $h = \max \{h_i, i = \overline{1, q-1}\}$. With these, from (24) and (25) we obtain for the remainders the estimation :

$$\left| r_{1,k}^{(n)}(f) \right| \leq \frac{M'''nh^4}{160}, \text{ and}$$

$$\left| \widetilde{r_{2,k}^{(n)}(f)} \right| \leq \frac{M'''}{160} \left(nh^4 + n^2 h^5 L + \frac{n^2 h^6}{3} M_2 \right), \forall k = \overline{1, q} \tag{26}$$

Continuing by induction for $m \geq 3$, from (23), (25) and (26) we obtain for $m \in \mathbb{N}$, $m > 2$, $k = \overline{1, q}$:

$$\begin{aligned}
\varphi_m(t_k) &= \frac{1}{2} \sum_{i=k-n}^k h_i [f(t_i, \widetilde{\varphi_{m-1}}(t_i) + \widetilde{r_{m-1,i}^{(n)}(f)}) + f(t_{i+1}, \widetilde{\varphi_{m-1}}(t_{i+1}) \\
& + \widetilde{r_{m-1,i+1}^{(n)}(f)})] - \frac{1}{12} \sum_{i=k-n}^k h_i^2 \left[\frac{\partial f}{\partial t}(t_{i+1}, \widetilde{\varphi_{m-1}}(t_{i+1}) + \widetilde{r_{m-1,i+1}^{(n)}(f)}) + \right. \\
& + \frac{\partial f}{\partial x}(t_{i+1}, \widetilde{\varphi_{m-1}}(t_{i+1}) + \widetilde{r_{m-1,i+1}^{(n)}(f)}) \cdot [f(t_{i+1}, \widetilde{\varphi_{m-2}}(t_{i+1}) + \\
& + \widetilde{r_{m-2,i+1}^{(n)}(f)}) - f(t_{i+1-n}, \widetilde{\varphi_{m-2}}(t_{i+1-n}) + \widetilde{r_{m-2,i+1-n}^{(n)}(f)})] - \\
& - \frac{\partial f}{\partial t}(t_i, \widetilde{\varphi_{m-1}}(t_i) + \widetilde{r_{m-1,i}^{(n)}(f)}) - \frac{\partial f}{\partial x}(t_i, \widetilde{\varphi_{m-1}}(t_i) + \widetilde{r_{m-1,i}^{(n)}(f)}) \cdot \\
& \cdot \{f(t_i, \widetilde{\varphi_{m-2}}(t_i) + \widetilde{r_{m-2,i}^{(n)}(f)}) - f(t_{i-n}, \widetilde{\varphi_{m-2}}(t_{i-n}) + \widetilde{r_{m-2,i-n}^{(n)}(f)}) \} \Big]
\end{aligned}$$

$$\begin{aligned}
+r_{m,k}^{(n)}(f) &= \frac{1}{2} \sum_{i=k-n}^k h_i [f(t_i, \widetilde{\varphi_{m-1}}(t_i)) + f(t_{i+1}, \widetilde{\varphi_{m-1}}(t_{i+1}))] - \\
&- \frac{1}{12} \sum_{i=k-n}^k h_i^2 \cdot [\frac{\partial f}{\partial t}(t_{i+1}, \widetilde{\varphi_{m-1}}(t_{i+1})) + \frac{\partial f}{\partial x}(t_{i+1}, \widetilde{\varphi_{m-1}}(t_{i+1})) \cdot \\
&\quad \cdot \{f(t_{i+1}, \widetilde{\varphi_{m-2}}(t_{i+1})) - f(t_{i+1-n}, \widetilde{\varphi_{m-2}}(t_{i+1-n}))\} - \\
&- \frac{\partial f}{\partial t}(t_i, \widetilde{\varphi_{m-1}}(t_i)) - \frac{\partial f}{\partial x}(t_i, \widetilde{\varphi_{m-1}}(t_i)) \cdot \{f(t_i, \widetilde{\varphi_{m-2}}(t_i)) - \\
&- f(t_{i-n}, \widetilde{\varphi_{m-2}}(t_{i-n}))\}] + r_{m,k}^{(n)}(f) = \widetilde{\varphi_m}(t_k) + r_{m,k}^{(n)}(f). \quad (27)
\end{aligned}$$

4.3 Main result

We can obtain a concise estimation for $\left| \widetilde{r_{m,k}^{(n)}}(f) \right|$. Indeed, for $m \geq 2$ and $k = \overline{1, q}$ we have:

$$\begin{aligned}
\left| \widetilde{r_{m,k}^{(n)}}(f) \right| &\leq \left| r_{m,k}^{(n)}(f) \right| + Lnh \cdot \left| \widetilde{r_{m-1,k}^{(n)}}(f) \right| + \\
&+ \frac{1}{12} \cdot 2nh^2(M_2(M+1) + M_1L) \left| \widetilde{r_{m-1,k}^{(n)}}(f) \right| = \left| r_{m,k}^{(n)}(f) \right| + \\
&+ [Lnh + \frac{1}{6} \cdot nh^2(M_2(M+1) + M_1L)] \cdot \left| \widetilde{r_{m-1,k}^{(n)}}(f) \right|.
\end{aligned}$$

and $\left| \widetilde{r_{1,k}^{(n)}}(f) \right| = \left| r_{1,k}^{(n)}(f) \right| \leq \frac{M'''nh^4}{160}$, $\forall k = \overline{1, q}$. Then we obtain

$$\begin{aligned}
\left| \widetilde{r_{2,k}^{(n)}}(f) \right| &\leq \left| r_{2,k}^{(n)}(f) \right| + [Lnh + \frac{nh^2}{6}(M_2(M+1) + M_1L)] \cdot \\
&\cdot \left| r_{1,k}^{(n)}(f) \right| \leq [1 + Lnh + \frac{nh^2}{6}(M_2(M+1) + M_1L)] \cdot \frac{M'''nh^4}{160}
\end{aligned}$$

and

$$\begin{aligned}
\left| \widetilde{r_{3,k}^{(n)}}(f) \right| &\leq \left| r_{3,k}^{(n)}(f) \right| + [Lnh + \frac{nh^2}{6}(M_2(M+1) + M_1L)] \cdot \\
&\cdot \left| \widetilde{r_{2,k}^{(n)}}(f) \right| \leq [1 + (Lnh + \frac{nh^2}{6}(M_2(M+1) + M_1L)) + \\
&+ (Lnh + \frac{nh^2}{6}(M_2(M+1) + M_1L))^2] \cdot \frac{M'''nh^4}{160}, \quad \forall k = \overline{1, q}.
\end{aligned}$$

By induction we obtain for $m \in \mathbb{N}, m \geq 2$, and $k = \overline{1, q}$:

$$\left| \widetilde{r_{m,k}^{(n)}}(f) \right| \leq [1 + (Lnh + \frac{nh^2}{6}(M_2(M+1) + M_1L)) + \dots +$$

$$\begin{aligned}
& + (Ln h + \frac{1}{6} \cdot n h^2 (M_2(M+1) + M_1 L))^{m-1}] \cdot \frac{M''' n h^4}{160} = \\
& = \frac{1 - [Ln h + \frac{1}{6} \cdot n h^2 (M_2(M+1) + M_1 L)]^m}{1 - [Ln h + \frac{1}{6} \cdot n h^2 (M_2(M+1) + M_1 L)]} \cdot \frac{M''' n h^4}{160}.
\end{aligned}$$

Since $h = O(\frac{\tau}{n})$ we can consider $\mu = \frac{hn}{\tau}$. If $L\mu\tau < \frac{3}{4}$ and $n \in \mathbb{N}$ is such that

$$Ln h + \frac{1}{6} \cdot n h^2 (M_2(M+1) + M_1 L) < 1,$$

then for $m \in \mathbb{N}, m \geq 2$, and $k = \overline{1, q}$, we have the estimation

$$\left| \widetilde{r_{m,k}^{(n)}}(f) \right| \leq \frac{M''' n h^4}{160(1 - [Ln h + \frac{1}{6} \cdot n h^2 (M_2(M+1) + M_1 L)])}. \quad (28)$$

In this way we got the sequence $(\widetilde{\varphi_m}(t_k))_{m \in \mathbb{N}^*}$, which approximates the sequence of successive approximations (19) on the knots $t_k, k = \overline{1, q}$, with the error : $|\varphi_m(t_k) - \widetilde{\varphi_m}(t_k)| = \left| \widetilde{r_{m,k}^{(n)}}(f) \right|$ which has the estimation (28).

Proposition 3 Consider the initial value problem (2) under the conditions of Theorem 1 and of Corollary 2. If $f \in C^3([-\tau, T] \times [a, \beta])$, $\Phi \in C^3[-\tau, 0]$, $L\mu\tau < \frac{3}{4}$ and the exact solution φ is approximated by the sequence $(\widetilde{\varphi_m}(t_k))_{m \in \mathbb{N}^*}$, on the knots $t_k, k = \overline{1, q}$, through the successive approximation method (19), combined with the quadrature rule (20), then the following error estimation holds

$$\begin{aligned}
|\varphi(t_k) - \widetilde{\varphi_m}(t_k)| & \leq \frac{M''' n h^4}{160(1 - [Ln h + \frac{1}{6} \cdot n h^2 (M_2(M+1) + LM_1)])} + \\
& + \frac{\tau^m \cdot L^m}{1 - \tau L} \cdot \|\varphi_0 - \varphi_1\|_{C[0, T]}, \quad \forall m \in \mathbb{N}^*, \forall k = \overline{1, q},
\end{aligned} \quad (29)$$

for $n \in \mathbb{N}$ large enough.

Proof. From Theorem 1, we have the estimation

$$|\varphi(t_k) - \varphi_m(t_k)| \leq \frac{\tau^m L^m}{1 - \tau L} \cdot \|\varphi_0 - \varphi_1\|_{C[0, T]}, \forall k = \overline{1, q},$$

$\forall m \in \mathbb{N}^*$. Since $|\varphi_m(t_k) - \widetilde{\varphi_m}(t_k)| = \left| \widetilde{r_{m,k}^{(n)}}(f) \right|$, from this estimation and from the inequality (28), we get the estimation (29). ■

Remark 3 In the equidistant case, $h_i = h = \frac{\tau}{n}, \forall i = \overline{1, n}$, after analogous calculus we obtain for $m \in \mathbb{N}, m \geq 2$, and $k = \overline{1, q}$ the estimation

$$\left| \widetilde{r_{m,k}^{(n)}}(f) \right| \leq \frac{1 - [\tau L + \frac{1}{6n^2}(\tau^2(M_2(M+1)) + LM_1)]^m}{1 - [\tau L + \frac{1}{6n^2}(\tau^2(M_2(M+1)) + LM_1)]} \cdot \frac{\tau^4 M'''}{160n^3}.$$

Finally, we obtain the main result of our paper :

Theorem 4 *In the conditions of Proposition 3 , if $h_i = \frac{\tau}{n}, \forall i = \overline{1, n}$, $\tau L < \frac{3}{4}$ and $n \in \mathbb{N}$ is such that $n > [\frac{2}{3}(M_2(M+1) + LM_1)]^{\frac{1}{2}}$ then for any $m \in \mathbb{N}, m \geq 2$ we have*

$$|\varphi(t_k) - \widetilde{\varphi}_m(t_k)| \leq \frac{\tau^m L^m}{1 - \tau L} \cdot \|\varphi_0 - \varphi_1\|_{C[0, T]} + \frac{\tau^4 M'''}{160n^3 (1 - [\tau L + \frac{1}{6n^2}(\tau^2(M_2(M+1) + LM_1))])}, \forall k = \overline{1, q}.$$

Proof. Follows from Proposition 3 and previous remark. ■

The above presented numerical method is also useful for first order ODE.

Remark 4 *If we use the trapezoidal quadrature rule, then we obtain the values $\widetilde{\varphi}_m(t_k)$ from the first sum in (23). The remainder estimation is*

$$\left| \widetilde{r}_{m,k}^{(n)}(f) \right| \leq \frac{1 - (Ln h)^m}{1 - Ln h} \cdot \frac{M'' n h^3}{12}$$

where, $M'' = M_2(M_0 + 1)^2 + M_1(M_0 + 1)$.

4.4 An example

To illustrate the algorithm we give an example in the conditions of Theorem 4. Here, $\tau = 1$, $T = 3$, $n = 10$, $q = 30$ and

$$\Phi(t) = \frac{1}{4}[\cos^2 \pi t + 0.25 \cdot \sin^2 2\pi t] + 0.1$$

$$f(t, x) = \frac{1}{20} \sin^2 \pi t + \frac{4x + 3}{10(4x + 5)} + 0.25832.$$

We have $\beta = 2$, $a = 0.1$, $m = 0.2$, $M = 0.5$, $\varphi(0) = \frac{7}{20}$, $L = \frac{4}{125}$, $M_1 = \frac{1}{20}\pi$, $M_2 = \frac{1}{10}\pi^2$ and $M''' \leq 24$.

For given $\varepsilon > 0$ we find the smallest $m \in \mathbb{N}$ such that

$$|\widetilde{\varphi}_m(t_k) - \widetilde{\varphi}_{m-1}(t_k)| < \varepsilon, \quad \forall k = \overline{1, q}.$$

In the following table there are the values $\widetilde{\varphi}_m(t_k)$, for $k = \overline{1, 30}$. For $\varepsilon = 10^{-12}$ we get $m = 8$.

t_k	$\widetilde{\varphi}_m(t_k)$	t_k	$\widetilde{\varphi}_m(t_k)$	t_k	$\widetilde{\varphi}_m(t_k)$
t_1	0.34795627524	t_{11}	0.35207611617	t_{21}	0.35211158318
t_2	0.34816606071	t_{12}	0.35208360617	t_{22}	0.35211127612
t_3	0.34849749458	t_{13}	0.35209051377	t_{23}	0.36211086721
t_4	0.34906707741	t_{14}	0.35209665275	t_{24}	0.35211053170
t_5	0.34986045867	t_{15}	0.35210167720	t_{25}	0.35211041254
t_6	0.35065567756	t_{16}	0.35210539510	t_{26}	0.35211056646
t_7	0.35123022656	t_{17}	0.35210798938	t_{27}	0.35211094266
t_8	0.35156856971	t_{18}	0.35210979766	t_{28}	0.35211140258
t_9	0.35178615824	t_{19}	0.35211100019	t_{29}	0.35211177352
t_{10}	0.35197711605	t_{20}	0.35211158784	t_{30}	0.35211191551

Acknowledgement

The authors are indebted to Professor Horea Oros from University of Oradea for his help in computer implementation of the algorithm.

References

- [1] Barnett N.S., Dragomir S.S., Perturbed trapezoidal inequality in terms of the third derivative and applications, *RGMIA Preprint* vol.4, no.2, (2001), 221-233.
- [2] Bica A., Numerical methods for approximating the solution of a delay integral equation from biomathematics solution of a delay , (submitted).
- [3] Bica A., Mangrău I., Mezei R., Numerical method for Volterra integral equations, *Proceedings of the 11-th Conference on Applied and Industrial Mathematics* (2003), 29-33, Oradea, Romania.
- [4] Guo D., Lakshmikantham V., Positive solutions of nonlinear integral equations arising in infectious diseases, *J. Math. Anal. Appl.* 134 (1988), 1-8.
- [5] Iancu C., *The analysis and data processing with spline functions*, PhD Thesis, "Babes-Bolyai" University, Cluj-Napoca, Romania, 1983 (in Romanian).
- [6] Iancu C., On the cubic spline of interpolation, *Seminar of Functional Analysis and Numerical Methods, Preprint* 4, Cluj-Napoca, (1981), 52-71.
- [7] Iancu C., A numerical method for approximating the solution of an integral equation from biomathematics, *Studia Univ. "Babes-Bolyai", Mathematica*, vol. XLIII, No.4, (1988), 37-45.
- [8] Ichida K., Yoshimoto F., Kiyono T., Curve Fitting by a Piecewise Cubic Polynomial, *Computing* 16, (1976), 329-338.
- [9] Precup R., Positive solutions of the initial value problem for an integral equation modelling infectious diseases, *Seminar of Fixed Point Theory, Preprint* 3, Cluj-Napoca, (1991), 25-30.
- [10] Precup R., Periodic solutions for an integral equation from biomathematics via Leray-Schauder principle, *Studia Univ. "Babes-Bolyai", Mathematica*, vol. XXXIX, 1, (1994), 47-58.
- [11] Rus I.A., *Principles and applications of the fixed point theory*, Ed. Dacia, Cluj-Napoca, 1979 (in Romanian).
- [12] Rus I.A., A delay integral equation from biomathematics, *Seminar of Differential Equations, Preprint* 3, Cluj-Napoca, (1989), 87-90.

- [13] Yoshimoto F., *Studies on data fitting with spline functions*, PhD Thesis, Kyoto University, Japan, 1977.
- [14] L.R. Williams, R.W. Legget, Nonzero solutions of nonlinear integral equations modeling infectious disease, *SIAM J. Math. Anal.*, 13 (1982), 112-121.

Set-valued Variational Inclusions with Fuzzy Mappings in Banach Spaces

A. H. Siddiqi

Department of Mathematical Sciences
King Fahd University of Petroleum and Minerals
P.O. Box 1745, Dhahran 31261, Saudi Arabia
E-mail: ahasan@kfupm.edu.sa

Rais Ahmad

Department of Mathematics
Aligarh Muslim University
Aligarh 202 002, India
E-mail: raisain@lycos.com

and

S. S. Irfan

Department of Matematicas
Aligarh Muslim University
Aligarh 202 002, India
E-mail: shakaib11@rediffmail.com

Abstract. In this paper, we study set-valued variational inclusions with fuzzy mappings in the setting of real Banach spaces. By using Nadler's Theorem and the resolvent operator technique for m -accretive mappings, we propose an iterative algorithm for computing the approximate solutions of this class of set-valued variational inclusions. We also discuss the existence of a solution of our problem without compactness assumption and prove the convergence of the iterative sequences generated by the proposed algorithm.

Keywords: Variational inclusion, Algorithm, Convergence, Resolvent operators, m -accretive mappings.

2000 AMS subject Classification: 49J40, 47S40, 47H06

1. INTRODUCTION

The theory of variational inequalities have had a great impact and influence in the development of almost all branches of pure and applied sciences. Variational inequalities have been generalized and extended in different directions using novel and innovative techniques. An important and useful generalization of variational inequalities is called the variational inclusion, which is mainly due to Hassouni and Moudafi [10]. Ansari [1] and Chang and Zhu [8], separately, introduced and studied a class of variational inequalities for fuzzy mappings. Several classes of variational and quasi-variational inequalities with fuzzy mappings are considered and studied by Chang [3], Chang and Huang [7], Huang [11], Irfan [13] and Lee et al [14]. Recently, Chang et al [5, 6] introduced and studied the following class of set-valued variational inclusion problems in the setting of Banach space E .

For a given m -accretive mapping $A : E \rightarrow 2^E$, a nonlinear mapping $N : E \times E \rightarrow E$, set-valued mappings $T, F : E \rightarrow CB(E)$ and a single-valued mapping $g : E \rightarrow E$, find $x \in E$, $w \in T(x)$ and $q \in F(x)$ such that

$$0 \in N(w, q) + A(g(x)), \quad (1)$$

where $CB(E)$ denotes the family of all nonempty closed and bounded subsets of E .

The purpose of this paper is to study the set-valued variational inclusion (1) with fuzzy mappings without compactness condition in real Banach spaces. By using the resolvent technique for m -accretive mappings, we establish the equivalence between the set-valued variational inclusions with fuzzy mappings and the resolvent equation in the setting of real Banach spaces. We use this equivalence and Nadler's Theorem [15] to propose an iterative algorithm for solving our class of set-valued variational inclusions with fuzzy mappings in real Banach spaces. The inequality of Petryshyn [16] is used to establish the existence of a solution of our inclusion without compactness assumption and prove the convergence of iterative sequences generated by the algorithm.

2. PRELIMINARIES

Throughout the paper, we assume that E is a real Banach space, E^* is the topological dual space of E , $CB(E)$ is the family of nonempty closed and bounded subsets of E , $D(., .)$ is the Hausdorff metric on $CB(E)$ defined by

$$\tilde{D}(A, B) = \max \left\{ \sup_{x \in A} d(x, B), \sup_{y \in B} d(A, y) \right\}, \quad A, B \in CB(E),$$

We denote by $\langle ., . \rangle$ the dual pairing between E and E^* and by $D(T)$ the domain of T . We also assume that $J : E \rightarrow 2^{E^*}$ is the normalized duality mapping defined by

$$J(x) = \{f \in E^* : \langle x, f \rangle = \|x\| \quad \|f\| \quad \text{and} \quad \|f\| = \|x\|\}, \quad x \in E.$$

Definition 1. [2] A set-valued mapping $A : D(A) \subset E \rightarrow 2^E$ is said to be

- (i) *accretive* if for any $x, y \in D(A)$ there exists $j(x - y) \in J(x - y)$ such that for all $u \in Ax$ and $v \in Ay$,

$$\langle u - v, j(x - y) \rangle \geq 0;$$

- (ii) *k-strongly accretive*, $k \in (0, 1)$, if for any $x, y \in D(A)$, there exists $j(x - y) \in J(x - y)$ such that for all $u \in Ax$ and $v \in Ay$,

$$\langle u - v, j(x - y) \rangle \geq k \|x - y\|^2;$$

- (iii) *m-accretive* if A is accretive and $(I + \rho A)(D(A)) = E$ for every (equivalently, for some) $\rho > 0$, where I is the identity mapping, (equivalently, if A is accretive and $(I + A)(D(A)) = E$).

Remark 1. If $E = E^* = H$ is a Hilbert space, then $A : D(A) \subset H \rightarrow 2^H$ is an m -accretive mapping if and only if $A : D(A) \subset H \rightarrow 2^H$ is a maximal monotone mapping, see for example [9].

Let E be a real Banach space. Let $\mathcal{F}(E)$ be a collection of fuzzy sets over E . A mapping $F : E \rightarrow \mathcal{F}(E)$ is said to be *fuzzy mapping*. For each $x \in E$, $F(x)$ (in the sequel, it will be denoted by F_x) is a fuzzy set on E and $F_x(y)$ is the *membership function of y in $F(x)$* .

A fuzzy mapping $F : E \rightarrow \mathcal{F}(E)$ is said to be *closed* if for each $x \in E$, the function $y \mapsto F_x(y)$ is upper semicontinuous, that is, for any given net $\{y_\alpha\} \subset E$ satisfying $y_\alpha \rightarrow y_0 \in E$ we have $\limsup_\alpha F_x(y_\alpha) \leq F_x(y_0)$.

For $B \in \mathcal{F}(E)$ and $\lambda \in [0, 1]$, the set $(B)_\lambda = \{x \in E : B(x) \geq \lambda\}$ is called a λ -*cut* set of B .

A closed fuzzy mapping $A : E \rightarrow \mathcal{F}(E)$ is said to satisfy condition $(*)$:

if there exists a function $a : E \rightarrow [0, 1]$ such that for each $x \in E$, $(A_x)_{a(x)}$ is a nonempty bounded subset of E .

It is clear that if A is a closed fuzzy mapping satisfying condition $(*)$, then for each $x \in E$, the set $(A_x)_{a(x)} \in CB(E)$.

In fact, let $\{y_\alpha\}_{\alpha \in \Gamma} \subset (A_x)_{a(x)}$ be a net and $y_\alpha \rightarrow y_0 \in E$. Then $(A_x)(y_\alpha) \geq a(x)$ for each $\alpha \in \Gamma$. Since A is closed, we have

$$A_x(y_0) \geq \limsup_{\alpha \in \Gamma} A_x(y_\alpha) \geq a(x).$$

This implies that $y_0 \in (A_x)_{a(x)}$ and so $(A_x)_{a(x)} \in CB(E)$.

Let $T, F : E \rightarrow \mathcal{F}(E)$ be two closed fuzzy mappings satisfying condition $(*)$. Then there exists two functions $a, b : E \rightarrow [0, 1]$ such that for each $x \in E$, we have $(T_x)_{a(x)}, (F_x)_{b(x)} \in CB(E)$. Therefore, we can define two set-valued mappings $\tilde{T}, \tilde{F} : E \rightarrow CB(E)$ by

$$\tilde{T}(x) = (T_x)_{a(x)}, \quad \tilde{F}(x) = (F_x)_{b(x)}, \quad \text{for each } x \in E.$$

In the sequel, $\tilde{T}$ and $\tilde{F}$ are called the *set-valued mappings induced by fuzzy mappings* T and F , respectively.

Let $N : E \times E \rightarrow E$ and $g : E \rightarrow E$ be the single-valued mappings and let $T, F : E \rightarrow \mathcal{F}(E)$ be fuzzy mappings. Let $a, b : E \rightarrow [0, 1]$ be given functions. Suppose that $A : E \rightarrow 2^E$ is an m -accretive mapping. We consider the following set-valued variational inclusion problem with fuzzy mappings in Banach spaces.

$$(SVIPFM) \quad \begin{cases} \text{Find } x, w, q \in E \text{ such that} \\ T_x(w) \geq a(x), \quad F_x(q) \geq b(x) \text{ and} \\ 0 \in N(w, q) + A(g(x)). \end{cases}$$

If $T, F : E \rightarrow CB(E)$ are classical set-valued mappings, we can define fuzzy mappings $T, F : E \rightarrow \mathcal{F}(E)$ by

$$x \rightarrow \mathcal{X}_{T(x)}, \quad x \rightarrow \mathcal{X}_{F(x)},$$

where $\mathcal{X}_{T(x)}$ and $\mathcal{X}_{F(x)}$ are characteristic functions of $T(x)$ and $F(x)$, respectively.

Taking $a(x) = b(x) = 1$ for all $x \in E$, then (SVIPFM) is equivalent to problem (1). For a suitable choice of mappings T, F, A, g, N and the space E , a number of known and new classes of variational inequalities and variational inclusions studied in [4, 5, 6, 13] can be obtained from (SVIPFM).

3. ITERATIVE ALGORITHM

Definition 2. Let $A : D(A) \subset E \rightarrow 2^E$ be an m -accretive mapping. For any $\rho > 0$, the mapping $R_\rho^A : E \rightarrow D(A)$ associated with A defined by

$$R_\rho^A(x) = (I + \rho A)^{-1}(x), \quad x \in E,$$

is called the *resolvent operator*.

Remark 2. We note that R_ρ^A is a single-valued and nonexpansive mapping, see for example [2].

We first transfer (SVIPFM) into a fixed point problem.

Theorem 1. (x, w, q) is a solution of (SVIPFM) if and only if (x, w, q) satisfies the following relation

$$g(x) = R_\rho^A[g(x) - \rho N(w, q)], \quad (2)$$

where $w \in \tilde{T}(x)$, $q \in \tilde{F}(x)$, and $\rho > 0$ is a constant.

Proof. By the definition of resolvent operator R_ρ^A associated with A , we have that (2) holds if and only if $w \in \tilde{T}(x)$ and $q \in \tilde{F}(x)$ such that

$$g(x) - \rho N(w, q) \in g(x) + \rho A(g(x))$$

if and only if

$$0 \in N(w, q) + A(g(x)).$$

Hence (x, w, q) is a solution of (SVIPFM) if and only if $w \in \tilde{T}(x)$ and $q \in \tilde{F}(x)$ are such that (2) holds.

We also need the following lemma.

Lemma 1. [12] Let $g : E \rightarrow E$ be a continuous and k -strongly accretive mapping. Then g maps E onto E .

We now invoke Theorem 1, Lemma 1 and Nadler's Theorem [15] to propose the following iterative algorithm.

Algorithm 1. Let $T, F : E \rightarrow \mathcal{F}(E)$ be two closed fuzzy mappings satisfying condition (*) and $\tilde{T}, \tilde{F} : E \rightarrow CB(E)$ be the set-valued mappings induced by the fuzzy mappings T and F , respectively. Let $N : E \times E \rightarrow E$ be a single-valued bifunction and $g : E \rightarrow E$ a continuous and k -strongly accretive mapping. For given $x_0 \in E$, $w_0 \in \tilde{T}(x_0)$ and $q_0 \in \tilde{F}(x_0)$, let

$$g(x_1) = R_\rho^A[g(x_0) - \rho N(w_0, q_0)].$$

Since $w_0 \in \tilde{T}(x_0) \in CB(E)$ and $q_0 \in \tilde{F}(x_0) \in CB(E)$, by Nadler [15], there exist $w_1 \in \tilde{T}(x_1)$ and $q_1 \in \tilde{F}(x_1)$ such that

$$\|w_0 - w_1\| \leq (1 + 1)\tilde{D}(\tilde{T}(x_0), \tilde{T}(x_1)),$$

$$\|q_0 - q_1\| \leq (1 + 1)\tilde{D}(\tilde{F}(x_0), \tilde{F}(x_1)).$$

Let

$$g(x_2) = R_\rho^A[g(x_1) - \rho N(w_1, q_1)].$$

Continuing the above process inductively, we can obtain sequences $\{x_n\}$, $\{w_n\}$ and $\{q_n\}$ satisfying

$$g(x_{n+1}) = R_\rho^A[g(x_n) - \rho N(w_n, q_n)],$$

$$w_n \in \tilde{T}(x_n), \quad \|w_n - w_{n+1}\| \leq \left(1 + \frac{1}{n+1}\right) \tilde{D}(\tilde{T}(x_n), \tilde{T}(x_{n+1})),$$

$$q_n \in \tilde{F}(x_n), \quad \|q_n - q_{n+1}\| \leq \left(1 + \frac{1}{n+1}\right) \tilde{D}(\tilde{F}(x_n), \tilde{F}(x_{n+1})),$$

for all $n = 0, 1, 2, \dots$

4. EXISTENCE AND CONVERGENCE RESULTS

In this section, we establish the existence of a solution of (SVIPFM) and prove the convergence of iterative sequences generated by Algorithm 1.

Definition 3. A set-valued mapping $T : E \rightarrow CB(E)$ is said to be ξ -Lipschitz continuous if for any $x, y \in E$

$$\tilde{D}(Tx, Ty) \leq \xi \|x - y\|,$$

where $\xi > 0$ is a constant.

Definition 4. A mapping $N : E \times E \rightarrow E$ is said to be Lipschitz continuous in the first argument if there exists a constant $\alpha > 0$ such that

$$\|N(x_1, \cdot) - N(x_2, \cdot)\| \leq \alpha \|x_1 - x_2\|, \quad \text{for all } x_1, x_2 \in E.$$

In a similar way, we can define the Lipschitz continuity of N in the second argument.

The following lemma plays an important role to prove our main result.

Lemma 2. [16] Let E be a real Banach space and $J : E \rightarrow 2^{E^*}$ be the normalized duality mapping. Then for any $x, y \in E$,

$$\|x + y\|^2 \leq \|x\|^2 + 2\langle y, j(x + y) \rangle, \quad \text{for all } j(x + y) \in J(x + y).$$

Theorem 2. Let E be a real Banach space. Let $T, F : E \rightarrow \mathcal{F}(E)$ be two closed fuzzy mappings satisfying condition $(\star)$ and $\tilde{T}, \tilde{F} : E \rightarrow CB(E)$ the set-valued mappings induced by the fuzzy mappings T, F , respectively. Let $A : E \rightarrow 2^E$ be an m -accretive mapping and $g : E \rightarrow E$ σ -Lipschitz continuous and k -strongly accretive. Assume that $\tilde{T}$ is μ -Lipschitz continuous, $\tilde{F}$ is γ -Lipschitz continuous and $N : E \times E \rightarrow E$ is Lipschitz continuous in both the arguments with constants α, β , respectively. If

$$\rho < \frac{2k - 1 - \sigma^2}{2(\beta\gamma + \alpha\mu)\sqrt{2k - 1}}, \quad k > \frac{1}{2} \quad (3)$$

then there exist $x \in E$, $w \in \tilde{T}(x)$, $q \in \tilde{F}(x)$ such that (x, w, q) is a solution of (SVIPFM) and $x_n \rightarrow x$, $w_n \rightarrow w$, $q_n \rightarrow q$ as $n \rightarrow \infty$.

Proof. Since $\tilde{T}$ is μ -Lipschitz continuous and $\tilde{F}$ is γ -Lipschitz continuous, it follows from Algorithm 1 that

$$\begin{aligned} \|w_n - w_{n+1}\| &\leq \left(1 + \frac{1}{n+1}\right) \tilde{D}(\tilde{T}(x_n), \tilde{T}(x_{n+1})) \leq \left(1 + \frac{1}{n+1}\right) \mu \|x_n - x_{n+1}\|, \\ \|q_n - q_{n+1}\| &\leq \left(1 + \frac{1}{n+1}\right) \tilde{D}(\tilde{F}(x_n), \tilde{F}(x_{n+1})) \leq \left(1 + \frac{1}{n+1}\right) \gamma \|x_n - x_{n+1}\|, \end{aligned} \quad (4)$$

for all $n = 0, 1, 2, \dots$. Let $z_{n+1} = g(x_n) - \rho(N(w_n, q_n))$ for all $n = 0, 1, 2, \dots$. Since g is σ -Lipschitz continuous, N is Lipschitz continuous in both the arguments with constants α and β , respectively, and by using Lemma 2 and (4), we have

$$\|z_{n+1} - z_n\|^2 = \|g(x_n) - g(x_{n-1}) - \rho[N(w_n, q_n) - N(w_{n-1}, q_{n-1})]\|^2$$

$$\begin{aligned}
&\leq \|g(x_n) - g(x_{n-1})\|^2 - 2\rho \langle (N(w_n, q_n) - N(w_{n-1}, q_{n-1})), j(z_{n+1}, z_n) \rangle \\
&\leq \sigma^2 \|x_n - x_{n-1}\|^2 + 2\rho \|N(w_n, q_n) - N(w_{n-1}, q_{n-1})\| \|z_{n+1} - z_n\| \\
&= \sigma^2 \|x_n - x_{n-1}\|^2 + 2\rho \|N(w_n, q_n) - N(w_n, q_{n-1}) + N(w_n, q_{n-1}) \\
&\quad - N(w_{n-1}, q_{n-1})\| \|z_{n+1} - z_n\| \\
&\leq \sigma^2 \|x_n - x_{n-1}\|^2 + 2\rho \{ \|N(w_n, q_n) - N(w_n, q_{n-1})\| \\
&\quad + \|N(w_n, q_{n-1}) - N(w_{n-1}, q_{n-1})\| \} \|z_{n+1} - z_n\| \\
&\leq \sigma^2 \|x_n - x_{n-1}\|^2 + 2\rho \{ \beta \|q_n - q_{n-1}\| + \alpha \|w_n - w_{n-1}\| \} \|z_{n+1} - z_n\| \\
&\leq \sigma^2 \|x_n - x_{n-1}\|^2 + 2\rho \left\{ \beta\gamma \left(1 + \frac{1}{n}\right) + \alpha\mu \left(1 + \frac{1}{n}\right) \right\} \|x_n - x_{n-1}\| \|z_{n+1} - z_n\| \\
&= \sigma^2 \|x_n - x_{n-1}\|^2 + 2\rho \left(1 + \frac{1}{n}\right) (\beta\gamma + \alpha\mu) \|x_n - x_{n-1}\| \|z_{n+1} - z_n\|.
\end{aligned} \tag{5}$$

Since g is k -strongly accretive, R_ρ^A is non-expansive and from Lemma 2, it follows that for any $j(x_n - x_{n-1}) \in J(x_n - x_{n-1})$,

$$\begin{aligned}
\|x_n - x_{n-1}\|^2 &= \| [R_\rho^A(z_n) - R_\rho^A(z_{n-1})] - [g(x_n) - x_n - (g(x_{n-1}) - x_{n-1})] \|^2 \\
&\leq \|R_\rho^A(z_n) - R_\rho^A(z_{n-1})\|^2 - 2 \langle g(x_n) - x_n - (g(x_{n-1}) - x_{n-1}), j(x_n - x_{n-1}) \rangle \\
&\leq \|z_n - z_{n-1}\|^2 - 2k \|x_n - x_{n-1}\|^2 + 2 \|x_n - x_{n-1}\|^2
\end{aligned}$$

which implies that

$$\|x_n - x_{n-1}\| \leq \frac{1}{\sqrt{2k-1}} \|z_n - z_{n-1}\|. \tag{6}$$

From (5) and (6), we have

$$\begin{aligned}
\|z_{n+1} - z_n\|^2 &\leq \frac{\sigma^2}{2k-1} \|z_n - z_{n-1}\|^2 + \frac{2\rho(1+1/n)(\beta\gamma + \alpha\mu)}{\sqrt{2k-1}} \|z_n - z_{n-1}\| \|z_{n+1} - z_n\| \\
&\leq \frac{\sigma^2}{2k-1} \|z_n - z_{n-1}\|^2 + \frac{\rho(1+1/n)(\beta\gamma + \alpha\mu)}{\sqrt{2k-1}} [\|z_n - z_{n-1}\|^2 + \|z_{n+1} - z_n\|^2]
\end{aligned}$$

Finally, we have

$$\|z_{n+1} - z_n\| \leq \theta_n \|z_n - z_{n-1}\| \tag{7}$$

where

$$\theta_n = \left[\frac{\sigma^2 + \rho(1+1/n)(\beta\gamma + \alpha\mu)\sqrt{2k-1}}{(2k-1)[1 - \rho(1+1/n)(\beta\gamma + \alpha\mu)/\sqrt{2k-1}]} \right]^{1/2}$$

Let

$$\theta = \left[\frac{\sigma^2 + \rho(\beta\gamma + \alpha\mu)\sqrt{2k-1}}{(2k-1)[1 - \rho(\beta\gamma + \alpha\mu)/\sqrt{2k-1}]} \right]^{1/2}.$$

Then we know that $\theta_n \rightarrow \theta$ as $n \rightarrow \infty$. From condition (3) it follows that $\theta < 1$. Hence $\theta_n < 1$ for n sufficiently large. Therefore (7) implies that $\{z_n\}$ is a Cauchy sequence in E . Since E is a Banach space, there exists $z \in E$ such that $z_n \rightarrow z$ as $n \rightarrow \infty$. From (4) and (6), we know that $\{x_n\}$, $\{w_n\}$ and $\{q_n\}$ are also Cauchy sequences in E . Therefore, there exist $x \in E$, $w \in E$ and $q \in E$ such that $x_n \rightarrow x$, $w_n \rightarrow w$ and $q_n \rightarrow q$ as $n \rightarrow \infty$. Note that $w_n \in \tilde{T}(x_n)$, we have

$$\begin{aligned} d(w, \tilde{T}(x)) &= \inf\{\|w - p\| : p \in \tilde{T}(x)\} \\ &\leq \|w - w_n\| + d(w_n, \tilde{T}(x)) \\ &\leq \|w - w_n\| + \tilde{D}(\tilde{T}(x_n), \tilde{T}(x)) \\ &\leq \|w - w_n\| + \mu \|x_n - x\| \rightarrow 0, \quad n \rightarrow \infty, \end{aligned}$$

which implies that $d(w, \tilde{T}(x)) = 0$. Since $\tilde{T}(x) \in CB(E)$, it follows that $w \in \tilde{T}(x)$. Similarly we can show that $q \in \tilde{F}(x)$. Since g , R_ρ^A , $N(.,.)$, $\tilde{T}$ and $\tilde{F}$ are continuous, it follows from Algorithm 1 that

$$g(x) = R_\rho^A[g(x) - \rho N(w, q)]$$

By Theorem 1, we know that (x, w, q) is a solution of set-valued variational inclusions with fuzzy mappings (1) in real Banach spaces. This completes the proof.

References

- [1] Q. H. Ansari, *Certain Problems Concerning Variational Inequalities*, Ph. D. Thesis, Aligarh Muslim University, Aligarh, India, 1988.
- [2] V. Barbu, *Nonlinear Semigroups and Differential Equations in Banach Spaces*, Noordhoff International Publishing, Leyden, The Netherlands, 1976.
- [3] S. S. Chang, Coincidence theorem and variational inequalities for fuzzy mappings, *Fuzzy Sets and Systems* 61, 359-368 (1994).
- [4] S. S. Chang, On the Mann and Ishikawa iterative approximation of solutions to variational inclusions with accretive type mappings, *Comput. Math. Appl.* 37 (9), 17-24 (1999).
- [5] S. S. Chang, Set-valued variational inclusions in Banach Spaces, *J. Math. Anal. Appl.* 248, 438-454 (2000).
- [6] S. S. Chang, Y. G. Cho, B. B. Lee and I. H. Jung, Generalized set-valued variational inclusions in Banach Spaces, *J. Math. Anal. Appl.* 246, 409-422 (2000).
- [7] S. S. Chang and N. J. Huang, Generalized complementarity problem for fuzzy mappings, *Fuzzy Sets and systems* 55 (2), 227-234 (1993).

- [8] S. S. Chang and Y. G. Zhu, On Variational inequalities for fuzzy mappings, *Fuzzy Sets and Systems* 32, 359-367 (1989).
- [9] K. Deimling, *Nonlinear Functional Analysis*, Springer-Verlag, Berlin, 1985.
- [10] A. Hassouni and A. Moudafi, A perturbed algorithm for variational inclusions, *J. Math. Anal. Appl.* 185, 706-712 (1994).
- [11] N. J. Huang, A new method for a class of nonlinear variational inequalities with fuzzy mappings, *Appl. Math. Lett.* 10 (6), 129-133 (1997).
- [12] N. J. Huang, A new class of generalized set-valued implicit variational inclusions in Banach spaces with an application, *Comput. Math. Appl.* 41, 937-943 (2001).
- [13] S. S. Irfan, *Existence Results and Solution Methods for Certain Variational Inequalities and Complementarity Problems*, Ph. D. Thesis, Aligarh Muslim University, Aligarh, India, (2002).
- [14] G. M. Lee, D. S. Kim, B. S. Lee and S. J. Cho, Generalized vector variational inequality and fuzzy extension, *Appl. Math. Lett.* 66, 47-51 (1993).
- [15] S. B. Nadler, Multi-valued contraction mappings, *Pacific J. Math.* 30, 475-488 (1969).
- [16] W. V. Petryshyn, A characterization of strictly convexity of Banach spaces and other uses of duality mappings, *J. Func. Anal.* 6, 282-291 (1970).

Acknowledgments. The first author would like to express his thanks to the Department of Mathematical Sciences, King Fahd University of Petroleum & Minerals, Dhahran, Saudi Arabia for providing excellent facilities for carrying out this work. In this research, second author was supported by the Chilean Government Agency CONICYT under FONDAP program.

On some properties of semilinear functional differential inclusions in abstract spaces

Candida Gori^a - Valeri Obukhovskii^{b *} - Marcello Ragni^a - Paola Rubbioni^c

^a Dipartimento di Matematica e Informatica, Università di Perugia

via L. Vanvitelli 1, 06123 Perugia, ITALY

^b Department of Mathematics, Voronezh University

Universitetskaya pl. 1, 394006 Voronezh, RUSSIA

^c INFM and Dipartimento di Matematica e Informatica, Università di Perugia

via L. Vanvitelli 1, 06123 Perugia, ITALY

Telephone: +39-075-5855042/5040; Fax: +39-075-5855024

e-mails: cgori@unipg.it - ingar@dipmat.unipg.it - valerio@math.vsu.ru - rubbioni@dipmat.unipg.it

Abstract

We consider local and global existence results for differential inclusions with infinite delay in a Banach space of the form $y'(t) \in Ay(t) + F(t, y_t)$, $t \geq \sigma$, $y_\sigma = \varphi \in \mathcal{B}$, properties of the translation multioperator along trajectories and existence of periodic solutions.

*The work of the author is supported by RFBR Grants 05-01-00100, 04-01-00081 and NATO Grant ICS.NR.CLG 981757.

AMS Subject Classification : 34K30, 34K13, 34G25

Key words and Phrases: *Functional differential inclusion, infinite delay, phase space, measure of noncompactness, topological degree, fixed point, periodic solutions.*

1 Introduction

Semilinear functional differential inclusions in a Banach space have been the object of intensive study by many researchers in recent years (see, e.g. [7], [12], [13], [16], [17], [18], [19]).

Several works related to this subject are concerned with the Cauchy problem

$$y'(t) \in Ay(t) + F(t, y_t), \quad t \geq \sigma \quad y_\sigma = \varphi \in \mathcal{B}$$

where y_t represents the "history" of the system, $\mathcal{B}$ is the abstract Hale-Kato phase space (see, e.g. [8], [10]) and A is the densely defined linear operator infinitesimal generator of a strongly continuous semigroup of linear operators e^{At} , $t \geq 0$.

Notice that inclusions of that type appear in a very natural way in the description of processes of controlled heat transfer (see, e.g. [2], [15]), in obstacle problems, in the study of hybrid systems with dry friction, in the control of a transmission line process and other problems (see [12] and references therein).

In this paper we prove the existence of local, global and periodic solutions for semilinear differential inclusions with infinite delay in a Banach space.

Here we continue the study we began in [7], but in a slightly modified setting, more suitable

to deal with the periodic problem. In particular, the regularity conditions posed on the right-hand side as well as the condensivity of integral and translation multioperators are expressed in terms of the Kuratowski measure of noncompactness.

An analogous problem for differential equations was studied by H. Henriquez in [9] under a stronger assumption of compactness on the semigroup generated from the linear part of the equation and by using the quotient space of the phase space.

About the periodic problem for differential inclusions without delay, there exists a description in the recent book of M. Kamenskii, V. Obukhovskii, P. Zecca ([12]), where topological methods in multivalued analysis are developed and applied.

The paper is organized as follows.

In Section 2, we give some definitions and preliminary results. Afterwards, in Section 3, we provide local and global existence results for our abstract Cauchy problem. In Section 4, under some additional hypotheses, we study the upper semicontinuous dependence of the solution set on the initial data as well as its topological properties. Further we apply the properties of the translation multioperator along the trajectories of the solutions to the study of the periodic problem. In particular, in order to use the corresponding topological degree theory, we give a sufficient condition under which the translation multioperator is condensing with respect to the Kuratowski measure of noncompactness of the phase space.

2 Preliminaries

Let X and Y be topological spaces and let us denote by $P(Y)$ the collection of all nonempty subsets of Y . A multivalued map (multimap) $\mathcal{F} : X \rightarrow P(Y)$ is said to be:

- (i) *upper semicontinuous* (u.s.c) if $\mathcal{F}^{-1}(V) = \{x \in X : \mathcal{F}(x) \subset V\}$ is an open subset of X for every open $V \subset Y$;
- (ii) *closed* if its graph $\Gamma_F = \{(x, y) : y \in F(x)\}$ is a closed subset of the space $X \times Y$.

Recall also the following notions (see, e.g. [1], [12]).

Let $\mathcal{E}$ be a Banach space and $(\mathcal{A}, \geq)$ a (partially) ordered set. A function $\beta : P(\mathcal{E}) \rightarrow \mathcal{A}$ is called a *measure of noncompactness* (MNC) in $\mathcal{E}$ if

$$\beta(\overline{\text{co}}\Omega) = \beta(\Omega)$$

for every $\Omega \in P(\mathcal{E})$.

A MNC is called:

- i) *monotone* if $\Omega_0, \Omega_1 \in P(\mathcal{E})$, $\Omega_0 \subset \Omega_1$ implies $\beta(\Omega_0) \leq \beta(\Omega_1)$;
- ii) *nonsingular* if $\beta(\{a\} \cup \Omega) = \beta(\Omega)$ for every $a \in \mathcal{E}$, $\Omega \in P(\mathcal{E})$;
- iii) *real* if $\mathcal{A} = [0, +\infty]$ with the natural ordering, and $\beta(\Omega) < +\infty$ for every bounded Ω .

If $\mathcal{A}$ is a cone in a Banach space we will say that the MNC β is *regular* if $\beta(\Omega) = 0$ is equivalent to the relative compactness of Ω .

As the example of the MNC possessing all these properties, we may consider the *Kuratowski MNC*

$$\alpha(\Omega) = \inf\{\delta > 0 : \Omega \text{ has a partition into a finite number of sets with diameter less than } \delta\}$$

and the *Hausdorff MNC*

$$\chi(\Omega) = \inf\{\varepsilon > 0 : \Omega \text{ has a finite } \varepsilon\text{-net}\}.$$

From the definitions it follows that the Kuratowski and Hausdorff MNCs are connected by the following relations:

$$\chi(\Omega) \leq \alpha(\Omega) \leq 2\chi(\Omega). \quad (1)$$

By means of the Kuratowski MNC we can define some important characteristics of a linear operator in a Banach space.

Definition 2.1. Let $L : \mathcal{E} \rightarrow \mathcal{E}$ be a bounded linear operator and B the unit ball of $\mathcal{E}$.

The number

$$\|L\|^{(\alpha)} \equiv \alpha(LB)$$

is said to be *the (α) -norm of the operator L* .

It is easy to see that the following properties hold:

$$\|L\|^{(\alpha)} \leq \|L\|$$

$$\alpha(L\Omega) \leq \|L\|^{(\alpha)} \alpha(\Omega) \quad (2)$$

Let $\mathcal{X}$ be a closed subset of a Banach space $\mathcal{E}$ and $K(\mathcal{E})$ [$Kv(\mathcal{E})$] denote the collection of all nonempty compact [compact convex] subsets of $\mathcal{E}$.

Definition 2.2. Let β be a real MNC in $\mathcal{E}$ and $0 \leq k < 1$. A multimap $\mathcal{F} : \mathcal{X} \rightarrow K(\mathcal{E})$ is said to be *(k, β) -condensing* or simply *β -condensing* if

$$\beta(\mathcal{F}(\Omega)) \leq k\beta(\Omega)$$

for every $\Omega \subseteq \mathcal{X}$.

The following fixed point theorem (see [12], Corollary 3.2, Proposition 3.5.1) will be useful in the forthcoming local existence result.

Theorem 2.1. *If $\mathcal{M}$ is a bounded closed convex subset of $\mathcal{E}$, and $\mathcal{F} : \mathcal{M} \rightarrow Kv(\mathcal{M})$ is a closed β -condensing multimap, where β is a real nonsingular and regular MNC in $\mathcal{E}$, then the fixed points set $Fix \mathcal{F} = \{x : x \in \mathcal{F}(x)\}$ is nonempty and compact.*

Everywhere in the following E will denote a separable Banach space with the norm $\|\cdot\|$.

Let us recall some notions (see, e.g. [5], [11], [12]).

A multifunction $\mathcal{G} : [c, d] \rightarrow K(E)$ is said to be *strongly measurable* if there exists a sequence $\{\mathcal{G}_n\}_{n=1}^{\infty}$ of step multifunctions such that

$$h(\mathcal{G}_n(t), \mathcal{G}(t)) \rightarrow 0$$

as $n \rightarrow \infty$ for μ -a.e. $t \in [c, d]$ where μ denotes a Lebesgue measure on $[c, d]$ and h is the Hausdorff metric on $K(E)$. Every strongly measurable multifunction $\mathcal{G}$ admits a *strongly measurable selection* $g : [c, d] \rightarrow E$, i.e., $g(t) \in \mathcal{G}(t)$ for a.e. $t \in [c, d]$.

By the symbol $L^1([c, d]; E)$ we will denote the space of all Bochner summable functions.

A multifunction $\mathcal{G} : [c, d] \rightarrow K(E)$ is said to be:

i) integrable provided it has a summable selection $g \in L^1([c, d]; E)$;

ii) integrably bounded if there exists a summable function $\omega(\cdot) \in L^1_+([c, d])$ such that

$$\|\mathcal{G}(t)\| := \sup\{\|g\| : g \in \mathcal{G}(t)\} \leq \omega(t)$$

for a.e. $t \in [c, d]$.

It is clear that strongly measurable and integrably bounded multifunction is integrable.

The set of all summable selections of the multifunction $\mathcal{G}$ will be denoted by $\mathcal{S}^1_{\mathcal{G}}$.

Put α_E the Kuratowski measure of noncompactness in the space E , the following result about α_E -estimates for a multivalued integral is the consequence of Theorem 4.2.3 of [12] and the estimates (1):

Lemma 2.1. *Let a multifunction $\mathcal{G} : [c, d] \rightarrow P(E)$ be integrable, integrably bounded and*

$$\alpha_E(\mathcal{G}(t)) \leq q(t)$$

for a.e. $t \in [c, d]$, where $q \in L_+^1[c, d]$.

Then

$$\alpha_E \left(\int_c^t \mathcal{G}(s) ds \right) \leq 2 \int_c^t q(s) ds$$

for all $t \in [c, d]$.

Consider an abstract operator $S : L^1([c, d]; E) \rightarrow C([c, d]; E)$ satisfying the following conditions (cf. [12]):

(S1) there exists $\zeta \geq 0$ such that

$$\|Sf(t) - Sg(t)\| \leq \zeta \int_c^t \|f(s) - g(s)\| ds$$

for every $f, g \in L^1([c, d]; E)$, $c \leq t \leq d$;

(S2) for any compact $\mathcal{K} \subset E$ and sequence $\{f_n\}_{n=1}^\infty \subset L^1([c, d]; E)$ such that $\{f_n(t)\}_{n=1}^\infty \subset \mathcal{K}$ for a.e. $t \in [c, d]$ the weak convergence $f_n \rightharpoonup f_0$ implies $Sf_n \rightarrow Sf_0$.

Of course, condition (S1) implies that the operator S satisfies the Lipschitz condition

$$(S1') \quad \|Sf - Sg\|_C \leq \zeta \|f - g\|_{L^1}$$

where the first norm is the usual sup-norm in the space $C([c, d]; E)$.

In the following we will suppose that

- (A) the linear operator $A : D(A) \subseteq E \rightarrow E$ is the densely defined infinitesimal generator of a strongly continuous semigroup of linear operators e^{At} , $t \geq 0$ which is immediately norm continuous, i.e. the function $t \rightarrow e^{At}$ is norm continuous from $(0, \infty)$ into $\mathcal{L}(E)$, the space of bounded linear operators in E (see [6]).

Remark 2.1. Note that the *Cauchy operator* $G : L^1([c, d]; E) \rightarrow C([c, d]; E)$ defined as

$$Gf(t) = \int_c^t e^{A(t-s)} f(s) ds$$

satisfies properties (S1) and (S2) (see [12], Lemma 4.2.1).

The sequence $\{f_n\}_{n=1}^\infty \subset L^1([c, d]; E)$ is said to be *semicompact* if:

- (i) it is integrably bounded: $\|f_n(t)\| \leq \omega(t)$ for a.e. $t \in [c, d]$ and every $n \geq 1$ where $\omega(\cdot) \in L^1_+[c, d]$
- (ii) the set $\{f_n(t)\}_{n=1}^\infty$ is relatively compact for a.e. $t \in [c, d]$.

Let us mention the following important property of semicompact sequences (see, e.g. [12], Proposition 4.2.1).

Lemma 2.2. *Every semicompact sequence is weakly compact in the space $L^1([c, d]; E)$.*

In the sequel we will need also the following properties of the operator S satisfying assumptions mentioned above (see [12], Theorem 5.1.1).

Lemma 2.3. *Let $S : L^1([c, d]; E) \rightarrow C([c, d]; E)$ be an operator satisfying the conditions (S1') and (S2). Then for every semicompact sequence $\{f_n\}_{n=1}^\infty \subset L^1([c, d]; E)$ the sequence $\{Sf_n\}_{n=1}^\infty$ is relatively compact in $C([c, d]; E)$ and, moreover, if $f_n \rightharpoonup f_0$ then $Sf_n \rightarrow Sf_0$.*

Now let σ be a real number and $a > 0$. For any function $y : (-\infty, \sigma + a] \rightarrow E$ and for every $t \in (-\infty, \sigma + a]$, y_t represents the function from $(-\infty, 0]$ into E defined by

$$y_t(\theta) = y(t + \theta), \quad -\infty < \theta \leq 0.$$

Throughout this paper we will employ the axiomatic definition of the phase space $\mathcal{B}$ introduced by Hale and Kato [8]. The space $\mathcal{B}$ will be considered as a linear topological space of functions mapping $(-\infty, 0]$ into E endowed with a seminorm $\|\cdot\|_{\mathcal{B}}$. We will assume that $\mathcal{B}$ satisfies the following set of axioms.

If $y : (-\infty, \sigma + a] \rightarrow E$ is continuous on $[\sigma, \sigma + a]$ and $y_\sigma \in \mathcal{B}$, then for every $t \in [\sigma, \sigma + a]$ we have

$$(B1) \quad y_t \in \mathcal{B};$$

$$(B2) \quad \text{the function } t \mapsto y_t \text{ is continuous};$$

$$(B3) \quad \|y_t\|_{\mathcal{B}} \leq K(t - \sigma) \sup_{\sigma \leq \tau \leq t} \|y(\tau)\| + M(t - \sigma) \|y_\sigma\|_{\mathcal{B}}$$

where $K, M : [0, +\infty) \rightarrow [0, +\infty)$ are independent of y , K is strictly positive and continuous, and M is locally bounded.

3 Existence results

We consider the following problem for systems governed by a semilinear functional differential inclusion

$$y'(t) \in Ay(t) + F(t, y_t), \quad t \geq \sigma \tag{3}$$

$$y_\sigma = \varphi \in \mathcal{B} \tag{4}$$

where the multivalued nonlinearity $F : [\sigma, \sigma + a] \times \mathcal{B} \rightarrow Kv(E)$ will satisfy the following hypotheses:

(F1) for every fixed $\psi \in \mathcal{B}$ the multifunction $F(\cdot, \psi) : [\sigma, \sigma + a] \rightarrow Kv(E)$ admits a strongly measurable selector;

(F2) for a.e. $t \in [\sigma, \sigma + a]$ the multimap $F(t, \cdot) : \mathcal{B} \rightarrow Kv(E)$ is u.s.c.;

(F3) for every nonempty bounded set $\Omega \subset \mathcal{B}$ there exists a number $\mu_\Omega \geq 0$ such that for every $\psi \in \Omega$

$$\|F(t, \psi)\| \leq \mu_\Omega$$

for a.e. $t \in [\sigma, \sigma + a]$;

(F4) there exists a function $k \in L^1_+([\sigma, \sigma + a])$ such that for every bounded $D \subset \mathcal{B}$

$$\alpha_E(F(t, D)) \leq k(t) \alpha_{\mathcal{B}}(D) ,$$

for a.e. $t \in [\sigma, \sigma + a]$ where $\alpha_{\mathcal{B}}$ is Kuratowski measure of noncompactness in the space $\mathcal{B}$ generated by the seminorm $\|\cdot\|_{\mathcal{B}}$.

For $0 < b \leq a$, let us denote by the symbol $\mathcal{C}((-\infty, \sigma + b]; E)$ the linear topological space of functions $y : (-\infty, \sigma + b] \rightarrow E$ such that $y_\sigma \in \mathcal{B}$ and the restriction $y|_{[\sigma, \sigma + b]}$ is continuous, endowed with a seminorm

$$\|y\|_{\mathcal{C}} = \|y_\sigma\|_{\mathcal{B}} + \|y|_{[\sigma, \sigma + b]}\|_{\mathcal{C}} .$$

Definition 3.1. A function $y \in \mathcal{C}((-\infty, \sigma + b]; E)$ ($0 < b \leq a$) is a *mild solution* of (3),

(4) if

$$(i) \ y_\sigma = \varphi$$

$$(ii) \ y(t) = e^{A(t-\sigma)}\varphi(0) + \int_\sigma^t e^{A(t-s)}f(s)ds, \quad t \in [\sigma, \sigma + b], \quad \text{where } f \in \mathcal{S}_{F(\cdot, y(\cdot))}^1.$$

Note that conditions under which a mild solution satisfies some regularity properties are described in Section 5.2.1 of [12].

Let X be any set of functions $x : (-\infty, \sigma + a] \rightarrow E$ such that $x_\sigma \in \mathcal{B}$ and $x|_{[\sigma, \sigma+a]}$ is continuous for all $x \in X$. We denote, for $t \in [\sigma, \sigma + a]$,

$$X_t = \{x_t \in \mathcal{B} : x \in X\}$$

$$X[\sigma, t] = \{x|_{[\sigma, t]} : x \in X\}.$$

The main result of this section is an existence result for problem (3), (4). For the proof we will need the following result of J.S.Shin ([20], Theorem 2.1):

Lemma 3.1. *For any $t \in [\sigma, \sigma + a]$ the following relation holds:*

$$\alpha_{\mathcal{B}}(X_t) \leq K(t - \sigma)\alpha_C(X[\sigma, t]) + M(t - \sigma)\alpha_{\mathcal{B}}(X_\sigma) .$$

Let us also mention the following useful property (see, e.g. [3]).

Lemma 3.2. *If $\Delta \subset C([c, d]; E)$ is a bounded equicontinuous set, then*

$$\alpha_C(\Delta) = \sup_{c \leq t \leq d} \alpha_E(\Delta(t)) .$$

Theorem 3.1. *Under assumptions (A), (B1)-(B3) and (F1)-(F4) there exist $h \in (0, a]$ and a mild solution y_* of (3), (4) on $(-\infty, \sigma + h]$.*

Proof. The proof follows the lines of Theorem 2 in [7].

Let us take $r > 0$.

Since the semigroup e^{At} is strongly continuous, there exists $h_1 \in (0, a]$ such that

$$\|(e^{A\tau} - I)\varphi(0)\| \leq \frac{r}{2}, \quad 0 \leq \tau \leq h_1$$

and there exists $C \in \mathbb{R}^+$ such that

$$\|e^{A\theta}\| \leq C, \quad 0 \leq \theta \leq a. \quad (5)$$

Let us take the continuous function $s : [0, a] \rightarrow \mathcal{B}$ defined as

$$s(t)(\theta) = \begin{cases} \varphi(t + \theta), & -\infty < \theta \leq -t \\ \varphi(0), & -t < \theta \leq 0, \end{cases}$$

and consider the set

$$Q = s([0, a]) \subset \mathcal{B},$$

which is compact.

Let $K^* > 0$ be the value

$$K^* = \max_{0 \leq t \leq a} K(t), \quad (6)$$

where $K(\cdot)$ is the function introduced in the axiom (B3).

Then, taking the closure W of the rK^* -neighbourhood of Q ,

$$W = \overline{Q_{rK^*}}$$

and the corresponding number μ_W (see condition (F3)) we may choose $h_2 \in (0, a]$ so that

$$C\mu_W h_2 \leq \frac{r}{2} .$$

Moreover, there exists $h_3 \in (0, a]$ such that

$$2K^*C \int_{\sigma}^{\sigma+h_3} k(\tau) d\tau < 1 \quad (7)$$

where K^* and C are as above, and $k(\cdot)$ is from the condition (F4).

We put

$$h = \min\{h_1, h_2, h_3\} . \quad (8)$$

For $\varphi \in \mathcal{B}$ given in (4), we consider

$$D(\varphi, h) = \{x \in C([\sigma, \sigma + h]; E) : x(\sigma) = \varphi(0)\}$$

which is a closed convex set.

Further, for any $x \in D(\varphi, h)$, we consider the function $x[\varphi] \in \mathcal{C}((-\infty, \sigma + h]; E)$

$$x[\varphi](t) = \begin{cases} \varphi(t - \sigma) & -\infty < t < \sigma \\ x(t) & \sigma \leq t \leq \sigma + h . \end{cases} \quad (9)$$

Then, for any $t \in [\sigma, \sigma + h]$,

$$x[\varphi]_t(\theta) = \begin{cases} \varphi(t - \sigma + \theta) & -\infty < \theta < \sigma - t \\ x(t + \theta) & \sigma - t \leq \theta \leq 0 . \end{cases} \quad (10)$$

We consider the map $j : [\sigma, \sigma + h] \times D(\varphi, h) \rightarrow \mathcal{B}$ defined by

$$j(t, x) = x[\varphi]_t$$

and recall ([7], Theorem 2) that $j(\cdot, x)$ is continuous and $j(t, \cdot)$ is Lipschitz continuous in the seminorm $\|\cdot\|_{\mathcal{B}}$ uniformly with respect to $t \in [\sigma, \sigma + h]$.

Now, we consider the multivalued superposition operator

$$\mathcal{P}_F : D(\varphi, h) \rightarrow P(L^1([\sigma, \sigma + h]; E))$$

defined as

$$\begin{aligned} \mathcal{P}_F(x) &= \mathcal{S}_{F(\cdot, j(\cdot, x))}^1 = \\ &= \{f \in L^1([\sigma, \sigma + h]; E) : f(s) \in F(s, j(s, x)) = F(s, x[\varphi]_s) \text{ for a.e. } s \in [\sigma, \sigma + h]\}. \end{aligned}$$

$\mathcal{P}_F$ is correctly defined and, using Lemma 4 in [7], one may verify that it is weakly closed.

We consider the integral multioperator

$$\Gamma : D(\varphi, h) \rightarrow P(D(\varphi, h))$$

defined by

$$\Gamma(x) = \left\{ z : z(t) = e^{A(t-\sigma)}\varphi(0) + \int_{\sigma}^t e^{A(t-s)}f(s)ds, \quad f \in \mathcal{P}_F(x) \right\}.$$

Of course, every mild solution $y \in \mathcal{C}((-\infty, \sigma + h]; E)$ of (3), (4) is determined by a fixed point x of Γ by means of

$$y(t) = x[\varphi](t).$$

The fact that Γ is a closed multioperator with compact convex values and that it transforms the ball $\bar{B}_r(\bar{\varphi})$ into itself, where $\bar{\varphi}(t) \equiv \varphi(0)$, $t \in [\sigma, \sigma + h]$, can be proved by the same method as in [7].

To prove the α -condensivity of Γ on bounded sets, we use the following important property of the integral multioperator.

Lemma 3.3. *For any bounded $\Omega \subset D(\varphi, h)$ the image $\Gamma(\Omega)$ is bounded equicontinuous.*

Proof. Let Ω be a bounded subset of $D(\varphi, h)$ and let N be a positive number such that $\|x\|_C \leq N$ for any $x \in \Omega$.

First of all, we prove that $\Gamma(\Omega)$ is bounded.

Let $y \in \Gamma(\Omega)$. By means of the definition of $\Gamma(\Omega)$, there exists $x \in \Omega$ such that

$$y(t) = e^{A(t-\sigma)}\varphi(0) + \int_{\sigma}^t e^{A(t-s)}f(s)ds, \quad t \in [\sigma, \sigma + h]$$

where $f \in \mathcal{P}_F(x)$, so, in particular, $f(s) \in F(s, x[\varphi]_s)$, a.e. $s \in [\sigma, t]$.

Let $\tilde{\Omega}$ be the set defined as

$$\tilde{\Omega} = \{x[\varphi]_t \in \mathcal{B} : x \in \Omega, t \in [\sigma, \sigma + h]\}.$$

It is a bounded set. In fact, for any $x \in \Omega$ and $t \in [\sigma, \sigma + h]$, from (B3) we get

$$\begin{aligned} \|x[\varphi]_t\|_{\mathcal{B}} &\leq K^* \sup_{\sigma \leq \tau \leq t} \|x(\tau)\| + M^* \|\varphi\|_{\mathcal{B}} \leq \\ &\leq K^* \|x\|_C + M^* \|\varphi\|_{\mathcal{B}} \leq K^* N + M^* \|\varphi\|_{\mathcal{B}} \end{aligned}$$

where K^* is from (6) and

$$M^* = \sup_{0 \leq t \leq a} M(t). \quad (11)$$

Therefore, by using (5) and by applying condition (F3) with $D = \tilde{\Omega}$, for every $t \in [\sigma, \sigma + h]$

we have the following estimation :

$$\|y(t)\| \leq C \|\varphi(0)\| + C \int_{\sigma}^t \|f(s)\| ds \leq C \|\varphi(0)\| + C \mu_{\tilde{\Omega}} h,$$

where C is from (5).

To prove the equicontinuity of $\Gamma(\Omega)$, for any $t_1, t_2 \in [\sigma, \sigma + h]$ with $t_1 < t_2$, we consider

$$y(t_2) - y(t_1) = J_1 + J_2 + J_3,$$

where

$$J_1 = \left(e^{A(t_2-\sigma)} - e^{A(t_1-\sigma)} \right) \varphi(0) ,$$

$$J_2 = \int_{\sigma}^{t_1} \left(e^{A(t_2-s)} - e^{A(t_1-s)} \right) f(s) ds$$

and

$$J_3 = \int_{t_1}^{t_2} e^{A(t_2-s)} f(s) ds .$$

Given arbitrary $\varepsilon > 0$, let us estimate every term J_i , $i = 1, 2, 3$. At first,

$$\|J_1\| = \|e^{A(t_1-\sigma)}\| \| \left(e^{A(t_2-t_1)} - I \right) \varphi(0) \| \leq C \| \left(e^{A(t_2-t_1)} - I \right) \varphi(0) \| < \varepsilon$$

provided $t_2 - t_1 < \delta_1$.

To estimate J_2 we may assume w.l.o.g. that $t_1 > \sigma$. Let us take $\kappa > 0$ such that $2C\mu_{\tilde{\Omega}}\kappa < \frac{\varepsilon}{2}$ and $\kappa < t_1 - \sigma$. By virtue of condition (A), there exists $\delta_2 > 0$ such that

$$\|e^{A(t_2-s)} - e^{A(t_1-s)}\| < \frac{\varepsilon}{2\mu_{\tilde{\Omega}}h}$$

for $t_2 - t_1 < \delta_2$; $s \leq t_1 - \kappa$, where h is from (8). Therefore

$$\begin{aligned} \|J_2\| &\leq \int_{\sigma}^{t_1-\kappa} \|e^{A(t_2-s)} - e^{A(t_1-s)}\| \|f(s)\| ds + \int_{t_1-\kappa}^{t_1} \|e^{A(t_2-s)} - e^{A(t_1-s)}\| \|f(s)\| ds \\ &\leq \frac{\varepsilon}{2\mu_{\tilde{\Omega}}h} \mu_{\tilde{\Omega}}(t_1 - \kappa - \sigma) + 2C\mu_{\tilde{\Omega}}\kappa < \varepsilon. \end{aligned}$$

At last, notice that $\|J_3\|$ is obviously estimated by $C\mu_{\tilde{\Omega}}(t_2 - t_1)$. □

Lemma 3.4. Γ is α_C -condensing on bounded subsets of $D(\varphi, h)$.

Proof. Let $\Omega \subset D(\varphi, h)$ be a bounded set and let $t \in [\sigma, \sigma + h]$ be fixed.

Of course,

$$\Gamma(\Omega)(\theta) \subset e^{A(\theta-\sigma)}\varphi(0) + \int_{\sigma}^{\theta} e^{A(\theta-\tau)} F(\tau, \Omega[\varphi]_{\tau}) d\tau . \quad (12)$$

First of all, we consider the MNC α_E of the integrand in the right-hand side of the above relation. From (5) and assumption (F4), we have

$$\alpha_E(e^{A(\theta-\tau)} F(\tau, \Omega[\varphi]_{\tau})) \leq C\alpha_E(F(\tau, \Omega[\varphi]_{\tau})) \leq Ck(\tau)\alpha_{\mathcal{B}}(\Omega[\varphi]_{\tau}) .$$

It is clear that $\Omega[\varphi][\sigma, \tau] = \Omega[\sigma, \tau]$ and $\Omega[\varphi]_{\sigma} = \varphi$, therefore, by applying Lemma 3.1, we have the following estimate:

$$\begin{aligned} \alpha_E(e^{A(\theta-\tau)} F(\tau, \Omega[\varphi]_{\tau})) &\leq Ck(\tau) [K(\tau - \sigma)\alpha_C(\Omega[\sigma, \tau]) + M(\tau - \sigma)\alpha_{\mathcal{B}}(\varphi)] \leq \\ &\leq CK^*k(\tau)\alpha_C(\Omega) \end{aligned}$$

where K^* is from (6).

Then, by means of Lemma 2.1 and by applying the above inequality, from (12) we have

$$\alpha_E(\Gamma(\Omega)(\theta)) \leq \alpha_E\left(\int_{\sigma}^{\theta} e^{A(\theta-\tau)} F(\tau, \Omega[\varphi]_{\tau}) d\tau\right) \leq 2CK^*\alpha_C(\Omega) \int_{\sigma}^{\theta} k(\tau) d\tau .$$

Since the right-hand part does not depend on t , by using Lemma 3.2 we obtain

$$\alpha_C(\Gamma(\Omega)) = \sup_{\sigma \leq t \leq \sigma+h} \alpha_E(\Gamma(\Omega)(t)) \leq 2K^*C \int_{\sigma}^{\sigma+h} k(\tau) d\tau \alpha_C(\Omega) .$$

Afterwards, since h is small enough to provide $2K^*C \int_{\sigma}^{\sigma+h} k(\tau) d\tau < 1$ (see (7)), we conclude that Γ is α_C -condensing. $\square$

As a direct consequence of Theorem 2.1, Γ has a fixed point. $\square$

Under a condition stronger than (F3), we can obtain a global existence result.

Theorem 3.2. *Assume that conditions (A), (B1)-(B3), (F1), (F2), (F4) are satisfied and that*

(F3') there exists a number $\mu \geq 0$ such that for every $\psi \in \mathcal{B}$ we have the estimate

$$\|F(t, \psi)\| \leq \mu(1 + \|\psi\|_{\mathcal{B}})$$

for a.e. $t \in [\sigma, \sigma + a]$.

Assume also that

(H1) $2K^*C \int_{\sigma}^{\sigma+a} k(\tau) d\tau < 1$ *where $K^*, C, k(\cdot)$ are from (6), (5), (F4) respectively.*

Then the set Σ_{φ} of all mild solutions of problem (3), (4) on $(-\infty, \sigma + a]$ is nonempty and compact.

Proof. The proof is analogous to that of Theorem 3 in [7] by using the Kuratowski MNC instead of the MNC adopted in the cited theorem. We just note that, by applying the same arguments of Lemma 3.4 under the further assumption (H1), the multifunction Γ is α_C -condensing on the whole interval $[\sigma, \sigma + a]$. $\square$

4 The translation multioperator and the periodic problem

In this section we study the periodic problem associated to (3), (4) by developing the method of the translation multioperator along the trajectories of solutions.

To this end, we need to investigate on the α -condensivity of the translation multioperator.

In fact, from the structure of the integral funnel and the theory of topological degree for condensing nonconvex-valued multimaps ([12], Section 3.4), the existence of periodic solutions follows.

We need some additional hypotheses on the phase space $\mathcal{B}$.

The first is the following:

(B4) there exists $l > 0$ such that

$$\|\varphi(0)\| \leq l\|\varphi\|_{\mathcal{B}}$$

for all $\varphi \in \mathcal{B}$.

Then, we need two hypotheses concerning continuous representatives in $\mathcal{B}$.

It is easy to see ([10], Proposition 1.2.1) that the space C_{00} of all continuous functions from $(-\infty, 0]$ into E with compact support is a subset of any space $\mathcal{B}$.

Our first assumption is nothing but hypothesis (C2) in [10]:

(D1) if a uniformly bounded sequence $\{\varphi_n\}_{n=1}^{+\infty} \subset C_{00}$ converges to a function φ uniformly on every compact subset of $(-\infty, 0]$, then $\varphi \in \mathcal{B}$ and $\lim_{n \rightarrow +\infty} \|\varphi_n - \varphi\|_{\mathcal{B}} = 0$.

The hypothesis (D1) yields that the Banach space $BC((-\infty, 0]; E)$ of bounded continuous functions is continuously imbedded into $\mathcal{B}$ (see [10], Proposition 7.1.1).

Our next assumption means that the seminorm $\|\cdot\|_{\mathcal{B}}$ is sensitive enough to distinguish continuous functions with compact support:

(D2) if $x \in C_{00}$ and $\|x\|_{BC} \neq 0$, then $\|x\|_{\mathcal{B}} \neq 0$.

This hypothesis implies that the space C_{00} endowed by $\|\cdot\|_{\mathcal{B}}$ is a normed space. We will denote it as $\mathcal{BC}_{00}$.

We recall that, under assumptions of existence of solutions, the translation multioperator maps every $\varphi \in \mathcal{B}$ into the set $\Sigma(\varphi)_{\sigma+h} \subset \mathcal{B}$ (see (10)), where $\Sigma(\varphi)$ is the set of all mild solutions of (3), (4).

By the definition, if $y : (-\infty, \sigma + h] \rightarrow E$ is a mild solution of the problem (3),(4) with the initial function φ in $\mathcal{BC}_0$, then $y_{\sigma+h} \in \mathcal{BC}_0$; therefore, we may consider the translation multioperator along the trajectories of the inclusion (3) as acting in the space $\mathcal{BC}_0$:

$$P_{\sigma+h} : \mathcal{BC}_0 \rightarrow P(\mathcal{BC}_0) .$$

Now, we assume that the multimap F is T -periodic in the first argument:

$$(F_T) \quad F(t+T, \psi) = F(t, \psi) \text{ for every } (t, \psi) \in \mathbb{R}^+ \times \mathcal{B} .$$

We consider on the space $\mathcal{BC}_0 \times C([\sigma, \sigma + T]; E)$ the subset

$$V = \{(\varphi, x) : \varphi(0) = x(\sigma)\}$$

which is closed by (B4), and we define an integral multioperator $G : V \rightarrow P(C([\sigma, \sigma + T]; E))$ by

$$G(\varphi, x) = \{z : z(t) = e^{A(t-\sigma)}\varphi(0) + \int_{\sigma}^t e^{A(t-s)}f(s)ds, \quad f \in L^1([\sigma, \sigma + T]; E),$$

$$f(s) \in F(s, x[\varphi]_s) \text{ for a.e. } s\}$$

where, as in Theorem 3.1,

$$x[\varphi]_s(\theta) = \begin{cases} \varphi(s - \sigma + \theta), & -\infty < \theta \leq \sigma - s \\ x(s + \theta), & \sigma - s < \theta \leq 0 . \end{cases}$$

It is clear that if $x \in G(\varphi, x)$, then the element $x[\varphi] \in \mathcal{C}((-\infty, \sigma + T]; E)$ (see (9)) is a mild solution of our Cauchy problem.

In the forthcoming Lemmas, we assume that the hypotheses of Theorem 3.2 on the interval $[\sigma, \sigma + T]$ and the further assumptions (B4), (D1), (D2) are fulfilled.

Lemma 4.1. *The multimap $\Sigma^\sigma : \mathcal{BC}_{00} \rightarrow K(C([\sigma, \sigma + T]; E))$ defined as*

$$\Sigma^\sigma(\varphi) = \{x \in D(\varphi, T) : x \in G(\varphi, x)\}$$

is upper semicontinuous.

Proof. Let us assume, to the contrary, that there exist $\varepsilon_0 > 0$ and sequences $\{\varphi_n\}_{n=1}^\infty \subset \mathcal{BC}_{00}$, $\|\varphi_n - \varphi_0\|_{\mathcal{B}} \rightarrow 0$, $\{x_n\}_{n=1}^\infty \subset C([\sigma, \sigma + T]; E)$, $x_n \in \Sigma^\sigma(\varphi_n)$, such that

$$x_n \notin W_{\varepsilon_0}(\Sigma^\sigma(\varphi_0)) , \quad n \geq 1 \quad (13)$$

where W_{ε_0} denotes the open ε_0 -neighborhood of $\Sigma^\sigma(\varphi_0)$.

It is clear by definition that

$$x_n \in G(\varphi_n, x_n) , \quad n \geq 1 ,$$

i.e.

$$x_n(t) = e^{A(t-\sigma)}\varphi_n(0) + \int_{\sigma}^t e^{A(t-s)}f_n(s)ds$$

where $f_n \in L^1([\sigma, \sigma + T]; E)$, $f_n(s) \in F(s, x_n[\varphi_n]_s)$, for a.e. $s \in [\sigma, t]$.

Further, the sequence $\{x_n\}_{n=1}^\infty$ is a priori bounded. In fact, for every $n \geq 1$, by applying condition (F3') and axiom (B3), we have the following estimations:

$$\begin{aligned} \|x_n(t)\| &\leq C\|\varphi_n(0)\| + C\mu \int_{\sigma}^t (1 + \|x_n[\varphi_n]_s\|_{\mathcal{B}})ds \leq \\ &\leq C\|\varphi_n(0)\| + C\mu T + C\mu \int_{\sigma}^t \left(K^* \sup_{\sigma \leq \tau \leq s} \|x_n(\tau)\| + M^*\|\varphi_n\|_{\mathcal{B}} \right) ds \leq \\ &\leq C\|\varphi_n(0)\| + C\mu T(1 + M^*\|\varphi_n\|_{\mathcal{B}}) + CK^*\mu \int_{\sigma}^t \sup_{\sigma \leq \tau \leq s} \|x_n(\tau)\| ds \end{aligned}$$

where C , K^* and M^* are defined by (5), (6) and (11) respectively.

Since the last expression does not decrease, we have

$$\sup_{\sigma \leq \tau \leq t} \|x_n(\tau)\| \leq C\|\varphi_n(0)\| + C\mu T(1 + M^*\|\varphi_n\|_{\mathcal{B}}) + CK^*\mu \int_{\sigma}^t \sup_{\sigma \leq \tau \leq s} \|x_n(\tau)\| ds.$$

Applying to the function $w_n(t) = \sup_{\sigma \leq \tau \leq t} \|x_n(\tau)\|$ the Gronwall-Bellmann inequality we obtain that

$$w_n(t) \leq N_n \exp(CK^*\mu(t - \sigma))$$

where $N_n = C\|\varphi_n(0)\| + C\mu T(1 + M^*\|\varphi_n\|_{\mathcal{B}})$.

From the convergence of the sequence $\{\varphi_n\}_{n=1}^{\infty}$ in the space $\mathcal{BC}_{00}$, which is normed, it follows that $\{\varphi_n\}_{n=1}^{\infty}$ is bounded in $\mathcal{BC}_{00}$; therefore, by using (B4), the sequence $\{\varphi_n(0)\}_{n=1}^{\infty}$ is also bounded in E , implying the desired boundedness of sequence $\{x_n\}_{n=1}^{\infty}$.

Besides, the sequence $\{x_n\}_{n=1}^{\infty}$ is equicontinuous. In fact, for any $t_1, t_2 \in [\sigma, \sigma + T]$ with $t_1 < t_2$ and for any $n \geq 1$, we have

$$x_n(t_2) - x_n(t_1) = I_{n1} + I_{n2} + I_{n3},$$

where

$$I_{n1} = (e^{A(t_2-\sigma)} - e^{A(t_1-\sigma)}) \varphi_n(0);$$

$$I_{n2} = \int_{\sigma}^{t_1} (e^{A(t_2-s)} - e^{A(t_1-s)}) f_n(s) ds;$$

and

$$I_{n3} = \int_{t_1}^{t_2} e^{A(t_2-s)} f_n(s) ds.$$

Given arbitrary $\varepsilon > 0$ and taking into account, by (B4), that the sequence $\{\varphi_n(0)\}_{n=1}^{\infty}$ is convergent and hence relatively compact, we conclude that the term I_{n1} may be estimated

uniformly with respect to n , i.e. there exists $\delta_1 > 0$ such that $\|I_{n1}\| < \varepsilon$, $n = 1, 2, \dots$ provided $t_2 - t_1 < \delta_1$.

Consider the sequence $\{x_n[\varphi_n]_s\}_{n=1}^\infty$. It is bounded in $\mathcal{BC}_{00}$, in fact, by applying axiom (B3), for every $n \geq 1$ the following estimation holds:

$$\begin{aligned} \|x_n[\varphi_n]_s\|_{\mathcal{B}} &\leq K^* \sup_{\sigma \leq \tau \leq s} \|x_n(\tau)\| + M^* \|\varphi_n\|_{\mathcal{B}} \leq \\ &\leq K^* \|x_n\|_C + M^* \|\varphi_n\|_{\mathcal{B}} \end{aligned}$$

where K^*, M^* are as above. Bearing in mind that sequences $\{x_n\}_{n=1}^\infty$ and $\{\varphi_n\}_{n=1}^\infty$ are bounded in $C([\sigma, \sigma + T]; E)$ and $\mathcal{BC}_{00}$ respectively, the conclusion follows.

Now, from $(F3')$ we have that

$$\|f_n(s)\| \leq \mu(1 + \|x_n[\varphi_n]_s\|_{\mathcal{B}}) \text{ for a.e. } s \in [0, t]$$

and hence the sequence $\{f_n\}_{n=1}^\infty$ is bounded and we may apply the same arguments as while the proof of Lemma 3.3 to obtain the estimates for I_{n2}, I_{n3} uniform with respect to n yielding the equicontinuity of the sequence $\{x_n\}_{n=1}^\infty$.

Now, for any $t \in [\sigma, \sigma + T]$, it is easy to see that

$$\{x_n(t)\}_{n=1}^\infty \subset e^{A(t-\sigma)}\{\varphi_n(0)\}_{n=1}^\infty + \int_\sigma^t e^{A(t-s)}F(s, \{x_n[\varphi_n]_s\}_{n=1}^\infty)ds$$

hence, by means of the properties of the MNC,

$$\alpha_E(\{x_n(t)\}_{n=1}^\infty) \leq \alpha_E(e^{A(t-\sigma)}\{\varphi_n(0)\}_{n=1}^\infty) + \alpha_E\left(\int_\sigma^t e^{A(t-s)}F(s, \{x_n[\varphi_n]_s\}_{n=1}^\infty)ds\right).$$

From the relative compactness of the sequence $\{\varphi_n(0)\}_{n=1}^\infty$ it follows that the first term of the right hand side vanishes. Estimating, by means of (2), $(F4)$ and Lemma 3.1,

$$\alpha_E\left(e^{A(t-s)}F(s, \{x_n[\varphi_n]_s\}_{n=1}^\infty)\right) \leq K^* Ck(s) \alpha_C(\{x_n\}_{n=1}^\infty)$$

and using Lemma 2.1 we get

$$\alpha_C(\{x_n\}_{n=1}^\infty) = \sup_{\sigma \leq t \leq \sigma+T} \alpha_E(\{x_n(t)\}_{n=1}^\infty) \leq 2K^*C \int_\sigma^{\sigma+T} k(\tau) d\tau \alpha_C(\{x_n\}_{n=1}^\infty) .$$

Now from hypothesis (H1) it follows that $\alpha_C(\{x_n\}_{n=1}^\infty) = 0$ and the sequence $\{x_n\}_{n=1}^\infty$ is relatively compact. We may assume, w.l.o.g., that $x_n \rightarrow x_0 \in C([\sigma, \sigma+T]; E)$.

Our aim is to show that $x_0 \in G(\varphi_0, x_0)$, i.e. $x_0 \in \Sigma^\sigma(\varphi_0)$, contrary to (13).

To this end, we consider the identity

$$x_n(t) = e^{A(t-\sigma)}\varphi_n(0) + \int_\sigma^t e^{A(t-s)}f_n(s)ds, \quad f_n(s) \in F(s, x_n[\varphi_n]_s) \text{ a.e. } s \in [0, t]. \quad (14)$$

From the convergence of $\{\varphi_n(0)\}_{n=1}^\infty$ to $\varphi_0(0)$, we have

$$e^{A(t-\sigma)}\varphi_n(0) \rightarrow e^{A(t-\sigma)}\varphi_0(0).$$

Now, let us estimate the MNC of the sequence $\{f_n(t)\}_{n=1}^\infty$ for a.e. $t \in [\sigma, \sigma+T]$.

By using assumption (F4) and Lemma 3.1, we have

$$\begin{aligned} \alpha_E(\{f_n(t)\}_{n=1}^\infty) &\leq \alpha_E(F(t, \{x_n[\varphi_n]_t\}_{n=1}^\infty)) \leq k(t)\alpha_{\mathcal{B}}(\{x_n[\varphi_n]_t\}_{n=1}^\infty) \leq \\ &\leq k(t) \left[K^* \alpha_C(\{x_n|_{[\sigma, t]}\}_{n=1}^\infty) + M^* \alpha_{\mathcal{BC}_{00}}(\{\varphi_n\}_{n=1}^\infty) \right] . \end{aligned}$$

Since the sequences $\{x_n\}_{n=1}^\infty$ and $\{\varphi_n\}_{n=1}^\infty$ converge, the last expression is equal to zero, therefore

$$\alpha_E(\{f_n(t)\}_{n=1}^\infty) = 0 \quad \text{for a.e. } t \in [\sigma, \sigma+T]. \quad (15)$$

Afterwards, by means of (F3') and (15), the sequence $\{f_n\}_{n=1}^\infty$ is semicompact. By applying Lemma 2.2, we deduce that, w.l.o.g., $\{f_n\}_{n=1}^\infty$ weakly converges to a function $f_0 \in L^1([\sigma, \sigma+T]; E)$.

By using axiom (B3) and the assumptions on F , with the same arguments as in Lemma 5.1.1 in [12], we have

$$f_0(t) \in F(t, x_0[\varphi_0]_t), \quad \text{a.e. } t.$$

Hence, by applying Remark 2.1 and Lemma 2.3, the integral term in (14) converges to $\int_{\sigma}^t e^{A(t-s)} f_0(s) ds$, which concludes the proof. $\square$

To formulate the next statement, let us recall (see e.g. [12]) that a nonempty space is said to be an R_δ -set if it can be represented as the intersection of a decreasing sequence of compact, contractible sets. It is clear that every R_δ -set is acyclic.

Lemma 4.2. *The set $\Sigma^\sigma(\varphi)$ is compact and, moreover, it is R_δ .*

Proof. For any $\varphi \in \mathcal{B}$ the set $\Sigma^\sigma(\varphi)$ is a priori bounded and it can be seen as the set $\text{Fix } G(\varphi, \cdot)$.

As proved in Lemma 3.4, the integral multioperator $G(\varphi, \cdot)$ is α -condensing, then the compactness of $\Sigma^\sigma(\varphi)$ follows from Proposition 3.5.1 of [12].

The fact that $\Sigma^\sigma(\varphi)$ is an R_δ -set can be proved following the lines of Theorem 5.3.1 in [12]. $\square$

Lemma 4.3. *The translation multioperator P_T is quasi R_δ -multimap, i.e. it can be represented as a composition of an u.s.c. multimap with R_δ values and a continuous map.*

Proof. We consider the multimap $\Pi : \mathcal{BC}_0 \rightarrow P(V)$ defined as

$$\Pi(\varphi) = \{\varphi\} \times \Sigma^\sigma(\varphi).$$

A multimap $\Sigma^\sigma(\varphi)$ can be regarded as a fixed point set of a condensing multioperator $G(\varphi, \cdot)$ depending on parameter and so it is u.s.c. by virtue of Proposition 3.5.2 of [12]. The set

$\Pi(\varphi)$ is an R_δ since so is $\Sigma^\sigma(\varphi)$. Further, the map $\mathcal{X} : V \rightarrow \mathcal{BC}_{00}$ defined as

$$\mathcal{X}(\varphi, x) = x[\varphi]_{\sigma+T} ,$$

is continuous from (B3).

The translation multioperator P_T can be regarded as the composition of Π and $\mathcal{X}$, i.e. $P_T = \mathcal{X} \circ \Pi$, hence it is quasi R_δ . $\square$

We will find conditions under which P_T will be α -condensing.

First of all, we suppose the following additional hypothesis holds

(A1) the semigroup e^{At} is uniformly continuous and α -decreasing:

$$\|e^{At}\|^{(\alpha)} \leq C_1 e^{-\gamma t}$$

where γ and C_1 are positive constants (see Definition 2.1).

By $\alpha_{\mathcal{BC}_{00}}$ we will denote the Kuratowski MNC generated in the space $\mathcal{BC}_{00}$ by the norm $\|\cdot\|_{\mathcal{B}}$.

Let $\Omega \subset \mathcal{BC}_{00}$ be a bounded set and $t \in [\sigma, \sigma + T]$.

We can apply Lemma 3.1 to the set $X = \Sigma(\Omega)$ and, bearing in mind that $\Sigma(\Omega)_t = P_t(\Omega)$ and $\Sigma(\Omega)_\sigma = \Omega$, we obtain the following estimate:

$$\alpha_{\mathcal{BC}_{00}}(P_t(\Omega)) \leq K(t - \sigma)\alpha_C(\Sigma(\Omega)[\sigma, t]) + M(t - \sigma)\alpha_{\mathcal{BC}_{00}}(\Omega) . \quad (16)$$

At first, let us estimate $\alpha_C(\Sigma(\Omega)[\sigma, t])$.

To this aim we need the following statement which may be verified using hypothesis (A1).

Lemma 4.4. *For any bounded $\Omega \subset \mathcal{BC}_{00}$ the set $\Sigma(\Omega)[\sigma, t]$ is bounded equicontinuous.*

From the Lemma above and Lemma 3.2, we get

$$\alpha_C(\Sigma(\Omega)[\sigma, t]) = \sup_{\sigma \leq \theta \leq t} \alpha_E(\Sigma(\Omega)(\theta)) .$$

Moreover, from the definition of the integral multioperator, for $\theta \in [\sigma, t]$ we have that

$$\Sigma(\Omega)(\theta) \subset e^{A(\theta-\sigma)}\Omega(0) + \int_{\sigma}^{\theta} e^{A(\theta-s)}F(s, P_s(\Omega)) ds . \quad (17)$$

The MNC of the first term of the previous sum may be estimated, by using (2) and (A1), in the following way:

$$\begin{aligned} \alpha_E(e^{A(\theta-\sigma)}\Omega(0)) &\leq \|e^{A(\theta-\sigma)}\|^{(\alpha)} \alpha_E(\Omega(0)) \leq \\ &\leq C_1 e^{-\gamma(\theta-\sigma)} \alpha_E(\Omega(0)) . \end{aligned}$$

Moreover, from (B4) we obtain

$$\alpha_E(\Omega(0)) \leq l \alpha_{\mathcal{BC}_{00}}(\Omega).$$

Therefore,

$$\alpha_E(e^{A(\theta-\sigma)}\Omega(0)) \leq C_1 e^{-\gamma(\theta-\sigma)} l \alpha_{\mathcal{BC}_{00}}(\Omega) . \quad (18)$$

On the other hand, by applying (2), (A1) and (F4), for the integrand we have the following estimate:

$$\begin{aligned} \alpha_E(e^{A(\theta-s)}F(s, P_s(\Omega))) &\leq \|e^{A(\theta-s)}\|^{(\alpha)} \alpha_E(F(s, P_s(\Omega))) \leq \\ &\leq C_1 e^{-\gamma(\theta-s)} k(s) \alpha_{\mathcal{BC}_{00}}(P_s(\Omega)). \end{aligned} \quad (19)$$

Taking into account the fact that the space $\mathcal{BC}_{00}$ is separable, hypothesis (B2), and following the lines of Theorem 4.2.4 in [12] we may demonstrate that the function $s \mapsto \alpha_{\mathcal{BC}_{00}}(P_s(\Omega))$ is summable. Therefore, by means of (17), (18), Lemma 2.1 and (19), we obtain

$$\alpha_E(\Sigma(\Omega)(\theta)) \leq C_1 e^{-\gamma(\theta-\sigma)} l \alpha_{\mathcal{BC}_{00}}(\Omega) + C_1 \int_{\sigma}^{\theta} e^{-\gamma(\theta-s)} k(s) \alpha_{\mathcal{BC}_{00}}(P_s(\Omega)) ds .$$

Hence, taking the supremum for $\theta \in [\sigma, t]$,

$$\alpha_C(\Sigma(\Omega)[\sigma, t]) \leq C_1 l \alpha_{\mathcal{BC}_{00}}(\Omega) + C_1 e^{-\gamma\sigma} \int_{\sigma}^t e^{\gamma s} k(s) \alpha_{\mathcal{BC}_{00}}(P_s(\Omega)) ds .$$

At last, from (16) we have

$$\begin{aligned} \alpha_{\mathcal{BC}_{00}}(P_t(\Omega)) &\leq K(t-\sigma) C_1 l \alpha_{\mathcal{BC}_{00}}(\Omega) + K(t-\sigma) C_1 e^{-\gamma\sigma} \int_{\sigma}^t e^{\gamma s} k(s) \alpha_{\mathcal{BC}_{00}}(P_s(\Omega)) ds + \\ &\quad + M(t-\sigma) \alpha_{\mathcal{BC}_{00}}(\Omega) = \\ &= [K(t-\sigma) C_1 l + M(t-\sigma)] \alpha_{\mathcal{BC}_{00}}(\Omega) + \\ &\quad + K(t-\sigma) C_1 e^{-\gamma\sigma} \int_{\sigma}^t e^{\gamma s} k(s) \alpha_{\mathcal{BC}_{00}}(P_s(\Omega)) ds . \end{aligned}$$

So, denoting $u(t) = \alpha_{\mathcal{BC}_{00}}(P_t(\Omega))$, we have the following integral inequality:

$$u(t) \leq R(t) u(\sigma) + Q(t) \int_{\sigma}^t e^{\gamma s} k(s) u(s) ds$$

where

$$R(t) = C_1 l K(t-\sigma) + M^*, \text{ with } M^* \text{ as in (11);}$$

$$Q(t) = C_1 e^{-\gamma\sigma} K(t-\sigma) .$$

Finally, by applying Theorem 17.1 of [4], we get the following Gronwall-type inequality for $u(t)$:

$$\begin{aligned} u(t) &\leq R(t)u(\sigma) + Q(t)u(\sigma) \int_{\sigma}^t R(s)e^{\gamma s}k(s)e^{\int_s^t Q(\tau)e^{\gamma\tau}k(\tau)d\tau} ds = \\ &= \left[R(t) + Q(t) \int_{\sigma}^t R(s)e^{\gamma s}k(s)e^{\int_s^t Q(\tau)e^{\gamma\tau}k(\tau)d\tau} ds \right] u(\sigma). \end{aligned}$$

Therefore, we obtain the following condition for the condensivity of P_T :

$$(H2) \quad L = R(T) + Q(T) \int_{\sigma}^T R(s)e^{\gamma s}k(s)e^{\int_s^T Q(\tau)e^{\gamma\tau}k(\tau)d\tau} ds < 1 .$$

Remark 4.1. In the particular case when $F(t, \psi)$ is completely u.s.c. in the second argument for a.e. t , i.e. $k(t) \equiv 0$, then condition (H2) reduces to

$$R(T) < 1 .$$

Now we may summarize our reasonings in the form of the following statement.

Theorem 4.1. *Let us suppose that assumptions (A1), (B1)-(B4), (D1)-(D2), (F1)-(F2), (F3'), (F4), (H1)-(H2) hold. Then the translation multioperator $P_{\sigma+T} : \mathcal{BC}_0 \rightarrow P(\mathcal{BC}_0)$ is α -condensing on bounded sets of $\mathcal{BC}_0$.*

Since the translation multioperator P_T is quasi R_{δ} (Lemma 4.3) the corresponding topological degree theory (see [14], [12]) may be applied to obtain fixed points of P_T which obviously are initial values of T -periodic mild solutions of the inclusion (3). In fact, we have the following general principle.

Theorem 4.2. *Suppose that conditions of Theorem 4.1 hold. Let $U \subset \mathcal{BC}_0$ be an open bounded set such that $\text{Fix } P_T \cap \partial U = \emptyset$ and $\deg(i - P_T, \bar{U}) \neq 0$. Then the differential inclusion (3) has a T -periodic mild solution.*

Corollary 4.1 *Suppose that conditions of Theorem 4.1 hold. Let $N \subset \mathcal{BC}_{00}$ be a convex bounded set such that $P_T(N) \subseteq N$. Then the differential inclusion (3) has a T -periodic mild solution.*

Notice that sufficient conditions for the existence of such invariant set N may be found, e.g. in [9].

References

- [1] R.R.Akhmerov, M.I.Kamenskii, A.S.Potapov, A.E.Rodkina, B.N.Sadovskii, *Measures of Non-compactness and Condensing Operators*, Birkhauser, Boston - Basel - Berlin, 1992.
- [2] G.Anichini, P.Zecca, Multivalued differential equations in Banach spaces. An application to control theory, *J. Optimiz. Theory and Appl.* 21, 477-486 (1977).
- [3] J.Banas, K.Goebel, *Measures of Noncompactness in Banach Spaces*, Marcel Dekker, New York - Basel, 1980.
- [4] D.Bainov, P.Simeonov, *Integral Inequalities and Applications*, Kluwer, Dordrecht-Boston-London, 1992.
- [5] C.Castaing, M.Valadier, *Convex Analysis and Measurable Multifunctions*, Lect. Notes in Math. 580, Springer-Verlag, Berlin - Heidelberg - New York, 1977.
- [6] K.J.Engel, R.Nagel, *One-Parameter Semigroups for Linear Evolution Equations*, Graduate Texts in Math. 194, Springer-Verlag, New York - Berlin - Heidelberg, 2000.
- [7] C.Gori, V.Obukhovskii, M.Ragni, P.Rubbioni, Existence and continuous dependence results for semilinear functional differential inclusions with infinite delay, *Nonlinear Anal. T.M.A.* 51, 5, 765-782 (2002).

- [8] J.K.Hale, J.Kato, Phase space for retarded equations with infinite delay, Funkcial. Ekvac. 21, 11-41 (1978).
- [9] H.R.Henriquez, Periodic solutions of quasi-linear partial functional differential equations with unbounded delay, Funkc. Ekvac. 37, 329-343 (1994).
- [10] Y.Hino, S.Murakami, T.Naito, *Functional Differential Equations with Infinite Delay*, Lect. Notes in Math. 1473, Springer-Verlag, Berlin - Heidelberg - New York, 1991.
- [11] S.Hu, N.S.Papageorgiou, *Handbook of Multivalued Analysis, vol. I: Theory*, Kluwer Acad. Publ., Dordrecht - Boston - London, 1997.
- [12] M.Kamenskii, V.Obukhovskii, P.Zecca, *Condensing Multivalued Maps and Semilinear Differential Inclusions in Banach Spaces*, De Gruyter Ser. Nonlinear Anal. Appl. 7, Walter de Gruyter, Berlin - New York, 2001.
- [13] A.G.Ibrahim, Topological properties of solution sets for functional differential inclusions governed by a family of operators, Port. Math. (N.S.) 58, No.3, 255-270 (2001).
- [14] V.V.Obukhovskii, On the topological degree for a class of noncompact multivalued mappings, Funktsionalnyi Analiz (Ulyanovsk) 23, 82-93 (1984) (in Russian).
- [15] V.V.Obukhovskii, Semilinear Functional-Differential Inclusions in a Banach Space and Controlled Parabolic Systems, Soviet J. Automat. Inform. Sci. 24, 71-79 (1991).
- [16] V.V.Obukhovskii, P.Rubbioni, On a controllability problem for systems governed by semilinear functional differential inclusions in Banach spaces, Topol. Meth. Nonlin. Anal. 15, 141-151 (2000).
- [17] V.Obukhovskii, P.Zecca, On some properties of dissipative functional differential inclusions in a Banach space, Topol. Meth. Nonlin. Anal. 15, 369-384 (2000).

- [18] N.S.Papageorgiou, On multivalued evolutions equations and differential inclusions in Banach spaces, *Comment. Math. Univ. San. Paul.* 36, 21-39 (1987).
- [19] N.S.Papageorgiou, On quasilinear differential inclusions in Banach spaces, *Bull. Pol. Acad. Sci. Math.* 35, 407-415 (1987).
- [20] J.S.Shin, An existence theorem for functional differential equations with infinite delay in a Banach space, *Funkcial. Ekvac.* 30, 19-29 (1987).

COMPLETELY GENERALIZED NONLINEAR VARIATIONAL INCLUSIONS WITH FUZZY SETVALUED MAPPINGS

Ravi. P. Agarwal¹, M. Firdosh Khan², Donal O'Regan³ and Salahuddin²

¹Department of Mathematical Sciences, Florida Institute of Technology, Melbourne, Florida 32901-6975, USA

²Department of Mathematics, Aligarh Muslim University, Aligarh-202002, India

³Department of Mathematics, National University of Ireland, Galway, Ireland.

Abstract : In this paper, we study a class of completely generalized nonlinear variational inclusions for fuzzy setvalued mappings and construct some general iterative algorithms. We prove the existence of solutions for completely generalized nonlinear variational inclusions for fuzzy setvalued mappings and the convergence of iterative sequences generated by algorithms.

Key words : Variational inclusions, fuzzy mappings, proximal operators, existence and convergence theorems, Hilbert spaces and Hausdorff metric.

AMS Subject Classification : 49J40, 47H10

1. INTRODUCTION

In recent years, fuzzy set theory introduced by Zadeh [14] in 1965 has emerged as an interesting and fascinating branch of pure and applied sciences. The applications of fuzzy set theory can be found in many branches of regional, physical, mathematical and engineering sciences including artificial intelligence, computer sciences, control engineering, management science, economics, transportation problems and operation research, see [5,11,14,15] and references therein. The concept of variational inequalities and complementarity problem for fuzzy mappings were introduced and studied by many authors (see [2,3,6,8,9,10,13,15]). Motivated and inspired by recent research in this field, in this paper we consider a class of completely generalized nonlinear variational inclusions for fuzzy set-valued mappings and its equivalence with a class of proximal operator equations for fuzzy mappings will be shown. Using this equivalence, the iterative algorithms for the class of inclusions will be developed and we will prove that the approximate solutions obtained by this iterative algorithms, converge to the exact solutions of the variational inclusion.

2. PRELIMINARIES

Let H be a real Hilbert space with inner product and norm define by $\langle u, u \rangle = \|u\|^2$, we denote the collection of all fuzzy sets on H by $\mathcal{F}(H) = \{\mu : H \rightarrow [0, 1]\}$. A mapping

$F : H \rightarrow \mathcal{F}(H)$ is said to be a fuzzy mapping. For each $x \in H$, $F(x)$ (denote it by F_x , in the sequel) is a fuzzy set on H and $F_x(y)$ is the membership function of y in F_x .

A fuzzy mapping $F : H \rightarrow \mathcal{F}(H)$ is said to be closed, if for any $x \in H$, the function $F_x(y)$ is upper semicontinuous with respect to y , (i.e., for any given point $y_0 \in H$ and any net $\{y_\alpha\} \subset H$, when $y_\alpha \rightarrow y_0$, we have $F_x(y_0) \geq \limsup_\alpha F_x(y_\alpha)$).

Let $E \in \mathcal{F}(H)$, $q \in [0, 1]$. Then the set

$$(E)_q = \{x \in H : E(x) \geq q\}$$

is called a q -cut set of E .

Let $M, S, T : H \rightarrow \mathcal{F}(H)$ be closed fuzzy mappings satisfying the following condition:

Condition (I) : There exist functions $a, b, c : H \rightarrow [0, 1]$ such that for all $x \in H$, we have $(Mx)_{a(x)}, (Sx)_{b(x)}, (Tx)_{c(x)} \in CB(H)$, where $CB(H)$ denote the family of all nonempty closed and bounded subsets of H . Therefore, we can define three multivalued mappings, $\tilde{M}, \tilde{S}, \tilde{T} : H \rightarrow CB(H)$ by

$$\tilde{M}(x) = (M_x)_{a(x)}, \quad \tilde{S}(x) = (S_x)_{b(x)}, \quad \tilde{T}(x) = (T_x)_{c(x)},$$

for each $x \in H$. In the sequel $\tilde{M}, \tilde{S}, \tilde{T}$ are called the multivalued mappings induced by the fuzzy mappings M, S and T , respectively. More precisely, let $\partial\phi$ denote the subdifferential of a proper convex and lower semicontinuous function $\phi : H \times H \rightarrow R \cup \{+\infty\}$. Let $a, b, c : H \rightarrow [0, 1]$ be given functions and $M, S, T : H \rightarrow \mathcal{F}(H)$ fuzzy mappings. Let $g, F, G, P : H \rightarrow H$ be single-valued mappings with $\text{Im}g \cap \text{dom}\partial\phi(\cdot, v) \neq \emptyset$, and we consider the following completely generalized nonlinear variational inclusion with fuzzy setvalued mappings for finding $u, x, y, z \in H$ such that $g(u) \cap \text{dom}\partial\phi \neq \emptyset$ and $M_u(x) \geq a(u)$, $S_u(y) \geq b(u)$, $T_u(z) \geq c(u)$,

$$\langle P(x) - (Fy - Gz), v - g(u) \rangle \geq \phi(g(u), u) - \phi(v, u), \quad \forall v \in H. \quad (2.1)$$

Inequality (2.1) is called the completely generalized nonlinear variational inclusion for fuzzy setvalued mappings.

It is clear that completely generalized nonlinear variational inclusions for fuzzy setvalued mappings (2.1) includes many kinds of variational inclusions and inequalities as special cases, such as those in [1,2,3,7,8,9,10,13].

Definition 2.1[5]. If $G : H \rightarrow 2^H$ (2^H denotes the power set of H) is a maximal monotone multivalued mappings, then for any $\eta > 0$, the mapping $J_\eta^G : H \rightarrow H$ defined by

$$J_\eta^G(u) = (I + \eta G)^{-1}(u), \quad \forall u \in H,$$

is said to be the proximal operator of index η of G , where I is the identity mapping on H . Furthermore, the proximal operator J_η^G is single-valued and nonexpansive i.e.,

$$\|J_\eta^G(u) - J_\eta^G(v)\| \leq \|u - v\|, \quad \forall u, v \in H.$$

Since the subdifferential $\partial\phi$ of a proper convex and lower semicontinuous function $\phi : H \rightarrow R \cup \{+\infty\}$ is a maximal monotone multivalued mappings, it follows that the proximal operator $J_\eta^{\partial\phi}$ of index η of $\partial\phi$ is given by

$$J_\eta^{\partial\phi}(u) = (I + \eta\partial\phi)^{-1}(u), \quad \forall u \in H.$$

Lemma 2.1[5]. A given $u, w \in H$ satisfies the inequality

$$\langle u - w, v - u \rangle + \eta\phi(v) - \eta\phi(u) \geq 0, \quad \forall v \in H$$

if and only if

$$u = J_\eta^{\partial\phi}(w),$$

where $J_\eta^{\partial\phi} = (I + \eta\partial\phi)^{-1}$ is the proximal operator and $\eta > 0$ is a constant.

We now consider the problem of finding $u, w, x, y, z \in H$, $M_u(x) \geq a(u)$, $S_u(y) \geq b(u)$, $T_u(z) \geq c(u)$ and

$$P(x) + \eta^{-1}R_\eta^{\partial\phi(\cdot, u)}(w) = Fy - Gz, \quad (2.2)$$

where $R_\eta^{\partial\phi} = I - J_\eta^{\partial\phi}$, I is the identity mapping and $\eta > 0$. Equation (2.2) is called the proximal operator equation for fuzzy setvalued mappings.

3. EQUIVALENCE RELATION AND ITERATIVE ALGORITHMS

In this section, we establish the equivalence between (2.1) and (2.2), and we develop the iterative algorithms.

Lemma 3.1. The set (u, x, y, z) is the solution for (2.1) if and only if there exist $x \in \tilde{M}(u)$, $y \in \tilde{S}(u)$, $z \in \tilde{T}(u)$ such that

$$g(u) = J_\eta^{\partial\phi(\cdot, u)}[g(u) - \eta(P(x) - (Fy - Gz))], \quad (3.1)$$

where $\eta > 0$ is a constant and $J_\eta^{\partial\phi(\cdot, u)} = (I + \eta\partial\phi)^{-1}(u)$ is the so-called proximal mapping on H .

Let $u \in H$, $x \in \tilde{M}(u)$, $y \in \tilde{S}(u)$, $z \in \tilde{T}(u)$ and from the definition of the proximal operator $J_{\eta}^{\partial\phi(\cdot, u)}$ of index η of $\partial\phi(\cdot, u)$ and relation (3.1), we have

$$\begin{aligned} g(u) &= J_{\eta}^{\partial\phi(\cdot, u)}[g(u) - \eta(P(x) - (Fy - Gz))] \\ &= (I + \eta\partial\phi(\cdot, u))^{-1}[g(u) - \eta(P(x) - (Fy - Gz))] \end{aligned}$$

and

$$g(u) - \eta(P(x) - (Fy - Gz)) \in g(u) + \eta\partial\phi(\cdot, u)(g(u)),$$

which gives

$$(Fy - Gz) - P(x) \in \partial\phi(\cdot, u)(g(u)).$$

From the definition of $\partial\phi(\cdot, u)$, we have

$$\phi(v, u) \geq \phi(g(u), u) + \langle (Fy - Gz) - P(x), v - g(u) \rangle, \quad \forall v \in H.$$

Thus u, x, y and z are solutions of problem (2.1).

From Lemma 3.1, we conclude that (2.1) is equivalent to the fuzzy fixed point problem

$$u \in N(u) \tag{3.2}$$

where

$$N(u) = u - g(u) + J_{\eta}^{\partial\phi(\cdot, u)}[g(u) - \eta(P(x) - (Fy - Gz))].$$

Based on (3.1) and (3.2), we have the following iterative algorithm.

Algorithm 3.1. Suppose $P, F, G, g : H \rightarrow H$ are single-valued mappings. Let $M, S, T : H \rightarrow \mathcal{F}(H)$ be fuzzy mappings satisfying condition (I) and $\tilde{M}, \tilde{S}, \tilde{T} : H \rightarrow CB(H)$ the fuzzy mappings induced by M, S, T respectively. For a given $u_0 \in H$, we take $x_0 \in \tilde{M}(u_0)$, $y_0 \in \tilde{S}(u_0)$ and $z_0 \in \tilde{T}(u_0)$ and let

$$u_1 = u_0 - g(u_0) + J_{\eta}^{\partial\phi(\cdot, u_0)}[g(u_0) - \eta(P(x_0) - (Fy_0 - Gz_0))] \quad .$$

where $\eta > 0$ is a constant. Since $x_0 \in \tilde{M}(u_0) \in CB(H)$, $y_0 \in \tilde{S}(u_0) \in CB(H)$ and $z_0 \in \tilde{T}(u_0) \in CB(H)$, by Nadler [12 pp.480], there exist $x_1 \in \tilde{M}(u_1)$, $y_1 \in \tilde{S}(u_1)$ and $z_1 \in \tilde{T}(u_1)$ such that

$$\|x_0 - x_1\| \leq (1 + 1)\hat{H}(\tilde{M}(u_0), \tilde{M}(u_1)),$$

$$\|y_0 - y_1\| \leq (1 + 1)\hat{H}(\tilde{S}(u_0), \tilde{S}(u_1)),$$

$$\|z_0 - z_1\| \leq (1 + 1)\hat{H}(\tilde{T}(u_0), \tilde{T}(u_1)),$$

where $\hat{H}(\cdot, \cdot)$ is the Hausdörff metric on $CB(H)$. Let

$$u_2 = u_1 - g(u_1) + J_\eta^{\partial\phi(\cdot, u_1)}[g(u_1) - \eta(P(x_1) - (Fy_1 - Gz_1))].$$

By induction we can obtain sequences $\{u_n\}$, $\{x_n\}$, $\{y_n\}$ and $\{z_n\}$ such that

$$x_n \in \tilde{M}(u_n), \quad \|x_n - x_{n+1}\| \leq (1 + (n + 1)^{-1})\hat{H}(\tilde{M}(u_n), \tilde{M}(u_{n+1})),$$

$$y_n \in \tilde{S}(u_n), \quad \|y_n - y_{n+1}\| \leq (1 + (n + 1)^{-1})\hat{H}(\tilde{S}(u_n), \tilde{S}(u_{n+1})),$$

$$z_n \in \tilde{T}(u_n), \quad \|z_n - z_{n+1}\| \leq (1 + (n + 1)^{-1})\hat{H}(\tilde{T}(u_n), \tilde{T}(u_{n+1})),$$

$$u_{n+1} = u_n - g(u_n) + J_\eta^{\partial\phi(\cdot, u_n)}[g(u_n) - \eta(P(x_n) - (Fy_n - Gz_n))], \quad (3.3)$$

where $\eta > 0$ is a constant.

Now, we show that (2.1) is equivalent to (2.2).

Theorem 3.1. The problem (2.1) has a solution $u \in H$ such that $x \in \tilde{M}(u)$, $y \in \tilde{S}(u)$, $z \in \tilde{T}(u)$ if and only if (2.2) has a solution set (u, x, y, z) , where

$$g(u) = J_\eta^{\partial\phi(\cdot, u)}(w) \quad (3.4)$$

and

$$w = g(u) - \eta(P(x) - (Fy - Gz)), \quad (3.5)$$

here $\eta > 0$ is a constant.

Proof. Let $u \in H$, $x \in \tilde{M}(u)$, $y \in \tilde{S}(u)$, $z \in \tilde{T}(u)$ be a solution of (2.1), then by Lemma 3.1,

$$g(u) = J_\eta^{\partial\phi(\cdot, u)}[g(u) - \eta(P(x) - (Fy - Gz))]. \quad (3.6)$$

Using $R_\eta^{\partial\phi(\cdot, u)} = I - J_\eta^{\partial\phi(\cdot, u)}$ and (3.6), we have

$$\begin{aligned}
& R_{\eta}^{\partial\phi(\cdot, u)}[g(u) - \eta(P(x) - (Fy - Gz))] \\
&= [g(u) - \eta(P(x) - (Fy - Gz))] - J_{\eta}^{\partial\phi(\cdot, u)}[g(u) - \eta(P(x) - (Fy - Gz))] \\
&= [g(u) - \eta(P(x) - (Fy - Gz))] - g(u) = -\eta[P(x) - (Fy - Gz)],
\end{aligned}$$

which implies that

$$P(x) + \eta^{-1}R_{\eta}^{\partial\phi(\cdot, u)}(w) = Fy - Gz,$$

where $w = g(u) - \eta(P(x) - (Fy - Gz))$, here $\eta > 0$ is a constant.

Conversely, let $u \in H$, $x \in \tilde{M}(u)$, $y \in \tilde{S}(u)$, $z \in \tilde{T}(u)$ be a solution of (2.2), then

$$\begin{aligned}
-\eta(Fy - Gz) + \eta P(x) &= -R_{\eta}^{\partial\phi(\cdot, u)}(w) \\
&= J_{\eta}^{\partial\phi(\cdot, u)}(w) - w, \\
g(u) &= J_{\eta}^{\partial\phi(\cdot, u)}(w).
\end{aligned} \tag{3.7}$$

Now from Lemma 2.1 and equation (3.7), we have

$$\begin{aligned}
0 &\leq \langle g(u) - w, v - g(u) \rangle + \eta\phi(v, u) - \eta\phi(g(u), u) \\
&= \eta\{\langle P(x) - (Fy - Gz), v - g(u) \rangle + \phi(v, u) - \phi(g(u), u)\}.
\end{aligned}$$

Note that $\eta > 0$ is a constant, so the above relation implies $u \in H$, $x \in \tilde{M}(u)$, $y \in \tilde{S}(u)$, $z \in \tilde{T}(u)$ a solution set of (2.1).

The (2.2) can be written as

$$R_{\eta}^{\partial\phi(\cdot, u)}(w) = \eta(Fy - Gz) - \eta P(x),$$

from which it follows that

$$w = g(u) - \eta(P(x) - (Fy - Gz)), \tag{3.8}$$

completing the proof of the Theorem 3.1.

This fixed point formulation for fuzzy mappings enables us to suggest the following Algorithm.

Algorithm 3.2. Let g, F, G, P, S, T and M be as in Algorithm 3.1. For any given $u_0 \in H$, compute the sequences $\{u_n\}$, $\{w_n\}$, $\{x_n\}$, $\{y_n\}$ and $\{z_n\}$ for fuzzy mappings by the iteration method

$$g(u_n) = J_\eta^{\partial\phi(\cdot, u_n)}(w_n),$$

$$\begin{aligned} x_n \in \tilde{M}(u_n), \quad \|x_n - x_{n+1}\| &\leq (1 + (n+1)^{-1})\hat{H}(\tilde{M}(u_n), \tilde{M}(u_{n+1})), \\ y_n \in \tilde{S}(u_n), \quad \|y_n - y_{n+1}\| &\leq (1 + (n+1)^{-1})\hat{H}(\tilde{S}(u_n), \tilde{S}(u_{n+1})), \\ z_n \in \tilde{T}(u_n), \quad \|z_n - z_{n+1}\| &\leq (1 + (n+1)^{-1})\hat{H}(\tilde{T}(u_n), \tilde{T}(u_{n+1})), \quad n \geq 0, \end{aligned} \quad (3.9)$$

$$w_{n+1} = g(u_n) - \eta(P(x_n) - (Fy_n - Gz_n)), \quad (3.10)$$

here $\eta > 0$ is a constant.

(ii) (2.2) may be written as

$$0 = -\eta^{-1}R_\eta^{\partial\phi(\cdot, u_n)}(w_n) - (P(x_n) - (Fy_n - Gz_n)),$$

which implies that

$$R_\eta^{\partial\phi(\cdot, u_n)}(w_n) = (1 - \eta^{-1})R_\eta^{\partial\phi(\cdot, u_n)}(w_n) - (P(x_n) - (Fy_n - Gz_n)),$$

that is

$$\begin{aligned} w_n &= J_\eta^{\partial\phi(\cdot, u_n)}(w_n) - (P(x_n) - (Fy_n - Gz_n)) + (1 - \eta^{-1})R_\eta^{\partial\phi(\cdot, u_n)}(w_n) \\ &= g(u_n) - (P(x_n) - F(y_n - Gz_n)) + (1 - \eta^{-1})R_\eta^{\partial\phi(\cdot, u_n)}(w_n). \end{aligned} \quad (3.11)$$

Using the fixed point formulation for fuzzy mappings, the following iterative algorithm is proposed.

Algorithm 3.3. Let g, F, G, P, M, S and T be as in Algorithm 3.1. For any given $u_0 \in H$, compute the sequences $\{u_n\}$, $\{w_n\}$, $\{x_n\}$, $\{y_n\}$ and $\{z_n\}$ for fuzzy mappings by the iterative schemes

$$g(u_n) = J_\eta^{\partial\phi(\cdot, u_n)}(w_n), \quad (3.12)$$

$$\begin{aligned}
x_n \in \tilde{M}(u_n), \quad \|x_n - x_{n+1}\| &\leq (1 + (n+1)^{-1})\tilde{H}(\tilde{M}(u_n), \tilde{M}(u_{n+1})), \\
y_n \in \tilde{S}(u_n), \quad \|y_n - y_{n+1}\| &\leq (1 + (n+1)^{-1})\tilde{H}(\tilde{S}(u_n), \tilde{S}(u_{n+1})), \\
z_n \in \tilde{T}(u_n), \quad \|z_n - z_{n+1}\| &\leq (1 + (n+1)^{-1})\tilde{H}(\tilde{T}(u_n), \tilde{T}(u_{n+1})), \\
w_{n+1} &= g(u_n) - (P(x_n) - (Fy_n - Gz_n)) + (1 - \eta^{-1})R_\eta^{\partial\phi(\cdot, u_n)}(w_n),
\end{aligned} \tag{3.13}$$

here $\eta > 0$ is a constant.

Now, we consider for each fixed $u \in H$, the sequences $\{\phi_n\}$ of proper convex lower semicontinuous $\phi_n : H \times H \rightarrow R \cup \{+\infty\}$, approximating $\{\phi\}$, then we have following more general iterative algorithms for fuzzy set-valued mappings.

Algorithm 3.4. For any given $u_0 \in H$, $x_0 \in \tilde{M}(u_0)$, $y_0 \in \tilde{S}(u_0)$, $z_0 \in \tilde{T}(u_0)$, compute the sequences $\{u_n\}$, $\{w_n\}$, $\{x_n\}$, $\{y_n\}$ and $\{z_n\}$ for fuzzy mappings by iterative schemes

$$g(u_n) = J_\eta^{\partial\phi_n(\cdot, u_n)}(w_n), \tag{3.14}$$

$$\begin{aligned}
x_n \in \tilde{M}(u_n), \quad \|x_n - x_{n+1}\| &\leq (1 + (n+1)^{-1})\hat{H}(\tilde{M}(u_n), \tilde{M}(u_{n+1})), \\
y_n \in \tilde{S}(u_n), \quad \|y_n - y_{n+1}\| &\leq (1 + (n+1)^{-1})\hat{H}(\tilde{S}(u_n), \tilde{S}(u_{n+1})), \\
z_n \in \tilde{T}(u_n), \quad \|z_n - z_{n+1}\| &\leq (1 + (n+1)^{-1})\hat{H}(\tilde{T}(u_n), \tilde{T}(u_{n+1})), \\
w_{n+1} &= g(u_n) - \eta(P(x_n) - (Fy_n - Gz_n)), \quad \eta \geq 0,
\end{aligned} \tag{3.15}$$

here $\eta > 0$ is a constant.

Algorithm 3.5. For any given $u_0 \in H$, compute the sequences $\{u_n\}$, $\{w_n\}$, $\{x_n\}$, $\{y_n\}$ and $\{z_n\}$ for fuzzy mappings by iterative schemes

$$g(u_n) = J_\eta^{\partial\phi_n(\cdot, u_n)}(w_n), \tag{3.16}$$

$$\begin{aligned}
x_n \in \tilde{M}(u_n), \quad \|x_n - x_{n+1}\| &\leq (1 + (n+1)^{-1})\tilde{H}(\tilde{M}(u_n), \tilde{M}(u_{n+1})), \\
y_n \in \tilde{S}(u_n), \quad \|y_n - y_{n+1}\| &\leq (1 + (n+1)^{-1})\tilde{H}(\tilde{S}(u_n), \tilde{S}(u_{n+1})), \\
z_n \in \tilde{T}(u_n), \quad \|z_n - z_{n+1}\| &\leq (1 + (n+1)^{-1})\tilde{H}(\tilde{T}(u_n), \tilde{T}(u_{n+1})), \\
w_{n+1} &= g(u_n) - (P(x_n) - (Fy_n - Gz_n)) + (I - \eta^{-1})R_\eta^{\partial\phi_n(\cdot, u_n)}(w_n),
\end{aligned} \tag{3.17}$$

here $\eta > 0$ is a constant and

$$R_\eta^{\partial\phi_n(\cdot, u_n)}(w_n) = I - J_\eta^{\partial\phi_n(\cdot, u_n)}(w_n), \text{ for each } u \in H.$$

4. MAIN RESULTS

In this section, we consider those conditions under which the solutions of (2.1) exists and for which the sequences of approximate solutions obtained by our iterative algorithms converge to the exact solution of (2.2). For this purpose, we need the following concepts.

Definition 4.1. A mapping $g : H \rightarrow H$ is said to be

(i) strongly monotone if there exists a constant $r > 0$ such that

$$\langle g(u_1) - g(u_2), u_1 - u_2 \rangle \geq r \|u_1 - u_2\|^2, \quad \forall u_1, u_2 \in H,$$

(ii) Lipschitz continuous if there exists a constant $s > 0$ such that

$$\|g(u_1) - g(u_2)\| \leq s \|u_1 - u_2\|, \quad \forall u_1, u_2 \in H.$$

Definition 4.2. Let $\tilde{S} : H \rightarrow CB(H)$ be a mapping

(i) $\tilde{S}$ is said to be relaxed Lipschitz continuous with respect to a mapping $F : H \rightarrow H$ if there exists a constant $\alpha > 0$ such that

$$\langle Fy_1 - Fy_2, u_1 - u_2 \rangle \leq -\alpha \|u_1 - u_2\|^2, \quad \forall u_i \in H, y_i \in \tilde{S}(u_i), i = 1, 2;$$

(ii) $\tilde{S}$ is said to be relaxed monotone with respect to a mapping $G : H \rightarrow H$ if there exists a constant $\beta > 0$ such that

$$\langle Gy_1 - Gy_2, u_1 - u_2 \rangle \geq -\beta \|u_1 - u_2\|^2, \quad \forall y_i \in \tilde{S}(u_i), u_i \in H, i = 1, 2.$$

Definition 4.3[11]. A sequence $\{\phi_n\}$ of convex, proper, lower semicontinuous function $\phi_n : H \rightarrow R \cup \{+\infty\}$ is said to be convergent to ϕ in the sense of Mosco if

(i) for every $u \in H$, we have

$$\phi(u) \leq \liminf_{n \rightarrow \infty} \phi_n(u_n),$$

for every sequence $\{u_n\}$ in H which converges weakly to u , and

(ii) there exists a sequence $\{u_n\}$ in H which converges strongly to u and satisfies

$$\phi(u) \geq \limsup_{n \rightarrow \infty} \phi_n(u_n).$$

Lemma 4.1[4]. If the sequence $\{\phi_n\}$ of convex, proper, lower semicontinuous function $\phi_n : H \rightarrow R \cup \{+\infty\}$ converges to ϕ in the sense of Mosco, then

$$J_{\eta}^{\phi_n}(u) \rightarrow J_{\eta}^{\phi}(u) \quad \text{as } n \rightarrow \infty, \quad \text{for any } u \in H,$$

where $\eta > 0$ is a constant.

Assumption 4.1. For each fixed $v \in H$, let $\phi(., v)$ be a proper, convex and lower semicontinuous function $\phi : H \times H \rightarrow R \cup \{+\infty\}$. Then for each $u, v, w \in H$, there exists a coefficient $\mu > 0$ such that

$$\|J_{\eta}^{\partial\phi(.,u)}(w) - J_{\eta}^{\partial\phi(.,v)}(w)\| \leq \mu\|u - v\|,$$

where $\eta > 0$ is a constant.

Theorem 4.1. Let $M, S, T : H \rightarrow \mathcal{F}(H)$ be the closed fuzzy mappings satisfying condition (I) and $\tilde{M}, \tilde{S}, \tilde{T} : H \rightarrow CB(H)$ the multivalued mappings induced by the fuzzy mappings M, S, T , respectively. Let $\tilde{M}, \tilde{S}, \tilde{T}$ be δ -Lipschitz continuous, γ -Lipschitz continuous, ϵ -Lipschitz continuous, respectively. Let $F, G, P : H \rightarrow H$ be Lipschitz continuous with corresponding coefficients ξ, ρ and σ , respectively. Let $g : H \rightarrow H$ be Lipschitz continuous and strongly monotone with corresponding coefficients $s \geq 0$ and $r \geq 0$, respectively. Let $\tilde{S}$ be a relaxed Lipschitz continuous with respect to F with constant $\alpha > 0$ and $\tilde{T}$ be a relaxed monotone with respect to G with constant $\beta > 0$. Let the sequence $\{\phi_n\}$ of proper, convex and lower semicontinuous function $\phi_n : H \times H \rightarrow R \cup \{+\infty\}$ converges to ϕ in the sense of Mosco and assume Assumption 4.1 holds (for $\{\phi_n\}$).

Suppose that there exists a constant $\eta > 0$ such that

$$\begin{aligned} & \left| \eta - \frac{\alpha - \beta - \sigma\delta(1 - k - \mu)}{(\xi\gamma + \rho\epsilon)^2 - \sigma^2\delta^2} \right| \\ & < \frac{\sqrt{(\alpha - \beta - \sigma\delta(1 - k - \mu))^2 - ((\xi\gamma + \rho\epsilon)^2 - \sigma^2\delta^2)(k + \mu)(2 - k - \mu)}}{(\xi\gamma + \rho\epsilon)^2 - \sigma^2\delta^2}, \\ & \alpha > \beta + \sigma\delta(1 - k - \mu) + \sqrt{((\xi\gamma + \rho\epsilon)^2 - \sigma^2\delta^2)(k + \mu)(2 - k - \mu)}, \end{aligned} \quad (4.1)$$

$$\sigma\delta < \xi\gamma + \rho\epsilon,$$

where $k = 2\sqrt{1 - 2r + s^2}$.

Then there exists a set of elements $u \in H$, $x \in \tilde{M}(u)$, $y \in \tilde{S}(u)$, $z \in \tilde{T}(u)$, which satisfies (2.2) and the sequences $\{u_n\}$, $\{w_n\}$, $\{x_n\}$, $\{y_n\}$ and $\{z_n\}$ generated by Algorithm

3.4 converge to u, w, x, y and z strongly in H , respectively.

Proof. From Algorithm 3.4, we have

$$\begin{aligned} \|w_{n+1} - w_n\| &\leq \|g(u_n) - g(u_{n-1}) + \eta(Fy_n - Fy_{n-1}) - \eta(Gz_n - Gz_{n-1})\| + \eta\|P(x_n) - P(x_{n-1})\| \\ &\leq \|u_n - u_{n-1} - (g(u_n) - g(u_{n-1}))\| + \eta\|P(x_n) - P(x_{n-1})\| \\ &\quad + \|u_n - u_{n-1} + \eta(Fy_n - Fy_{n-1}) - \eta(Gz_n - Gz_{n-1})\|. \end{aligned} \quad (4.2)$$

By Lipschitz continuity and the strong monotonicity of the operator g , we obtain

$$\|u_n - u_{n-1} - (g(u_n) - g(u_{n-1}))\|^2 \leq (1 - 2r + s^2)\|u_n - u_{n-1}\|^2. \quad (4.3)$$

Again from the $\hat{H}$ -Lipschitz continuity of the fuzzy operators $\tilde{M}, \tilde{S}, \tilde{T}$ and the Lipschitz continuity of the operators P, F, G , we have

$$\begin{aligned} \|P(x_n) - P(x_{n-1})\| &\leq \sigma\|x_n - x_{n-1}\| \\ &\leq \sigma(1 + (n+1)^{-1})\hat{H}(\tilde{M}(u_n), \tilde{M}(u_{n-1})) \\ &\leq \sigma\delta(1 + n^{-1})\|u_n - u_{n-1}\|, \end{aligned} \quad (4.4)$$

$$\begin{aligned} \|F(y_n) - F(y_{n-1})\| &\leq \xi\|y_n - y_{n-1}\| \\ &\leq \xi(1 + (n+1)^{-1})\hat{H}(\tilde{S}(u_n), \tilde{S}(u_{n-1})) \\ &\leq \gamma\xi(1 + n^{-1})\|u_n - u_{n-1}\|, \end{aligned} \quad (4.5)$$

$$\begin{aligned} \|G(z_n) - G(z_{n-1})\| &\leq \rho\|z_n - z_{n-1}\| \\ &\leq \rho(1 + (n+1)^{-1})\hat{H}(\tilde{T}(u_n), \tilde{T}(u_{n-1})) \\ &\leq \rho\epsilon(1 + n^{-1})\|u_n - u_{n-1}\|. \end{aligned} \quad (4.6)$$

Further, since $\tilde{S}$ is relaxed Lipschitz continuous with respect to the operator F and the operator $\tilde{T}$ is relaxed monotone with respect to the operator G , we have

$$\begin{aligned} \|u_n - u_{n-1} + \eta(Fy_n - Fy_{n-1}) - \eta(Gz_n - Gz_{n-1})\|^2 &\leq \|u_n - u_{n-1}\|^2 \\ &\quad + 2\eta\langle Fy_n - Fy_{n-1}, u_n - u_{n-1} \rangle - 2\eta\langle Gz_n - Gz_{n-1}, u_n - u_{n-1} \rangle \end{aligned}$$

$$\begin{aligned}
& +\eta^2\|(Fy_n - Fy_{n-1}) - (Gz_n - Gz_{n-1})\|^2 \\
& \leq [1 - 2\eta(\alpha - \beta) + \eta^2(1 + n^{-1})^2(\xi\gamma + \rho\epsilon)^2]\|u_n - u_{n-1}\|^2.
\end{aligned} \tag{4.7}$$

From (4.2)-(4.7), we obtain

$$\begin{aligned}
\|w_{n+1} - w_n\| & \leq \{\sqrt{1 - 2r + s^2} + \eta\sigma\delta(1 + n^{-1}) + \sqrt{1 - 2\eta(\alpha - \beta) + \eta^2(1 + n^{-1})^2(\xi\gamma + \rho\epsilon)^2}\} \\
& \quad \|u_n - u_{n-1}\| \\
& = \left\{\frac{k}{2} + \eta\sigma\delta(1 + n^{-1}) + \Pi_n(\eta)\right\}\|u_n - u_{n-1}\|,
\end{aligned} \tag{4.8}$$

where $k = 2\sqrt{1 - 2r + s^2}$ and $\Pi_n(\eta) = \sqrt{1 - 2\eta(\alpha - \beta) + \eta^2(1 + n^{-1})^2(\xi\gamma + \rho\epsilon)^2}$.

From (3.14) and (4.3), we have

$$\begin{aligned}
\|u_n - u_{n-1}\| & = \|u_n - u_{n-1} - (g(u_n) - g(u_{n-1})) + J_\eta^{\partial\phi_n(\cdot, u_n)}(w_n) - J_\eta^{\partial\phi_{n-1}(\cdot, u_{n-1})}(w_{n-1})\| \\
& \leq \|u_n - u_{n-1} - (g(u_n) - g(u_{n-1}))\| + \|J_\eta^{\partial\phi_n(\cdot, u_n)}(w_n) - J_\eta^{\partial\phi_n(\cdot, u_n)}(w_{n-1})\| \\
& \quad + \|J_\eta^{\partial\phi_n(\cdot, u_n)}(w_{n-1}) - J_\eta^{\partial\phi_n(\cdot, u_{n-1})}(w_{n-1})\| + \|J_\eta^{\partial\phi_n(\cdot, u_{n-1})}(w_{n-1}) - J_\eta^{\partial\phi_{n-1}(\cdot, u_{n-1})}(w_{n-1})\| \\
& \leq \frac{k}{2}\|u_n - u_{n-1}\| + \|w_n - w_{n-1}\| + \mu\|u_n - u_{n-1}\| + \varepsilon_n,
\end{aligned}$$

since $J_\eta^{\partial\phi(\cdot, u)}$ is nonexpansive and Assumption 4.1 holds (with $\{\phi_n\}$), which implies that

$$\|u_n - u_{n-1}\| \leq \frac{1}{1 - \frac{k}{2} - \mu} [\|w_n - w_{n-1}\| + \varepsilon_n], \tag{4.9}$$

where

$$\varepsilon_n = \|J_\eta^{\partial\phi_n(\cdot, u_{n-1})}(w_{n-1}) - J_\eta^{\partial\phi_{n-1}(\cdot, u_{n-1})}(w_{n-1})\|.$$

Combining (4.8)-(4.9), we have

$$\begin{aligned}
\|w_{n+1} - w_n\| & \leq \left\{\frac{\frac{k}{2} + \eta\sigma\delta(1 + n^{-1}) + \Pi_n(\eta)}{1 - \frac{k}{2} - \mu}\right\}\|w_n - w_{n-1}\| \\
& \quad + \left\{\frac{\frac{k}{2} + \eta\sigma\delta(1 + n^{-1}) + \Pi_n(\eta)}{1 - \frac{k}{2} - \mu}\right\}\varepsilon_n \\
& \leq \theta_n\|w_n - w_{n-1}\| + \theta_n\varepsilon_n,
\end{aligned} \tag{4.10}$$

where $\theta_n = \frac{\frac{k}{2} + \eta\sigma\delta(1 + n^{-1}) + \Pi_n(\eta)}{1 - \frac{k}{2} - \mu}$. Let

$$\theta = \frac{\frac{k}{2} + \eta\sigma\delta + \Pi(\eta)}{1 - \frac{k}{2} - \mu} \quad \text{with} \quad \Pi(\eta) = \sqrt{1 - 2\eta(\alpha - \beta) + \eta^2(\xi\gamma + \rho\epsilon)^2}.$$

Then $\theta_n \rightarrow \theta$ as $n \rightarrow \infty$. It follows from (4.1) that $\theta < 1$. Hence $\theta_n < 1$ for n sufficiently large. Since $\varepsilon_n \rightarrow 0$ as $n \rightarrow \infty$, it follows from (4.10) that $\{w_n\}$ is a Cauchy sequence in H , that is $w_{n+1} \rightarrow w \in H$ as $n \rightarrow \infty$.

From (4.9), we see that $\{u_n\}$ is also Cauchy sequence in H , so there exists $u \in H$, such that $u_{n+1} \rightarrow u$ as $n \rightarrow \infty$.

Also from (4.4)-(4.6), we have

$$\begin{aligned} \|x_n - x_{n-1}\| &\leq (1 + n^{-1})\hat{H}(\tilde{M}(u_n), \tilde{M}(u_{n-1})) \leq (1 + n^{-1})\delta\|u_n - u_{n-1}\|, \\ \|y_n - y_{n-1}\| &\leq (1 + n^{-1})\hat{H}(\tilde{S}(u_n), \tilde{S}(u_{n-1})) \leq (1 + n^{-1})\gamma\|u_n - u_{n-1}\|, \\ \|z_n - z_{n-1}\| &\leq (1 + n^{-1})\hat{H}(\tilde{T}(u_n), \tilde{T}(u_{n-1})) \leq (1 + n^{-1})\epsilon\|u_n - u_{n-1}\|. \end{aligned}$$

It follows that $\{x_n\}$, $\{y_n\}$ and $\{z_n\}$ are also Cauchy sequences in H . Since H is complete, we may let $x_n \rightarrow x$, $y_n \rightarrow y$, $z_n \rightarrow z$ as $n \rightarrow \infty$. Further we have

$$\begin{aligned} d(x, \tilde{M}(u)) &\leq \|x - x_n\| + d(x_n, \tilde{M}(u)) \\ &\leq \|x - x_n\| + \hat{H}(\tilde{M}(u_n), \tilde{M}(u)) \\ &\leq \|x - x_n\| + \delta\|u_n - u\| \rightarrow 0 \quad \text{as } n \rightarrow \infty, \end{aligned}$$

where $d(x, \tilde{M}(u)) = \inf\{\|u - p\| : p \in \tilde{M}(u)\}$, and so we have $d(x, \tilde{M}(u)) = 0$. Hence we must have $x \in \tilde{M}(u)$. In a similar way, we can show $y \in \tilde{S}(u)$ and $z \in \tilde{T}(u)$.

Using the continuity of operators $P, F, G, M, S, T, \phi, J^{\partial\phi}$ and Algorithm 3.4, we have

$$\begin{aligned} w &= g(u) - \eta(P(x) - (Fy - Gz)) \\ &= J_{\eta}^{\partial\phi(\cdot, u)}(w) - \eta(P(x) - (Fy - Gz)) \in H, \quad \eta > 0 \text{ is a constant.} \end{aligned}$$

From Theorem 3.1, we see that $u, w \in H$, such that $x \in \tilde{M}(u)$, $y \in \tilde{S}(u)$ and $z \in \tilde{T}(u)$ are solution of (2.2) and consequently, $w_{n+1} \rightarrow w$, $u_{n+1} \rightarrow u$, $x_{n+1} \rightarrow x$, $y_{n+1} \rightarrow y$ and $z_{n+1} \rightarrow z$ strongly in H . This completes the proof.

References

1. R. Ahmad, K. R. Kazmi and Salahuddin, *Completely generalized nonlinear variational inclusions involving relaxed Lipschitz and relaxed monotone mappings*, Non-linear Anal. Forum, **5**(2000), 61-69.
2. R. Ahmad, Salahuddin and S. Husain, *Generalized nonlinear variational inclusions with fuzzy mappings*, Indian J. Pure Appl. Math. **32**(6)(2001), 943-947.
3. R. Ahmad, Salahuddin and S. S. Irfan, *Generalized quasi-complementarity problems with fuzzy set valued mappings*, Appearing in Int. J. Fuzzy Systems, September (2003).
4. H. Attouch, *Variational convergence for function and operators*, Pitman, London, 1984.
5. H. Brezis, *Operateurs Maximaux Monotones*, North Holland, Amsterdam, 1973.
6. S. S. Chang, *Coincidence theorems and variational inequalities for fuzzy mappings*, Fuzzy Sets and Systems, **61**(1994), 359-368.
7. S. S. Chang and N. J. Huang, *Generalized complementarity problems for fuzzy mappings*, Fuzzy Sets and Systems, **55**(1993), 227-234.
8. S. S. Chang and Y. G. Zhu, *On variational inequalities for fuzzy mappings*, Fuzzy Sets and Systems, **32**(1989), 359-367.
9. X. P. Ding and J. Y. Park, *A new class of generalized nonlinear implicit quasi-variational inclusions with fuzzy mappings*, J. Comput. Appl. Math., **138**(2002), 243-257.
10. N. J. Huang, *A general class of nonlinear variational inclusions for fuzzy mappings*, Indian J. Pure Appl. Math., **29**(9)(1998), 957-964.
11. U. Mosco, *Implicit variational problems and quasi-variational inequalities*, Lecture Notes in Maths, Springer-Verlag, Berlin, **543**(1976), 83-126.
12. Jr. S. B. Nadler, *Multivalued contraction mappings*, Pacific J. Math., **30**(1969), 475-488.
13. M. A. Noor, *Variational inequalities for fuzzy mappings*, Fuzzy Sets and System, **56**(1993), 309-312.
14. L. A. Zadeh, *Fuzzy sets*, Information and control, **8**(1965), 338-353.
15. H. J. Zimmerman, *Fuzzy sets theory and its applications*, Kluwer Academic Publishers, Boston, 1988.

Some Direct Results For The Iterative Combinations Of The Second Kind Beta Operators

Zoltán Finta

Babeş - Bolyai University, Department of Mathematics and Computer Science
1, M. Kogălniceanu St., 400084 Cluj-Napoca, Romania
e-mail : fzoltan@math.ubbcluj.ro

Vijay Gupta

School of Applied Sciences, Netaji Subhas Institute of Technology
Sector 3 Dwarka, Azad Hind Fauj Marg, New Delhi 110045, India
e-mail : vijay@nsit.ac.in

Abstract : In the present paper, we study the second kind beta operators introduced by D. D. Stancu [8] and prove some local and global direct results for the iterative combinations of these Stancu beta operators.

2000 **AMS Subject Classification :** 41A30 41A36

Key Words and Phrases : linear positive operators, iterative combinations, order of approximation, modulus of smoothness of order m , K - functional, Ditzian - Totik modulus of second order.

1 Introduction

The second kind beta operators L_n associating for $n \in N = \{1, 2, \dots\}$, are defined by

$$L_n(f, x) \equiv (L_n f)(x) = \frac{1}{B(nx, n+1)} \int_0^\infty f(t) \frac{t^{nx-1}}{(1+t)^{nx+n+1}} dt \quad (1)$$

The operators (1) were introduced by D. D. Stancu [8]. Recently U. Abel [1] estimated the complete asymptotic expansion for these operators. Another beta approximating operators of second kind have been studied in [2], [3] and [4].

Alternatively we may rewrite (1) as

$$L_n(f, x) = \int_0^\infty K_n(x, t) f(t) dt, \quad (2)$$

where

$$K_n(x, t) = \frac{1}{B(nx, n+1)} \cdot \frac{t^{nx-1}}{(1+t)^{nx+n+1}}$$

It is easily verified the operators L_n are linear positive operators and preserve the linear functions : $L_n(1, x) = 1$ and $L_n(t, x) = x$. Therefore we can take the iterative combinations of these operators easily.

It has been observed that the order of approximation by these operators L_n is $O(n^{-1})$, $n \rightarrow \infty$. On the other hand C. A. Micchelli [7] offered an approach for improving the order of approximation of Bernstein polynomials. It is interesting that by considering the iterative combinations due to C. A. Micchelli [7], the higher order of approximation may be achieved.

We consider here the iterative combinations of the Stancu beta operators. The iterative combinations $L_{n,k}(f, x)$ of the operators L_n are defined as follows:

$$L_{n,k}(f, x) \equiv (L_{n,k}f)(x) = [I - (I - L_n)^k](f, x) = \sum_{r=1}^k (-1)^{r+1} \binom{k}{r} L_n^r(f, x) \quad (3)$$

where L_n^r denotes the r -th iterate (superposition) of the operator L_n .

In the present paper we prove an asymptotic formula and an error estimate in terms of higher order integral modulus of smoothness, for the iterative combinations of these Stancu beta operators of second kind. Further, in the last section we give global direct results in terms of Ditzian - Totik modulus of smoothness.

2 Lemmas

In this section we present certain lemmas which are necessary to prove the main results of next two section.

Lemma 1 *Let the function $\mu_{n,m}(x)$, $m \in N \cup \{0\}$, be defined as*

$$\mu_{n,m}(x) = L_n((t-x)^m, x).$$

Then

$$\mu_{n,0}(x) = 1, \quad \mu_{n,1}(x) = 0, \quad \mu_{n,2}(x) = \frac{x(x+1)}{n-1}$$

and

$$\mu_{n,m}(x) = O\left(n^{-[(m+1)/2]}\right), \quad n \rightarrow \infty$$

where $0 \leq x < \infty$ and $m \in N \cup \{0\}$.

Proof. See [8, p. 234, Theorem 1].

For each $m \in N \cup \{0\}$ the m -th order moment $\mu_{n,m}^{\{q\}}(x)$ for the operators L_n^q is defined by

$$\mu_{n,m}^{\{q\}}(x) = L_n^q((t-x)^m, x) \quad (4)$$

For $q = 1$, $\mu_{n,m}^{\{1\}}(x) \equiv \mu_{n,m}(x)$.

Lemma 2 *Let γ and δ be two positive numbers. Then for any $m \in N$ there exists a constant $C_1 = C_1(m) > 0$ such that*

$$\left\| \int_{|t-x| \geq \delta} K_n(x, t) t^\gamma dt \right\|_{C[a,b]} \leq C_1 n^{-m}$$

Proof. It follows easily by using Lemma 1.

Lemma 3 *There holds the following relation*

$$\mu_{n,m}^{\{q+1\}}(x) = \sum_{j=0}^m \binom{m}{j} \sum_{i=0}^{m-j} \frac{1}{i!} D^i (\mu_{n,m-j}^{\{q\}}(x)) \mu_{n,i+j}(x), \quad D \equiv \frac{d}{dx}.$$

Proof. By (4), we have

$$\begin{aligned} \mu_{n,m}^{\{q+1\}}(x) &= L_n(L_n^q((t-x)^m, u), x) \\ &= \sum_{j=0}^m \binom{m}{j} L_n((u-x)^j L_n^q((t-u)^{m-j}, u), x) \\ &= \sum_{j=0}^m \binom{m}{j} L_n\left(\sum_{i=0}^{m-j} \frac{(u-x)^{i+j}}{i!} D^i(\mu_{n,m-j}^{\{q\}}(x)), x\right). \end{aligned}$$

Now making use of Lemma 1, the conclusion of this lemma follows immediately.

Lemma 4 *We have*

$$\mu_{n,m}^{\{q\}}(x) = O\left(n^{-[(m+1)/2]}\right), \quad n \rightarrow \infty \quad (5)$$

where $0 \leq x < \infty$ and $m \in N \cup \{0\}$.

Proof. For $q = 1$ the above result follows from Lemma 1. We shall prove the result (5) by the principle of mathematical induction. Let the result be true for q and we shall prove it for $q + 1$. Because $\mu_{n,m-j}^{\{q\}}(x) = O\left(n^{-[(m-j+1)/2]}\right)$ and $\mu_{n,m-j}^{\{q\}}(x)$ is a polynomial in x of degree $\leq m - j$, it is obviously seen that $D^i\left(\mu_{n,m-j}^{\{q\}}(x)\right) = O\left(n^{-[(m-j+1)/2]}\right)$. Applying Lemma 3, we get

$$\begin{aligned} \mu_{n,m}^{\{q+1\}}(x) &= \sum_{j=0}^m \binom{m}{j} \sum_{i=0}^{m-j} \frac{1}{i!} O\left(n^{-[(m-j+1)/2] + [(i+j+1)/2]}\right) \\ &= O\left(\sum_{j=0}^m \sum_{i=0}^{m-j} n^{-[(m+i+1)/2]}\right), \end{aligned}$$

which implies equation (5). This completes the proof of lemma.

By direct application of Lemma 1, Lemma 3 and Lemma 4, we have

$$L_{n,k}((t-x)^i, x) = O(n^{-k}), \quad n \rightarrow \infty,$$

where $L_{n,k}$ is the iterative combination defined by (3).

3 Local Approximation

In this section we present a Voronovskaja type asymptotic formula and an error estimate in terms of higher order modulus of smoothness, using the technique of linear approximating method, namely Steklov mean's.

Throughout this section let

$$M_\gamma[0, \infty) = \{ f : f \text{ be locally integrable on } (0, \infty) \text{ and } f(t) = O(t^\gamma), t \rightarrow \infty \text{ for some } \gamma > 0 \}.$$

For $f \in M_\gamma[0, \infty)$ we define the norm $\|\cdot\|_\gamma$ on $M_\gamma[0, \infty)$ by

$$\|f\|_\gamma = \sup_{0 < x < \infty} |f(x)| x^{-\gamma}.$$

Theorem 1 *Let $f \in M_\gamma[0, \infty)$. If the $2k$ -th derivative of f exists at a fixed point $x \in [0, \infty)$ then*

$$\lim_{n \rightarrow \infty} n^k [L_{n,k}(f, x) - f(x)] = \sum_{j=2}^{2k} Q(j, k, x) \cdot \frac{f^{(j)}(x)}{j!} \quad (6)$$

where $Q(j, k, x)$ are certain polynomials in x . Further, if $f^{(2k-1)}$ exists and is absolutely continuous over $[0, b]$ and $f^{(2k)} \in L_\infty[0, b]$ then, for any $[c, d] \subset (0, b)$, there holds

$$\|L_{n,k}f - f\|_{C[c,d]} \leq C_2 n^{-k} \left\{ \|f\|_\gamma + \|f^{(2k)}\|_{L_\infty[0,b]} \right\}, \quad (7)$$

where C_2 is a constant independent of f and n .

Proof. First by Taylor's expansion, we have

$$\begin{aligned} n^k [L_{n,k}(f, x) - f(x)] &= \\ &= n^k \sum_{j=1}^{2k} \frac{f^{(j)}(x)}{j!} L_{n,k}((t-x)^j, x) + n^k \sum_{r=1}^k (-1)^k \binom{k}{r} L_n^r(\varepsilon(t, x)(t-x)^{2k}, x) \\ &= E_1 + E_2, \end{aligned}$$

where $\varepsilon(t, x) \rightarrow 0$ as $t \rightarrow x$ and $|\varepsilon(t, x)| \leq Mt^\gamma$, $M > 0$. Using (5), we have

$$E_1 = n^k \sum_{j=2}^{2k} \frac{f^{(j)}(x)}{j!} L_{n,k}((t-x)^j, x) = \sum_{j=2}^{2k} \frac{f^{(j)}(x)}{j!} Q(j, k, x) + o(1).$$

If $\phi_\delta(t)$ denotes the characteristic function of the interval $(x-\delta, x+\delta)$ then

$$\begin{aligned} |E_2| &\leq n^k \sum_{r=1}^k \binom{k}{r} L_n^r(|\varepsilon(t, x)|(t-x)^{2k} \phi_\delta(t), x) + \\ &+ n^k \sum_{r=1}^k \binom{k}{r} L_n^r((t-x)^{2k}(1-\phi_\delta(t)), x) = E_3 + E_4. \end{aligned}$$

Now we estimate E_3 . Since $\varepsilon(t, x) \rightarrow 0$ as $t \rightarrow x$ therefore for given $\varepsilon > 0$ there exists a $\delta > 0$ such that $|\varepsilon(t, x)| < \varepsilon$ whenever $|t - x| < \delta$. Using Lemma 4, we get

$$E_3 \leq \left(\sup_{|t-x|<\delta} |\varepsilon(t, x)| \right) n^k \sum_{r=1}^k \binom{k}{r} L_n^r((t-x)^{2s}, x) < \varepsilon \cdot C_3$$

Next applying Lemma 2, for arbitrary $p > 0$, we obtain

$$E_4 \leq n^k \sum_{r=1}^k \binom{k}{r} L_n^r(Mt^\gamma(t-x)^{2k}(1-\phi_\delta(t)), x) < C_4 n^{-p} = o(1)$$

Due to the arbitrariness of $\varepsilon > 0$ we conclude that $E_2 \rightarrow 0$ as $n \rightarrow \infty$.

Combining the estimates of E_1, E_2 , we get the required estimate (6).

To prove (7), we may write

$$\begin{aligned} L_{n,k}(f, x) - f(x) &= L_{n,k}(f(t), x) - f(x) \\ &= L_{n,k}(\phi(t)(f(t) - f(x)), x) + L_{n,k}((1 - \phi(t))(f(t) - f(x)), x) \\ &= E_5 + E_6, \end{aligned}$$

where $\phi(t)$ denotes the characteristic function of the closed interval $[0, b]$. Proceeding along the lines of the proof of E_4 , for all $x \in [c, d]$, we have

$$E_6 \leq C_5 n^{-k} \|f\|_\gamma.$$

Again, for $t \in [0, b]$ and $x \in [c, d]$, we have, by given hypothesis

$$f(t) - f(x) = \sum_{i=1}^{2k-1} \frac{f^{(i)}(x)}{i!} (t-x)^i + \frac{1}{(2k-1)!} \int_x^t (t-w)^{2k-1} f^{(2k)}(w) dw$$

Thus

$$\begin{aligned} E_5 &= \sum_{i=1}^{2k-1} \frac{f^{(i)}(x)}{i!} L_{n,k}(\phi(t)(t-x)^i, x) + \\ &+ \frac{1}{(2k-1)!} L_{n,k} \left(\phi(t) \int_x^t (t-w)^{2k-1} f^{(2k)}(w) dw, x \right) \\ &= \sum_{i=1}^{2k-1} \frac{f^{(i)}(x)}{i!} \{L_{n,k}((t-x)^i, x) + L_{n,k}((\phi(t)-1)(t-x)^i, x)\} + \\ &+ \frac{1}{(2k-1)!} L_{n,k} \left(\phi(t) \int_x^t (t-w)^{2k-1} f^{(2k)}(w) dw, x \right) \\ &= \sum_{i=1}^{2k-1} \frac{f^{(i)}(x)}{i!} \{E_7 + E_8\} + E_9. \end{aligned}$$

By (6), we get $E_7 = O(n^{-k})$, uniformly for $x \in [c, d]$. Also as in the estimate of E_4 , we have $E_8 = O(n^{-k})$. Finally, by (6) and Lemma 4, we get

$$E_9 = \|f^{(2k)}\|_{L_\infty[0,b]} \cdot O(n^{-k})$$

Combining the estimates E_7, E_8, E_9 , we obtain

$$\|E_5\|_{C[c,d]} \leq C_6 \left\{ \sum_{i=1}^{2k-1} \|f^{(i)}\|_{C[c,d]} + \|f^{(2k)}\|_{L_\infty[0,b]} \right\}.$$

Finally, using the interpolation property due to S. Goldberg and V. Meir [6], we obtain the required result (7), which completes the proof of the theorem.

We now define the linear approximating function viz. Steklov's mean, which is the main tool to prove our next theorem.

Let $f \in M_\gamma[0, \infty)$, $m \in N$. Then the Steklov's mean $f_{\eta,m}$ of m -th order corresponding to f , for sufficiently small $\eta > 0$ is defined by

$$f_{\eta,m} = \eta^{-m} \int_{-\eta/2}^{\eta/2} \int_{-\eta/2}^{\eta/2} \cdots \int_{-\eta/2}^{\eta/2} \left\{ f(t) + (-1)^{m-1} \Delta_u^m f(t) \right\} dt_1 dt_2 \dots dt_m,$$

where $u = \sum_{i=1}^m t_i$, $t \in [a, b]$ and $\Delta_\eta^m f(t)$ is the m -th forward difference with step length η . We assume throughout this section that $0 < a_1 < a_2 < b_2 < b_1 < \infty$. It is easily verified (see e.g. [9]) that :

1. $f_{\eta,m}$ has continuous derivatives up to order m on $[a_1, b_1]$
2. $\|f_{\eta,m}^{(r)}\|_{C[a_1,b_1]} \leq C_7 \eta^{-r} \omega_r(f, \eta, a_1, b_1)$, $r = 1, 2, \dots, m$
3. $\|f - f_{\eta,m}\|_{C[a_2,b_2]} \leq C_8 \omega_m(f, \eta, a_1, b_1)$
4. $\|f_{\eta,m}\|_{C[a_2,b_2]} \leq C_9 \|f\|_\gamma$
5. $\|f_{\eta,m}^{(m)}\|_{C[a_2,b_2]} \leq C_{10} \|f\|_\gamma$,

where C_i 's, $i \in \{7, 8, \dots, 10\}$ are certain constants independent of f and η and $\omega_m(f, \eta, a, b)$ is the modulus of smoothness of order m corresponding to f :

$$\omega_m(f, \eta, a, b) = \sup_{a_1 \leq x \leq b_1} \sup_{0 \leq h \leq \eta} |\Delta_h^m f(x)|.$$

Theorem 2 *Let $f \in M_\gamma[0, \infty)$. Then, for sufficiently large n , there holds*

$$\|L_{n,k}f - f\|_{C[a_2,b_2]} \leq C_{11} \left\{ \omega_{2k}(f, n^{-1/2}, a_1, b_1) + n^{-k} \|f\|_\gamma \right\},$$

where C_{11} is an absolute constant.

Proof. By linearity property, we have

$$\begin{aligned} \|L_{n,k}f - f\|_{C[a_2,b_2]} &\leq \|L_{n,k}(f - f_{\eta,2k})\|_{C[a_2,b_2]} + \|L_{n,k}f_{\eta,2k} - f_{\eta,2k}\|_{C[a_2,b_2]} + \\ &+ \|f_{\eta,2k} - f\|_{C[a_2,b_2]} = H_1 + H_2 + H_3. \end{aligned}$$

Applying property 3. of Steklov's mean, we get

$$H_3 \leq C_{12} \omega_{2k}(f, \eta, a_1, b_1)$$

Making use of Theorem 1, we have

$$H_2 \leq n^{-k} C_{13} \sum_{j=2}^{2k} \|f_{\eta,2k}^{(j)}\|_{C[a_2,b_2]} \leq C_{14} \left\{ \|f_{\eta,2k}\|_{C[a_2,b_2]} + \|f_{\eta,2k}^{(2k)}\|_{C[a_2,b_2]} \right\},$$

where we have used the interpolation property due to S. Goldberg and V. Meir [6]. Now, by using properties 4. and 5. of Steklov's mean, we get

$$H_2 \leq C_{15} \|f\|_\gamma.$$

Setting a^*, b^* satisfying $a_1 < a^* < a_2 < b_2 < b^* < b_1$ and let $\xi(t)$ be the characteristic function of the closed interval $[a^*, b^*]$, then, by Lemma 2 and property 3. of Steklov's mean, we obtain

$$\begin{aligned} H_1 &= \|L_{n,k}(f(t) - f_{\eta,2k})\|_{C[a_2,b_2]} \\ &\leq \|L_{n,k}(\xi(t)(f(t) - f_{\eta,2k}))\|_{C[a_2,b_2]} + \|L_{n,k}((1 - \xi(t))(f(t) - f_{\eta,2k}))\|_{C[a_2,b_2]} \\ &\leq C_{16} \|f - f_{\eta,2k}\|_{C[a^*,b^*]} + C_{17} n^{-m} \|f\|_\gamma \\ &\leq C_{18} \omega_{2k}(f, \eta, a_1, b_1) + C_{19} \|f\|_\gamma. \end{aligned}$$

Choosing $m \geq k$ and $\eta = n^{-1/2}$ in the above estimate we get the required result.

4 Global Approximation

In this section we establish direct global approximation theorems for Stancu beta operators. Let $C_B[0, \infty)$ be the space of all real valued continuous bounded functions f on $[0, \infty)$ endowed with the norm $\|f\| = \sup_{0 \leq x < \infty} |f(x)|$. Our results are given with the aid of Ditzian - Totik modulus of smoothness of second order defined by

$$\omega_\varphi^2(f, \delta) = \sup_{0 \leq h \leq \delta} \sup_{x \pm h\varphi(x) \in [0, \infty)} |f(x + h\varphi(x)) - 2f(x) + f(x - h\varphi(x))|,$$

where $\varphi(x) = \sqrt{x(1+x)}$, $0 \leq x < \infty$. The corresponding K -functional is

$$K_{2,\varphi}(f, \delta^2) = \inf_{g \in W_\infty^2} \left\{ \|f - g\| + \delta^2 \|\varphi^2 g''\| \right\},$$

where $W_\infty^2(\varphi) = \{ g \in C_B[0, \infty) : g' \in AC_{loc}[0, \infty), \varphi^2 g'' \in C_B[0, \infty) \}$ denote the weighted Sobolev space.

Our first theorem in this section is :

Theorem 3 *Let $f \in C_B[0, \infty)$. Then*

$$\|L_n f - f\| \leq C_{20} \omega_\varphi^2 \left(f, \frac{1}{\sqrt{n-1}} \right), \quad n = 2, 3, \dots$$

where $C_{20} > 0$ is an absolute constant.

Proof. Let $g \in W_\infty^2(\varphi)$. By Taylor's formula, we have

$$g(t) = g(x) + g'(x) (t - x) + \int_x^t (t - u) g''(u) du.$$

In view of Lemma 1, we get

$$L_n(g, x) - g(x) = L_n \left(\int_x^t (t - u) g''(u) du, x \right)$$

Hence

$$\begin{aligned} |L_n(g, x) - g(x)| &\leq L_n \left(\left| \int_x^t \frac{|t-u|}{\varphi^2(u)} \cdot \varphi^2(u) |g''(u)| du \right|, x \right) \\ &\leq L_n \left(\left| \int_x^t \frac{|t-u|}{u(1+u)} du \right|, x \right) \cdot \|\varphi^2 g''\| \end{aligned}$$

Furthermore, in view of [5, p. 141, (9.6.2)] and Lemma 1, we have

$$\begin{aligned}
 |L_n(g, x) - g(x)| &\leq L_n \left(\frac{(t-x)^2}{x} \cdot \left(\frac{1}{1+x} + \frac{1}{1+t} \right), x \right) \|\varphi^2 g''\| \\
 &= \left\{ \frac{1}{x(1+x)} L_n((t-x)^2, x) + \frac{1}{x} L_n \left(\frac{(t-x)^2}{1+t}, x \right) \right\} \|\varphi^2 g''\| \\
 &= \left\{ \frac{1}{n-1} + \frac{1}{x} L_n \left(\frac{(t-x)^2}{1+t}, x \right) \right\} \|\varphi^2 g''\| \quad (8)
 \end{aligned}$$

By direct computation we obtain

$$\begin{aligned}
 L_n \left(\frac{(t-x)^2}{1+t}, x \right) &= L_n \left(\frac{t^2}{1+t} - 2x \frac{t}{1+t} + x^2 \frac{1}{1+t}, x \right) \\
 &= \frac{x(nx+1)}{nx+n+1} - 2x \frac{nx}{nx+n+1} + x^2 \frac{n+1}{nx+n+1} \\
 &= \frac{x(1+x)}{nx+n+1}
 \end{aligned}$$

Hence, by (8), we have

$$\begin{aligned}
 |L_n(g, x) - g(x)| &\leq \left(\frac{1}{n-1} + \frac{1}{n} \cdot \frac{x(1+x)}{nx+n+1} \right) \|\varphi^2 g''\| \\
 &= \left(\frac{1}{n-1} + \frac{1+x}{nx+n+1} \right) \|\varphi^2 g''\| \\
 &\leq \left(\frac{1}{n-1} + \frac{1}{n} \right) \|\varphi^2 g''\| \leq \frac{2}{n-1} \|\varphi^2 g''\|.
 \end{aligned}$$

Also, L_n is a contraction, i.e.

$$\|L_n f\| \leq \|f\|, \quad f \in C_B[0, \infty). \quad (9)$$

Thus

$$\begin{aligned}
 |L_n(f, x) - f(x)| &\leq |L_n(f - g, x) - (f - g)(x)| + |L_n(g, x) - g(x)| \\
 &\leq 2 \|f - g\| + \frac{2}{n-1} \|\varphi^2 g''\|
 \end{aligned}$$

Taking the infimum on the right - hand side over all $g \in W_\infty^2(\varphi)$, we get

$$|L_n(f, x) - f(x)| \leq 2 K_{2,\varphi} \left(f, \frac{1}{n-1} \right) \quad (10)$$

By [5, p. 11, Theorem 2.1.1], there exists an absolute constant $C_{21} > 0$ such that

$$K_{2,\varphi}\left(f, \frac{1}{n-1}\right) \leq C_{21} \omega_{\varphi}^2\left(f, \frac{1}{\sqrt{n-1}}\right), \quad n > 1.$$

Hence, by (10)

$$\|L_n f - f\| \leq C_{20} \omega_{\varphi}^2\left(f, \frac{1}{\sqrt{n-1}}\right), \quad n = 2, 3, \dots,$$

which was to be proved.

The next theorem is :

Theorem 4 *Let $f \in C_B[0, \infty)$. Then*

$$\|L_{n,k} f - f\| \leq 2^k k K_{2,\varphi}\left(f, \frac{1}{n-1}\right), \quad n = 2, 3, \dots$$

Proof. We have

$$\begin{aligned} L_{n,k}(f, x) - f(x) &= \\ &= \sum_{r=1}^k (-1)^{r+1} \binom{k}{r} [L_n^r(f, x) - f(x)] + \sum_{r=1}^k (-1)^{r+1} \binom{k}{r} f(x) - f(x) \\ &= \sum_{r=1}^k (-1)^{r+1} \binom{k}{r} [L_n^r(f, x) - f(x)] + \sum_{r=0}^k (-1)^r \binom{k}{r} f(x) \\ &= \sum_{r=1}^k (-1)^{r+1} \binom{k}{r} [L_n^r(f, x) - f(x)]. \end{aligned} \tag{11}$$

On the other hand, by (9), we obtain

$$\|L_n^r f - f\| \leq r \|L_n f - f\|, \quad r = 1, 2, \dots$$

Hence, by (11) and (10),

$$\begin{aligned} |L_{n,k}(f, x) - f(x)| &\leq \sum_{r=1}^k \binom{k}{r} |L_n^r(f, x) - f(x)| \leq \sum_{r=1}^k \binom{k}{r} \|L_n^r f - f\| \\ &\leq \sum_{r=1}^k \binom{k}{r} r \|L_n f - f\| = \|L_n f - f\| \sum_{r=1}^k k \binom{k-1}{r-1} \\ &= 2^{k-1} \cdot k \|L_n f - f\| \leq 2^k \cdot k K_{2,\varphi}\left(f, \frac{1}{n-1}\right). \end{aligned}$$

This completes the proof of Theorem 4 .

Corollary 1 *There exists an absolute constant $C_{22} = C_{22}(k) > 0$ such that*

$$\|L_{n,k}f - f\| \leq C_{22} \omega_{\varphi}^2\left(f, \frac{1}{\sqrt{n-1}}\right),$$

for all $f \in C_B[0, \infty)$ and $n = 2, 3, \dots$

Proof. It follows from Theorem 4 and [5, p. 11, Theorem 2.1.1].

References

- [1] U. Abel, Asymptotic approximation with Stancu beta operators, *Rev. Anal. Numér. Théorie Approx.* 27 (1), 5 - 13 (1998).
- [2] J. A. Adell, J. de la Cal, On a Bernstein - type operator associated with the inverse Polya - Eggenberger distribution, *Rend. Circolo Matem. Palermo, Ser. II* 33, 143 - 154 (1993).
- [3] J. A. Adell, J. de la Cal, Limiting properties of inverse beta and generalized Bleimann - Butzer - Hahn operators, *Math. Proc. Cambridge Philos. Soc.* 114 (3), 489 - 498 (1993).
- [4] J. A. Adell, J. de la Cal, M. San Miguel, Inverse beta and generalized Bleimann - Butzer - Hahn operators, *J. Approx. Theory* 76 (1), 54 - 64 (1994).
- [5] Z. Ditzian, V. Totik, *Moduli of Smoothness*, Springer - Verlag, New York Berlin Heidelberg, 1987.
- [6] S. Goldberg, V. Meir, Minimum moduli of ordinary differential operators, *Proc. London Math. Soc.* 23 (3), 1 - 15 (1971).
- [7] C. A. Micchelli, The saturation class and iterates of the Bernstein polynomials, *J. Approx. Theory* 8, 1 - 18 (1973).
- [8] D. D. Stancu, On the beta approximating operators of second kind, *Rev. Anal. Numér. Théorie Approx.* 24 (1 - 2), 231 - 239 (1995).
- [9] A. F. Timan, *Theory of Approximation of Functions of a Real Variable*, Hindustan Publ. Corp., Delhi, 1966.

Instructions to Contributors
Journal of Concrete and Applicable Mathematics

A quarterly international publication of Eudoxus Press, LLC, of TN.

Editor in Chief: George Anastassiou

Department of Mathematical Sciences
 University of Memphis
 Memphis, TN 38152-3240, U.S.A.

1. Manuscripts hard copies in triplicate, and in English, should be submitted to the Editor-in-Chief:

Prof. George A. Anastassiou
Department of Mathematical Sciences
The University of Memphis
Memphis, TN 38152, USA.
Tel. 001. 901.678.3144
e-mail: ganastss@memphis.edu.

Authors may want to recommend an associate editor the most related to the submission to possibly handle it.

Also authors may want to submit a list of six possible referees, to be used in case we cannot find related referees by ourselves.

2. Manuscripts should be typed using any of TEX, LaTeX, AMS-TEX, or AMS-LaTeX and according to EUDOXUS PRESS, LLC. LATEX STYLE FILE. (VISIT www.msci.memphis.edu/~ganastss/jcaam/ to save a copy of the style file.) They should be carefully prepared in all respects. Submitted copies should be brightly printed (not dot-matrix), double spaced, in ten point type size, on one side high quality paper 8(1/2)x11 inch. Manuscripts should have generous margins on all sides and should not exceed 24 pages.

3. Submission is a representation that the manuscript has not been published previously in this or any other similar form and is not currently under consideration for publication elsewhere. A statement transferring from the authors (or their employers, if they hold the copyright) to Eudoxus Press, LLC, will be required before the manuscript can be accepted for publication. The Editor-in-Chief will supply the necessary forms for this transfer. Such a written transfer of copyright, which previously was assumed to be implicit in the act of submitting a manuscript, is necessary under the U.S. Copyright Law in order for the publisher to carry through the dissemination of research results and reviews as widely and effectively as possible.

4. The paper starts with the title of the article, author's name(s) (no titles or degrees), author's affiliation(s) and e-mail addresses. The affiliation should comprise the department, institution (usually university or company), city, state (and/or nation) and mail code.

The following items, 5 and 6, should be on page no. 1 of the paper.

5. An abstract is to be provided, preferably no longer than 150 words.
 6. A list of 5 key words is to be provided directly below the abstract. Key words should express the precise content of the manuscript, as they are used for indexing purposes. The main body of the paper should begin on page no. 1, if possible.

7. All sections should be numbered with Arabic numerals (such as: 1. INTRODUCTION).

Subsections should be identified with section and subsection numbers (such as 6.1. Second-Value Subheading).

If applicable, an independent single-number system (one for each category) should be used to label all theorems, lemmas, propositions, corollaries, definitions, remarks, examples, etc. The label (such as Lemma 7) should be typed with paragraph indentation, followed by a period and the lemma itself.

8. Mathematical notation must be typeset. Equations should be numbered consecutively with Arabic numerals in parentheses placed flush right, and should be thusly referred to in the text [such as Eqs.(2) and (5)]. The running title must be placed at the top of even numbered pages and the first author's name, et al., must be placed at the top of the odd numbered pages.

9. Illustrations (photographs, drawings, diagrams, and charts) are to be numbered in one consecutive series of Arabic numerals. The captions for illustrations should be typed double space. All illustrations, charts, tables, etc., must be embedded in the body of the manuscript in proper, final, print position. In particular, manuscript, source, and PDF file version must be at camera ready stage for publication or they cannot be considered.

Tables are to be numbered (with Roman numerals) and referred to by number in the text. Center the title above the table, and type explanatory footnotes (indicated by superscript lowercase letters) below the table.

10. List references alphabetically at the end of the paper and number them consecutively. Each must be cited in the text by the appropriate Arabic numeral in square brackets on the baseline.

References should include (in the following order):

initials of first and middle name, last name of author(s)

title of article,

name of publication, volume number, inclusive pages, and year of publication.

Authors should follow these examples:

Journal Article

1. H.H.Gonska, Degree of simultaneous approximation of bivariate functions by Gordon operators, (journal name in italics) *J. Approx. Theory*, 62,170-191(1990).

Book

2. G.G.Lorentz, (title of book in italics) *Bernstein Polynomials* (2nd ed.), Chelsea, New York, 1986.

Contribution to a Book

3. M.K.Khan, Approximation properties of beta operators, in (title of book in italics) *Progress in Approximation Theory* (P.Nevai and A.Pinkus, eds.), Academic Press, New York, 1991, pp.483-495.

11. All acknowledgements (including those for a grant and financial support) should occur in one paragraph that directly precedes the References section.

12. Footnotes should be avoided. When their use is absolutely necessary, footnotes should be numbered consecutively using Arabic numerals and should be typed at the bottom of the page to which they refer. Place a line above the footnote, so that it is set

off from the text. Use the appropriate superscript numeral for citation in the text.

13. After each revision is made please again submit three hard copies of the revised manuscript, including in the final one. And after a manuscript has been accepted for publication and with all revisions incorporated, manuscripts, including the TEX/LaTeX source file and the PDF file, are to be submitted to the Editor's Office on a personal-computer disk, 3.5 inch size. Label the disk with clearly written identifying information and properly ship, such as:

Your name, title of article, kind of computer used, kind of software and version number, disk format and files names of article, as well as abbreviated journal name.

Package the disk in a disk mailer or protective cardboard. Make sure contents of disks are identical with the ones of final hard copies submitted!

Note: The Editor's Office cannot accept the disk without the accompanying matching hard copies of manuscript. No e-mail final submissions are allowed! The disk submission must be used.

14. Effective 1 Nov. 2005 the journal's page charges are \$8.00 per PDF file page. Upon acceptance of the paper an invoice will be sent to the contact author. The fee payment will be due one month from the invoice date. The article will proceed to publication only after the fee is paid. The charges are to be sent, by money order or certified check, in US dollars, payable to Eudoxus Press, LLC, to the address shown in the Scope and Prices section.

No galleys will be sent and the contact author will receive one(1) electronic copy of the journal issue in which the article appears.

15. This journal will consider for publication only papers that contain proofs for their listed results.

TABLE OF CONTENTS, JOURNAL OF CONCRETE AND APPLICABLE
MATHEMATICS, VOL.4, NO.2, 2006

POWER SERIES OF FUZZY NUMBERS WITH CROSS PRODUCT AND
APPLICATIONS TO FUZZY DIFFERENTIAL EQUATIONS,
A.BAN, B.BEDE,125

ON A DELAY INTEGRAL EQUATION IN BIOMATHEMATICS,
A.BICA, C.IANCU,153

SET-VALUED VARIATIONAL INCLUSIONS WITH FUZZY MAPPINGS
IN BANACH SPACES,
A.SIDDIQI, R.AHMAD, S.IRFAN,171

ON SOME PROPERTIES OF SEMILINEAR FUNCTIONAL DIFFERENTIAL
INCLUSIONS IN ABSTRACT SPACES, C.GORI, V.OBUKHOVSKII,
M.RAGNI, P.RUBBIONI,183

COMPLETELY GENERALIZED NONLINEAR VARIATIONAL INCLUSIONS
WITH FUZZY SETVALUED MAPPINGS, R.AGARWAL, M.KHAN,
D.O'REGAN, SALAHUDDIN,215

SOME DIRECT RESULTS FOR THE ITERATIVE COMBINATIONS OF THE
SECOND KIND BETA OPERATORS,
Z.FINTA, V.GUPTA,229

VOLUME 4,NUMBER 3 JULY 2006

ISSN:1548-5390 PRINT,1559-176X ONLINE



**JOURNAL
OF CONCRETE
AND APPLICABLE
MATHEMATICS**

EUDOXUS PRESS,LLC

SCOPE AND PRICES OF THE JOURNAL
Journal of Concrete and Applicable Mathematics

A quartely international publication of **Eudoxus Press, LLC**

Editor in Chief: George Anastassiou
Department of Mathematical Sciences,
University of Memphis
Memphis, TN 38152, U.S.A.
ganastss@memphis.edu

The main purpose of the "Journal of Concrete and Applicable Mathematics" is to publish high quality original research articles from all subareas of Non-Pure and/or Applicable Mathematics and its many real life applications, as well connections to other areas of Mathematical Sciences, as long as they are presented in a Concrete way. It welcomes also related research survey articles and book reviews. A sample list of connected mathematical areas with this publication includes and is not restricted to: Applied Analysis, Applied Functional Analysis, Probability theory, Stochastic Processes, Approximation Theory, O.D.E, P.D.E, Wavelet, Neural Networks, Difference Equations, Summability, Fractals, Special Functions, Splines, Asymptotic Analysis, Fractional Analysis, Inequalities, Moment Theory, Numerical Functional Analysis, Tomography, Asymptotic Expansions, Fourier Analysis, Applied Harmonic Analysis, Integral Equations, Signal Analysis, Numerical Analysis, Optimization, Operations Research, Linear Programming, Fuzzyness, Mathematical Finance, Stochastic Analysis, Game Theory, Math. Physics aspects, Applied Real and Complex Analysis, Computational Number Theory, Graph Theory, Combinatorics, Computer Science Math. related topics, combinations of the above, etc. In general any kind of Concretely presented Mathematics which is Applicable fits to the scope of this journal. Working Concretely and in Applicable Mathematics has become a main trend in many recent years, so we can understand better and deeper and solve the important problems of our real and scientific world. "Journal of Concrete and Applicable Mathematics" is a peer-reviewed International Quarterly Journal. We are calling for papers for possible publication. The contributor should send three copies of the contribution to the editor in-Chief typed in TEX, LATEX double spaced. [See: Instructions to Contributors]

Journal of Concrete and Applicable Mathematics(JCAAM)
ISSN:1548-5390 PRINT, 1559-176X ONLINE.

is published in January, April, July and October of each year by
EUDOXUS PRESS, LLC,
1424 Beaver Trail Drive, Cordova, TN 38016, USA,
Tel. 001-901-751-3553
anastassioug@yahoo.com
<http://www.EudoxusPress.com>.

Annual Subscription Current Prices: For USA and Canada, Institutional: Print \$250, Electronic \$220, Print and Electronic \$310. Individual: Print \$77, Electronic \$60, Print & Electronic \$110. For any other part of the world add \$25 more to the above prices for Print.

Single article PDF file for individual \$8. Single issue in PDF form for individual \$25.
The journal carries page charges \$8 per page of the pdf file of an article, payable upon acceptance of the article within one month and before publication.
No credit card payments. Only certified check, money order or international check in US dollars are acceptable.
Combination orders of any two from JoCAAA, JCAAM, JAFA receive 25% discount, all three receive 30% discount.

Copyright©2004 by Eudoxus Press, LLC all rights reserved. JCAAM is printed in USA.

JCAAM is reviewed and abstracted by AMS Mathematical Reviews, MATHSCI, and Zentralblatt MATH.

It is strictly prohibited the reproduction and transmission of any part of JCAAM and in any form and by any means without the written permission of the publisher. It is only allowed to educators to Xerox articles for educational purposes. The publisher assumes no responsibility for the content of published papers.

JCAAM IS A JOURNAL OF RAPID PUBLICATION

Editorial Board Associate Editors

Editor in -Chief:
George Anastassiou
Department of Mathematical Sciences
The University Of Memphis
Memphis, TN 38152, USA
tel. 901-678-3144, fax 901-678-2480
e-mail ganastss@memphis.edu
www.msci.memphis.edu/~ganastss/jcaam
Areas: Approximation Theory,
Probability, Moments, Wavelet,
Neural Networks, Inequalities, Fuzzyness.

Associate Editors:

1) Ravi Agarwal
Florida Institute of Technology
Applied Mathematics Program
150 W. University Blvd.
Melbourne, FL 32901, USA
agarwal@fit.edu
Differential Equations, Difference
Equations, inequalities

2) Shair Ahmad
University of Texas at San Antonio
Division of Math. & Stat.
San Antonio, TX 78249-0664, USA
SAHMAD@utsa.edu
Differential Equations, Mathematical
Biology

3) Drumi D. Bainov
Medical University of Sofia
P.O. Box 45, 1504 Sofia, Bulgaria
drumibainov@yahoo.com
Differential Equations, Optimal Control,
Numerical Analysis, Approximation Theory

4) Carlo Bardaro
Dipartimento di Matematica & Informatica
Universita' di Perugia
Via Vanvitelli 1
06123 Perugia, ITALY
tel. +390755855034, +390755853822,
fax +390755855024
bardaro@unipg.it ,
bardaro@dipmat.unipg.it
Functional Analysis and Approximation Th.,

19) Rupert Lasser
Institut fur Biomathematik & Biomertie, GSF
-National Research Center for environment and
health
Ingolstaedter landstr.1
D-85764 Neuherberg, Germany
lasser@gsf.de
Orthogonal Polynomials, Fourier Analysis,
Mathematical Biology

20) Alexandru Lupas
University of Sibiu
Faculty of Sciences
Department of Mathematics
Str. I. Ratiu nr. 7
2400-Sibiu, Romania
lupas@ulbsibiu.ro
Classical Analysis, Inequalities,
Special Functions, Umbral Calculus,
Approximation Th., Numerical Analysis
and Methods

21) Ram N. Mohapatra
Department of Mathematics
University of Central Florida
Orlando, FL 32816-1364
tel. 407-823-5080
ramm@mail.ucf.edu
Real and Complex analysis, Approximation Th.,
Fourier Analysis, Fuzzy Sets and Systems

22) Rainer Nagel
Arbeitsbereich Funktionalanalysis
Mathematisches Institut
Auf der Morgenstelle 10
D-72076 Tuebingen
Germany
tel. 49-7071-2973242
fax 49-7071-294322
rana@fa.uni-tuebingen.de
evolution equations, semigroups, spectral th.,
positivity

23) Panos M. Pardalos
Center for Appl. Optimization
University of Florida
303 Weil Hall
P.O. Box 116595
Gainesville, FL 32611-6595

Summability, Signal Analysis, Integral
Equations,
Measure Th., Real Analysis

5) Francoise Bastin
Institute of Mathematics
University of Liege
4000 Liege
BELGIUM
f.bastin@ulg.ac.be
Functional Analysis, Wavelets

6) Paul L. Butzer
RWTH Aachen
Lehrstuhl A für Mathematik
D-52056 Aachen
Germany
tel. 0049/241/80-94627 office,
0049/241/72833 home,
fax 0049/241/80-92212
Butzer@rwth-aachen.de
Approximation Th., Sampling Th., Signals,
Semigroups of Operators, Fourier Analysis

7) Yeol Je Cho
Department of Mathematics Education
College of Education
Gyeongsang National University
Chinju 660-701
KOREA
tel. 055-751-5673 Office,
055-755-3644 home,
fax 055-751-6117
yjcho@nongae.gsnu.ac.kr
Nonlinear operator Th., Inequalities,
Geometry of Banach Spaces

8) Sever S. Dragomir
School of Communications and Informatics
Victoria University of Technology
PO Box 14428
Melbourne City M.C
Victoria 8001, Australia
tel 61 3 9688 4437, fax 61 3 9688 4050
sever.dragomir@vu.edu.au,
sever@sci.vu.edu.au
Math. Analysis, Inequalities, Approximation
Th.,
Numerical Analysis, Geometry of Banach
Spaces,
Information Th. and Coding

9) A.M. Fink
Department of Mathematics
Iowa State University
Ames, IA 50011-0001, USA

tel. 352-392-9011
pardalos@ufl.edu
Optimization, Operations Research

24) Svetlozar T. Rachev
Dept. of Statistics and Applied Probability
Program
University of California, Santa Barbara
CA 93106-3110, USA
tel. 805-893-4869
rachev@pstat.ucsb.edu
AND
Chair of Econometrics and Statistics
School of Economics and Business Engineering
University of Karlsruhe
Kollegium am Schloss, Bau II, 20.12, R210
Postfach 6980, D-76128, Karlsruhe, Germany
tel. 011-49-721-608-7535
rachev@lsoe.uni-karlsruhe.de
Mathematical and Empirical Finance,
Applied Probability, Statistics and Econometrics

25) Paolo Emilio Ricci
Universita' degli Studi di Roma "La Sapienza"
Dipartimento di Matematica-Istituto
"G. Castelnuovo"
P.le A. Moro, 2-00185 Roma, ITALY
tel. ++39 0649913201, fax ++39 0644701007
riccip@uniroma1.it, Paoloemilio.Ricci@uniroma1.it
Orthogonal Polynomials and Special functions,
Numerical Analysis, Transforms, Operational
Calculus,
Differential and Difference equations

26) Cecil C. Rousseau
Department of Mathematical Sciences
The University of Memphis
Memphis, TN 38152, USA
tel. 901-678-2490, fax 901-678-2480
ccrousse@memphis.edu
Combinatorics, Graph Th.,
Asymptotic Approximations,
Applications to Physics

27) Tomasz Rychlik
Institute of Mathematics
Polish Academy of Sciences
Chopina 12, 87100 Torun, Poland
T.Rychlik@impan.gov.pl
Mathematical Statistics, Probabilistic
Inequalities

28) Bl. Sendov
Institute of Mathematics and Informatics
Bulgarian Academy of Sciences
Sofia 1090, Bulgaria

tel.515-294-8150
fink@math.iastate.edu
Inequalities, Ordinary Differential
Equations

10) Sorin Gal
Department of Mathematics
University of Oradea
Str. Armatei Romane 5
3700 Oradea, Romania
galso@uoradea.ro
Approximation Th., Fuzzyness, Complex
Analysis

11) Jerome A. Goldstein
Department of Mathematical Sciences
The University of Memphis,
Memphis, TN 38152, USA
tel. 901-678-2484
jgoldste@memphis.edu
Partial Differential Equations,
Semigroups of Operators

12) Heiner H. Gonska
Department of Mathematics
University of Duisburg
Duisburg, D-47048
Germany
tel. 0049-203-379-3542 office
gonska@informatik.uni-duisburg.de
Approximation Th., Computer Aided
Geometric Design

13) Dmitry Khavinson
Department of Mathematical Sciences
University of Arkansas
Fayetteville, AR 72701, USA
tel. (479) 575-6331, fax (479) 575-8630
dmitry@uark.edu
Potential Th., Complex Analysis, Holomorphic
PDE, Approximation Th., Function Th.

14) Virginia S. Kiryakova
Institute of Mathematics and Informatics
Bulgarian Academy of Sciences
Sofia 1090, Bulgaria
virginia@diogenes.bg
Special Functions, Integral Transforms,
Fractional Calculus

15) Hans-Bernd Knoop
Institute of Mathematics
Gerhard Mercator University
D-47048 Duisburg
Germany
tel. 0049-203-379-2676

bSENDOV@BAS.BG
Approximation Th., Geometry of Polynomials,
Image Compression

29) Igor Shevchuk
Faculty of Mathematics and Mechanics
National Taras Shevchenko
University of Kyiv
252017 Kyiv
UKRAINE
shevchuk@univ.kiev.ua
Approximation Theory

30) H.M. Srivastava
Department of Mathematics and Statistics
University of Victoria
Victoria, British Columbia V8W 3P4
Canada
tel. 250-721-7455 office, 250-477-6960 home,
fax 250-721-8962
harimsri@math.uvic.ca
Real and Complex Analysis, Fractional Calculus
and Appl.,
Integral Equations and Transforms, Higher
Transcendental
Functions and Appl., q-Series and q-Polynomials,
Analytic Number Th.

31) Ferenc Szidarovszky
Dept. Systems and Industrial Engineering
The University of Arizona
Engineering Building, 111
PO Box 210020
Tucson, AZ 85721-0020, USA
szidar@sie.arizona.edu
Numerical Methods, Game Th., Dynamic Systems,
Multicriteria Decision making,
Conflict Resolution, Applications
in Economics and Natural Resources
Management

32) Gancho Tachev
Dept. of Mathematics
Univ. of Architecture, Civil Eng. and Geodesy
1 Hr. Smirnenski blvd
BG-1421 Sofia, Bulgaria
Approximation Theory

33) Manfred Tasche
Department of Mathematics
University of Rostock
D-18051 Rostock
Germany
manfred.tasche@mathematik.uni-rostock.de
Approximation Th., Wavelet, Fourier Analysis,
Numerical Methods, Signal Processing,

knoop@math.uni-duisburg.de
Approximation Theory, Interpolation

16) Jerry Koliha
Dept. of Mathematics & Statistics
University of Melbourne
VIC 3010, Melbourne
Australia
koliha@unimelb.edu.au
Inequalities, Operator Theory,
Matrix Analysis, Generalized Inverses

17) Mustafa Kulenovic
Department of Mathematics
University of Rhode Island
Kingston, RI 02881, USA
kulenm@math.uri.edu
Differential and Difference Equations

18) Gerassimos Ladas
Department of Mathematics
University of Rhode Island
Kingston, RI 02881, USA
gladas@math.uri.edu
Differential and Difference Equations

Image Processing, Harmonic Analysis

34) Chris P. Tsokos
Department of Mathematics
University of South Florida
4202 E. Fowler Ave., PHY 114
Tampa, FL 33620-5700, USA
profcpt@math.usf.edu, profcpt@chuma1.cas.usf.edu
Stochastic Systems, Biomathematics,
Environmental Systems, Reliability Th.

35) Lutz Volkmann
Lehrstuhl II fuer Mathematik
RWTH-Aachen
Templergraben 55
D-52062 Aachen
Germany
volkm@math2.rwth-aachen.de
Complex Analysis, Combinatorics, Graph Theory

On the uniform exponential stability of evolution families in terms of the admissibility of an Orlicz sequence space

Ciprian Preda,¹ Sever S. Dragomir² and Constantin Chilarescu³

¹ *Department of Electrical Engineering, University of California
Los Angeles, CA 90095, U.S.A
e-mail: preda@ee.ucla.edu*

² *School of Computer Science and Mathematics Victoria University
PO Box 14428, Melbourne City, MC 8001, Australia
e-mail: sever@matilda.vu.edu.au*

³ *Department of Mathematics, West University of Timisoara
4 Parvan Blv., 300223, Timisoara, Romania
e-mail: cchilarescu@rectorat.wt.ro*

A characterization of the exponential stability for evolution families in terms of the admissibility of some pairs of sequence spaces, is given. It is used the method of "test function" using a very well-known class of sequence spaces (the so-called Orlicz sequence spaces). Thus, are extended to the case of the abstract evolution families some results due to Coffman, Schäffer, Przyluski, Rolewicz.

1

¹Mathematics Subject Classification: 34D05, 47D06, 93D20.

Key words and phrases: evolution families, exponential stability, Orlicz sequence spaces, admissibility.

1 Introduction

Over the past ten years, the asymptotic theory of linear evolution equations in Banach spaces has witnessed an explosive development. A number of long-standing open problems have recently been solved and the theory seems to have obtained a certain degree of maturity. An important role in the early development of qualitative theory of differential systems, was played by the paper "Die stabilitätsfrage bei differentialgleichungen" [13], where Perron gave a characterization of exponential stability of the solutions to the linear differential equations

$$\frac{dx}{dt} = A(t)x, \quad t \in [0, +\infty), \quad x \in \mathbb{R}^n$$

where $A(t)$ is a matrix bounded continuous function, in terms of the existence of bounded solutions of the equations $\frac{dx}{dt} = A(t)x + f(t)$, where f is a continuous bounded function on $\mathbb{R}_+$. After these seminal researches of O. Perron, relevant results concerning the extension of Perron's problem in the more general framework of infinite-dimensional Banach spaces were obtained, in their pioneering monographs, by M. G. Krein, J. L. Daleckij, J. L. Massera and J. J. Schäffer. In the present there are different important characterizations of exponential stability or dichotomy for a strongly continuous, exponentially bounded evolution family, which can be found for instance in the papers due to N. van Minh [11,12], Y. Latushkin[2,7,8], P. Randolph [8], P. Preda[10,15], R. Schnaubelt [11,17], S. Montgomery-Smith[7]. For the case of discrete-time systems analogous results were firstly obtained by Ta Li in 1934 [see 18]. In Ta Li's paper, can be found the same principal concern as in Perron's work, but using discrete arguments. This approach was later developed for the discrete-time systems in the infinite-dimensional case by Ch.V. Coffman and J.J. Schäffer in 1967 [3] and D. Henry in 1981[5]. More recently we refer the reader to the the papers due to A. Ben-Artzi[2], I. Gohberg[2], M. Pinto[19], J. P. La Salle[8]. Applications of this "discrete-time theory" to stability theory of linear infinite-dimensional continuous-time systems have been presented by Przyłuski and Rolewicz in [16]. The first aim of this paper is to give a characterization of the exponential stability for an abstract evolution family (not necessary provided by a differential system) in terms of the admissibility of a Orlicz sequence space, so this work

fit into the context initiated by Ta Li. Roughly speaking, this means that an evolution family, acting on a Banach space X , is uniformly exponentially stable if and only if the corresponding inhomogeneous difference equation with the inhomogeneous term from $l^\Phi(X)$ admits a solution in $l^\Phi(X)$, where $l^\Phi(X) = \{f : \mathbb{N} \mapsto X : (\|f(n)\|)_{n \in \mathbb{N}} \text{ is in the Orlicz sequence space } l^\Phi\}$. Also, we note that the technique used in this work does not require any continuity hypothesis about the evolution families (as the strongly continuity or partial strongly continuity, assumptions required in most of the above cited works). Classical examples of Orlicz sequence spaces are l^p -spaces, $p \in [1, \infty)$, and there is also presented in the paper other example of a Orlicz sequence spaces which is not l^p (also, for well-chosen Young functions Φ other various examples of Orlicz sequence space can be found). So, this approach generalize the equivalence between (l^p, l^p) -admissibility and the exponential stability of evolution families using the admissibility of a Orlicz sequence space. Also, roughly speaking, we note that using this discrete-time theory we obtain here the uniform exponential stability of the continuous-time evolution families, fact which is very useful. Moreover, this approach does break a new ground and it can bring other useful situations (as Example 2.2), adding some nice twist to the subject concerning the connection between admissibility and exponential stability of evolution families.

2 Preliminaries

First, let us fix some standard notation. For X a Banach space we will denote by $l^p(X)$ and $l^\infty(X)$ the normed spaces,

$$l^p(X) = \{f : \mathbb{N} \rightarrow X : \sum_{n=0}^{\infty} \|f(n)\|^p < \infty\}, \quad p \in [1, \infty),$$

$$l^\infty(X) = \{f : \mathbb{N} \rightarrow X : \sup_{n \in \mathbb{N}} \|f(n)\| < \infty\}.$$

We note that $l^p(X), l^\infty(X)$ are Banach spaces endowed with the respectively norms

$$\|f\|_p = \left(\sum_{n=0}^{\infty} \|f(n)\|^p \right)^{1/p};$$

$$\|f\|_\infty = \sup_{n \in \mathbb{N}} \|f(n)\|.$$

For the simplicity of notations we denote by $l^p = l^p(\mathbf{R}), l^\infty = l^\infty(\mathbf{R})$.

For more convenience we will list in the next the definition of a Orlicz sequence space. Let $\varphi : \mathbb{R}_+ \rightarrow \mathbb{R}_+$ be a function which is non-decreasing, with $\varphi(t) > 0$, for each $t > 0$. Define

$$\Phi(t) = \int_0^t \varphi(s) ds$$

A function Φ of this form is called an Young function. For $f : \mathbb{N} \rightarrow \mathbb{R}$ a real sequence and the Young function Φ we define

$$m^\Phi(f) = \sum_{k=0}^{\infty} \Phi(|f(k)|).$$

The set l^Φ of all f for which there exists a $j > 0$ that $m^\Phi(jf) < \infty$ is easily checked to be a linear space. Equipped with the Luxemburg norm

$$\|f\|_\Phi = \inf \left\{ j > 0 : m^\Phi\left(\frac{1}{j}f\right) \leq 1 \right\}$$

the space $(l^\Phi, \|\cdot\|_\Phi)$ becomes a Banach space.

Remark 2.1. It is easy to check that $\chi_{\{0, \dots, n\}} \in l^\Phi$ and $\|\chi_{\{0, \dots, n\}}\|_\Phi = \frac{1}{\Phi^{-1}(\frac{1}{n+1})}$, for all $n \in \mathbb{N}$, where in general χ_A denotes the characteristic (indicator) function of the set A .

Example 2.1. Classical examples of Orlicz sequence spaces are the well-known l^p -spaces, for all $p \in [1, \infty)$, (i.e. $l^\Phi = l^p$ if $\Phi(t) = t^p$).

Remark 2.2 If $(l^\Phi, \|\cdot\|) = (l^p, \|\cdot\|)$ then $\lim_{t \rightarrow 0} \frac{\Phi(t)}{t^p} = 1$.

Proof. If $(l^\Phi, \|\cdot\|) = (l^p, \|\cdot\|)$ then $\|\chi_{\{0, \dots, n\}}\|_\Phi = \|\chi_{\{0, \dots, n\}}\|_p$, for all $n \in \mathbb{N}$, which is equivalent with:

$$\Phi^{-1}\left(\frac{1}{n+1}\right) = \left(\frac{1}{n+1}\right)^{\frac{1}{p}}, \quad \text{for all } n \in \mathbb{N}.$$

Let $x \in (0, 1]$ and $m = \left[\frac{1}{x}\right] \in \mathbb{N}^*$, where $[a]$ denotes the greatest integer which is less or equal than a . Using the fact that Φ^{-1} is nondecreasing we

have that

$$\left(\frac{1}{m+1}\right)^{\frac{1}{p}} = \Phi^{-1}\left(\frac{1}{m+1}\right) \leq \Phi^{-1}(x) \leq \Phi^{-1}\left(\frac{1}{m}\right) = \left(\frac{1}{m}\right)^{\frac{1}{p}}$$

which implies that

$$\left[\frac{1}{\left(\left[\frac{1}{x}\right] + 1\right)x}\right]^{\frac{1}{p}} \leq \frac{\Phi^{-1}(x)}{x^{\frac{1}{p}}} \leq \left[\frac{1}{x\left[\frac{1}{x}\right]}\right]^{\frac{1}{p}}, \text{ for all } x \in (0, 1].$$

Hence

$$\lim_{x \rightarrow 0} \frac{\Phi^{-1}(x)}{x^{\frac{1}{p}}} = 1$$

and

$$\lim_{u \rightarrow 0} \frac{\Phi(u)}{u^p} = \lim_{u \rightarrow 0} \frac{1}{\left[\frac{\Phi^{-1}(\Phi(u))}{(\Phi(u))^{\frac{1}{p}}}\right]^p} = 1$$

Example 2.2. Consider $\varphi, \Phi : \mathbb{R}_+ \rightarrow \mathbb{R}_+$

$$\varphi(t) = \sum_{m=1}^{\infty} \frac{\sqrt[m]{t}}{m^2}, \quad \Phi(t) = \int_0^t \varphi(s) ds = \sum_{m=1}^{\infty} \frac{t^{1+\frac{1}{m}}}{m(m+1)}$$

We pretend that $l^\Phi \neq l^p$, for all $p \in [1, \infty)$. Indeed, $\lim_{t \rightarrow 0} \frac{\Phi(t)}{t} = 0$ and $\lim_{t \rightarrow 0} \frac{\Phi(t)}{t^p} = \infty$, for all $p \in (1, \infty)$ and using the above remark, our claim follows easily.

In what follows, if l^Φ is a Orlicz sequence space we denote by

$$l^\Phi(X) = \{f : \mathbb{N} \mapsto X : (||f(n)||)_{n \in \mathbb{N}} \text{ is in } l^\Phi\}.$$

Remark 2.3. $l^\Phi(X)$ is a Banach space endowed with the norm

$$||f||_{l^\Phi(X)} = || ||f(\cdot)|| ||_\Phi.$$

Remark 2.4. For any Orlicz sequence space l^Φ we have that:

i) $l^\Phi \subset l^\infty$;

$$\text{ii)} \|f\|_{\infty} \leq \frac{1}{\|\chi_{\{0\}}\|_{\Phi}} \|f\|_{\Phi}$$

Proposition 2.1. *If Φ is an Young function of the Orlicz sequence space l^{Φ} then the followings statements hold:*

i) The map $a_{\Phi} : \mathbb{N} \rightarrow \mathbb{R}_{+}^{}$ given by $a_{\Phi}(n) = (n+1)\Phi^{-1}(\frac{1}{n+1})$ is a nondecreasing one.*

ii) $\sum_{k=0}^n |f(k)| \leq a_{\Phi}(n) \|f\|_{\Phi}$, for all $n \geq 0, f \in l^{\Phi}$.

Proof.

i) First let us prove that the map $b : \mathbb{R}_+^* \rightarrow \mathbb{R}_+^*$, given by $b(u) = \frac{\Phi(u)}{u}$ is nondecreasing. If $0 < u_1 \leq u_2$ then

$$\begin{aligned} \frac{\Phi(u_1)}{u_1} &= \frac{1}{u_1} \int_0^{u_1} \varphi(s) ds = \frac{1}{u_1} \int_0^{u_2} \varphi\left(\frac{u_1}{u_2}v\right) \frac{u_1}{u_2} dv = \\ &= \frac{1}{u_2} \int_0^{u_2} \varphi\left(\frac{u_1}{u_2}v\right) dv \leq \frac{1}{u_2} \int_0^{u_2} \varphi(v) dv = \frac{\Phi(u_2)}{u_2}. \end{aligned}$$

In order to prove that a_Φ is nondecreasing let $n \in \mathbb{N}$. It follows that $0 < w_2 := \Phi^{-1}\left(\frac{1}{n+2}\right) \leq \Phi^{-1}\left(\frac{1}{n+1}\right) := w_1$ and so $b(w_2) \leq b(w_1)$. Having in mind that

$$b(w_1) = \frac{1}{a_\Phi(n+1)} \quad \text{and} \quad b(w_2) = \frac{1}{a_\Phi(n)},$$

it results that a_Φ is a nondecreasing function.

ii) Consider $f \in l^\Phi, n \in \mathbb{N}^*, c > 0$ such that $m^\Phi\left(\frac{1}{c}f\right) \leq 1$. Then we have that

$$\Phi\left(\frac{1}{c(n+1)} \sum_{k=0}^n |f(k)|\right) \leq \frac{1}{n+1} \sum_{k=0}^n \Phi\left(\frac{1}{c}|f(k)|\right) \leq \frac{1}{n+1},$$

and so

$$\sum_{k=0}^n |f(k)| \leq (n+1) \Phi^{-1}\left(\frac{1}{n+1}\right) c$$

which implies that

$$\sum_{k=0}^n |f(k)| \leq (n+1) \Phi^{-1}\left(\frac{1}{n+1}\right) \|f\|_\Phi = a_\Phi(n) \|f\|_\Phi,$$

for all $n \in \mathbb{N}, f \in l^\Phi$.

Remark 2.5. Using a simple translation argument we may state that

$$\sum_{k=n_0}^{n_0+n} |f(k)| \leq a_\Phi(n) \|f\|_\Phi,$$

for all $n_0, n \in \mathbb{N}, f \in L^\Phi$.

Definition 2.1. A family of bounded linear operators acting on X denoted by $\mathcal{U} = \{U(t, s)\}_{t \geq s \geq 0}$ is called an evolution family if the following statements hold:

- $e_1)$ $U(t, t) = I$ (where I is the identity operator on X), for all $t \geq 0$;
- $e_2)$ $U(t, s) = U(t, r)U(r, s)$, for all $t \geq r \geq s \geq 0$;
- $e_3)$ there exist $M > 0, \omega > 0$ such that

$$\|U(t, s)\| \leq Me^{\omega(t-s)} \quad \text{for all } t \geq s \geq 0.$$

Definition 2.2. The evolution family $\mathcal{U} = \{U(t, s)\}_{t \geq s \geq 0}$ is called uniformly exponentially stable (u.e.s) if there exist two strictly positive constants N, ν such that the following statement hold:

$$\|U(t, s)\| \leq Ne^{-\nu(t-s)}.$$

If l^Φ is a Orlicz sequence space we give the following:

Definition 2.3. l^Φ is said to be admissible to the evolution family $\mathcal{U}$ if for all $f \in l^\Phi(X)$ the application $x_f : \mathbb{N} \rightarrow X$ defined by

$$x_f(n) = \sum_{k=0}^n U(n, k)f(k) \quad \text{lies in } l^\Phi(X).$$

3. The main result

Let l^Φ be a Orlicz sequence space.

Lemma 3.1. *If l^Φ is admissible to $\mathcal{U}$ then there is $K > 0$ such that*

$$\|x_f\|_{l^\Phi(X)} \leq K\|f\|_{l^\Phi(X)}$$

Proof. We define the linear operator, acting on $l^\Phi(X)$, denoted by $T : l^\Phi(X) \rightarrow l^\Phi(X)$, and given by

$$(Tf)(m) = \sum_{k=0}^m U(m, k)f(k).$$

Consider $\{g_n\}_{n \in \mathbb{N}} \subset l^\Phi(X)$, $g \in l^\Phi(X)$, $h \in l^\Phi(X)$ such that

$$g_n \xrightarrow{l^\Phi(X)} g, \quad Tg_n \xrightarrow{l^\Phi(X)} h$$

Then

$$\begin{aligned} \|(Tg_n)(m) - (Tg)(m)\| &\leq \sum_{k=0}^m \|U(m, k)(g_n(k) - g(k))\| \leq \\ &\leq \sum_{k=0}^m Me^{\omega m} \|g_n(k) - g(k)\| \leq Me^{\omega m} a_\Phi(m) \|g_n - g\|_{l^\Phi(X)}, \end{aligned}$$

for all $m, n \in \mathbb{N}$.

It follows, using again the Remark 2.4, that $Tg = h$, and hence T is closed and so, by the Closed-Graph Theorem it is also bounded.

So we obtain that

$$\|x_f\|_{l^\Phi(X)} = \|Tf\|_{l^\Phi(X)} \leq \|T\| \|f\|_{l^\Phi(X)}, \text{ for all } f \in l^\Phi(X) \text{ as required.}$$

Lemma 3.2. *Let $g : \{(t, t_0) \in \mathbb{R}^2 : t \geq t_0 \geq 0\} \rightarrow \mathbb{R}_+$ be a function such that the following properties hold.*

1) $g(t, t_0) \leq g(t, s)g(s, t_0)$, for all $t \geq s \geq t_0 \geq 0$;

2) $\sup_{0 \leq t_0 \leq t \leq t_0+1} g(t, t_0) < \infty$;

3) *there exist a sequence $h : \mathbb{N} \rightarrow \mathbb{R}_+$, $\lim_{n \rightarrow \infty} h(n) = 0$ and $g(m+n, n) \leq h(m)$, for all $m, n \in \mathbb{N}$.*

Then there exist two constants $N, \nu > 0$ such that

$$g(t, t_0) \leq Ne^{-\nu(t-t_0)} , \quad \text{for all } t \geq t_0 \geq 0$$

Proof. Let $a = \sup_{0 \leq t_0 \leq t \leq t_0+1} g(t, t_0)$, $m_0 = \min \left\{ m \in \mathbb{N}^* : h(m) \leq \frac{1}{e} \right\}$.

Conditions 1) and 2) imply that $\sup_{0 \leq t_0 \leq t \leq t_0+2m_0} g(t, t_0) \leq a^{2m_0}$.

Fix $t_0 \geq 0, t \geq t_0 + 2m_0, m = \left\lfloor \frac{t}{m_0} \right\rfloor, n = \left\lfloor \frac{t_0}{m_0} \right\rfloor$ where $[s]$ is the largest integer

equal or less than $s \in \mathbf{R}$. One can see that $m_0 m \leq t < m_0(m+1)$, $m_0 n \leq t_0 < m_0(n+1)$, $m \geq n+2$, and so,

$$\begin{aligned} g(t, t_0) &\leq g(t, m_0 m) g(m_0 m, m_0(n+1)) g(m_0(n+1), t_0) \leq \\ &\leq a^{4m_0} \prod_{k=n+2}^m g(m_0 k, m_0(k-1)) \leq a^{4m_0} \prod_{k=n+2}^m h(m_0) \\ &\leq a^{4m_0} e^{-(m-n-1)} \leq a^{4m_0} e^{-\frac{t-t_0}{m_0}+2}. \end{aligned}$$

If we note that

$$g(t, t_0) \leq a^{2m_0} \leq a^{2m_0} e^2 e^{-\frac{t-t_0}{m_0}}, \quad \text{for all } t_0 \geq 0, t \in [t_0, t_0 + 2m_0]$$

we obtain easily that

$$g(t, t_0) \leq N e^{-\nu(t-t_0)}, \quad \text{for all } t \geq t_0 \geq 0, \quad \text{where}$$

$$N = \max\{a^{4m_0} e^2, a^{2m_0} e^2\}, \quad \nu = \frac{1}{m_0}.$$

Theorem 3.1. $\mathcal{U}$ is u.e.s. if and only if there exists a Orlicz sequence space l^Φ admissible to $\mathcal{U}$.

Proof.

Necessity. It follows easily from (Definition 2.3) that the space l^1 is admissible to $\mathcal{U}$.

Sufficiency. Let $m \in \mathbf{N}$, $x \in X$, and $f : \mathbf{N} \rightarrow X$, $f = \chi_{\{m\}} x$. It is easy to verify that $f \in l^\Phi(X)$ and $\|f\|_\Phi = \|\chi_{\{0\}}\|_\Phi \|x\|$ and

$$(x_f)(k) = \sum_{j=0}^k U(k, j) f(j) =$$

$$\begin{cases} U(k, m)x & , \quad k \geq m \\ 0 & , \quad k < m \end{cases}$$

and so $\|U(k, m)x\| \leq \|x_f\|_\infty \leq \frac{1}{\|\chi_{\{0\}}\|_\Phi} \|x_f\|_\Phi \leq K \frac{1}{\|\chi_{\{0\}}\|_\Phi} \|f\|_\Phi = K \|x\|$,
for all $k \geq m$.

Let $n_0, m \in \mathbf{N}, x \in X, f : \mathbf{N} \rightarrow X$, given by

$$f(n) = \begin{cases} U(n, n_0)x & , \quad n \in \{n_0, \dots, n_0 + m\} \\ 0 & , \quad n \notin \{n_0, \dots, n_0 + m\} \end{cases}$$

Then $f \in l^\Phi(X)$, $\|f\|_{l^\Phi(X)} \leq K \frac{1}{\Phi^{-1}(\frac{1}{m+1})} \|x\|$. It follows that

$$(x_f)(n) = \sum_{k=0}^n U(n, k)f(k) = \begin{cases} 0 & , \quad n < n_0 \\ (n - n_0 + 1)U(n, n_0)x & , \quad n \in \{n_0, \dots, n_0 + m\} \\ (m + 1)U(m, n_0)x & , \quad n \geq n_0 + m + 1 \end{cases}$$

and so

$$\begin{aligned} \frac{(m+1)(m+2)}{2} \|U(m+n_0, n_0)x\| &= \sum_{n=n_0}^{n_0+m} (n - n_0 + 1) \|U(m+n_0, n_0)x\| \\ &\leq K \sum_{n=n_0}^{n_0+m} (n - n_0 + 1) \|U(n, n_0)x\| = K \sum_{n=n_0}^{n_0+m} \|x_f(n)\| \leq K a_\Phi(m) \|x_f\|_{l^\Phi(X)} \\ &\leq K a_\Phi(m) \|f\|_{l^\Phi(X)} \leq K^3 \|x\| (m+1). \end{aligned}$$

We obtain that

$$\|U(m+n_0, n_0)\| \leq \frac{2K^3}{m+2} \|x\| \quad , \quad \text{for all } m, n_0 \in \mathbf{N}$$

By Lemma 3.2 it results that there exist two constants $N, \nu > 0$ such that

$$\|U(t, t_0)\| \leq N e^{-\nu(t-t_0)} \quad , \quad \text{for all } t \geq t_0 \geq 0 \quad .$$

3 References

- [1] A. Ben-Artzi, I. Gohberg, Dichotomies of systems and invertibility of linear ordinary differential operators, *Oper. Theory Adv. Appl.*, **56**, 1992, 90-119.
- [2] C. Chicone, Y. Latushkin, Evolution semigroups in dynamical systems and differential equations, Mathematical Surveys and Monographs, vol. 70, Providence, RO: American Mathematical Society, 1999.

- [3] C. V. Coffman, J.J. Schäffer, Dichotomies for linear difference equations, *Math. Annalen* **172**, 1967, 139-166.
- [4] J.L. Daleckij, M.G. Krein, Stability of differential equations in Banach space, Amer. Math. Soc, Providence, R.I. 1974
- [5] D. Henry, Geometric theory of semilinear parabolic equations, Springer Verlag, New-York, 1981
- [6] J. P. La Salle, The stability and control of discrete processes, Springer Verlag, Berlin, 1990.
- [7] Y. Latushkin, S. Montgomery-Smith, Evolutionary semigroups and Lyapunov theorems in Banach spaces, *J. Funct. Anal.*, **127** (1995), 173-197.
- [8] Y. Latushkin, T. Randolph, Dichotomy of differential equations on Banach spaces and an algebra of weighted composition operators, *Integral Equations Operator Theory*, 23 (1995), 472-500.
- [9] J.L. Massera, J.J. Schäffer, Linear Differential Equations and Function Spaces, Academic Press, New York, 1966.
- [10] M. Megan, P. Preda, Admissibility and dichotomies for linear discrete-time systems in Banach spaces, *An. Univ. Tim., Seria St. Mat.*, vol. **26** (1), 1988, 45-54.
- [11] N. van Minh, F. Răbiger, R. Schnaubelt, Exponential stability, exponential expansiveness and exponential dichotomy of evolution equations on the half-line, *Int. Eq. Op. Theory*, **32** (1998), 332-353.
- [12] N. van Minh, On the proof of characterizations of the exponential dichotomy, *Proc. Amer. Math. Soc.*, **127** (1999), 779-782.
- [13] O. Perron, Die stabilitätsfrage bei differentialgleichungen, *Math. Z.*, **32** (1930), 703-728.
- [14] M. Pinto, Discrete dichotomies, *Computers Math. Applic.*, vol. **28**, 1994, 259-270.

- [15] P.Preda, On a Perron condition for evolutionary processes in Banach spaces, *Bull. Math. de la Soc. Sci. Math. de la R.S. Roumanie*, T **32** (80), no.1, 1988, 65-70.
- [16] K.M. Przyluski, S. Rolewicz, On stability of linear time-varying infinite-dimensional discrete-time systems, *Systems Control Lett.* **4**, 1994, 307-315.
- [17] R. Schnaubelt, Sufficient conditions for exponential stability and dichotomy of evolution equations, *Forum Mat.*, **11** (1999), 543-566.
- [18] L.Ta, Die stabilitätsfrage bei differenzengleichungen, *Acta Math.* **63**, 1934, 99-141.
- [19] A. C. Zaanen, *Integration*, North-Holland, Amsterdam, 1967.

Generalized Quasilinearization Technique for the Second Order Differential Equations with General Mixed Boundary Conditions

Bashir Ahmad^a and Rahmat A. Khan^b

^aDepartment of Mathematics, Faculty of Science, King Abdulaziz University,
P.O.Box 80203, Jeddah-21589, Saudi Arabia

E-mail: *bashir_qau@yahoo.com*

^bDepartment of Mathematics, Quaid-i-Azam University, Islamabad, Pakistan

E-mail: *rahmat_alipk@yahoo.com*

Abstract

The generalized quasilinearization method for a second order nonlinear differential equation with general type of boundary conditions is developed. A monotone sequence of approximate solutions converging uniformly and rapidly to a solution of the problem is obtained. Some special cases have also been discussed.

Keywords and Phrases: Quasilinearization, mixed boundary conditions, rapid convergence.

AMS Subject Classifications (2000): 34A45, 34B15.

1. Introduction

The method of generalized quasilinearization introduced by Lakshmikantham [4,5] has been effectively applied to find a sequence of approximate solutions of the nonlinear initial and boundary value problems, see, for example, [6-9]. In reference [2], the generalized quasilinearization method was discussed for a nonlinear Robin problem.

In this paper, we develop a generalized quasilinearization technique for a second order nonlinear differential equation subject to general mixed boundary conditions and obtain a sequence of approximate solutions converging monotonically and rapidly to a solution of the problem. Some interesting observations have also been presented.

2. Some Basic Results

Consider the nonlinear problem

$$-u''(t) = f(t, u), \quad t \in J = [0, \pi], \quad (1)$$

$$p_o u(0) - q_o u'(0) = c_1, \quad p_o u(\pi) + q_o u'(\pi) = c_2,$$

where $f \in C[J \times R, R]$, $p_o, q_o, c_1, c_2 \in R$ with $p_o, q_o > 0$ ($2p_o > q_o$) and $c_1, c_2 \geq 0$.

We know that the linear problem

$$-u''(t) = \lambda u(t), \quad t \in J,$$

$$p_o u(0) - q_o u'(0) = 0, \quad p_o u(\pi) + q_o u'(\pi) = 0,$$

has a trivial solution for all real values of λ except for those values of λ which are the roots of the equation

$$\tan \sqrt{\lambda} \pi = \frac{2\sqrt{\lambda} p_o q_o}{(q_o^2 \lambda - p_o^2)}.$$

Consequently, for all such values of λ for which the homogeneous linear problem has a trivial solution, the corresponding non homogeneous problem

$$-u''(t) - \lambda u(t) = \sigma(t), \quad t \in J,$$

$$p_o u(0) - q_o u'(0) = c_1, \quad p_o u(\pi) + q_o u'(\pi) = c_2,$$

has a unique solution

$$u(t) = \int_0^\pi G_\lambda(t, s) \sigma(s) ds + c_2 \frac{d}{ds} G_\lambda(t, s)|_{s=\pi} - c_1 \frac{d}{ds} G_\lambda(t, s)|_{s=0},$$

where $G_\lambda(t, s)$ is the Green's function of the associated homogeneous problem and is given by

$$\begin{aligned} G_\lambda(t, s) &= \frac{1}{2p_o q_o \lambda \cos \sqrt{\lambda} \pi + (p_o^2 \sqrt{\lambda} - q_o^2 \lambda^{\frac{3}{2}}) \sin \sqrt{\lambda} \pi} \\ &\times \begin{cases} \frac{(p_o \sin \sqrt{\lambda} t + q_o \sqrt{\lambda} \cos \sqrt{\lambda} t)}{[q_o \sqrt{\lambda} \cos \sqrt{\lambda}(\pi - s) + p_o \sin \sqrt{\lambda}(\pi - s)]^{-1}}, & \text{if } 0 \leq t \leq s \leq \pi, \\ \frac{(p_o \sin \sqrt{\lambda} s + q_o \sqrt{\lambda} \cos \sqrt{\lambda} s)}{[q_o \sqrt{\lambda} \cos \sqrt{\lambda}(\pi - t) + p_o \sin \sqrt{\lambda}(\pi - t)]^{-1}}, & \text{if } 0 \leq s \leq t \leq \pi \quad (\lambda > 0), \end{cases} \\ G_\lambda(t, s) &= \frac{1}{p_o(2q_o + p_o \pi)} \begin{cases} (q_o + p_o t)(q_o + p_o(\pi - s)), & \text{if } 0 \leq t \leq s \leq \pi, \\ (q_o + p_o s)(q_o + p_o(\pi - t)), & \text{if } 0 \leq s \leq t \leq \pi \quad (\lambda = 0), \end{cases} \\ G_\lambda(t, s) &= \frac{-1}{2p_o q_o (-\lambda) \cosh \sqrt{-\lambda} \pi + (p_o^2 \sqrt{-\lambda} + q_o^2 (-\lambda)^{\frac{3}{2}}) \sinh \sqrt{-\lambda} \pi} \\ &\times \begin{cases} \frac{(p_o \sinh \sqrt{-\lambda} t + q_o \sqrt{-\lambda} \cosh \sqrt{-\lambda} t)}{[q_o \sqrt{-\lambda} \cosh \sqrt{-\lambda}(\pi - s) + p_o \sinh \sqrt{-\lambda}(\pi - s)]^{-1}}, & \text{if } 0 \leq t \leq s \leq \pi, \\ \frac{(p_o \sinh \sqrt{-\lambda} s + q_o \sqrt{-\lambda} \cosh \sqrt{-\lambda} s)}{[q_o \sqrt{-\lambda} \cosh \sqrt{-\lambda}(\pi - t) + p_o \sinh \sqrt{-\lambda}(\pi - t)]^{-1}}, & \text{if } 0 \leq s \leq t \leq \pi \quad (\lambda < 0). \end{cases} \end{aligned}$$

We observe that $G_\lambda \geq 0$ for $\lambda < 1/4$. Now, we state the following lemma (for proof, see [2]).

Lemma 2.1. (Comparison Result) Let $\lambda < 1/4$, $\sigma(t) \geq 0$ and $u \in C^2(J)$. Then the nonhomogeneous problem

$$-u''(t) - \lambda u(t) = \sigma(t), \quad t \in J,$$

$$p_o u(0) - q_o u'(0) = c_1, \quad p_o u(\pi) + q_o u'(\pi) = c_2,$$

has a nonnegative solution.

We shall say that $\alpha(t) \in C^2(J)$ is a lower solution of (1) if

$$-\alpha''(t) \leq f(t, \alpha(t)), \quad t \in J,$$

$$p_o \alpha(0) - q_o \alpha'(0) \leq c_1, \quad p_o \alpha(\pi) + q_o \alpha'(\pi) \leq c_2.$$

Similarly, $\beta \in C^2(J)$ is an upper solution of (1) if

$$-\beta''(t) \geq f(t, \beta(t)), \quad t \in J,$$

$$p_o \beta(0) - q_o \beta'(0) \geq c_1, \quad p_o \beta(\pi) + q_o \beta'(\pi) \geq c_2.$$

Now, we state some theorems (without proof) which can be proved using the standard arguments [2, 9].

Theorem 2.2. Let $f \in C^2[J \times R, R]$ be nonincreasing in u for each $t \in J$ and α, β are lower and upper solutions of (1) respectively. Further, there exists a constant $L \geq 0$ such that

$$f(t, u_1) - f(t, u_2) \leq L(u_1 - u_2),$$

whenever $u_1 \geq u_2$. Then $\alpha(t) \leq \beta(t)$ on J .

Theorem 2.3. Let α, β be lower and upper solutions of (1) respectively such that $\alpha(t) \leq \beta(t)$ on J and $f \in C^2[J \times R, R]$. Then there exists a solution $u(t)$ of (1) such that $\alpha(t) \leq u(t) \leq \beta(t)$, $t \in J$.

3. Generalized Quasilinearization Technique

Theorem 3.1. Assume that

(A₁) α and $\beta \in C^2(J)$ are lower and upper solutions of (1) respectively such that $\alpha(t) \leq \beta(t)$ on J .

(A₂) $\frac{\partial f}{\partial u}(t, u)$, $\frac{\partial^2 f}{\partial u^2}(t, u)$ exist and are continuous such that $\frac{\partial f}{\partial u}(t, u) < 1/4$ and $\frac{\partial^2}{\partial u^2}(f(t, u) + \phi(t, u)) \geq 0$, where $\phi \in C[J \times R, R]$ such that $\frac{\partial^2 \phi}{\partial u^2}(t, u) \geq 0$ for every $(t, u) \in J \times R$.

Then there exists a monotone sequence $\{w_n\}$ of solutions converging uniformly and quadratically to the unique solution of the problem.

Proof. We define

$$F(t, u) = f(t, u) + \phi(t, u), \quad t \in J.$$

Clearly $F(t, u) \in C[J \times R, R]$. In view of (A_2) and the generalized mean value theorem, we have

$$F(t, u) \geq F(t, v) + F_u(t, v)(u - v) - \phi(t, u),$$

where $\alpha \leq v \leq u \leq \beta$ on J ($u, v \in R$). Define

$$g(t, u, v) = F(t, v) + F_u(t, v)(u - v) - \phi(t, u), \quad (2)$$

and observe that

$$g(t, u, v) \leq F(t, u), \quad g(t, u, u) = F(t, u). \quad (3)$$

Also, it can easily be seen that $g(t, u, v)$ is nonincreasing in u for each fixed (t, v) . Further, using the mean value theorem repeatedly, we have

$$g(t, u, v_1) - g(t, u, v_2) \geq \frac{\partial^2}{\partial u^2} F(t, \xi_1)(\xi - v_2)(v_1 - v_2),$$

($v_2 \leq \xi_1 \leq \xi$, $v_2 \leq \xi \leq v_1$, $t \in J$), which in view of (A_2) implies that $g(t, u, v)$ is nondecreasing in v for each fixed (t, u) . Also, $g(t, u, v)$ satisfies Lipschitz condition:

$$g(t, u, v_1) - g(t, u, v_2) = (F_u(t, v) - \phi_u(t, \eta))(u_1 - u_2) \leq f_u(t, u)(u_1 - u_2) \leq L(u_1 - u_2),$$

where $L > 0$, and $u_1 \geq u_2$.

Now, we set $w_o = \alpha$ and consider

$$-u''(t) = g(t, u(t), w_o(t)), \quad t \in J, \quad (4)$$

$$p_o u(0) - q_o u'(0) = c_1, \quad p_o u(\pi) + q_o u'(\pi) = c_2.$$

Using (A_1) , (3) and the nondecreasing nature of $g(t, u, v)$ in v , we get

$$-w_o''(t) \leq f(t, w_o(t)) = g(t, w_o(t), w_o(t)), \quad t \in J,$$

$$p_o w_o(0) - q_o w_o'(0) \leq c_1, \quad p_o w_o(\pi) + q_o w_o'(\pi) \leq c_2,$$

and

$$-\beta''(t) \geq f(t, \beta(t)) \geq g(t, w_o, \beta), \quad t \in J$$

$$p_o \beta(0) - q_o \beta'(0) \geq c_1, \quad p_o \beta(\pi) + q_o \beta'(\pi) \geq c_2,$$

which imply that w_o and β are lower and upper solution of (4) respectively. Hence, by Theorem 2.3, there exists a solution w_1 of (4) such that $w_o \leq w_1 \leq \beta$ on J .

Now, consider the problem

$$-u''(t) = g(t, u(t), w_1(t)), \quad t \in J, \quad (5)$$

$$p_o u(0) - q_o u'(0) = c_1, \quad p_o u(\pi) + q_o u'(\pi) = c_2.$$

From

$$\begin{aligned} -w_1''(t) &= g(t, w_1(t), w_o(t)) \leq g(t, w_1, w_1), \quad t \in J, \\ p_o w_1(0) - q_o w_1'(0) &= c_1, \quad p_o w_1(\pi) + q_o w_1'(\pi) = c_2, \end{aligned}$$

and

$$\begin{aligned} -\beta''(t) &\geq f(t, \beta(t)) \geq g(t, w_o, \beta), \quad t \in J, \\ p_o \beta(0) - q_o \beta'(0) &\geq c_1, \quad p_o \beta(\pi) + q_o \beta'(\pi) \geq c_2, \end{aligned}$$

it follows that w_1 and β are lower and upper solutions of (5) respectively. Again, by Theorem 2.3, there exists a solution w_2 of (5) such that $w_1 \leq w_2 \leq \beta$ on J . Continuing this process successively, we obtain a monotone sequence $\{w_n\}$ of solutions on J satisfying

$$w_o \leq w_1 \leq w_2 \leq w_3 \leq \dots \leq w_{n-1} \leq w_n \leq \beta,$$

where the element w_n of the sequence is the solution of the problem

$$\begin{aligned} -u''(t) &= g(t, u(t), w_{n-1}(t)), \quad t \in J, \\ p_o u(0) - q_o u'(0) &= c_1, \quad p_o u(\pi) + q_o u'(\pi) = c_2. \end{aligned}$$

Since the sequence $\{w_n\}$ is monotone, it has a pointwise limit w . To show that w is in fact a solution of (1), we note that

$$w_n(t) = \int_0^\pi G_0(t, s) \sigma_n(s) ds + c_2 \frac{d}{ds} G_o(t, s)|_{s=\pi} - c_1 \frac{d}{ds} G_o(t, s)|_{s=0},$$

($G_0(t, s)$ is given in section 2) is a solution of the following problem

$$\begin{aligned} -u''(t) &= \sigma_n(t), \quad t \in J, \\ p_o u(0) - q_o u'(0) &= c_1, \quad p_o u(\pi) + q_o u'(\pi) = c_2, \end{aligned}$$

where

$$\sigma_n(t) = F(t, w_{n-1}) + F_u(t, w_{n-1})(w_n - w_{n-1}) - \phi(t, w_n).$$

Since σ_n is continuous for each $n \in N$ and $\alpha \leq w_n \leq \beta$ ($n = 1, 2, 3, \dots$), it follows that $\{\sigma_n\}$ is bounded in $C(J)$. Also, $\lim_{n \rightarrow \infty} \sigma_n(t) = f(t, w)$. Thus, the sequence $\{w_n\}$ is bounded in $C^2(J)$ and hence $\{w_n\} \nearrow w$ uniformly on J . Passing onto the limit $n \rightarrow \infty$, we get

$$w = \int_0^\pi G_0(t, s) f(s, w) ds + (c_2 + c_1) \frac{q_o}{(2q_o + p_o \pi)}.$$

This shows that w is a solution of (1).

Now we show that the convergence is quadratic. For that, we set $e_n = u - w_n$, $n =$

1, 2, 3... and note that $e_n \geq 0$, $p_o e_n(0) - q_o e'_n(0) = 0$ and $p_o e_n(\pi) + q_o e'_n(\pi) = 0$. Using the definition of $g(t, u, v)$ and the generalized mean value theorem, we get

$$\begin{aligned} -e''_n(t) &= -u''(t) + w''_n(t) \\ &= \left[\frac{\partial F}{\partial u}(t, \xi_1) - \frac{\partial F}{\partial u}(t, u_{n-1}) \right] (u - u_{n-1}) \\ &\quad - \phi_u(t, \xi_2) (u - u_n) + \frac{\partial F}{\partial u}(t, u_{n-1}) (u - u_n) \\ &= \frac{\partial^2 F}{\partial u^2}(t, \xi_3) (\xi_1 - u_{n-1}) (u - u_{n-1}) + \left[\frac{\partial F}{\partial u}(t, u_{n-1}) - \phi_u(t, \xi_2) \right] (u - u_n) \\ &\leq C e_{n-1}^2 + \frac{\partial f}{\partial u}(t, u_{n-1}) e_n \leq C e_{n-1}^2 + a_n(t) e_n, \end{aligned}$$

where $w_{n-1}(x) \leq \xi_3 \leq \xi_1 \leq u$, $u_n \leq \xi_2 \leq u$ and $a_n = \frac{\partial f}{\partial u}(t, u_{n-1})$. Since $\lim_{n \rightarrow \infty} a_n(t) = f_u(t, w) < 1/4$, therefore we can choose $\lambda < 1/4$ and $n_o \in N$ such that for $n \geq n_o$, $a_n(t) < \lambda$, $t \in J$. Thus, for some $b_n(t) \leq 0$, the error function $e_n(t)$ satisfies the problem

$$\begin{aligned} -e''_n(t) - \lambda(t) e_n(t) &= (a_n(t) - \lambda) e_n(t) + c e_{n-1}^2 + b_n(t), \quad t \in J, \\ p_o e_n(0) - q_o e'_n(0) &= 0, \quad p_o e_n(\pi) + q_o e'_n(\pi) = 0, \end{aligned}$$

whose solution is

$$e_n(t) = \int_0^\pi G_\lambda(t, s) [(a_n(s) - \lambda) e_n(s) + c e_{n-1}^2 + b_n(s)] ds, \quad t \in J.$$

Since $a_n(t) - \lambda \leq 0$, $e_n = w - w_n \geq 0$, $b_n \leq 0$, therefore

$$(a_n(t) - \lambda) e_n(t) + b_n(t) + c e_{n-1}^2 \leq c e_{n-1}^2.$$

Consequently, we have

$$e_n(t) \leq c \int_0^\pi G_\lambda(t, s) e_{n-1}^2(s) ds, \quad n \geq n_o,$$

which, on taking the maximum, can be written as

$$\|e_n\| \leq \delta \|e_{n-1}\|^2,$$

where δ provides a bound on $c \int_0^\pi G_\lambda(t, s) ds$ and

$$\|u\| = \max\{|u(t)|, t \in J\}.$$

Theorem 3.2. (Rapid Convergence) Assume that

(B₁) $\alpha, \beta \in C^2(J)$ are lower and upper solutions of (1) respectively such that $\alpha \leq \beta$ on J .

(B₂) $\frac{\partial^i f}{\partial x^i}(t, u)$, $i = 1, 2, 3, \dots, k$ exist and are continuous on $\Omega = \{(t, u) \in J \times R\}$ such that $\frac{\partial f}{\partial u}(t, u) < \frac{1}{4}$, $\frac{\partial^k}{\partial u^k}(f(t, u) + \phi(t, u)) \geq 0$ for some function $\phi \in C^{0, (k-1)}[J \times R, R]$ such that $\frac{\partial^k \phi}{\partial u^k}(t, u) \geq 0$.

Then there exists a monotone sequence $\{w_n\}$ of solutions converging uniformly and rapidly with the order of convergence k ($k \geq 2$) to a solution of (1).

Proof. Define

$$F(t, u) = f(t, u) + \phi(t, u), \quad t \in J.$$

Using (B₂) and generalized mean value theorem, we have

$$f(t, u) \geq \sum_{i=0}^{k-1} \frac{\partial^i F}{\partial u^i}(t, v) \frac{(u-v)^i}{(i)!} - \phi(t, u).$$

Now, we set

$$\begin{aligned} K(t, u, v) &= \sum_{i=0}^{k-1} \frac{\partial^i F}{\partial u^i}(t, v) \frac{(u-v)^i}{(i)!} - \phi(t, u) \\ &= \sum_{i=0}^{k-1} \frac{\partial^i f}{\partial u^i}(t, v) \frac{(u-v)^i}{(i)!} - \frac{\partial^k \phi}{\partial u^k}(t, \xi) \frac{(u-v)^k}{(k)!}, \end{aligned}$$

where $v \leq \xi \leq u$ and $\alpha \leq v \leq u \leq \beta$ on J .

Observe that

$$K(t, u, v) \leq F(t, u), \quad K(t, u, u) = F(t, u). \quad (6)$$

Fix $w_o = \alpha$ and consider the problem

$$-u''(t) = K(t, u(t), w_o(t)), \quad t \in J, \quad (7)$$

$$p_o u(0) - q_o u'(0) = c_1, \quad p_o u(\pi) + q_o u'(\pi) = c_2.$$

Using (B₁) and (6), we get

$$-w_o''(t) \leq f(t, w_o(t)) = K(t, w_o(t), w_o(t)), \quad t \in J,$$

$$p_o w_o(0) - q_o w_o'(0) \leq c_1, \quad p_o w_o(\pi) + q_o w_o'(\pi) \leq c_2,$$

and

$$-\beta''(t) \geq f(t, \beta(t)) \geq K(t, w_o, \beta), \quad t \in J,$$

$$p_o \beta(0) - q_o \beta'(0) \geq c_1, \quad p_o \beta(\pi) + q_o \beta'(\pi) \geq c_2,$$

which imply that w_o and β are lower and upper solution of (7) respectively. Hence, by Theorem 2.3, there exists a solution w_1 of (7) such that $w_o \leq w_1 \leq \beta$ on J . Similarly, we can show that the following problem

$$-u''(t) = K(t, u(t), w_1(t)), \quad t \in J, \quad (8)$$

$$p_o u(0) - q_o u'(0) = c_1, \quad p_o u(\pi) + q_o u'(\pi) = c_2,$$

has a solution w_2 such that $w_1 \leq w_2 \leq \beta$ on J , where w_1 and β are lower and upper solution of (8) respectively.

Continuing this process successively, we obtain a monotone sequence $\{w_n\}$ of solutions satisfying

$$w_o \leq w_1 \leq w_2 \leq w_3 \leq \dots \leq w_{n-1} \leq w_n \leq \beta$$

on J , where the element w_n of the sequence is the solution of the problem

$$-u''(t) = K(t, u(t), w_{n-1}(t)), \quad t \in J,$$

$$p_o u(0) - q_o u'(0) = c_1, \quad p_o u(\pi) + q_o u'(\pi) = c_2.$$

Employing the procedure used in the preceding theorem, we can show that the sequence $\{w_n\}$ converges to the unique solution w of the boundary value problem (1).

To show that the convergence of the sequence is of order k ($k \geq 2$), we set $e_n = w - w_n$, $a_n = w_{n+1} - w_n$, $n = 1, 2, 3, \dots$. Clearly, $a_n \geq 0$, $e_n \geq 0$, $e_n - a_n = e_{n+1}$, $a_n \leq e_n$ and $a_n^k \leq e_n^k$. Using the mean value theorem repeatedly, we have

$$\begin{aligned} -e_n''(t) &= w_n''(t) - w''(t) \\ &= \sum_{i=0}^{k-1} \frac{\partial^i f}{\partial u^i}(t, w_{n-1}) \frac{(e_{n-1}^i - a_{n-1}^i)}{(i)!} \\ &\quad + \frac{\partial^k f}{\partial u^k}(t, \zeta(t)) \frac{e_{n-1}^k}{k!} + \frac{\partial^k \phi}{\partial u^k}(t, \zeta(t)) \frac{a_{n-1}^k}{k!} \\ &\leq \left(\sum_{i=0}^{k-1} \frac{\partial^i f}{\partial u^i}(t, w_{n-1}) \frac{1}{(i)!} \sum_{j=0}^{i-1} e_{n-1}^{j-1} a_{n-1}^{i-j-1} \right) e_n \\ &\quad + \left[\frac{\partial^k f}{\partial u^k}(t, \zeta(t)) + \frac{\partial^k \phi}{\partial u^k}(t, \zeta(t)) \right] \frac{e_{n-1}^k}{k!} \\ &\leq q_n(t) e_n + N e_{n-1}^k, \end{aligned}$$

where

$$q_n(t) = \sum_{i=1}^{k-1} \frac{\partial^i f}{\partial u^i}(t, w_{n-1}) \frac{1}{(i)!} \sum_{j=0}^{i-1} e_{n-1}^{i-1-j} a_{n-1}^j,$$

and $N > 0$ provides bound for $\frac{\partial^k f}{\partial u^k}(t, w_{n-1}(t))$ on Ω . As $\lim_{n \rightarrow \infty} q_n(t) = f_u(t, w) < 1/4$, we can choose $\lambda < 1/4$ and $n_o \in \mathbb{N}$ such that for $n \geq n_o$, $q_n(t) < \lambda$, we have

$$-e_n''(t) - \lambda(t) e_n(t) \leq (q_n(t) - \lambda) e_n(t) + N e_{n-1}^k \leq N e_{n-1}^k,$$

$$p_o e_n(0) - q_o e_n'(0) = 0, \quad p_o e_n(\pi) + q_o e_n'(\pi) = 0,$$

whose solution is

$$e_n(t) = \int_0^\pi G_\lambda(t, s) N e_{n-1}^k ds, \quad t \in J.$$

Taking maximum over $[0, \pi]$, we obtain

$$\|e_n\| \leq C \|e_{n-1}\|^k,$$

where C provides a bound on $N \int_0^\pi G_\lambda(t, s) ds$.

4. Concluding Remarks

In this paper, we have developed a generalized quasilinearization method with rapid convergence for a somewhat more general type of boundary value problem. Some interesting cases can be obtained by assigning particular values to the parameters involved in the boundary conditions and these are given below:

- (i) The results of references [10] and [3] can be recorded as a special case by taking $p_o = 1$, $q_o = 0$, $c_1 = 0$, $c_2 = 0$ for $\phi = Mu^2$ and $\phi = Mu^k$ respectively.
- (ii) For $p_o = 0$, $q_o = 1$, $c_1 = 0$, $c_2 = 0$, we recover the results presented in reference [1].
- (iii) The results of reference [2] appear as a special case by taking $p_o = 1$, $q_o = 1$, $c_1 = 0$, $c_2 = 0$.

References

- [1] Bashir Ahmad, J.J. Nieto and N. Shahzad, The Bellman-Kalaba-Lakshmikantham quasilinearization method for Neumann problems, *J. Math. Anal. Appl.*, 257(2001) 356-363.
- [2] Bashir Ahmad, Rahmat A. Khan and S. Sivasundaram, Generalized quasilinearization method for nonlinear boundary value problems, *Dynamic Systems Appl.*, 11(2002) 359-370.
- [3] Albert Cabada, J.J. Nieto and Rafael Pita-da-veige, A note on rapid convergence of approximate solutions for an ordinary Dirichlet problem, *Dyna. Contin. Discrete Impuls. Syst.*, 4(1998) 23-30.
- [4] V. Lakshmikantham, An extension of the method of quasilinearization, *J. Optim. Theory Appl.*, 82(1994) 315-321.
- [5] V. Lakshmikantham, Further improvement of generalized quasilinearization, *Nonlinear Anal.*, 27(1996) 315-321.
- [6] V. Lakshmikantham, S. Leela and F.A. McRae, Improved generalized quasilinearization method, *Nonlinear Anal.*, 24(1995) 1627-1637.
- [7] V. Lakshmikantham and N. Shahzad, Further generalization of generalized quasilinearization method, *J. Appl. Math. Stoch. Anal.*, 7(1994), 545-552.
- [8] V. Lakshmikantham, N. Shahzad and J.J. Nieto, Method of generalized quasilinearization for periodic boundary value problems, *Nonlinear Anal.*, 27(1996) 143-151.

- [9] V. Lakshmikantham and A.S. Vatsala, Generalized Quasilinearization for Non-linear Problems, Kluwer Academic Publishes, Dordrecht, 1998.
- [10] J.J. Nieto, Generalized quasilinearization method for a second order ordinary differential equation with Dirichlet boundary conditions, Proc. Amer. Math. Soc., 125(1997) 2599-2604.

Existence and Construction of Finite Tight Frames

Peter G. Casazza, Manuel T. Leon*

Department of Mathematics

University of Missouri

Columbia, MO 652511

pete@math.missouri.edu

February 13, 2006

Abstract

The space of finite tight frames of M vectors in $\mathbb{R}^N$ with prescribed norms $\{b_j\}_{j=1}^M$ and frame constant A corresponds to the first N columns of matrices in $\mathbf{O}(M)$ (the orthogonal group) with the property that the norms of the first N elements of their rows equal the values $a_j = b_j/\sqrt{A}$. Then $a_j^2 \leq 1$ and $\sum_{j=1}^{j=M} a_j^2 = N$.

For any sequence $\{a_j\}_{j=1}^M$ of positive real numbers such that $a_j^2 \leq 1$ and $\sum_{j=1}^{j=M} a_j^2 = N$, correspond to the first N rows of matrices in $\mathbf{O}(M)$ (the orthogonal group) with the property that the norms of the first N elements of their rows equal the values a_j . we prove existence of such frames. The proof is constructive, giving an embedding of $\mathbf{O}(N) \times \mathbf{O}(M)$ into the space of such frames. In addition, for any such $\{a_j\}_{j=1}^M$, all solutions can be factorised into $\frac{M*(M-1)}{2} - \frac{N*(N-1)}{2}$ parameters and one matrix $R_N \in \mathbf{O}(N)$, and each different solution provides a different embedding of $\mathbf{O}(N)$ into the space of such frames. Replacing R_N by any other element of $\mathbf{O}(M)$, $\frac{N*(N-1)}{2}$ different solutions are obtained. A simple and fast algorithm providing particular solutions is given. In particular, Parseval frames correspond to sequences $\{a_j\}_{j=1}^M$ such that $a_j = \sqrt{\frac{N}{M}}$. The results provide an independent proof of some of the existence results in [CKLT02]. All proofs are constructive. A MATLAB toolbox implementing all results is freely distributed by the authors.

*P.G. Casazza and M. T. Leon were supported by NSF DMS 0102686

1 Introduction

Frames for Hilbert spaces play a fundamental role in a variety of important areas including signal/image processing [Cas00], multiple antenna coding [GKV98], perfect reconstruction filter banks [DGK02], quantum theory [EF01], [Š] and much more. For applications, tight frames are preferred since they allow simple reconstruction formulas. A common problem in applications of frames is to find a tight frame $\{\varphi_j\}_{j=1}^M$ for an N -dimensional Hilbert space H_N with $\{\|\varphi_j\|\}_{j=1}^M$ prescribed in advance. In [CKLT02] the authors give necessary and sufficient conditions on $\{a_j\}_{j=1}^M$ so that there exists a tight frame $\{\varphi_j\}_{j=1}^M$ for H_N with $\|\varphi_j\| = a_j$, for all $j = 1, \dots, M$. The main tool used in [CKLT02] is the notion of *frame potentials* introduced by Benedetto and Fickus [BF02]. However, this is an existence proof while for applications we need an exact representation for the frame. In this paper we will give an algorithm for constructing a family of tight frame vectors $\{\varphi_j\}_{j=1}^M$ for H_N with $\|\varphi_j\| = a_j$, for all $j = 1, \dots, M$. This provides an independent proof of some of the results in [CKLT02] while at the same time allowing exact constructions for the necessary frames. A MATLAB toolbox implementing all these results is freely distributed by the authors.

2 Definitions

Throughout the paper we let $N < M$ be fixed positive integers. First we recall the well known fact that finite normalized tight frames correspond to the first N columns of matrices in $O(M)$, the $M \times M$ orthonormal matrices. Background definitions and a detailed proof of this fact is found in [GK01] and [BF02].

For the purposes of this paper it suffices to recall that $\{\varphi_j\}_{j=1}^M$ is a (finite) tight frame in $\mathbb{R}^N$ with frame constant A if for every vector $y \in \mathbb{R}^N$

$$Sy = \sum_{j=1}^M \langle y, \varphi_j \rangle \varphi_j = Ay$$

The "analysis frame operator" L of the frame is given by

$$Ly = \{\langle y, \varphi_j \rangle\}_{j=1}^M$$

with adjoint operator L^* given by

$$\{a_n\}_{n=1}^M \longrightarrow \sum_{j=1}^M a_j \varphi_j.$$

We have the diagram

$$\mathbb{R}^N \xrightarrow{L} l^2(M) \xrightarrow{L^*} \mathbb{R}^M \xrightarrow{L} l^2(M)$$

Then the frame operator S and grammian operator G equal $S = L^*L$ and $G = LL^*$, where L is an $M \times N$ matrix and L^* is an $N \times M$ matrix given by

$$L = \begin{pmatrix} \varphi_1^* \\ \vdots \\ \varphi_M^* \end{pmatrix}, \quad L^* = (\varphi_1 | \dots | \varphi_M)$$

where φ_j^* are row vectors and φ_j are column vectors. The tight frame condition gives $L^*L = AI_{N \times N}$ and the diagonal of the grammian is $\{\|\varphi_1\|^2, \dots, \|\varphi_M\|^2\}$.

Now extend $\frac{L}{\sqrt{A}}$ to a matrix $U \in \mathbf{O}(M)$,

$$U = \begin{pmatrix} \frac{\varphi_1}{\sqrt{A}} & | & \cdot \\ \vdots & | & \vdots \\ \frac{\varphi_M}{\sqrt{A}} & | & \cdot \end{pmatrix}$$

In [CKLT02] it is proved that for $\{a_j \geq 0\}_{j=1}^{j=M}$ the condition $\sum_{j=1}^{j=M} a_j^2 \geq Na_1^2$ (suppose a_1 is the largest a_j) is sufficient for the existence of a tight frame $\{\varphi_j\}_{j=1}^{j=M}$ with $\|\varphi_j\| = a_j$. The associated matrix $U \in \mathbf{O}(M)$ satisfies the condition

$$N = \sum_{j=1}^{j=M} \left\| \frac{\varphi_j}{\sqrt{A}} \right\|^2, \quad \text{and} \quad \left\| \frac{\phi_j}{\sqrt{A}} \right\|^2 \leq 1,$$

where A is the tight frame bound of $\{\varphi_j\}_{j=1}^{j=M}$.

The main result in this paper is then:

Theorem 2.1. *For a sequence $\{a_j\}_{j=1}^M$ of positive real numbers such that $a_j^2 \leq 1$ and $\sum_{j=1}^{j=M} a_j^2 = N$ a matrix $U \in \mathbf{O}(M)$ exists such that, in squared norms, the first N columns are*

$$\begin{pmatrix} a_1^2 & | & \cdot \\ \vdots & | & \vdots \\ a_M^2 & | & \cdot \end{pmatrix}.$$

We proceed to obtain such matrices. Their construction provides an independent proof of some of the results on existence of tight frames in [CKLT02].

Any sequence $\{a_j\}_{j=1}^M$ of positive real numbers such that $a_j^2 \leq 1$ and $\sum_{j=1}^{j=M} a_j^2 = N$ will be called *admissible*. For $R_M \in \mathbf{O}(M)$ let $r_{left,j}$ denote the first N terms of its j th row. Then $\sum_{j=1}^{j=M} \|r_{left,j}\|^2 = N$. We will say that R_M is a *solution* for the *admissible* sequence $\{a_j\}_{j=1}^M$ if $\|r_{left,j}\| = a_j$ for $j = 1, \dots, M$.

3 Factorisation of $\mathbf{O}(M)$

The following factorisation is similar, with minor changes, to that described in [Vai93], p. 747. Every matrix in $\mathbf{O}(M)$ is obtained as a product of Givens rotations $\theta(t, j, k) \in \mathbf{O}(M)$, $j < k$, where

$$\theta(t, j, k) = \begin{pmatrix} I_{j-1, j-1} & 0 & 0 & 0 & 0 \\ 0 & \cos(t) & 0 & \sin(t) & 0 \\ 0 & 0 & I_{M-j-k-2, M-j-k-2} & 0 & 0 \\ 0 & -\sin(t) & 0 & \cos(t) & 0 \\ 0 & 0 & 0 & 0 & I_{k-1, k-1} \end{pmatrix}$$

It is clear that

$$\theta(t, j, k)^{-1} = \theta(-t, j, k)$$

To begin with, every $R_2 \in \mathbf{O}(2)$ is obtained as

$$R_2 = \begin{pmatrix} \cos(t) & \sin(t) \\ -\sin(t) & \cos(t) \end{pmatrix} \begin{pmatrix} \mu & 0 \\ 0 & 1 \end{pmatrix}, \mu = \pm 1$$

and thus can be stored with one parameter t and a sign ± 1 (taking account of the determinant).

Let $R_M \in \mathbf{O}(M)$, and we will see how to factor R_M . Let

$$t_M^{M-1} = \pi/2 \text{ if } R_M(M, M) = 0,$$

and

$$t_M^{M-1} = \arg(R_M(M, M) - iR_M(M-1, M)) \text{ otherwise.} \quad (1)$$

Then, if $T = \theta(t_M^{M-1}, M-1, M)R_M$, we have that $T(M-1, M) = 0$ and $T(M, M) \geq 0$. That is,

$$T = \left(\begin{array}{c|c} \dots & T(1, M) \\ \dots & \dots \\ \dots & T(M-2, M) \\ \dots & 0 \\ \dots & T(M, M) \end{array} \right)$$

Now repeat the process on the $M-2$ term of M th column of T , thus obtaining t_M^{M-2} . The matrix

$$S = \theta(t_M^{M-2}, M-2, M)T = \theta(t_M^{M-2}, M-2, M)\theta(t_M^{M-1}, M-1, M)R_M$$

is such that $S(M-2, M) = S(M-1, M) = 0$ and $S(M, M) \geq 0$.

$$S = \left(\begin{array}{c|c} \dots & S(1, M) \\ \dots & \dots \\ \dots & S(M-3, M) \\ \dots & 0 \\ \dots & 0 \\ \dots & S(M, M) \end{array} \right)$$

Iterating the process we obtain $t_M^1, \dots, t_M^{M-1}$ such that if

$$Q_M = \theta(t_M^1, 1, M)\theta(t_M^2, 2, M) \dots \theta(t_M^{M-2}, M-2, M)\theta(t_M^{M-1}, M-1, M)R_M$$

then

$$Q_M(1, M) = Q_M(2, M) = \dots = Q_M(M-1, M) = 0$$

and $Q_M(M, M) \geq 0$. Since $Q_M \in \mathbf{O}(M)$, we now have $Q_M(M, M) = 1$. Also and since $Q_M(M, M) = 1$, it follows that

$$Q_M(M, 1) = Q_M(M, 2) = \dots = Q_M(M, M-1) = 0$$

This means that

$$Q_M = \begin{pmatrix} R_{M-1} & 0 \\ 0 & 1 \end{pmatrix} \quad (2)$$

for some matrix R_{M-1} , and since $Q_M \in \mathbf{O}(M)$, we have $R_{M-1} \in \mathbf{O}(M-1)$.

Therefore,

$$R_M = \theta(-t_M^{M-1}, M-1, M)\theta(-t_M^{M-2}, M-1, M) \dots \theta(-t_M^2, 2, M)\theta(-t_M^1, 1, M)Q_M$$

for some matrix Q_M as in (2).

Any choice of $M-1$ parameters $t_M^1, \dots, t_M^{M-1}$ will induce embeddings

$$\theta_{t_M^1, \dots, t_M^{M-1}} : \mathbf{O}(M-1) \longrightarrow \mathbf{O}(M) \quad (3)$$

taking $R_{M-1} \in \mathbf{O}(M-1)$ into

$$\theta(t_M^{M-1}, M-1, M)\theta(t_M^{M-2}, M-1, M) \dots \theta(t_M^2, 2, M)\theta(t_M^1, 1, M)Q_M$$

for Q_M as in (2). Multiplication by any of the matrices $\theta(t, j, k)$ will only introduce changes in rows j and k . Parameters defined in (1) are defined up to integer multiples of 2π . Therefore for different (up to integer multiples of 2π) $M-1$ parameters $s_M^1, \dots, s_M^{M-1}$ we obtain different embeddings $\theta_{s_M^1, \dots, s_M^{M-1}}$.

The process can be iterated for R_{M-1} . And composing such embeddings (3) we obtain embeddings

$$\mathbf{O}(N) \longrightarrow \mathbf{O}(N+1) \longrightarrow \dots \mathbf{O}(M-1) \longrightarrow \mathbf{O}(M)$$

The resulting parametrisation and the embeddings of $\mathbf{O}(N)$ into $\mathbf{O}(M)$ will play a significant role and we record them as follows.

Lemma 3.1. *If $N < M$, then*

(a) *Every matrix in $\mathbf{O}(M)$ is parametrised by $(M-1) + (M-2) + \dots + 2 + 1 = \frac{M(M-1)}{2}$ parameters and one sign.*

(b) *Every matrix in $\mathbf{O}(M)$ is parametrised by $(M-1) + (M-2) + \dots + (N+1) = \frac{M(M-1)}{2} - \frac{N(N-1)}{2}$ parameters and a matrix $R_N \in \mathbf{O}(N)$.*

(c) Both of these parametrisations are unique (up to integer multiples of 2π for the real parameters).

Proof. Parts (a) and (b) are the well known representation of matrices in $\mathbf{O}(M)$ described above. For part (c), the sign of the decomposition is the determinant of the matrix. The real parameters t in (3) are defined up to integer multiples of 2π . $\square$

The factorisation reformulates, for storage and computation purposes, the well known fact that as a Lie Group, $\mathbf{O}(M)$ has dimension $\frac{M(M-1)}{2}$ and two components, corresponding to the two possible values of the determinant.

Whenever we refer to a matrix in $\mathbf{O}(M)$ by giving the $M * (M-1)/2$ parameters and a sign we refer to the decomposition of the Lemma 3.1. For later reference we now record both factorisations of arbitrary $R_M \in \mathbf{O}(M)$.

$$\begin{aligned} R_M = & \theta(t_M^{M-1}, M-1, M) \theta(t_M^{M-2}, M-1, M) \dots \theta(t_M^2, 2, M) \theta(t_M^1, 1, M) \\ & \theta(t_{M-1}^{M-2}, M-2, M-1) \theta(t_{M-1}^{M-3}, M-2, M-1) \dots \theta(t_{M-1}^2, 2, M-1) \theta(t_{M-1}^1, 1, M-1) \\ & \dots \theta(t_3^2, 2, 3) \theta(t_3^1, 1, 3) \begin{pmatrix} \cos(t_2^1) & \sin(t_2^1) \\ -\sin(t_2^1) & \cos(t_2^1) \end{pmatrix} \begin{pmatrix} \mu & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} I_{M-2, M-2} \\ \mu =_{-}^{+} 1 \end{pmatrix}, \end{aligned} \quad (4)$$

and if any $Q_N \in \mathbf{O}(M)$ is given by

$$\begin{aligned} Q_N = & \theta(t_N^{N-1}, N-1, N) \theta(t_N^{N-2}, N-2, N) \dots \theta(t_N^2, 2, N) \theta(t_N^1, 1, N-1) \\ & \dots \theta(t_3^2, 2, 3) \theta(t_3^1, 1, 3) \begin{pmatrix} \cos(t_2^1) & \sin(t_2^1) \\ -\sin(t_2^1) & \cos(t_2^1) \end{pmatrix} \begin{pmatrix} \mu & 0 \\ 0 & 1 \end{pmatrix} \begin{pmatrix} I_{M-2, M-2} \\ \mu =_{-}^{+} 1 \end{pmatrix}, \end{aligned}$$

then

$$Q_N = \begin{pmatrix} R_N & 0 \\ 0 & I_{M-N, M-N} \end{pmatrix} \quad (5)$$

where $R_N \in \mathbf{O}(N)$, and

$$\begin{aligned} R_M = & \theta(t_M^{M-1}, M-1, M) \theta(t_M^{M-2}, M-1, M) \dots \theta(t_M^2, 2, M) \theta(t_M^1, 1, M) \dots \\ & \theta(t_{N+1}^N, N, N+1) \theta(t_{N+1}^{N-1}, N-1, N+1) \dots \theta(t_{N+1}^1, 1, N+1) Q_N = \Psi Q_N \end{aligned} \quad (6)$$

where

$$\begin{aligned} \Psi = & \theta(t_M^{M-1}, M-1, M) \theta(t_M^{M-2}, M-2, M) \dots \theta(t_M^2, 2, M) \theta(t_M^1, 1, M) \\ & \dots \theta(t_{N+1}^N, N, N+1) \theta(t_{N+1}^{N-1}, N-1, N+1) \dots \theta(t_{N+1}^1, 1, N+1) \end{aligned}$$

4 Existence and Construction Results

Lemma 4.1. *Let $R_M \in \mathbf{O}(M)$ be a solution for the admissible sequence $\{a_j\}_{j=1}^M$, and let R_M be factorised as in (6) by means of $\frac{M(M-1)}{2} - \frac{N(N-1)}{2}$ parameters and a matrix $R_N \in \mathbf{O}(N)$ as in (5). For any matrix $S_N \in \mathbf{O}(N)$, let*

$$P_N = \begin{pmatrix} S_N & 0 \\ 0 & I_{M-N, M-N} \end{pmatrix}$$

Then

$$S_M = \theta(t_M^{M-1}, M-1, M) \theta(t_M^{M-2}, M-2, M) \dots \theta(t_M^2, 2, M) \theta(t_M^1, 1, M) \dots \\ \theta(t_{N+1}^N, N, N+1) \theta(t_{N+1}^{N-1}, N-1, N+1) \dots \theta(t_{N+1}^1, 1, N+1) P_N = \Psi P_N$$

is a solution for the admissible sequence $\{a_j\}_{j=1}^M$. This provides an embedding of $\mathbf{O}(N)$ into the space of solutions for the admissible sequence $\{a_j\}_{j=1}^M$.

Proof. We will show that the values $\{a_j\}_{j=1}^M$ depend on the parameters

$$t_M^{M-1}, t_M^{M-2}, \dots, t_{N+1}^2, t_{N+1}^1 \quad (7)$$

and are independent of the particular choice of $R_M \in \mathbf{O}(N)$.

Let r_l be the l th row of R_N for $l = 1, \dots, N$. Then $\langle r_l, r_k \rangle = \delta_{l,k}$. Let x_j be the first N terms of the j th row of R_M for $j = 1, \dots, M$. Multiplication by a Givens rotation will replace the first N terms of any row by a linear combination of $\{r_l\}_{l=1}^M$. At the end of the process

$$x_j = \sum_{l=1}^{l=N} \lambda_l^j r_l$$

where the coefficients λ_l depend only on the $\frac{M(M-1)}{2} - \frac{N(N-1)}{2}$ parameters in (7) and are independent of the r_l given by R_N . Since the vectors r_l are orthonormal,

$$\|x_j\|^2 = \langle x_j, x_j \rangle = \left\langle \sum_{l=1}^{l=N} \lambda_l^j r_l, \sum_{l=1}^{l=N} \lambda_l^j r_l \right\rangle = \sum_{l=1}^{l=N} (\lambda_l^j)^2$$

Thus under replacement of R_N by another matrix in $\mathbf{O}(N)$ the values of $\|x_j\|^2$ do not change. Again from (6) it follows that the parameters in (7) induce an embedding of $\mathbf{O}(N)$ into the space of solutions for the admissible sequence $\{a_j\}_{j=1}^M$. $\square$

Remark 4.2. The equations giving λ_l^j in terms of (7) are in practical terms unmanageable. They are given in terms of the $\frac{M(M-1)}{2} - \frac{N(N-1)}{2}$ parameters.

Theorem 4.3. *Let $N < M$ be fixed. If $\{a_j\}_{j=1}^M$ is an admissible sequence, then*

(a) There is a solution for $\{a_j\}_{j=1}^M$

(b) Moreover, there is an embedding from $\mathbf{O}(N) \times \mathbf{O}(M - N)$

$$\mathbf{O}(N) \times \mathbf{O}(M - N)$$

into the space of all solutions for $\{a_j\}_{j=1}^M$. The embedding will take the pair of matrices $R_N \in \mathbf{O}(N)$ and $R_{M-N} \in \mathbf{O}(M - N)$ into

$$\Psi \begin{pmatrix} R_N & 0 \\ 0 & R_{M-N, M-N} \end{pmatrix}$$

where the matrix Ψ is determined by the *admissible* sequence $\{a_j\}_{j=1}^M$.

Proof. The proof is just an algorithm. Let $\Psi = I_{M, M}$. We start with the matrix $O \in \mathbf{O}(M)$ given by

$$O = \begin{pmatrix} R_N & 0 \\ 0 & R_{M-N} \end{pmatrix}$$

where R_N and R_{M-N} are arbitrary matrices in $\mathbf{O}(N)$ and $\mathbf{O}(M - N)$ respectively. The fact that there is no other requirement on R_N and R_{M-N} will provide the clue for the proof of the second claim of the lemma. We will construct matrices $O^n \in \mathbf{O}(M)$ for $n = 1, \dots, M - 1$. We will denote

$$\begin{aligned} o_{left, j}^n &= \text{first } N \text{ terms of } j\text{th row of } O^n \text{ and} \\ o_{right, j}^n &= \text{last } M - N \text{ terms of } j\text{th row of } O^n \end{aligned}$$

Once the values $a_1^2, \dots, a_{M-1}^2$ have been attained the problem is solved since each O^n is an orthogonal matrix. The *solution* will be O^{M-1} . Let

$$t_1 = \arg(a_1 + i\sqrt{1 - a_1^2})$$

and rename

$$\Psi = \theta(t_1, 1, M)\Psi.$$

Let

$$O^1 = \theta(t_1, 1, M)O$$

Then $O^1 = \Psi O$, and

$$\|o_{left, 1}^1\|^2 = a_1^2, \text{ and } \|o_{left, M}^1\|^2 = 1 - a_1^2.$$

Also $o_{left, M}^1$ remains orthogonal to $o_{left, j}^1$ for $j \geq 2$ and $o_{right, M}^1$ remains orthogonal to $o_{right, j}^1$ for $j \leq M - 1$. The values of the squared norms of $o_{left, j}^1$

and $o_{right,j}^1$ are described by the diagram:

$$\left(\begin{array}{c|c} a_1^2 & 1 - a_1^2 \\ 1 & 0 \\ 1 & 0 \\ \vdots & \vdots \\ 1 & 0 \\ \hline 0 & 1 \\ \vdots & \vdots \\ 0 & 1 \\ 1 - a_1^2 & a_1^2 \end{array} \right)$$

Now one of the following applies:

(CASE I) $1 \leq a_1^2 + a_2^2$, or

(CASE II) $1 > a_1^2 + a_2^2$

In CASE I, take t_2 a solution of

$$\sin(t_2)^2 = \frac{1 - a_1^2}{a_1^2} \quad \text{and} \quad O^2 = \theta(t_2, 2, M)O^1$$

Then

$$\begin{aligned} \|o_{left,2}^2\|^2 &= \|\cos(t_2)o_{left,2}^1 + \sin(t_2)o_{left,M}^1\|^2 \stackrel{(1)}{=} \|\cos(t_2)o_{left,2}^1\|^2 \\ &+ \|\sin(t_2)o_{left,M}^1\|^2 = \cos(t_2)^2 \|o_{left,2}^1\|^2 + \sin(t_2)^2 \|o_{left,M}^1\|^2 \stackrel{(2)}{=} \cos(t_2)^2 \\ &+ \sin(t_2)^2(1 - a_1^2) = 1 - \sin(t_2)^2(a_1^2) = 1 - \frac{1 - a_2^2}{a_1^2}a_1^2 = a_2^2 \end{aligned}$$

where equality (1) holds because $o_{left,2}^1$ and $o_{left,M}^1$ are orthogonal and equality (2) holds because $o_{left,2}^1$ is a unitary vector (in R_N). The values of the squared norms of $o_{left,j}^2$ and $o_{right,j}^2$ are described by the diagram:

$$\left(\begin{array}{c|c} a_1^2 & 1 - a_1^2 \\ a_2^2 & 1 - a_2^2 \\ 1 & 0 \\ \vdots & \vdots \\ 1 & 0 \\ \hline 0 & 1 \\ \vdots & \vdots \\ 0 & 1 \\ 2 - a_1^2 - a_2^2 & a_1^2 + a_2^2 - 1 \end{array} \right)$$

Now for $j = 3, \dots, N$, we have $o_{left,j}^2 = o_{left,j}^1$ (they come from R_N) so they are unitary and they remain orthogonal to $o_{left,M}^2$ since the latter has been obtained as a linear combination of $o_{left,j}^1$ for $j \leq 2$.

Similarly, for $j = N+1, \dots, M-1$, we have $o_{right,j}^2 = o_{right,j}^1$ (they come from R_M) so they are unitary and remain orthogonal to $o_{right,M}^2$ since the latter is just a scalar multiple of $o_{right,M}^1$.

In CASE II, take t_2 a solution of

$$\sin(t_2)^2 = \frac{a_2^2}{1-a_1^2} \quad \text{and let } O^2 = \theta(t_2, M-1, M)O^1$$

Then

$$\begin{aligned} \|o_{right,M-1}^2\|^2 &= \|\cos(t_2)o_{right,M-1}^1 + \sin(t_2)o_{right,M}^1\|^2 \stackrel{(1)}{=} \|\cos(t_2)o_{right,M-1}^1\|^2 \\ &\quad + \|\sin(t_2)o_{right,M}^1\|^2 = \cos(t_2)^2 \|o_{right,M-1}^1\|^2 + \sin(t_2)^2 \|o_{right,M}^1\|^2 \\ &\stackrel{(2)}{=} \cos(t_2)^2 + \sin(t_2)^2 a_1^2 = 1 - \sin(t_2)^2 + \sin(t_2)^2 a_1^2 \\ &= 1 + \sin(t_2)^2 (a_1^2 - 1) = 1 + \frac{a_2^2}{1-a_1^2} (a_1^2 - 1) = 1 - a_2^2 \end{aligned}$$

where equality (1) holds because $o_{right,M-1}^1$ and $o_{right,M}^1$ are orthogonal and equality (2) holds since $o_{right,M-1}^1$ is a unitary vector (in R_{M-N}). The values of the squared norms of $o_{left,j}^2$ and $o_{right,j}^2$ are described by the diagram:

$$\left(\begin{array}{c|c} a_1^2 & 1-a_1^2 \\ 1 & 0 \\ 1 & 0 \\ \cdots & \cdots \\ 1 & 0 \\ \hline 0 & 1 \\ \cdots & \cdots \\ 0 & 1 \\ a_2^2 & 1-a_2^2 \\ 1-a_1^2-a_2^2 & a_1^2+a_2^2 \end{array} \right)$$

Now for $j = 2, \dots, N$, we have $o_{left,j}^2 = o_{left,j}^1$ (they come from R_{M-N}). Hence, they are unitary and they remain orthogonal to $o_{left,M}^2$ since the latter has been obtained as a linear combination of $o_{left,j}^1$ for $j < 2$.

Similarly, for $j = N+1, \dots, M-2$, we have $o_{right,j}^2 = o_{right,j}^1$ (they come from R_M). So they are unitary and remain orthogonal to $o_{right,M}^2$ since the latter has been obtained as a linear combination of $o_{right,j}^1$ for $j > M-2$.

In either case, by renaming

$$\Psi = \theta(t_2, m, M)\Psi$$

where m is given by the CASE used, we have

$$O^2 = \Psi O$$

After k applications of CASE I and l applications of CASE II we obtain the matrix O^{k+l+1} . The values of the squared norms of $o_{left,j}^{k+l+1}$ and $o_{right,j}^{k+l+1}$ are described (up to reordering of $\{a_j\}_{j=1}^M$) by the diagram:

$$\left(\begin{array}{c|c} \begin{array}{c} a_1^2 \\ a_2^2 \\ \dots \\ a_{k+1}^2 \\ 1 \\ \dots \\ 1 \\ \hline 0 \\ \dots \\ 0 \\ \dots \\ a_{k+l+1}^2 \\ \dots \\ a_{k+3}^2 \\ a_{k+2}^2 \\ k+1-a_1^2-\dots-a_{k+l+1}^2 \end{array} & \begin{array}{c} 1-a_1^2 \\ 1-a_2^2 \\ \dots \\ 1-a_{k+1}^2 \\ 0 \\ \dots \\ 0 \\ \hline 1 \\ \dots \\ 1 \\ \dots \\ 1-a_{k+l+1}^2 \\ \dots \\ 1-a_{k+3}^2 \\ 1-a_{k+2}^2 \\ a_1^2+\dots+a_{k+l+1}^2-k \end{array} \end{array} \right)$$

Now the two cases correspond to :

(CASE I) $k+1 \leq a_1^2 + \dots + a_{k+l+1}^2 + a_{k+l+2}^2$, or

(CASE II) $k+1 > a_1^2 + \dots + a_{k+l+1}^2 + a_{k+l+2}^2$

In CASE I, take t_{k+l+2} a solution of

$$\sin(t_{k+l+2})^2 = \frac{1 - a_{k+l+2}^2}{a_1^2 + \dots + a_{k+l+1}^2 - k}$$

and

$$O^{k+l+2} = \theta(t_{k+l+2}, k+2, M) O^{k+l+1}$$

In CASE II, take t_{k+l+2} a solution of

$$\sin(t_{k+l+2})^2 = \frac{a_{k+l+2}^2}{k+1-a_1^2-\dots-a_{k+l+1}^2}$$

and

$$O^{k+l+2} = \theta(t_{k+l+2}, M-l-1, M) O^{k+l+1}$$

In either case, by renaming

$$\Psi = \theta(t_{k+l+2}, m, M)\Psi$$

where m is given by the CASE used, we have

$$O^{t_{k+l+2}} = \Psi O.$$

Now it remains to prove that CASE I eventually takes place whenever $k < N - 1$ and when $k = N - 1$ only CASE II can apply.

Suppose $k < N - 1$. Then $k + 1 \leq N - 1$. At the same time

$$a_1^2 + \cdots + a_{k+l+1}^2 + a_{M-1}^2 \geq N - 1$$

giving

$$k + 1 \leq N - 1 \leq a_1^2 + \cdots + a_{k+l+1}^2 + a_{M-1}^2$$

so CASE I will eventually apply .

If $k = N - 1$ and CASE I applies at some step,

$$k + 1 = N \leq a_1^2 + \cdots + a_{k+l+2}^2 \leq N$$

This could only happen if $N = a_1^2 + \cdots + a_{k+l+2}^2$, so $a_{k+l+3} = \cdots = a_{M-1} = a_M = 0$. In this case O^{k+l+1} is a *solution* for $\{a_j\}_{j=1}^M$. It is possible to reorder the sequence $\{a_j\}$ so that in the algorithm first take place $N - 1$ times CASE I and then It is clear from the construction that

$$O^{M-1} = \Psi O.$$

Finally, the parameters $t_1, \dots, t_{M-1}$ and the matrix Ψ were determined by the norms $a_1^2, \dots, a_{M-1}^2$ and so is Ψ . Replacement of O by any other

$$O' = \begin{pmatrix} R'_N & 0 \\ 0 & R'_{M-N} \end{pmatrix}$$

would have lead to another solution

$$O'^{M-1} = \Psi O'$$

It is clear that multiplication by a unitary matrix is an injection. The claim is thus proved. $\square$

It must be pointed out that the embedding of $O(N) \times O(M - M)$ into $O(M)$ does not fill out the space of *solutions*. Using alternative algorithms it is possible to construct *solutions* which do not belong to the image of the embedding, that is, can not be factorised as

$$\psi O = \psi \begin{pmatrix} R_N & 0 \\ 0 & R_{M-N} \end{pmatrix}$$

References

- [BF02] J.J. Benedetto and M. Fickus. Finite normalized tight frames. 2002. Preprint.
- [Cas00] P.G. Casazza. The art of frame theory. *Taiwan Jour. of Math.*, 4(2):129-201, 2000. To appear.
- [CKLT02] P.G. Casazza, J. Kovačević, M.T. Leon, and J.C. Tremain. Custom built tight frames. 2002. Preprint.
- [DGK02] P.L. Dragotti, V.K. Goyal, and J. Kovačević. Filter bank frame expansions with erasures. *IEEE Trans. Inform. Th., special issue in Honor of Aaron D. Wyner*, 46(6):1439-1450, 2002, Invited Paper.
- [EF01] Y. Eldar and G.D. Forney. Optimal tight frames and quantum measurement. 2001. Preprint.
- [GK01] V.K. Goyal and J. Kovačević. quantized frame expansions with erasures. *Appl. Comput. Harmon. Anal.*, 10(3):203-233, 2001. To appear.
- [GKV98] V.K. Goyal, J. J. Kovačević, and M. Vetterli. Multiple description transform coding: Robustness to erasures using tight frame expansions. *Proc. IEEE Int. Symp. on Inform. Th.*, 1998. Cambridge, MA.
- [Vai93] P.P. Vaidyanathan. *Multirate systems and filters banks*. Prentice Hall, Englewood Cliffs, N.J., 1993.
- [Š] E. Šoljanin. Tight frames in quantum information theory. *DIMACS Workshop on Source Coding and Harmonic Analysis*. Rutgers, New Jersey, May 2002, <http://dimacs.rutgers.edu/Workshops/Modern/abstracts.html>.

Semigroups and Genetic Repression

Genni Fragnelli

Dipartimento di Ingegneria dell'Informazione,
Università degli Studi di Siena,
Via Roma 56, 53100 Siena, Italy
fragnelli@dii.unisi.it

Abstract

In this article we start from a model presented in [5], [16] and in [21, Introduction]. We explain why this model is unrealistic and propose a system of modified equations. Our approach to the study of these equations is based on semigroup techniques. In particular we use the theory developed in [9] and [10]. The basic idea is to rewrite the equations that express the biological model as a delay equation with nonautonomous past on a suitable state space and then to study their well-posedness and stability.

2000 Mathematics Subject Classification: 34G10, 35B40, 47A10, 47D06, 47N60.

Key words and phrases: Genetic repression, partial differential equations, semigroups, stability.

1 Introduction

Systems of delayed reaction diffusion equations have been used frequently in modeling genetic repression, (see, e.g., J. Wu [21, Introduction]). The study of these equations goes back to the 60's with B.C. Goodwin, F. Jacob and J. Monod (see, e.g., [12], [13] or [14]).

In particular, Goodwin suggested that time delays caused by the processes of transcription and translation as well as spatial diffusion of reactants could play a role in the behavior of this system. The later studies of these models included either time delays (see, e.g., [2], [15] or [20]) or the spatial diffusion (see, e.g. [17]).

In this paper we study the one-dimensional model presented, e.g., in [5], [16], and in [21, Introduction] which includes spatial diffusion and time delay. According to this model the repressor in the cytoplasm at time t and position x depends on the mRNA (messenger ribo nuclein acid) that was at time $t - r$ at the *same* position x , where $r > 0$. This assumption, however, is unrealistic. For this reason we present a system of modified equations, which take into account the diffusion in the past of the mRNA contained in the cytoplasm. In particular we modify the equations associated to this genetic repression in such

a way and then study them with the theory developed in [9] and in [10]. In particular in Section 2 we explain this model following [16] or [21, Introduction] and then justify the modifications of these equations. In Section 3 we consider the model without delay. In Section 4 we rewrite the biological model as a delay equation with nonautonomous past, studying its well-posedness. The last section is dedicated to a stability analysis.

2 The Biological Model

2.1 The Classical Equations

In this section we briefly explain the biological model and rewrite it as a system of equations following, e.g., [5], [16] and [21, Introduction].

The eucaryotic cell Ω consists of two compartments where the most important chemical reactions happen. Such compartments are enclosed within the cell wall, unpermeable to the mRNA and to the repressor, and separated by the permeable nuclear membrane. The first compartment ω is the nucleus where mRNA is produced. The second compartment, denoted by $\Omega \setminus \omega$, is the cytoplasm in which the ribosomes are randomly dispersed. The process of translation and the production of the repressor occurs here.

We denote by u_i and v_i the concentrations of mRNA and of the repressor, respectively, in ω if $i = 1$ and in $\Omega \setminus \omega$ if $i = 2$. These two species interact to control each other's production. In the nucleus ω , mRNA is transcribed from the gene at a rate depending on the concentration of the repressor v_1 . The mRNA leaves ω and enters the cytoplasm $\Omega \setminus \omega$ where it diffuses and reacts with ribosomes. Through the delayed process of translation, a sequence of enzymes is produced which in turn produce a repressor v_2 . This repressor comes back to ω where it inhibits the production of u_1 . This process can be written, in one dimension, as the following system of equations.

$$\left\{ \begin{array}{ll} \frac{du_1(t)}{dt} = h(v_1(t + r_1)) - b_1 u_1(t) + a_1(u_2(t, 0) - u_1(t)), & t \geq 0, \\ \frac{dv_1(t)}{dt} = -b_2 v_1(t) + a_2(v_2(t, 0) - v_1(t)), & t \geq 0, \\ \frac{\partial u_2(t, x)}{\partial t} = D_1 \frac{\partial^2 u_2(t, x)}{\partial x^2} - b_1 u_2(t, x), & t \geq 0, x \in (0, 1], \\ \frac{\partial v_2(t, x)}{\partial t} = D_2 \frac{\partial^2 v_2(t, x)}{\partial x^2} - b_2 v_2(t, x) + c_0 u_2(t + r_2, x), & t \geq 0, x \in (0, 1], \end{array} \right. \quad (1)$$

with boundary conditions

$$\left\{ \begin{array}{ll} \frac{\partial u_2(t, 0)}{\partial x} = -\beta_1(u_2(t, 0) - u_1(t)), & t \geq 0, \\ \frac{\partial v_2(t, 0)}{\partial x} = -\beta_1^*(v_2(t, 0) - v_1(t)), & t \geq 0, \\ \frac{\partial u_2(t, 1)}{\partial x} = \frac{\partial v_2(t, 1)}{\partial x} = 0, & t \geq 0, \end{array} \right. \quad (2)$$

and initial conditions

$$\begin{cases} u_1(s) = f_1(s), \\ v_1(s) = g_1(s), \\ u_2(s, x) = f_2(s, x), \\ v_2(s, x) = g_2(s, x), \end{cases} \quad (3)$$

for $x \in (0, 1]$ and $s \in [-r_i, 0]$, $i = 1, 2$. As in [16], the interval $(0, 1]$ corresponds to the cytoplasm $\Omega \setminus \omega$ and the nucleus ω is localized in 0. The constants b_i are *kinetic rates of decay*, a_i are *rates of transfer* between ω and $\Omega \setminus \omega$ and are directly proportional to the concentration gradient. The constants D_i are the diffusivity coefficients, and the constant c_0 is the *production rate* for the repressor. The function h is a decreasing function and represents the production of mRNA. It is of the form $h(x) = \frac{1}{1 + kx^\rho}$, where k is a kinetic constant and ρ is the Hill coefficient. The delay $-r_1 \geq 0$ is the transcription time, i.e., the time necessary to the transcription reaction, and $-r_2 \geq 0$ is the translation time. The constants β_1 and β_1^* are the constants of Fick's law (see, e.g., [1, Chapter VI] or [19, Chapter V]). We have to underline the fact that all biological constants are positive.

2.2 A better model

According to model (1), the variation of the repressor $v_2(t, x)$ at time t and position x depends on the mRNA which was at time $t + r_2$ at the *same* position x . Clearly, this is unrealistic because the mRNA is subject to a diffusion process in the time interval $[t + r_2, t]$, where we recall that $r_2 < 0$.

To include such a phenomenon in our model, we suppose, for simplicity, that this migration is given by a diffusion of the form $e^{t\Delta_D}$, where $\Delta_D := \frac{d^2}{dx^2}$ is the Laplacian with Dirichlet boundary conditions.

To be more precise, we take the Laplacian Δ_D with domain

$$D(\Delta_D) := \{f \in W^{2,1}[0, 1] : f(0) = f(1) = 0\}$$

on the Banach space $L^1[0, 1]$.

Then the evolution family $\mathcal{U} := (U(t, s))_{-1 \leq t \leq s \leq 0}$ solving the corresponding Cauchy problem (see [9, Example 6.1]) is

$$U(t, s) := T(s - t), \quad -1 \leq t \leq s \leq 0, \quad (4)$$

where $(T(t))_{t \geq 0} = (e^{t\Delta_D})_{t \geq 0}$ is the heat semigroup on $L^1[0, 1]$. Observe that the growth bound of the evolution family $(U(t, s))_{-1 \leq t \leq s \leq 0}$ and of the heat semigroup $(T(t))_{t \geq 0}$ coincide and are negative, i.e.

$$-\pi^2 = \omega_0(\mathcal{U}) = \omega_0(T(\cdot)) < 0. \quad (5)$$

Thus, assuming that the mRNA in the cytoplasm is subject to a diffusion in the past of the form $e^{t\Delta_D}$, the term $u_2(t + r_2, x)$ must be modified. Let

$\tilde{u}_2(t + r_2, x)$ be the modification of $u_2(t + r_2, x)$ governed by $(U(t, s))_{-1 \leq t \leq s \leq 0}$, i.e.

$$\begin{aligned} \tilde{u}_2(t + r_2, x) : &= \begin{cases} U(r_2, 0)u_2(t + r_2, x), & 0 \leq r_2 + t, \\ U(r_2, t + r_2)f_2(t + r_2, x), & r_2 + t \leq 0, \end{cases} \\ &= \begin{cases} T(-r_2)u_2(t + r_2, x), & 0 \leq r_2 + t, \\ T(t)f_2(t + r_2, x), & r_2 + t \leq 0. \end{cases} \end{aligned} \quad (6)$$

Then, the system (1) becomes

$$\begin{cases} \frac{du_1(t)}{dt} = h(v_1(t + r_1)) - b_1u_1(t) + a_1(u_2(t, 0) - u_1(t)), & t \geq 0, \\ \frac{dv_1(t)}{dt} = -b_2v_1(t) + a_2(v_2(t, 0) - v_1(t)), & t \geq 0, \\ \frac{\partial u_2(t, x)}{\partial t} = D_1 \frac{\partial^2 u_2(t, x)}{\partial x^2} - b_1u_2(t, x), & t \geq 0, x \in (0, 1], \\ \frac{\partial v_2(t, x)}{\partial t} = D_2 \frac{\partial^2 v_2(t, x)}{\partial x^2} - b_2v_2(t, x) + c_0\tilde{u}_2(t + r_2, x), & t \geq 0, x \in (0, 1]. \end{cases} \quad (7)$$

Remark 2.1. If Δ_D is substituted by Δ_N , i.e. the Neumann Laplacian with domain

$$D(\Delta_N) := \{f \in W^{2,1}[0, 1] : f'(0) = f'(1) = 0\},$$

then for the growth bound of the associated heat semigroup $(T(t))_{t \geq 0}$ we only have $\omega_0(\mathcal{T}) \leq 0$.

3 Semigroup without Delay

We want to study the system (7) with the theory developed in [9] and [10]. To this aim we have to rewrite the genetic repression as a delay equation of the form

$$(NDE) \quad \begin{cases} \dot{W}(t) &= \mathcal{B}W(t) + \Phi \widetilde{W}_t, & t \geq 0 \\ W(0) &= y \in X \\ \widetilde{W}_0 &= f \in L^p([-1, 0], X), \end{cases}$$

on some Banach space X , where $(\mathcal{B}, D(\mathcal{B}))$ is the generator of a strongly continuous semigroup $(S(t))_{t \geq 0}$ on X , the **delay operator** $\Phi : W^{1,p}([-1, 0], X) \rightarrow X$ is a linear operator, $p \geq 1$ and the **modified history function** $\widetilde{W}_t : [-1, 0] \rightarrow X$ (see [9, Definition 3.2]) is given by

$$\widetilde{W}_t(\tau) := \begin{cases} \widetilde{U}(\tau, 0)W(t + \tau) & \text{for } 0 \leq t + \tau, \\ \widetilde{U}(\tau, t + \tau)f(t + \tau) & \text{for } t + \tau \leq 0 \end{cases} \quad (8)$$

for some evolution family $(\widetilde{U}(t, s))_{-1 \leq t \leq s \leq 0}$.

In the definition of the modified history function $\widetilde{W}_t$ two time variables t and τ appear. The variable t can be interpreted as the *absolute time* and τ as the *relative time*. The meaning of (NDE) is that if we start with the history function f , this function is shifted to the left by $-t$, and for values greater than

$-t$ the value of the solution is given by the delay operator Φ applied to the shifted function.

In order to rewrite (7) as a delay equation with nonautonomous past we proceed as follows.

As a first step we ignore the terms with delay, i.e. we consider the system

$$\begin{cases} \frac{du_1(t)}{dt} = -b_1 u_1(t) + a_1(u_2(t, 0) - u_1(t)), & t \geq 0, \\ \frac{dv_1(t)}{dt} = -b_2 v_1(t) + a_2(v_2(t, 0) - v_1(t)), & t \geq 0, \\ \frac{\partial u_2(t, x)}{\partial t} = D_1 \frac{\partial^2 u_2(t, x)}{\partial x^2} - b_1 u_2(t, x), & t \geq 0, x \in (0, 1], \\ \frac{\partial v_2(t, x)}{\partial t} = D_2 \frac{\partial^2 v_2(t, x)}{\partial x^2} - b_2 v_2(t, x), & t \geq 0, x \in (0, 1], \end{cases} \quad (9)$$

with initial conditions (3).

To rewrite (9) as an ordinary differential equation of the form

$$\begin{cases} \dot{W}(t) = \mathcal{B}W(t), & t \geq 0, \\ W(0) = y \in X, \end{cases} \quad (10)$$

for a linear operator $(\mathcal{B}, D(\mathcal{B}))$ on a Banach space X , we take as X the space $X := \mathbb{R}^2 \times (L^1[0, 1])^2$ and as $\mathcal{B}$ the operator

$$\mathcal{B} := \begin{pmatrix} -b_1 - a_1 & 0 & a_1 \delta_0 & 0 \\ 0 & -b_2 - a_2 & 0 & a_2 \delta_0 \\ 0 & 0 & D_1 \Delta - b_1 & 0 \\ 0 & 0 & 0 & D_2 \Delta - b_2 \end{pmatrix}, \quad (11)$$

with domain

$$D(\mathcal{B}) := \left\{ \begin{pmatrix} x \\ y \\ f \\ g \end{pmatrix} \in \mathbb{R}^2 \times (W^{2,1}[0, 1])^2 : L \begin{pmatrix} f \\ g \end{pmatrix} = \begin{pmatrix} x \\ y \end{pmatrix} \text{ and } f'(1) = g'(1) = 0 \right\}$$

where $\delta_0 f := f(0)$ for $f \in C[0, 1]$, $\Delta := \frac{d^2}{dx^2}$ and the operator $L : (W^{2,1}[0, 1])^2 \rightarrow \mathbb{R}^2$ is defined by

$$L \begin{pmatrix} f \\ g \end{pmatrix} = \begin{pmatrix} \frac{f'(0)}{\beta_1} + f(0) \\ \frac{g'(0)}{\beta_1^*} + g(0) \end{pmatrix}. \quad (12)$$

Remark 3.1. The Laplacian Δ in the definition of the operator matrix $\mathcal{B}$ has "maximal" domain $W^{2,1}[0, 1]$ and thus is different from the Dirichlet Laplacian Δ_D that governs the diffusion in the past.

Let now

$$W(t) := (u)_1(t) v_1(t) u_2(t) v_2(t). \quad (13)$$

Then the ordinary differential equation (10) and the system (9) are "equivalent". Thus, to prove the existence of a solution of (9) is sufficient to prove that $(\mathcal{B}, D(\mathcal{B}))$ is the generator of a strongly continuous semigroup $(S(t))_{t \geq 0}$. In

this case the solution of (10) is given by $W(t) = S(t)y$. To this aim we rewrite $\mathcal{B}$ in the form

$$\begin{aligned} \mathcal{B} = \mathcal{B}_0 + \Theta - \Pi &:= \begin{pmatrix} 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & D_1\Delta & 0 \\ 0 & 0 & 0 & D_2\Delta \end{pmatrix} + \begin{pmatrix} 0 & 0 & a_1\delta_0 & 0 \\ 0 & 0 & 0 & a_2\delta_0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{pmatrix} \\ &\quad - \begin{pmatrix} b_1 + a_1 & 0 & 0 & 0 \\ 0 & b_2 + a_2 & 0 & 0 \\ 0 & 0 & b_1 & 0 \\ 0 & 0 & 0 & b_2 \end{pmatrix}, \end{aligned}$$

with domains $D(\mathcal{B}_0) = D(\mathcal{B})$, $D(\Theta) = \mathbb{R}^2 \times (W^{1,1}[0, 1])^2$ and $\Pi \in \mathcal{L}(X)$.

Proposition 3.2. *The operator $(\mathcal{B}_0, D(\mathcal{B}_0))$ generates an analytic semigroup.*

Proof. Let

$$B_m := \begin{pmatrix} D_1\Delta & 0 \\ 0 & D_2\Delta \end{pmatrix}$$

with the maximal domain $D(D_m) = (W^{2,1}[0, 1])^2$ and $B_0 := B_m|_{\text{Ker } L}$. Then, by [8, Chapter VI, Section 4], B_0 generates an analytic positive semigroup $(S_0(t))_{t \geq 0}$.

Applying [6, Corollary 2.8] to the operator $(\mathcal{B}_0, D(\mathcal{B}_0))$, the assertion follows. $\square$

The next definition will be needed to state Lemma 3.4.

Definition 3.3 (see [8], Definition III.2.1). Let $A : D(A) \subset X \rightarrow X$ be a linear operator on the Banach space X . An operator $B : D(B) \subset X \rightarrow X$ is called A -bounded if $D(A) \subseteq D(B)$ and if there exist constants $a, b \in \mathbb{R}_+$ such that

$$\|Bx\| \leq a\|Ax\| + b\|x\| \quad (14)$$

for all $x \in D(A)$. The A -bound of B is

$$a_0 := \inf\{a \geq 0 : \text{there exists } b \in \mathbb{R}_+ \text{ such that (14) holds}\}.$$

Lemma 3.4. *The operator Θ is $\mathcal{B}_0$ -bounded with $\mathcal{B}_0$ -bound $a_0 = 0$.*

Proof. Obviously, $D(\mathcal{B}_0) \subseteq D(\Theta)$ and for

$$\mathcal{V} = \begin{pmatrix} x \\ y \\ f \\ g \end{pmatrix} \in D(\mathcal{B}_0),$$

we have

$$\|\Theta\mathcal{V}\| = \left\| \begin{pmatrix} a_1 f(0) \\ a_2 g(0) \\ 0 \\ 0 \end{pmatrix} \right\| \leq \|a_1 f(0)\| + \|a_2 g(0)\|.$$

Since

$$\|\mathcal{B}_0\mathcal{V}\| = \left\| \begin{pmatrix} 0 \\ 0 \\ f'' \\ g'' \end{pmatrix} \right\|,$$

it suffices to prove that for arbitrary small $a, b > 0$ there exist constants $c, d \in \mathbb{R}_+$, such that $\|f(0)\| \leq a\|f''\|_1 + b\|f\|_1$ and $\|g(0)\| \leq c\|f''\|_1 + d\|f\|_1$.

Using the fundamental theorem of calculus and the fact that the operator $\frac{d}{dx}$ is $\frac{d^2}{dx^2}$ -bounded with $\frac{d^2}{dx^2}$ -bound 0 (see [8, Example III.2.2]), the assertion follows. $\square$

The following proposition is an easy consequence of Proposition 3.2 and Lemma 3.4.

Proposition 3.5. *The operator $(\mathcal{B}, D(\mathcal{B}))$ generates a positive analytic semigroup $(S(t))_{t \geq 0}$ on X .*

Proof. By the previous proposition, the operator Θ and hence also $\Theta + \Pi$ is $\mathcal{B}_0$ -bounded with $\mathcal{B}_0$ -bound $a_0 = 0$. Using [8, Theorem III.2.10], one has that $\mathcal{B}_0 + \Theta + \Pi$ generates an analytic semigroup. Moreover, by [6, Proposition 5.2], this semigroup is positive as well. $\square$

4 The Genetic Repression as a Delay Equation with Nonautonomous Past

As we mentioned at the beginning of Section 3, we want to rewrite (7) as a delay equation with nonautonomous past, i.e. in the form (NDE) (see the previous section). To this end and for the sake of simplicity, assume $r_1 = r_2 = -1$ and, instead of $h(x) = \frac{1}{1+kx^\rho}$, consider its linearization, i.e. the function $\bar{h} : \mathbb{R} \rightarrow \mathbb{R}$ given by

$$\bar{h}(x) = -\frac{k\rho(v_1^e)^{\rho-1}}{(1+k(v_1^e)^\rho)^2}x, \quad x \in \mathbb{R}. \quad (15)$$

Here v_1^e is such that $(u_1^e, v_1^e, u_2^e, v_2^e)^T$ is one of the steady-state solutions of (7) (see [16, Section 5]). Hence, instead of (7) we consider the simplified and linearized system

$$\begin{cases} \frac{du_1(t)}{dt} = cv_1(t-1) - b_1u_1(t) + a_1(u_2(t,0) - u_1(t)), & t \geq 0, \\ \frac{dv_1(t)}{dt} = -b_2v_1(t) + a_2(v_2(t,0) - v_1(t)), & t \geq 0, \\ \frac{\partial u_2(t,x)}{\partial t} = D_1 \frac{\partial^2 u_2(t,x)}{\partial x^2} - b_1u_2(t,x), & t \geq 0, x \in (0,1], \\ \frac{\partial v_2(t,x)}{\partial t} = D_2 \frac{\partial^2 v_2(t,x)}{\partial x^2} - b_2v_2(t,x) + c_0\tilde{u}_2(t-1,x), & t \geq 0, x \in (0,1], \end{cases} \quad (16)$$

where $c := -\frac{k\rho(v_1^\varepsilon)^{\rho-1}}{(1+k(v_1^\varepsilon)^\rho)^2}$. Let $\Phi : W^{1,p}([-1, 0], X) \rightarrow X$ be the delay operator given by

$$\Phi := \begin{pmatrix} 0 & c\delta_{-1} & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & c_0\delta_{-1} & 0 \end{pmatrix}, \quad (17)$$

where $\delta_{-1}f = f(-1)$ for $f \in W^{1,p}([-1, 0], \mathbb{R}) \rightarrow \mathbb{R}$ and take

$$\tilde{U}(t, s) := (I) dIdU(t, s)Id = (I) dIdT(s - t)Id \quad \text{for } -1 \leq t \leq s \leq 0. \quad (18)$$

Recall that $(T(t))_{t \geq 0}$ denotes the heat semigroup on $L^1[-1, 0]$. Then it is easy to prove that the system (16) is equivalent to

$$\begin{cases} \dot{W}(t) = \mathcal{B}W(t) + \Phi \tilde{W}_t, & t \geq 0 \\ W(0) = y \in X, \\ \tilde{W}_0 = f \in L^p([-1, 0], X), \end{cases} \quad (19)$$

where

$$f = (f)_1 g_1 f_2 g_2, \quad (20)$$

$W(t)$ is as in (13) and the modified history function $\tilde{W}_t : [-1, 0] \rightarrow X$ is defined by

$$\tilde{W}_t(\tau) := \begin{cases} \tilde{U}(\tau, 0)W(t + \tau) & \text{for } 0 \leq t + \tau, \\ \tilde{U}(\tau, t + \tau)f(t + \tau) & \text{for } t + \tau \leq 0. \end{cases}$$

Here $\tilde{\mathcal{U}} := (\tilde{U}(t, s))_{-1 \leq t \leq s \leq 0}$ is the evolution family defined in (18). Thanks to the equivalence between (16) and (19), to prove the existence of a solution of (16) it is sufficient to study the well-posedness of (NDE).

Definition 4.1. 1. We call a function $W : \mathbb{R} \rightarrow X$ a *classical solution* of (NDE) if

- (i) $W \in C(\mathbb{R}, X) \cap C^1(\mathbb{R}_+, X)$,
- (ii) $W(t) \in D(\mathcal{B})$, $\tilde{W}_t \in D(\Phi)$, $t \geq 0$
- (iii) W satisfies (NDE) for all $t \geq 0$.

2. We call (NDE) **well-posed** if

- (i) for every $\begin{pmatrix} y \\ f \end{pmatrix}$ in a dense subspace $\mathcal{S} \subseteq X \times L^p(\mathbb{R}_-, X)$ there is a unique (classical) solution $u(y, f, \cdot)$ of (NDE) and
- (ii) the solutions depend continuously on the initial values, i.e., if a sequence $\begin{pmatrix} y_n \\ f_n \end{pmatrix}$ in $\mathcal{S}$ converges to $\begin{pmatrix} y \\ f \end{pmatrix} \in \mathcal{S}$, then $u(y_n, f_n, t)$ converges to $u(y, f, t)$ uniformly for t in compact intervals of $\mathbb{R}_-$.

The next two results are concerned with the well-posedness of (19) and as a consequence the existence of the solutions of (16).

Take

$$\mathcal{C} := \begin{pmatrix} \mathcal{B} & \Phi \\ 0 & \mathcal{G} \end{pmatrix}, \quad (21)$$

with domain

$$D(\mathcal{C}) := \left\{ \begin{pmatrix} y \\ f \end{pmatrix} \in D(\mathcal{B}) \times D(\mathcal{G}) : f(0) = y \right\} \quad (22)$$

on the product space $\mathcal{E} := X \times L^p([-1, 0], X)$. Here the operator $\mathcal{G}$ is the matrix

$$\mathcal{G} := \begin{pmatrix} \frac{d}{d\sigma} & 0 & 0 & 0 \\ 0 & \frac{d}{d\sigma} & 0 & 0 \\ 0 & 0 & G & 0 \\ 0 & 0 & 0 & \frac{d}{d\sigma} \end{pmatrix},$$

with domain

$$D(\mathcal{G}) := (W^{1,p}[-1, 0])^2 \times D(G) \times W^{1,p}([-1, 0], L^1[0, 1]),$$

$\frac{d}{d\sigma}$ is the weak derivative and $(G, D(G))$ is the closure of

$$Af = f' + f'', \quad (23)$$

for f in an appropriate subspace of $D(G)$ (for details see [11, Proposition 3.1]). The following theorem follows by [9, Theorem 4.3].

Theorem 4.2. *For the function $\bar{h}$ given by (15), the operator $(\mathcal{C}, D(\mathcal{C}))$ generates a strongly continuous semigroup $(\mathcal{T}(t))_{t \geq 0}$.*

As an immediate consequence of [9, Theorem 4.5], we obtain the next theorem.

Theorem 4.3. *If the function $\bar{h}$ is given by (15), then the delay equation (19) is well-posed. Thus the classical solutions u of (16) are given by*

$$u(t) = \begin{cases} \pi_1 \left(\mathcal{T}(t) \begin{pmatrix} y \\ f \end{pmatrix} \right), & t \geq 0, \\ f(t), & t \leq 0, \end{cases}$$

for every $\begin{pmatrix} y \\ f \end{pmatrix} \in D(\mathcal{C})$. Here π_1 is the projection onto the first component X of $\mathcal{E}$.

5 Stability

In this section our goal is to study the stability of the solutions of the system (16), in particular we want to find conditions such that the solutions decay exponentially. In this context *positivity* will be very helpful. First of all we have

to observe that the function $\bar{h}$ in (15) is negative, which implies that the delay operator Φ is negative. Then we consider the system

$$\begin{cases} \frac{du_1(t)}{dt} = -cv_1(t-1) - b_1u_1(t) + a_1(u_2(t,0) - u_1(t)), & t \geq 0, \\ \frac{dv_1(t)}{dt} = -b_2v_1(t) + a_2(v_2(t,0) - v_1(t)), & t \geq 0, \\ \frac{\partial u_2(t,x)}{\partial t} = D_1 \frac{\partial^2 u_2(t,x)}{\partial x^2} - b_1u_2(t,x), & t \geq 0, x \in (0,1], \\ \frac{\partial v_2(t,x)}{\partial t} = D_2 \frac{\partial^2 v_2(t,x)}{\partial x^2} - b_2v_2(t,x) + c_0\tilde{u}_2(t-1,x), & t \geq 0, x \in (0,1], \end{cases} \quad (24)$$

where c is the constant in (16). Let Φ^+ be the delay operator associated to this modified system, i.e.

$$\Phi^+ := \begin{pmatrix} 0 & -c\delta_{-1} & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & c_0\delta_{-1} & 0 \end{pmatrix}. \quad (25)$$

Hence Φ^+ is positive, thus also the semigroup $\mathcal{T}^+ := (\mathcal{T}^+(t))_{t \geq 0}$ generated by

$$\mathcal{C}^+ := \begin{pmatrix} \mathcal{B} & \Phi^+ \\ 0 & \mathcal{G} \end{pmatrix},$$

with $D(\mathcal{C}^+) = D(\mathcal{C})$ is positive (see [7, Theorem 3.2]).

Observe that

$$|\mathcal{T}(t)| \leq \mathcal{T}^+(t)$$

for all $t \geq 0$ and hence

$$\|\mathcal{T}(t)\| \leq \|\mathcal{T}^+(t)\|$$

for all $t \geq 0$ (see [18, Proposition II.4.1]). Thus, if we want to find conditions implying the exponential decay of the semigroup $(\mathcal{T}(t))_{t \geq 0}$, it is sufficient to find these conditions for the semigroup $(\mathcal{T}^+(t))_{t \geq 0}$. To this aim we first estimate the spectral bound of $\mathcal{C}$.

For each $\lambda \in \mathbb{C}$ with $\Re \lambda > \omega_0(\mathcal{U})$, where $\mathcal{U}$ denotes the evolution family associated to the heat semigroup (see (4)), we consider the bounded operator $E_\lambda : X \rightarrow L^p([-1,0], X)$ defined as

$$E_\lambda := \begin{pmatrix} \epsilon_\lambda & 0 & 0 & 0 \\ 0 & \epsilon_\lambda & 0 & 0 \\ 0 & 0 & \tilde{\epsilon}_\lambda & 0 \\ 0 & 0 & 0 & \hat{\epsilon}_\lambda \end{pmatrix}.$$

Here the bounded operators $\epsilon_\lambda : \mathbb{R} \rightarrow W^{1,p}([-1,0], L^1[0,1])$, $\tilde{\epsilon}_\lambda : L^1[0,1] \rightarrow L^p([-1,0], L^1[0,1])$ and $\hat{\epsilon}_\lambda : L^1[0,1] \rightarrow L^p([-1,0], L^1[0,1])$ are defined as

$$(\epsilon_\lambda x)(s) := e^{\lambda s} x, \quad \text{for } s \in [-1,0], x \in \mathbb{R},$$

$$(\hat{e}_\lambda f)(s) := e^{\lambda s} f, \quad \text{for } s \in [-1, 0], f \in L^1[0, 1],$$

and

$$(\tilde{e}_\lambda f)(s) := e^{\lambda s} U(s, 0) f = e^{\lambda s} T(-s) f, \quad \text{for } s \in [-1, 0], f \in L^1[0, 1],$$

respectively, for $(T(t))_{t \geq 0}$ as in (4).

The following corollary follows by the positivity of $(\mathcal{T}^+(t))_{t \geq 0}$ and by [7, Theorem 4.1].

Corollary 5.1. *For the spectral bounds of $\mathcal{G}$, $\mathcal{B} + \Phi^+ E_0$ and $\mathcal{C}^+$ the next relation holds*

$$\text{if } s(\mathcal{G}) \text{ and } s(\mathcal{B} + \Phi^+ E_0) < 0, \text{ then } s(\mathcal{C}^+) < 0. \quad (26)$$

For the proof of this corollary it is sufficient to observe (see [7, Theorem 4.1]) that we can rewrite $\mathcal{C}^+$ as

$$\mathcal{C}^+ := \begin{pmatrix} \mathcal{B} & 0 \\ 0 & \mathcal{G}_0 \end{pmatrix} \begin{pmatrix} Id & 0 \\ -E_0 & Id \end{pmatrix} + \begin{pmatrix} 0 & \Phi^+ \\ 0 & 0 \end{pmatrix},$$

where the operator $\mathcal{G}_0$ is the following matrix

$$\mathcal{G}_0 := \begin{pmatrix} \frac{d}{d\sigma} & 0 & 0 & 0 \\ 0 & \frac{d}{d\sigma} & 0 & 0 \\ 0 & 0 & G_0 & 0 \\ 0 & 0 & 0 & \frac{d}{d\sigma} \end{pmatrix},$$

with domain

$$D(\mathcal{G}_0) := (W^{1,p}[-1, 0])^2 \times D(G_0) \times W^{1,p}([-1, 0], L^1[0, 1]).$$

Here $(G_0, D(G_0))$ is the operator defined by

$$G_0 f := G f \text{ and } D(G_0) := D(G) \cap \{f(0) = 0\},$$

where $(G, D(G))$ is given by (23).

Remark 5.2. The negativity of the spectral bound of $\mathcal{G}$, $s(\mathcal{G})$, follows by (5) and by the fact that $\frac{d}{d\sigma}$ is the generator of the nilpotent translation semigroup $(T_l(t))_{t \geq 0}$ on $L^1[0, 1]$ (see, e.g., [3, Lemma 1.1.12]).

Thus the problem reduces to find conditions such that $s(\mathcal{B} + \Phi^+ E_0) < 0$. To this purpose, we first calculate

$$\mathcal{B} + \Phi^+ E_0 = \begin{pmatrix} -b_1 - a_1 & -c & a_1 \delta_0 & 0 \\ 0 & -b_2 - a_2 & 0 & a_2 \delta_0 \\ 0 & 0 & D_1 \Delta - b_1 & 0 \\ 0 & 0 & c_0 \delta_{-1} \tilde{e}_0 & D_2 \Delta - b_2 \end{pmatrix},$$

$$D(\mathcal{B} + \Phi^+ E_0) = D(\mathcal{B}),$$

where we recall

$$D(\mathcal{B}) := \left\{ \begin{pmatrix} x \\ y \\ f \\ g \end{pmatrix} \in \mathbb{R}^2 \times (W^{2,1}[0,1])^2 : L \begin{pmatrix} f \\ g \end{pmatrix} = \begin{pmatrix} x \\ y \end{pmatrix} \text{ and } f'(1) = g'(1) = 0 \right\}.$$

Hence the matrix $\mathcal{B} + \Phi^+ E_0$ is of the form

$$\begin{pmatrix} A & B \\ 0 & D_m \end{pmatrix}$$

on X . The operator B is a diagonal matrix and A and D_m are triangular matrices, i.e.

$$B := \begin{pmatrix} a_1 \delta_0 & 0 \\ 0 & a_2 \delta_0 \end{pmatrix},$$

$$A := \begin{pmatrix} -b_1 - a_1 & -c \\ 0 & -b_2 - a_2 \end{pmatrix},$$

and

$$D_m := \begin{pmatrix} D_1 \Delta - b_1 & 0 \\ c_0 \delta_{-1} \tilde{c}_0 & D_2 \Delta - b_2 \end{pmatrix}$$

with $D(B) := (W^{2,1}[0,1])^2$, $D(A) := \mathbb{R}^2$ and $D(D_m) := \left\{ \begin{pmatrix} f \\ g \end{pmatrix} \in (W^{2,1}[0,1])^2 : f'(1) = g'(1) = 0 \right\}$, respectively.

Now let $D \subset D_m$ with $D(D) = \text{Ker} L$, where, as above, $L : (W^{2,1}[0,1])^2 \rightarrow \mathbb{R}^2$ is defined by

$$L \begin{pmatrix} f \\ g \end{pmatrix} := \begin{pmatrix} \frac{f'(0)}{\beta_1} + f(0) \\ \frac{g'(0)}{\beta_1^*} + g(0) \end{pmatrix}.$$

Then for the matrix

$$\mathcal{D}_m := \begin{pmatrix} A & B \\ 0 & D_m \end{pmatrix},$$

with domain $D(\mathcal{D}_m) = \mathbb{R}^2 \times D(D_m)$ we have

$$\mathcal{B} + \Phi^+ E_0 \subseteq \mathcal{D}_m,$$

and the operator $\mathcal{B} + \Phi^+ E_0$ is *one-sided coupled* (see [7, Definition 1.1]). For the spectral bound of $\mathcal{B} + \Phi^+ E_0$ the following proposition holds.

Proposition 5.3. *The spectral bounds of the operators $\mathcal{B} + \Phi^+ E_0$, D and $A + BL_0$ satisfy*

$$s(\mathcal{B} + \Phi^+ E_0) < 0 \Leftrightarrow s(D) < 0 \text{ and } s(A + BL_0) < 0,$$

where $L_0 := (L|_{\ker D_m})^{-1} : \mathbb{R}^2 \rightarrow \ker(D_m) \subseteq (L^1[0,1])^2$.

The proof of this proposition follows again by [7, Theorem 4.1], rewriting $\mathcal{B} + \Phi^+ E_0$ as

$$\mathcal{B} + \Phi^+ E_0 := \begin{pmatrix} A & 0 \\ 0 & D \end{pmatrix} \begin{pmatrix} Id & 0 \\ -L_0 & 0 \end{pmatrix} + \begin{pmatrix} 0 & B \\ 0 & 0 \end{pmatrix}.$$

Now, let C_1 be the operator $C_1 \subseteq D_1 \Delta - b_1$ with domain

$$D(C_1) := \{f \in (W^{2,1}[0, 1], \mathbb{R}) : f'(1) = 0 \text{ and } f'(0) = -\beta_1 f(0)\}$$

and C_2 the operator $C_2 \subseteq D_2 \Delta - b_2$ with domain

$$D(C_2) := \{f \in (W^{2,1}[0, 1], \mathbb{R}) : f'(1) = 0 \text{ and } f'(0) = -\beta_1^* f(0)\}.$$

Proposition 5.4. *The spectral bound of the operator D satisfies the following property*

$$s(D) < 0 \Leftrightarrow s(C_1) < 0 \text{ and } s(C_2) < 0.$$

The next theorem gives a stability result for $(\mathcal{T}^+(t))_{t \geq 0}$. It follows by [10, Theorem 4.1] and by Corollary 5.1.

Theorem 5.5. *Assume that $s(C_1)$, $s(C_2)$, and $s(A + BL_0)$ are negative. Then $(\mathcal{T}^+(t))_{t \geq 0}$ and hence $(\mathcal{T}(t))_{t \geq 0}$ as well are uniformly exponentially stable.*

5.1 A Non Constant Diffusion in the Past

Assume now that the diffusion in the past of the mRNA presented in the cytoplasm is *not constant*, but is governed by the operators $A(t) := a(t)\Delta_D$, where $0 < a(\cdot) \in C(\mathbb{R}_-)$ and Δ_D is the Dirichlet Laplacian on $L^1[0, 1]$ as before. In this case the evolution family associated to these operators is given by

$$U(t, s) = e^{(\int_t^s a(\sigma) d\sigma) \Delta_D}. \quad (27)$$

Since the norm of the evolution family is

$$\|U(t, s)\| = e^{(\int_t^s a(\sigma) d\sigma) \lambda_0}, \quad (28)$$

where λ_0 denotes the largest eigenvalue of Δ_D , we can compute directly the growth bound of $(U(t, s))_{t \leq s \leq 0}$.

Proposition 5.6 (see [4], Example 5). *The growth bound of the evolution family $(U(t, s))_{-1 \leq t \leq s \leq 0}$ is given by*

$$\omega_0(\mathcal{U}) = \inf_{h \geq 0} \sup_{s+h \leq 0} \left(\frac{1}{h} \int_s^{s+h} a(\sigma) d\sigma \right) \lambda_0,$$

By Theorem 4.2 one has again that there exists a solution of (NDE) , but in this case the operator G defined in (23) is

$$Gf := f' + a(\cdot)f'', \quad \text{a.e.} \quad (29)$$

for appropriate f . For the stability, proceeding as before, we can find sufficient conditions in order to obtain that the solutions of (4.2) decay exponentially.

6 Acknowledgements

The author is supported by the Italian Programma di Ricerca di Rilevante interesse Nazionale "Analisi e controllo di equazioni di evoluzione nonlineari" (cofin 2000).

Part of this work has been written while the author was a Ph.D. student at the Department of Mathematics at the University of Tübingen and she takes the opportunity to thank Elena Babini, Klaus J. Engel, Rainer Nagel and Susanna Piazzera for many helpful discussions.

References

- [1] D.J. Aidley, P.E. Stanfield, *ION Channels. Molecules in Action*, Cambridge University Press, 2000.
- [2] H.T. Banks, J.M. Mahaffy, *Global asymptotic stability of certain models for protein synthesis and repression*, Quart. Appl. Math. **36**, 209–221 (1978).
- [3] A. Bátkai, S. Piazzera, *Semigroups for delay equations*, Research Notes in Mathematics **10**, 2005.
- [4] S. Brendle, R. Nagel, *Partial functional differential equations with nonautonomous past*, Disc. Cont. Dyn. Syst. **8**, 953–966 (2002).
- [5] S. Busenberg, J.M. Mahaffy, *Interaction of spatial diffusion and delays in models of genetic control by repression*, J. Math. Biol. **22**, 313–333 (1985).
- [6] V. Casarino, K.J. Engel, R. Nagel, G. Nickel, *A semigroup approach to boundary feedback systems*, Discrete Contin. Dyn. Syst. **12**, no. 4, 761–772 (2005).
- [7] K.J. Engel, *Positivity and stability for one-sided coupled operator matrices*, Positivity **1**, 103–124 (1997).
- [8] K.J. Engel, R. Nagel, *One-Parameter Semigroups for Linear Evolution Equations*, Springer-Verlag, 2000.
- [9] G. Fragnelli, G. Nickel, *Partial functional differential equations with nonautonomous past in L^p -phase spaces*, Diff. Int. Equ. **16**, 327–348 (2003).
- [10] G. Fragnelli, *A spectral mapping theorem for semigroups solving of PDEs with nonautonomous past*, Abstract and Applied Analysis, no. 16, 933–951 (2003).
- [11] G. Fragnelli, *Classical solutions for PDEs with nonautonomous past in L^p -spaces*, Bulletin of the Belgian Mathematical Society **11**, no. 1, 133–148 (2004).
- [12] B.C. Goodwin, *Temporal organization in cells*, Academic Press, New York, 1963.

- [13] B.C. Goodwin, *Oscillatory behavior of enzymatic control processes*, Adv. Enzyme Reg. **13**, 425–439 (1965).
- [14] F. Jacob, J. Monod, *On the regulation of gene activity*, Cold Spring harbor Symp. Quant. Biol. **26**, 389–401 (1961).
- [15] J.M. Mahaffy, *Periodic solutions for certain protein synthesis*, J. Math. Anal. Appl. **74**, 72–105 (1980).
- [16] J.M. Mahaffy, C.V. Pao, *Models of genetic control by repression with time delays and spatial effects*, J. Math. Biol. **20**, 39–58 (1984).
- [17] J.M. Mahaffy, C.V. Pao, *Qualitative analysis of a coupled reaction-diffusion model in biology with time delays*, J. Math. Anal. Appl. **109**, 355–371 (1985).
- [18] R. Nagel (ed.), *One-parameter Semigroups of Positive Operators*, Lect. Notes in Math **1184**, Springer Verlag, 1986.
- [19] V. Taglietti, C. Casella, *Elementi di Fisiologia e Biofisica della Cellula*, La Goliardica Pavese, 1991.
- [20] J.J. Tyson, *Periodic enzyme synthesis and oscillatory repression: Why is the period of oscillation close to the cell cycle time?*, J. Theor. Biol. **103**, 313–328 (1983).
- [21] J. Wu, *Theory and Applications of Partial Functional Differential Equations*, Springer-Verlag, 1996.

The irregular nesting problem involving triangles and rectangles

Loris Faina
Dipartimento di Matematica
Via L. Vanvitelli, 1
06123 Perugia (ITALY)
Fax # 39 75 585 5024
e-mail: faina@unipg.it

Abstract

This paper introduces a new geometrical method for a general two-dimensional irregular nesting problem; an efficient algorithm is derived. The validity of this algorithm is testified by several statistical tests, whose results include a numerical estimate of the asymptotic performance bound and the percentage of wastage of area utilization.

AMS Subject Classification: 90b06, 90c15, 90c42, 90c90

Key words and Phrases: Irregular Nesting, Asymptotic Performance, Method of Cones, Statistic Algorithm

1 Introduction

Cutting and Packing problems have applications accross the whole business activity, from design to manufacturing and distribution. These problems have a common logical structure that can be synthetized as follows. There are two groups of basic data, whose elements define geometric bodies in one or more dimensions: the stock of large objects, and the list of small items; the cutting and packing processes realize patterns being geometric combinations of small items assigned to large objects without overlaps. The residual pieces, that means figures occurring in patterns not belonging to small items, are usually treated as *trim loss*. The objective of most solution techniques is to minimize the wasted material.

The great majority of researchers have focused their attention on the case where both the objects and the items are rectangular and the orientation of the items is restricted (orthogonally oriented). Many of these papers are surveyed in Dyckhoff and Finke [10] and Dowsland and Dowsland [8]; for a wide discussion about the state of the research see also Faina [12] for the two-dimensional case and Faina [13] for the three-dimensional case.

However, there are many situations where either the pieces or the containing region are irregular in shape (see, for instance, textile, metal, or footwear industry). Problems involving irregular pieces comprise the most difficult class of packing problems (see Dowsland

and Dowsland [7] for a survey). Human intelligence is very able to face these problems. Indeed, it has been proved by several studies comparing layouts obtained by automatic processes with those produced by a human expert that the computer generated solutions are inferior in terms of minimizing the waste. However, in some industry the expertise required to beat the best packages often takes several years to acquire and in others the time required to produce a manual solution may be many times that of the automatic process.

Although the primary objective is to minimize the waste, often the cutting time or other constraints like minimizing the movements of the cutter between the cutting operations have a great relevance. Therefore good automatic computerised solutions are always welcome.

1.1 State of the research

There are two main approaches for facing the irregular nesting problem. The first produces one or a series of different layouts based on piece orderings and placement policies where the algorithm optimizes locally by means of shifts and rotations (see, for instance, Terno et al. [32]). The second approach is to produce an initial layout, which may be feasible but non-optimal or infeasible, and then to use small changes in order to improve it. This second approach may incorporate a metaheuristic technique such as simulated annealing or tabu search in order to allow non-improving moves (see, for instance, Reeves [29]).

In the spirit of first approach, the method developed by Amaral et al. [2] is noteworthy. Their method is based on a sophisticated process in order to find a suitable position for the next piece. Pieces are ordered in decreasing order of their areas and two different placement policies are used for large and small pieces, the rationale being that small pieces are used to fill holes left between the larger pieces in the layout. This reflects the strategy used by experts working on the problem manually.

Stoyan and Pankratov [31] find a way of packing identical copies of a polygon in a rectangular sheet using regular patterns allowing just two orientations.

Some researchers have considered an interesting alternative in which the irregular pieces are nested inside other more regular shapes, and these simpler shapes are then packed into the available area. The most popular shape for nesting is the rectangle. Freeman and Shapira [15] studied the problem of enclosing a piece, whose orientation is not restricted, in a rectangle of minimal area.

If many of the pieces are far from rectangular in shape a more effective option is to nest more than one piece together, as was developed by Adamowics and Albano [1].

An alternative to the use of rectangles is to nest all the pieces into identical polygons which can be used to tile the plane. Dori and Ben-Bassat [6] and Karoupi and Loftus [21] based their packing algorithms on tessellations using hexagons.

In the spirit of second approach, Ismail and Hon [18] studied the blanks nesting problem, but using only two pieces and allowing 180° degree of rotation only, using a genetic algorithm. Jain et al. [20], using simulated annealing, studied the same problem but where more than two pieces are involved, and these may take on any orientation. In both

these methods, the fitness value of each solution is a combination of the area utilisation and an overlap penalty.

Many other authors used simulated annealing, genetic algorithm or tabu search for facing the irregular nesting problem. We quote Lutfiyya et al. [23], Marques et al. [24], Oliveira and Ferreira [26], Dowsland and Dowsland [9], Blazewicz et al. [4], George et al. [17], Jacobs [19].

Milenkovic et al. [25] and Li and Milenkovic [22] considered the problem of planning motions of the pieces in an already existing layout to improve the layout in some fashion: increase efficiency by shortening the length, open space for new pieces without increasing the length, or eliminate overlap among pieces in an overlapping layout. This type of improvement strategy is also used by Stoyan et al. [30].

The method used to determine the geometry of the situation is essential in order to implement any of these algorithms. Depending on the strategy used, it is necessary to find the distance between two pieces along a given vector, the set of feasible positions for one piece with respect to another, or the existence of an area of overlap between any two pieces. The complexity of the calculations involved depends on the types of pieces considered and on the accuracy required.

For Qu and Sanders [28] pieces are approximated by unit squares on a rectangular grid. Preparata and Shamos [27] used basic trigonometric formulae to do these calculations, like the intersection of two lines or intersection of two polygons.

Things are even worse when non-convex shapes are involved. One possibility is to cut the shape into a number of convex pieces, as suggested by Cuninghame-Green [5].

One of the most used concepts in determining the interaction between two pieces is the concept of the no-fit polygon. This was first introduced by [1] as an efficient way of determining the minimum enclosure for a cluster of two or more pieces.

An important practical problem linked to the irregular cutting problem is to cut out a given shape or design from a given piece of parent material. Bhadury and Chandrasekaran [3] studied the problem to find a sequence of guillotine cuts to cut out a convex polygon P_{in} from a convex polygon $P_{out} \supset P_{in}$, such that the total length of the cutting sequence is minimized.

1.2 The aim of this paper

We study the problem of nesting triangles and rectangles on a strip with a fixed width and virtually infinite length, minimizing the length required. This is the first stage of a study whose goals is to nest convex polygons of any number of edges.

The limited number of sides permitted here is motivated by the exceptional accuracy of the calculations which allows us, in some cases, to get an optimal solution. Indeed, a peculiar characteristic of our algorithm *cone_{2d}* is the capability to solve puzzles and this gives many assurances on the validity of the algorithm.

We use a strategy which is an evolution of the geometrical method introduced in [12, 13, 14] for rectangular pieces, and that has been theoretically proved to lead to an optimal solution. This new geometrical procedure generates a particular finite class of placings to which a statistical global optimization algorithm, based on simulated annealing,

is applied. In some relevant cases, this reduction scheme is proved to lead to an optimal solution (see Section 3.1).

The pieces are considered one at a time and no overlap is tolerated in any stage of the process. In each stage of the allocation process, there are some free zones of conical shape where the pieces can be nested in such a way that no check for overlapping is necessary (with great saving of time!).

We consider any number of different pieces, and these may take any orientation; no a priori ranking rule is necessary.

After generating a starting feasible layout, the algorithm *cone_{2d}* uses small changes in order to improve it, incorporating a fully statistical technique based on simulated annealing.

We do not have a mathematical proof for the optimality of this algorithm except for the capability to solve puzzles. However, we believe that there is a substantial difference, on a principle level, between developing an algorithm whose layout is as close as possible to an optimal solution (and the capability to solve puzzles is the minimum request in this sense) and the description of a mere heuristic procedure which, at most, incorporates some metaheuristic technique, but having no care at all of the theoretical validity of the used algorithms. For instance, simulated annealing is an algorithm which has deep theoretical foundations, but in most of the papers quoted in the reference section which assert to use simulated annealing we see no traces of a proof or, at least, a discussion about the convergence of the algorithm presented and/or any relations with an optimal solution.

For sure it is difficult to prove the asymptotic convergence of simulated annealing to an optimal solution of an irregular nesting problem. Indeed we did not get it; but, at least, we prove the asymptotic convergence of the algorithm *cone_{2d}* to an optimal solution each time the items involved form a perfect puzzle.

In experimental sciences, the results, for being valid, should be at least reproducible; in mathematics, the results should be at least theoretically proved; operational research is in between these two philosophiae, and we believe that it is our duty to prove as much as possible, giving motivations for the unproved and reasons for the conjectures.

We will show the validity of the algorithm *cone_{2d}* by reporting the results of many tests based on randomly generated initial set of conditions, valuated with several parameters like the worst-case performance ratio, the asymptotic performance bound and the percentage of wastage of area utilisation. In some cases, a graphic representation of the solution is also presented, in order to show the difficult link between the parameters and the validity of the algorithm *cone_{2d}* (see figure 9).

2 Statement of the Problem

Let S be a rectangular strip, with a fixed width w and unbounded length; consider S embedded into a two dimensional Cartesian frame, in such a way that the lower left side corner of S coincides with the origin and each side of S is parallel to a reference frame axis, i.e. $S = \{(x, y) : 0 \leq y \leq w, x \geq 0\}$.

Let Q be the set of all closed triangular and rectangular items contained in S . Given

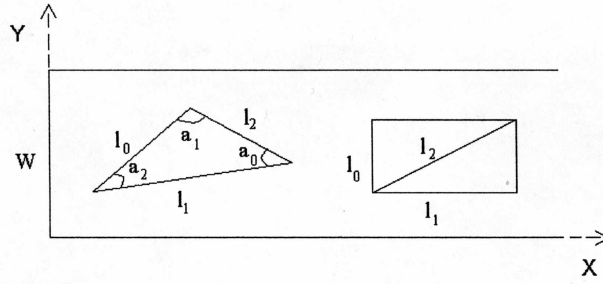


Figure 1: A bit of notations about the items.

an item $q \in Q$, we denote by $l_0(q)$ and $l_1(q)$ its width and length respectively, and by $l_2(q)$ its diagonal, if q is a rectangle; otherwise, we denote by $l_0(q), l_1(q), l_2(q)$ the three edge lengths of q and by $a_0(q), a_1(q), a_2(q)$ the angles opposed to $l_0(q), l_1(q)$, and $l_2(q)$ respectively. For the sake of shortness, when there will be no risk of confusion, we will write l_i, a_i instead of $l_i(q), a_i(q)$ (see figure 1).

Let $L = \{q_1, \dots, q_n\} \subset Q$ be a list of n items. A geometrical combination of the items of L is called a feasible placing P if the interior of the items do not overlap.

Let $\Pi(L, S)$ be the set of all the feasible placings P of L contained in S . If $P \in \Pi(L, S)$, then the length of the strip S containing P is called the length of the placing P , and it is denoted by $len(P)$.

The two-dimensional irregular nesting problem consists in finding a feasible placing $P_{opt}^L \in \Pi(L, S)$ which has the minimal length, i.e.

$$(2.1) \quad len(P_{opt}^L) = \min_{P \in \Pi(L, S)} len(P).$$

Clearly, this is a natural generalization of the regular two dimensional nesting problem involving only rectangles, and therefore it is NP-hard (see [16]); this means that, in general, optimal solutions are not known. From the statement of the problem (2.1) we derive that no restrictions are made neither on the orientation of the items nor on the number of different items.

2.1 The Method of the Cones

For placing an item, any nesting algorithm, sooner or later, should face the following questions: Which are the possible positions where the item could be located? Does the item fit into the strip? Is the item overlapped to the already placed items?

Mainly, there are two leading strategies: to store a lot of informations about some special positions on the strip, which permit to reduce the problem of overlapping; or to locate approximately an initial position of the item and then to apply some strategy for eliminating the possible overlaps. Both the strategies have their strenghts and weaknesses: the first strategy implies the storage of a lot of positions to avoid the penalty of a poor placement, but every position implies a lot of computations to do; the second strategy has the heavy duty to check and eliminate the possible overlaps; but even a simple *good* routine for checking the overlapping between two triangles is everything except easy.

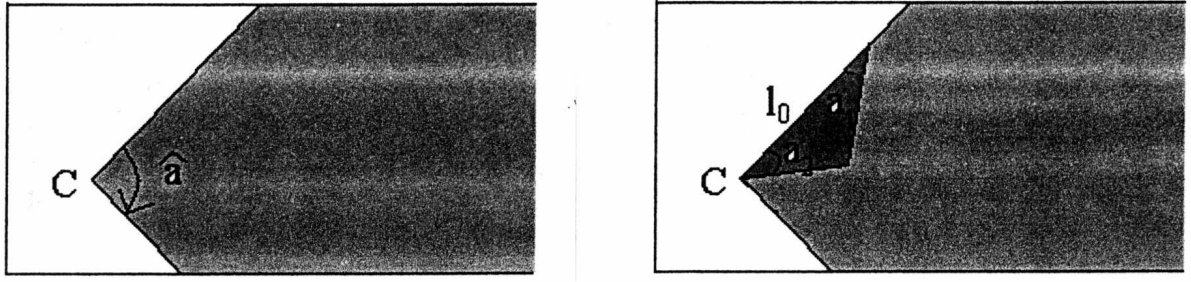


Figure 2: left: Description of a conical zone; right: Placement of an item into a conical zone.

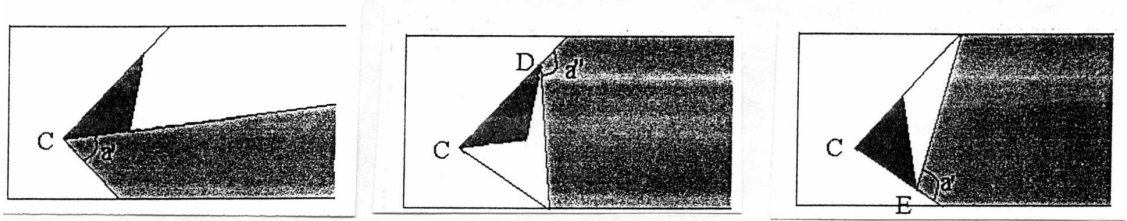


Figure 3: The conical zone used for the placing of an item is replaced with two new conical zones: case of a triangular item.

We cannot say that one strategy is better than the other. In the papers [12, 13], we developed algorithms for the regular cutting and packing problems in the spirit of the first strategy. We operate this choice also in the present two dimensional irregular nesting problem.

The strategy of our method is to determine a certain number of zones where it is possible to place an item without troubling about overlaps. Since these zones will have a conical shape, we will call this geometrical method the *method of the cones*. The method of the cones develops through three stages: the generation of the conical zones; the description of the strategy for placing an item inside such conical zones; and, each time a new item is placed, a reset policy for revising all the conical zones in order to preserve their properties.

Assume to have stored some conical zones; we describe now what happens when an item is placed in a particular conical zone. Suppose first to have a triangular item. Let C be the vertex and $\hat{a}$ be the amplitude of the conical zone (see figure 2-left), and suppose that the angle a_2 of the triangle q is smaller than $\hat{a}$. Then, it is possible to place the triangle in the conical zone; we put the corner relative to the angle a_2 on the point C and lean the side l_0 on the lower side (with respect to the clockwise orientation) of the cone (see figure 2-right). If the triangle remains into the strip, we find an admissible placing for the triangular item q . Now two new conical zones are naturally determined: the first conical zone has the vertex in C and amplitude equal to $a' = \hat{a} - a_2$ (see figure 3-left); the second conical zone has the vertex in D and amplitude equal to a'' (see figure 3-centre). It is to underline that the amplitude of the second conical zone is, in general, smaller than $\pi - a_1$, because we should guarantee that no other items occupy that conical zone. In the case $\hat{a} = a_2$, the first conical zone has vertex in E and amplitude equal to a' (see figure 3-right). In the same spirit, figure 4 indicate how to operate with a rectangular item.

Since the conical zones have to be truncated for fitting into the strip S , some tuning

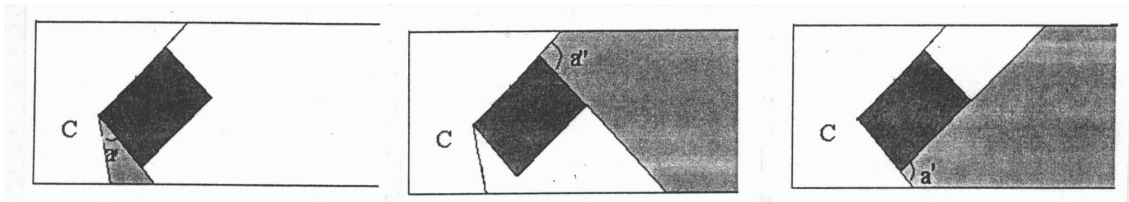


Figure 4: Generation of two new conical zones: case of a rectangular item.

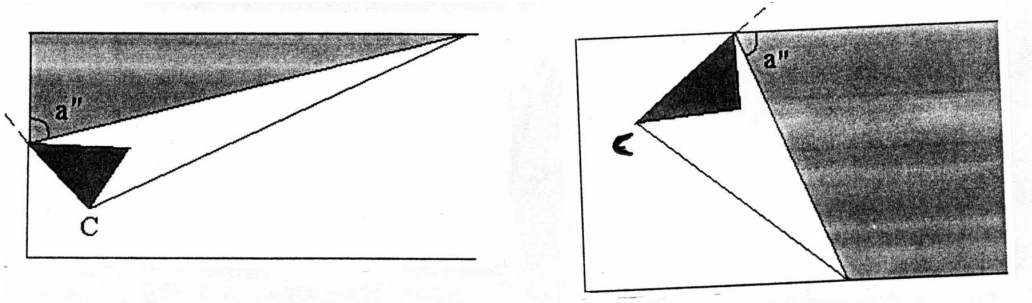


Figure 5: The amplitude of the new conical zones has to be tuned when their vertices lie on the border of the strip.

is necessary when the vertices of the new conical zones lie on the border of the strip (see figure 5).

Clearly each time we introduce a new item, all the conical zones already stored should be revised; first we check if the new item intersects the interior of a conical zone and, in the positive case, we reduce the amplitude of that conical zone in the spirit of figure 6.

Since for the choice of the sequence of conical zones where to check for the placing of the current item the algorithm *cone_{2d}* uses a metaheuristic technique, we postpone the complete description of the method of the cones to the subsection 3.1.

3 The base of the nesting algorithm: the simulated annealing

Simulated annealing algorithms are based on the analogy between the simulation of the annealing of solids and the problem of solving large combinatorial optimization problems.

These algorithms continuously transform the current configuration into another one by a small perturbation. The new proposed configuration, called *transition*, is accepted

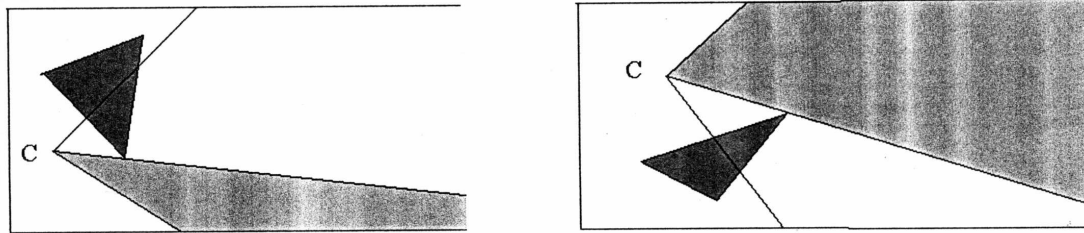


Figure 6: Reduction of the amplitude of the other conical zones after the placing of a new item.

criterion is strongly guided by a control parameter which is lowered step by step, and which reduces the number of transitions accepted as the algorithm proceeds.

It can be proved (see for example Faina [12] for an essential bibliography) that if any new transition from a starting configuration has the same probability to be generated then this process leads to an optimal solution with probability 1 as the control parameter goes to zero. It is important to underline that the solution obtained by simulated annealing does not depend on the initial configuration.

This framework solves the problem from a theoretical point of view; but asymptotic convergence is attained possibly in infinite time. Therefore, in any implementation of simulated annealing, asymptotic convergence can be only approximated, and the simulated annealing turns into a heuristic algorithm.

Indeed, the number of transitions for each value c_k of the control parameter, for instance, must be finite and $\lim_{k \rightarrow +\infty} c_k$ can only be approximated in a finite number of values c_k . Due to these approximations the algorithm is not guaranteed to find a global minimum with probability 1.

3.1 The nesting algorithm *cone_{2d}*

In the notation introduced in this section, a feasible placing $P \in \Pi(L, S)$ is considered a configuration and the length of the placing P , $len(P)$, is the cost functional.

Given a list L , the items are placed one by one following a certain order that could be as simple as that in which the items appear in the list (the initial ordering of the list is not important, since simulated annealing does not depend on the initial configuration).

At the very beginning, we have only one conical zone with vertex in the origin and amplitude equal to $\frac{\pi}{2}$. In general, we have to consider six different positions for valuating the possibility to locate a triangular item into a conical zone, and two different positions in the case of a rectangular item. Note that an item having at least a side not greater than w can be always placed into a conical zone with vertex lying on the x -axis and amplitude equal to $\frac{\pi}{2}$.

The first item is located, choosing randomly one among all the feasible positions of the item inside the initial conical zone. Then, two new conical zones are generated; we introduce always an additional conical zone with vertex in $(h_{max}, 0)$ and amplitude equal to $\frac{\pi}{2}$, where h_{max} is the length reached by the items already placed. This is a safety zone, which prevents the halt of the algorithm (see, for example, figure 7).

By induction, suppose to have located the i -th item; the geometrical method of the cones locates at most $(n + 1)$ conical zones where the $(i + 1)$ -th item could be located and a conical zone is chosen at random, until the $(i + 1)$ -th item is placed; then we reduce the already stored conical zones in order to eliminate overlaps with the $(i + 1)$ -th item; and so on.

In this way we obtain the initial configuration. Then, we operate a small perturbation to this initial configuration by modifying a little the order of the items, for instance by changing the position of any two items -at random- and by constructing a new configuration as shown above. Let $\Pi^*(L, S) \subset \Pi(L, S)$ be the set of all the finite feasible placings that can be obtained from the method of the cones. Since the transition mechanism depends only on the order of the items, and this order is perturbed at random, any new

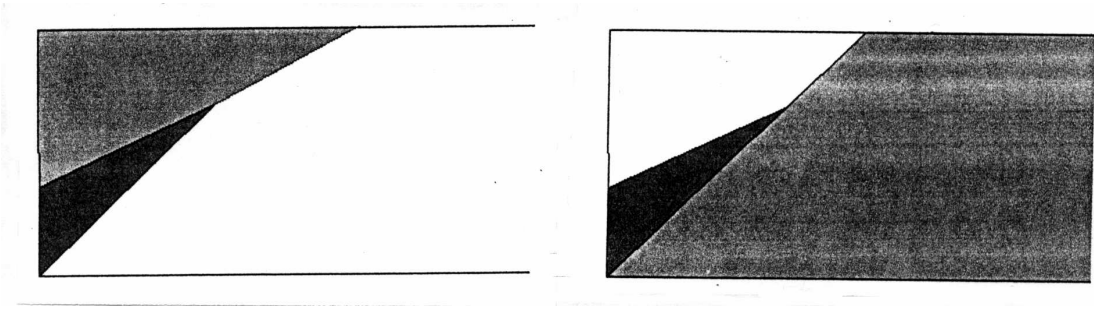


Figure 7: This is a simple case showing the necessity of the *safety zone*. The only two available conical zones (the shadowed areas) have the amplitude smaller than $\frac{\pi}{2}$, and therefore no rectangle can be placed at this stage.

configuration has the same probability of being generated. This property guarantees that simulated annealing will converge to a global optimal solution of the problem

$$(3.2) \quad \min_{P \in \Pi^*(L, S)} \text{len}(P).$$

In general, an optimal solution of (3.2) has a length bigger than an optimal solution of (2.1). It has proven difficult to find a relation between problem (3.2) and (2.1). However, if the items of L form a perfect puzzle, then it is always possible to find a suitable order for the items in such a way that the items' puzzle is represented by a placing of $\Pi^*(L, S)$. This property holds because, in the case of a puzzle, the items are located in the positions and with the orientations which are peculiar of the method of the cones.

Thus, in these particular but relevant cases, the two problems are equivalent, and therefore our algorithm converges asymptotically to an optimal solution of problem (2.1).

3.2 The implementation of the algorithm *cone_{2d}*

For the implementation of the algorithm we should specify the following parameters:

- the initial value of the control parameter, c_0 ;
- the final value of the control parameter, c_f (stop criterion);
- the number of new transitions generated for each value of the control parameter;
- a rule for changing the current value of the control parameter, c_k , into the next one, c_{k+1} .

A choice for these parameters is referred to as a '*cooling schedule*'.

The initial value c_0 is determined in such a way that virtually all transitions are accepted. The stop criterion which determines the final value of the control parameter, consists in fixing the number of values of c_k , say it F_k , for which the algorithm can be executed, or by terminating if consecutive decrements of the control parameter are identical for a given number of times, (a '*stalemate*' condition), whichever occur first.

For controlling the number of transitions generated, we have to face two conflicting requirements: the first is that for each value of the control parameter a minimal number of transitions should be accepted; the second is that, since transitions are accepted with

decreasing probability as $c_k \rightarrow 0$, we should aim to prevent the generation of large numbers of transitions for low values of c_k . A good compromise is to keep fixed the number of transitions generated, say it L_k , choosing a quite big number depending on the size of the problem, and to fix a maximum number of accepted transitions for each value of c_k , say it M_k . In this way, for values of c_k quite far from 0, the number of transitions generated will probably be much smaller than L_k (since a lot of transitions are accepted), while for small value of c_k this number will more likely be closer to L_k , assuring that a sufficient number of transitions are accepted. Thus, the number of accepted transitions is substantial for each value of c_k , and there is little waste of computational time for the initial steps.

For decrementing the control parameter, we use the rule

$$c_{k+1} = \alpha c_k \quad \text{for } k \geq 0,$$

where α is a constant smaller than but close to 1.

The cooling schedules depends on the number of items to place and on the objective to achieve. Indeed, heavier cooling schedules (this means that c_0, F_k, L_k, M_k are big) gives better and almost optimal solutions of problem (3.2) but with a larger computational time, and vice versa. The aim of the following tests was to present cooling schedules which guarantee a good quality of the solutions in a reasonable computational time.

4 Numerical Tests

The computer codes for the algorithm *cone_{2d}* is written in the programming language C, under the LINUX operative system. The computational time is computed through a virtual timer counter, which increments both when the routine executes and when the system is executing on behalf of the routine. This choice means that the computational time measures the time spent by the routine both in user and kernel space. The platform used is a PENTIUM III 450 MHz.

The evaluation of the algorithm *cone_{2d}* was performed by means of a self-comparison test on the basis of randomly generated initial set of conditions. Let *cone_{2d}*(L) be the final placing obtained by the algorithm *cone_{2d}* for a given list of items L . The main performance measure of the nesting algorithm is the asymptotic performance bound, which characterizes the behaviour of the ratio of $\text{len}(\text{cone}_{2d}(L))$ over $\text{len}(P_{opt}^L)$; if there are constants α and β such that for any list of items L ,

$$(4.3) \quad \text{len}(\text{cone}_{2d}(L)) \leq \alpha \cdot \text{len}(P_{opt}^L) + \beta D_{min}^L,$$

where $D_{min}^L = \max_{b_i \in L} \{\text{len}(P_{opt}^{\{b_i\}})\}$, then α is called an asymptotic performance bound for the algorithm *cone_{2d}*.

The number D_{min}^L represents the minimal length necessary for placing each single item; it represents a necessary correction to the formula (4.3) in the case of relatively small lists of items, for avoiding an unfair big value of α even in the case of quasi-optimal solutions (see, for example, figure 8).

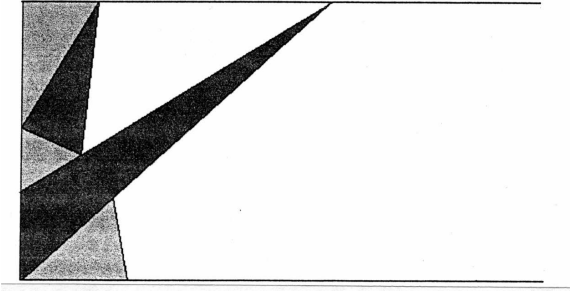


Figure 8: Optimal solution with a big percentage of wastage of area utilization.

Actually, since we do not have an estimate of $\text{len}(P_{opt}^L)$, we use a different form of formula (4.3) in the numerical estimations

$$\text{len}(\text{cone}_{2d}(L)) \leq \alpha^* \cdot \text{len}(A(L)) + D_{min}^L,$$

where $A(L)$ is the lower bound of $\text{len}(P_{opt}^L)$, equal to the length such that $A(L) \cdot w = \sum_{b_i \in L} \{\text{area of } b_i\}$ ($\beta = 1$); therefore α^* is called an approximate asymptotic performance bound.

We introduce another two parameters for a more complete evaluation: the worst-case performance ratio r_w

$$r_w = \frac{\text{len}(\text{cone}_{2d}(L))}{\text{len}(A(L))},$$

and the percentage of wastage of area utilization $A_{\%}$,

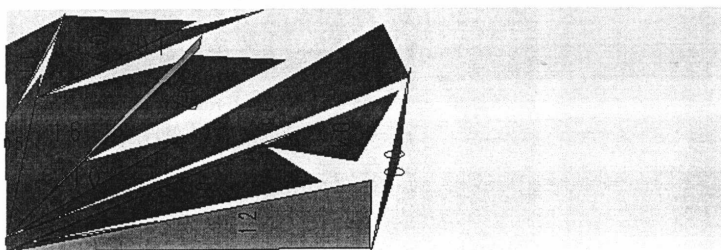
$$A_{\%} = \frac{\text{len}(\text{cone}_{2d}(L)) - \text{len}(A(L))}{\text{len}(\text{cone}_{2d}(L))} \cdot 100.$$

A comparison between r_w and α^* explains better the role of D_{min}^L . Clearly, the two value get closer as soon as the list of items becomes larger and larger.

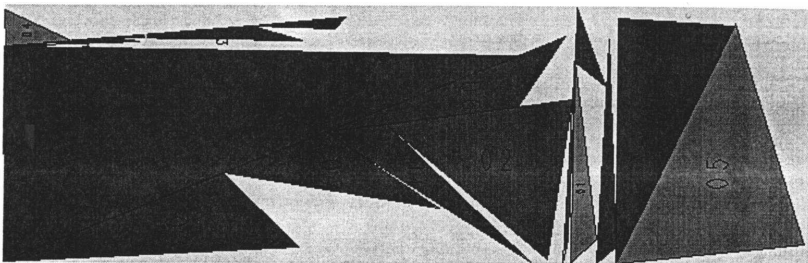
Furthermore, it is to underline that $A_{\%}$ does not coincide with the true wastage of area in comparison with the optimal solution, but it is, in general, much bigger. This means that a small value of $A_{\%}$ implies a good performance of the algorithm; but a relatively big value of $A_{\%}$ (or r_w) should not be necessarily interpreted as a failure of the algorithm. Some example reported in figure 9 help understanding better this situation.

The numerical tests have been performed as follows:

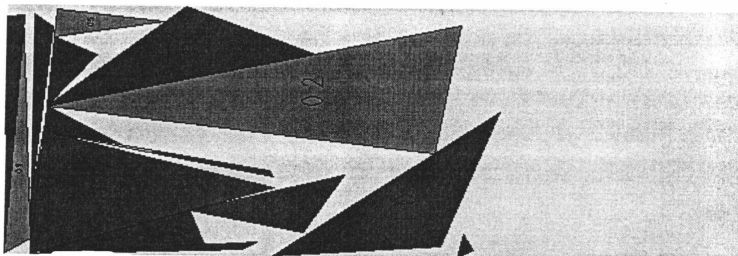
- each test has been executed on the basis of 100 runs;
- in each test, we have:
 - a fixed number of items (between 8 and 1000) whose parameters are uniformly generated random numbers: in the case of a rectangle, the two sides are between $\frac{w}{10}$ and w ; in the case of a triangle, one side is between $\frac{w}{10}$ and w , another side is between $\frac{w}{10}$ and $2w$ and the angle between them has an amplitude between 0.2 and $\frac{\pi}{2}$;
 - a fixed percentage of triangles and rectangles in the list of items;
 - each run has as output the numbers α^* , r_w , $A_{\%}$, and the computational time needed t ;



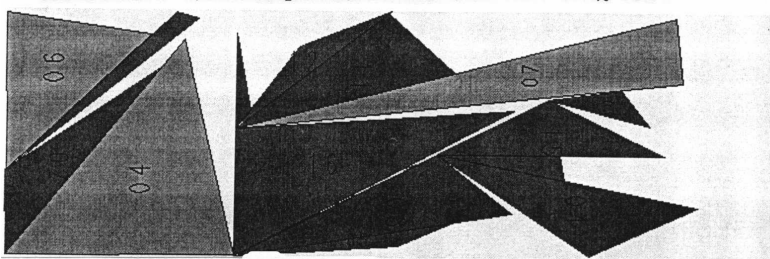
items=16, % triang=100, $\alpha^* = 0.16$, $r_w = 1.43$, $A_{\%} = 30$



items=16, % triang=100, $\alpha^* = 0.55$, $r_w = 1.43$, $A_{\%} = 30$



items=16, % triang=100, $\alpha^* = 0.59$, $r_w = 1.52$, $A_{\%} = 34$



items=16, % triang=100, $\alpha^* = 0.65$, $r_w = 1.42$, $A_{\%} = 29$

Figure 9: Some outputs of the algorithm *cone_{2d}*; the big values of $A_{\%}$ and r_w do not correctly express the quality of the solution obtained.

Test type	Number of runs	c_0	α	Max n. conf. generated (L_k)	Max n. conf. accepted (M_k)	Max n. decr. control param. (F_k)	Stalemate
<i>A</i>	100	50	0.7	16000	1000	30	20
<i>B</i>	100	10	0.9	12000	1000	20	10
<i>C</i>	100	10	0.9	16000	1000	30	20
<i>D</i>	100	10	0.9	12000	1000	30	10
<i>E</i>	100	10	0.9	12000	1000	20	10
<i>F</i>	100	50	0.7	12000	800	10	10
<i>G</i>	100	1	0.9	100	10	1	1
<i>P</i>	1	10	0.8	20000	1000	20	5

- the output of each test returns the average value of α^* , r_w , $A\%$ and t , plus their best/worst performances.

The Table 1 contains the cooling schedules of all the test performed.

5 Remarks on the Numerical Tests

The tests have the purpose to value the algorithm *cone_{2d}* with respect to very general conditions, where the lists of items contain every kind of pieces, from that very big, through that very thin and long, to that very small. These conditions are very difficult; the algorithm *cone_{2d}* has no shortcuts of any kind.

The complexity of these problems, mainly measured by the number of stored positions and feasible placings for each item, start to become uncontrollable as soon as the number of items exceeds 16; notwithstanding the results up to 128 items are still reasonable from all points of view (see Table 2).

Concerning numerical results for instances with many triangles, it should be said that the number of triangles influences the size of feasible region, since for each triangle and for each cone there are six possible placements. Therefore such instances are more difficult to be solved compared to instances with the same number of items but more rectangles, and this fact may explain the reduced performance of the algorithm. In the case of assence of triangles, the algorithm *cone_{2d}* is compared with the algorithm *zone_{2d}* developed by the author in [14]; *zone_{2d}* treats only rectangles, therefore it is an important comparison for *cone_{2d}*. These results are shown in Table 4 (see [14] for the values related to *zone_{2d}*); clearly, *zone_{2d}* is better than *cone_{2d}*, but *cone_{2d}* is not that bad. The results for 64 and 128 items show better the penalization of *cone_{2d}* that "does not know it has no triangles!".

Figure 10 shows an example of perfect puzzle solved by *cone_{2d}*; from the dimensions of the items involved, one can extrapolate the complexity of the puzzle, even in this case of 9 items only.

Finally, Table 3 reports the results of *cone_{2d}* stopped after very few iterations; this separate test shows the quality of the initial placings of the algorithm and the quality's limitations of the final placings of *cone_{2d}* in the case of very large list of items (within a computational time of few seconds).

Table 2: Tests for the algorithm $cone_{2d}$.

Number of items	% of Triangles	Test type	α^*		Perf. Ratio		$A\%$		Comp. Time in sec.	
			mean	min/max	mean	min/max	mean	min/max	mean	min/max
8	100	C	0.65	0.10/1.08	1.51	1.26/2.60	33.05	21.09/61.59	9.18	6.32/11.80
	50	A	0.78	0.43/1.03	1.26	1.06/1.59	20.69	6.26/37.09	46.17	8.59/72.04
	50	B	0.81	0.45/1.12	1.30	1.10/1.59	22.74	9.16/37.30	5.23	3.06/6.80
	0	C	0.79	0.64/1.02	1.10	1.06/1.25	9.82	6.15/20.32	7.26	5.26/8.67
16	100	C	1.00	0.59/1.56	1.52	1.27/2.10	33.76	21.61/52.46	22.29	18.26/29.25
	100	D	1.03	0.71/1.67	1.55	1.32/2.10	35.39	24.62/52.42	11.35	8.39/14.23
	50	D	1.01	0.79/1.26	1.31	1.12/1.55	23.49	11.07/35.86	11.9	8.28/17.55
	0	D	0.94	0.84/1.06	1.11	1.06/1.21	10.64	6.45/17.47	10.32	8.19/12.86
32	100	C	1.17	0.98/1.30	1.44	1.30/1.56	30.81	23.08/36.14	37.42	30.42/47.37
	100	F	1.20	1.05/1.37	1.46	1.34/1.68	31.87	25.56/41.56	20.94	16.10/24.89
	50	E	1.10	0.99/1.24	1.27	1.13/1.42	21.57	11.82/29.57	42.21	31.21/54.83
	0	E	1.03	0.98/1.08	1.12	1.07/1.17	10.87	7.29/15.08	38.30	29.84/49.41
64	100	E	1.30	1.17/1.45	1.44	1.32/1.58	30.72	24.28/36.99	116.6	91.9/144.3
	50	E	1.17	1.07/1.28	1.27	1.18/1.39	21.31	5.25/28.47	148.1	107.2/190.7
	50	F	1.19	1.11/1.30	1.28	1.20/1.40	22.39	16.96/28.63	99.8	77.98/127.0
	0	E	1.08	1.05/1.12	1.13	1.10/1.17	11.84	9.25/14.80	137.2	108.0/159.8
128	100	E	1.36	1.29/1.46	1.44	1.36/1.54	30.66	26.83/35.11	415.6	356.2/490.4
	50	E	1.22	1.17/1.28	1.27	1.22/1.34	21.50	18.25/25.56	568.6	466.4/698.6
	50	F	1.23	1.17/1.32	1.28	1.22/1.36	22.38	18.33/26.94	332.9	285.0/392.0
	0	E	1.12	1.09/1.16	1.44	1.12/1.18	12.99	10.61/15.81	533.7	452.9/612.9

Table 3: Tests for the algorithm $cone_{2d}$ stopped after few iterations.

Number of items	% of Triangles	Test type	α^*		Perf. Ratio		$A\%$		Comp. Time in sec.	
			mean	min/max	mean	min/max	mean	min/max	mean	min/max
32	100	G	1.71	1.38/2.24	1.99	1.66/2.59	49.52	39.93/61.46	0.02	0.02/0.04
	50	G	1.42	1.14/1.79	1.59	1.30/1.98	36.95	23.21/49.63	0.03	0.02/0.04
	0	G	1.18	1.08/1.29	1.26	1.16/1.39	21.13	14.25/28.35	0.02	0.02/0.04
64	100	G	1.80	1.58/2.39	1.94	1.71/2.50	48.29	41.63/60.78	0.08	0.07/0.11
	50	G	1.48	1.30/1.68	1.57	1.40/1.77	36.46	28.72/43.57	0.11	0.09/0.14
	0	G	1.24	1.16/1.42	1.29	1.20/1.48	22.78	17.08/32.59	0.09	0.08/0.11
128	100	G	1.86	1.67/2.20	1.93	1.75/2.28	48.24	43.00/56.21	0.29	0.24/0.36
	50	G	1.51	1.38/1.69	1.57	1.43/1.75	36.32	30.11/43.05	0.39	0.35/0.46
	0	G	1.28	1.21/1.35	1.30	1.24/1.38	23.45	19.36/27.62	0.34	0.31/0.39
1000	100	G	1.85	1.77/1.92	1.86	1.78/1.93	46.50	43.89/48.28	15.91	14.85/17.21
	50	G	1.50	1.44/1.57	1.51	1.45/1.58	34.08	31.08/36.91	21.30	20.02/22.76
	0	G	1.32	1.29/1.37	1.33	1.30/1.37	24.95	23.12/27.25	19.46	18.63/20.66

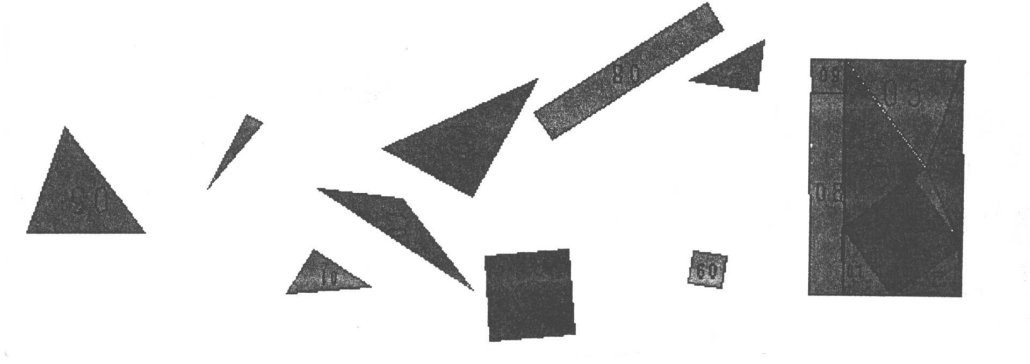


Figure 10: Example of perfect puzzle solved by $cone_{2d}$ in 30 seconds. We use the cooling schedule P .

Table 4: Comparison between $zone_{2d}$ and $cone_{2d}$ (only rectangular items!). We report just the mean values of the parameter measured; the first value is referred to $zone_{2d}$ and the second to $cone_{2d}$.

Number of items	α^*	$A\%$	Comp. Time in sec.
	mean	mean	mean
8	0.74/0.79	9.69/9.82	5.34/7.26
16	0.87/0.94	6.74/10.64	9.13/10.32
32	0.95/1.03	5.32/10.87	18.75/38.30
64	1.00/1.08	5.96/11.84	40.10/137.2
128	1.03/1.12	5.95/12.99	91.28/533.7

6 Conclusions

In this paper we propose a new geometrical method which reduces, in some cases, the two dimensional irregular nesting problem (involving triangles and rectangles) to a finite reduction scheme. The algorithm derived, $cone_{2d}$, is able to solve puzzles. In the case of absence of triangles, $cone_{2d}$ is compared with a fully regular algorithm (only for rectangular items), $zone_{2d}$, developed by the author in [14]. This comparison gives a "lower bound" for the algorithm $cone_{2d}$; the reasonable gap between each couple of parameter in Table 4 denotes that the increased complexity of the strategy for the algorithm $cone_{2d}$ due to the addition of triangles does not affect too much the quality of its final solutions, except for the computational time.

References

- [1] Adamowicz, M. and Albano, A., "Nesting two-dimensional shapes in rectangular modules", *Computer Aided Design*, vol. 8, n. 1, (1976), 27-33.
- [2] Amaral, C. and Bernardo, J. and Jorge, J., "Marker making using automatic placement of irregular shapes for the garment industry", *Computers and Graphics*, vol. 14, n. 1, (1990), 41-46.
- [3] Bhadury, J. and Chandrasekaran, R., "Stock cutting to minimize cutting length", *European Journal of Operational Research*, vol. 88, (1996), 69-87.
- [4] Blazewicz, J. and Hawryluk, P. and Walkowiak, R., "Using a tabu search approach for solving the two dimensional irregular cutting problem", *Annals of Operations Research*, vol. 41, (1993), 313-327.
- [5] Cuninghame-Green, R., "Cut out waste!", *O. R. Insight*, vol. 5, n. 3, (1992), 4-7.
- [6] Dori, D. and Ben-Bassat, M., "Efficient nesting of congruent convex figures", *Communications of the ACM*, vol. 27, n. 3, (1984), 228-235.
- [7] Dowsland, K. A. and Dowsland, W. B., "Solution approaches to irregular nesting problems", *European Journal of Operational Research*, vol. 84, (1995), 506-521.
- [8] Dowsland, K. A. and Dowsland, W. B., "Packing problems", *European Journal of Operational Research*, vol. 56, (1992), 2-14.
- [9] Dowsland, K. A. and Dowsland, W. B., "Heuristic approaches to irregular cutting problems", Working paper EBMS/1993/13, European Business Management school, UC Swansea, UK, (1993) .
- [10] Dyckhoff, H. and Finke, U., *Cutting and Packing in Production and Distribution*, Springer-Verlag, Berlin, (1992).

- [11] H. Dyckhoff, G. Scheithauer and J. Terno, "Cutting and packing", in: M. Dell'Amico, F. Maffioli and S. Martello Eds., "Annotated Bibliographies in Combinatorial Optimization", John Wiley & Sons, Chichester, 1997, pp. 393-412.
- [12] Faina, L., "An application of simulated annealing to the cutting stock problem", *European Journal of Operational Research* 144, 542-556, (1999).
- [13] Faina, L., "A global optimization algorithm for the three dimensional packing problem", *European Journal of Operational Research* 126, 340-354, (2000).
- [14] Faina, L., "A statistical approach to the cutting and packing problems", to appear.
- [15] Freeman, H. and Shapira, R., "Determining the minimum area encasing rectangle for an arbitrary closed curve", *Communications of the ACM*, vol. 81, n. 7, (1975), 409-413.
- [16] M. R. Gary, D. S. Johnson, "Computer and Intractability: A guide to the Theory of NP-Completeness", Freeman, San Francisco, CA, (1979).
- [17] George, J. A. and George, J. M. and Lamar, B. W., "Packing different-sized circles into a rectangular container", *European Journal of Operational Research*, vol. 84, (1995), 693-712.
- [18] Ismail, H. S. and Hon, K. K. B., "New approaches for the nesting of two-dimensional shapes for press tool design", *International J. of Production Research*, vol. 30, n. 4, (1992), 825-837.
- [19] Jacobs, S., "On genetic algorithms for the packing of polygons", *European Journal of Operational Research*, vol. 88, (1996), 165-181.
- [20] Jain, P. and Fenyves, P. and Richter, R., "Optimal blank nesting using simulated annealing", *Transaction of the ASME*, vol. 114, (1992), 160-165.
- [21] Karoupi, F. and Loftus, M., "Accommodating diverse shapes within hexagonal pavers", *International J. of Production Research*, vol. 29, n. 8, (1991), 1507-1519.
- [22] Li, Z. and Milenkovic, V., "Compaction and separation algorithms for non-convex polygons and their applications", *European Journal of Operational Research*, vol. 84, (1995), 539-561.
- [23] Lutfiyya, H. and McMillin, B. and Poshyammon, D. A. P. and Dagli, C., "Composite stock cutting through simulated annealing", *Mathematical Computer Modelling*, vol. 16, n. 1, (1992), 57-74.
- [24] Marques, V. M. M. and Bispo, C. F. G. and Sentieiro, J. J. S., "A system for the compactation of two-dimensional irregular shapes based on simulated annealing" in: *IECON-91 (IEEE)*, (1991), 1911-1916.

- [25] Milenkovic, V. and Daniels, K. and Li, Z., "Placement and compaction of non-convex polygons for clothing manufacture", in: *Proceedings of the 4th Canadian Conference on Computational Geometry*, St. John's, Newfoundland, August 10-14, (1992), 236-243.
- [26] Oliveira, J. F. C. and Ferreira, J. A. S., "Algorithms for nesting problems", in: Vidal, R. V. V. (Ed.) *Applied Simulated Annealing*, Lecture Notes in Economics and Mathematical Systems, vol. 306, Springer-Verlag, Berlin, (1993), 255-274.
- [27] Preparata, F. P. and Shamos, M. I., *Computational Geometry – An Introduction*, Springer-Verlag, Berlin, (1985).
- [28] Qu, W. and Sanders, J. L., "A nesting algorithm for irregular parts and factors affecting trim losses", *International J. of Production Research*, vol. 25, n. 3, (1987), 381-397.
- [29] Reeves, C. (Ed.), *Modern Heuristic Techniques for Combinatorial Problems*, Basil Blackwell, Oxford, (1993).
- [30] Stoyan, Y. G. and Novozhilova, M. V. and Kartashov, A. V., "Mathematical model and method of searching for a local extremum for the non-convex oriented polygons allocation problem", Working paper, Institute for Problems in Machinery, Ukrainian Academy of Sciences, (1993).
- [31] Stoyan, Y. G. and Pankratov, A. V., "Regular packing of congruent polygons on the regular sheet", *European Journal of Operational Research*, vol. 113, (1999), 653-675.
- [32] Terno, J. and Lindemann, R. and Scheithauer, G., *Zuschnittprobleme und ihre Praktische Loesung*, Harri Deutsch-Verlag, (1987).

Weak convergence and Prokhorov's Theorem for measures in a Banach space

Maria Cristina Isidori

Dipartimento di Matematica e Informatica, Università degli Studi di Perugia

Via Vanvitelli 1, 06123 - PERUGIA, ITALY

Fax: +39-75-5855024

e-mail: isidori@dipmat.unipg.it

Abstract

We prove the sufficient part of Prokhorov's Theorem in the vector setting.

AMS Subject Classification : 46G10, 28B05, 46B09

Key words and Phrases: *Weak convergence, semivariation, Rybakov control, prokhorov's Theorem.*

1 Introduction

As it is well known the Prokhorov's Theorem has a very important role in the study of cylindrical measures and in general in probability theory ([14]).

Vector versions of the Prokhorov's Theorem have been studied in the literature ([5], [7], [13], [14], [18]) in relation to several problems and applications (marginal problem, existence of the projective limit). However the authors either consider vector measures ranging in the positive cone of a Banach lattice or, even in the case of more general target spaces (as in [13], [14]), deal with special families of Radon measures.

In this paper we shall achieve one implication of Prokhorov's Theorem for a sequence of abstract measures in $bv(\Omega, \Sigma, X)$, where X is a general Banach space, along the ideas of Billingsley [2]. Indeed, we consider weak convergence in terms of the integral of continuous functions. The integral that is usually considered in the classic case ([2]), can be lifted to the vector case, since only continuous bounded integrands are taken into consideration (see for instance [18]).

In [4] Brooks and Martellotti show that in this particular class of integrands, the integral can be equivalently defined in the monotone sense, according to the scalar definition due to De Giorgi and Letta ([9]): this equivalence easily yields the validity of the vector version of Alexandroff Theorem ([12] IV.9.15).

The proof is rather long and elaborated, and goes along the scalar one as given in [2]: however one has to replace upper and lower limits by Cauchy nets, and also needs suitable extension results in the vector setting as those given by [17] [20].

The author wishes to thank Professor Anna Martellotti for the interesting and useful suggestions and the referees for their valuable suggestions.

2 Notations and Preliminaries

Let Ω be a separable metric space, Σ the Borel σ -algebra of Ω . Let d be the metric on Ω .

We denote by $C_b(\Omega)$ the set of all continuous and bounded functions $f : \Omega \rightarrow \mathbb{R}$.

If P_n and P are finite real measures, we say that P_n *weakly converges* to P , and we write $P_n \rightharpoonup P$, if for every $f \in C_b(\Omega)$

$$\lim_n \int_{\Omega} f dP_n = \int_{\Omega} f dP.$$

Weak convergence

A set $A \in \Sigma$ whose boundary ∂A satisfies $P(\partial A) = 0$ is called a *P-continuity set*.

Let X be a Banach space, X^* its dual and X_1^* the unit ball of X^* . Let $m : \Sigma \rightarrow X$ be a finitely additive measure, or simply a vector measure.

Definition 2.1 A vector measure m is said to be *strongly bounded* whenever given a sequence $(A_n)_n$ of pairwise disjoint members of Σ the limit $\lim_n m(A_n)$ is zero.

Theorem 2.2 ([16]) *A strongly bounded vector measure m is bounded.*

From Theorem 2.2 it follows that all countably additive measures, defined on σ -algebras, are strongly bounded.

Definition 2.3 The *variation* of m is the set function $|m| : \Sigma \rightarrow [0, +\infty]$ defined by

$$|m|(E) = \sup_{\pi} \sum_{A \in \pi} \|m(A)\|$$

where the supremum is taken over all partitions π of E into a finite number of pairwise disjoint elements of Σ .

If $|m|(\Omega) < +\infty$, then m is of *bounded variation* (b.v.).

We denote by $ca(\Omega, \Sigma, X)$ the set of all countably additive measures $m : \Sigma \rightarrow X$ and by $bvca(\Omega, \Sigma, X)$ the subset of all countably additive measures of bounded variation.

Definition 2.4 The *semivariation* of m is the function $\|m\| : \Sigma \rightarrow [0, +\infty]$ defined by

$$\|m\|(E) = \sup\{|x^*m|(E) : x^* \in X^*, \|x^*\| \leq 1\}.$$

If $\|m\|(\Omega) < +\infty$, then m is of *bounded semivariation*.

One can easily prove the following relationship:

$$\|m(E)\| \leq \|m\|(E) \leq |m|(E), \quad (1)$$

for every $E \in \Sigma$.

Definition 2.5 A set $A \in \Sigma$ will be called an *m-continuity set* if $\|m\|(\partial A) = 0$.

Definition 2.6 A finitely additive bounded measure m is *regular* if for every $E \in \Sigma$ and for every $\varepsilon > 0$ there exist an open set G and a closed set F with $F \subseteq E \subseteq G$ such that $\|m\|(G \setminus F) < \varepsilon$.

Lemma 2.7 ([8]) *Since Ω is a metrizable space, every countably additive measure is regular.*

Definition 2.8 Let m and $\nu : \Sigma \rightarrow [0, +\infty]$ be positive subadditive functions. We say that m is *absolutely continuous with respect to ν* , and we write $m \ll \nu$, if for every $\varepsilon > 0$ there exists $\delta(\varepsilon) > 0$ such that for every $E \in \Sigma$ with $\nu(E) < \delta$ it is $m(E) < \varepsilon$.

Definition 2.9 Let m and $\nu : \Sigma \rightarrow [0, +\infty]$ be positive subadditive functions. We say that m is *equivalent* to ν if they are absolutely continuous with respect to each other.

Definition 2.10 A positive subadditive function $\lambda : \Sigma \rightarrow [0, +\infty]$ is a *control measure* for m if λ is equivalent to $\|m\|$.

Theorem 2.11 ([3]) *If m is a strongly bounded vector measure, then m admits a control measure, which is finitely additive.*

Theorem 2.12 ([11]) *If m is a strongly bounded vector measure then there exists $x^* \in X^*$ with $\|x^*\| = 1$ such that $|x^*m|$ is a control measure for m .*

Weak convergence

In this case we say that $\lambda = |x^*m|$ is a *Rybakov control* for m .

Definition 2.13 ([4]) Let $f : \Omega \rightarrow [0, +\infty]$ be a measurable function.

f is $(\sim)$ -integrable with respect to m if $\|m\|(\{\omega \in \Omega : f(\omega) > t\})$ is Lebesgue integrable. In this case, as it is shown in [4], the X -valued vector function

$t \mapsto m(\omega \in \Omega : f(\omega) > t)$ turns out to be Bochner integrable on $[0, +\infty]$. In this case we set

$$\widetilde{\int_{\Omega}} f d\|m\| = \int_0^{+\infty} \|m\|(\{\omega \in \Omega : f(\omega) > t\}) dt,$$

and

$$\widehat{\int_{\Omega}} f dm = \int_0^{+\infty} m(\{\omega \in \Omega : f(\omega) > t\}) dt.$$

Also in [4] it is proved that if f is $(\sim)$ -integrable then, for every $A \in \Sigma$, $f1_A$ is $(\sim)$ -integrable.

Therefore we can set

$$\widehat{\int_A} f dm = \widehat{\int_{\Omega}} f 1_A dm,$$

and

$$\widetilde{\int_A} f d\|m\| = \widetilde{\int_{\Omega}} f 1_A d\|m\|.$$

These integrals satisfy the following properties:

$$(2.13.1) \quad \widehat{\int_{A \cup B}} f dm = \widehat{\int_A} f dm + \widehat{\int_B} f dm, \text{ where } A, B \in \Sigma, A \cap B = \emptyset;$$

$$(2.13.2) \quad \text{If } f = k \text{ in } A \text{ then } \widehat{\int_A} f dm = k \cdot m(A);$$

$$(2.13.3) \quad \text{If } f \leq g \text{ then } \widetilde{\int_{\Omega}} f d\|m\| \leq \widetilde{\int_{\Omega}} g d\|m\|;$$

$$(2.13.4) \quad \text{If } A \subset B \text{ then } \widetilde{\int_A} f d\|m\| \leq \widetilde{\int_B} f d\|m\|.$$

If f takes values in $\mathbb{R}$ we say that f is $(\sim)$ -integrable if f^+ and f^- are $(\sim)$ -integrable and we set

$$\widehat{\int_{\Omega}} f dm = \widehat{\int_{\Omega}} f^+ dm - \widehat{\int_{\Omega}} f^- dm,$$

and

$$\widetilde{\int_{\Omega}} f d\|m\| = \widetilde{\int_{\Omega}} f^+ d\|m\| - \widetilde{\int_{\Omega}} f^- d\|m\|.$$

Theorem 3.2 in [4] shows that every bounded function is $(\widetilde{\cdot})$ -integrable.

Definition 2.14 Let m_k and m be countably additive vector measures. We say that m_k *weakly converges* to m , and we write $m_k \rightharpoonup m$, if for every $f \in C_b(\Omega)$ and for every $x^* \in X^*$ it is

$$\lim_k x^* \left(\widehat{\int_{\Omega}} f dm_k \right) = x^* \left(\widehat{\int_{\Omega}} f dm \right).$$

If $\mathcal{A}$ is an algebra, we denote by $\sigma(\mathcal{A})$ the σ -algebra generated by $\mathcal{A}$.

Theorem 2.15 ([20]) *Let $\mu : \mathcal{A} \rightarrow X$ be a countably additive set function, where $\mathcal{A}$ is an algebra, such that for every monotone sequence $(A_n)_n \in \mathcal{A}$ there exists $\lim_n \mu(A_n)$. Then μ can be extended to $\sigma(\mathcal{A})$ in a countably additive way.*

Definition 2.16 A family of sets $\mathcal{L}$ is a *lattice* if $\mathcal{L}$ is closed under finite joints and meets, and $\emptyset \in \mathcal{L}$.

Let $\mathcal{L}$ be a lattice and $\lambda : \mathcal{L} \rightarrow \mathcal{X}$ be a function. We say that λ is *strongly additive* if the following hold:

- 1) $\lambda(A) + \lambda(B) = \lambda(A \cup B) + \lambda(A \cap B)$ with $A, B \in \mathcal{L}$;
- 2) $\lambda(\emptyset) = 0$.

Theorem 2.17 ([17]) *Every strongly additive set function λ on $\mathcal{L}$ has a unique extension to an additive set function defined on the ring generated by $\mathcal{L}$.*

Weak convergence

We recall that a subnet of a net $g : \Lambda \rightarrow X$ is the composition $g \circ \varphi$, where M is directed and $\varphi : M \rightarrow \Lambda$ is increasing and cofinal in Λ (i.e. for each $\lambda \in \Lambda$, there is some $\mu \in M$ such that $\lambda \leq \varphi(\mu)$).

3 The Prokhorov's Theorem for tight families in $bvca(\Omega, \Sigma, X)$

According to [2] we set the following definition:

Definition 3.1 A family $\Pi \subseteq bvca(\Omega, \Sigma, X)$ is *weakly sequentially compact* if for every sequence $(m_k)_k \subseteq \Pi$ there exists a subsequence $(m_{k_j})_j$ and a countably additive vector measure β such that $m_{k_j} \rightharpoonup \beta$.

Definition 3.2 A family $\Pi \subseteq bvca(\Omega, \Sigma, X)$ is *uniformly tight* if for every $\varepsilon > 0$ there exists a compact set $K_\varepsilon \subset \Omega$ such that

$$|m|(K_\varepsilon) > |m|(\Omega) - \varepsilon,$$

for every $m \in \Pi$.

Definition 3.3 A family $\Pi \subseteq bvca(\Omega, \Sigma, X)$ is *uniformly sequentially tight* if for every $(m_k)_k \subseteq \Pi$ there exists a subsequence $(m_{k_j})_j$ uniformly tight.

Theorem 3.4 (Prokhorov) *Let $\Pi \subseteq bvca(\Omega, \Sigma, X)$ be a family of measures such that*

- 1) Π is equibounded in the sense of the variations;
- 2) Π is uniformly sequentially tight;
- 3) for every compact subset S of Ω there exists a weakly sequentially compact set W_S in X such that $\{m(S) : m \in \Pi\} \subset W_S$.

Then Π is weakly sequentially compact.

Proof: Step 1 Fix $(m_k)_k$ in Π . We shall construct the weak limit β on Σ .

In order to do this, first we shall define a subsequence of indexes for which the measures and their variations simultaneously converge. This construction will coincide with that in [2] for the sequence $(|m_k|)_k$.

Since Π is uniformly sequentially tight, there exists a sequence of compact sets K_u with $K_1 \subset K_2 \subset \dots$ and a subsequence $(m_{k_j})_j$ such that

$$|m_{k_j}|(\Omega \setminus K_u) < \frac{1}{u},$$

for every $u \in \mathbb{N}$ and $j \in \mathbb{N}$. We will continue to denote (m_{k_j}) by (m_k) .

Let $D = \{x_j : j \in \mathbb{N}\}$ be dense in Ω and let S be the family $\{B(x_j, \rho) : j \in \mathbb{N}, \rho \in \mathbb{Q}^+\}$.

Let $\mathcal{H}$ consist of $\emptyset$ and of the finite unions of sets of the form $\overline{A} \cap K_u$, for $A \in \mathcal{S}$ and $u \geq 1$.

$\mathcal{H}$ is at most countable and closed with respect to the finite unions and every set of $\mathcal{H}$ is compact. We set

$$\mathcal{H}' = \{H_1 \cap H_2, H_i \in \mathcal{H}, i = 1, 2\}.$$

We define the function $\psi : \mathcal{H} \times \mathcal{H} \rightarrow \mathcal{H}'$ by $\psi(H_1, H_2) = H_1 \cap H_2$. Then ψ is onto and so $\mathcal{H}'$ is countable. Let $H \in \mathcal{H}'$. Then H is compact and by assumption $\{m_k(H), k \in \mathbb{N}\} \subset W_H$, for some relatively compact set W_H . Therefore there exists a subsequence of $(m_k(H))_k$ which weakly converges to an element of X . By a diagonal argument we can determine indexes k_j such that the following two limits exist:

$$\alpha(H) = (w) - \lim_j m_{k_j}(H); \quad (2)$$

$$\sigma(H) = \lim_j |m_{k_j}|(H). \quad (3)$$

Weak convergence

We assume that the subsequence $(m_{k_j})_j$ is weakly convergent.

We have

$$\|\alpha(H)\| \leq \lim_j |m_{k_j}|(H) = \sigma(H) \quad (4)$$

Claim 1) *Let $G \subset \Omega$ be an open set. Then $\{\alpha(H) : H \in \mathcal{H}, H \subset G\}$ is a strongly Cauchy net.*

Proof: Fix G open in Ω . We set

$$\Lambda_G = \{H \in \mathcal{H} : H \subseteq G\}.$$

The set Λ_G is non empty by definition of $\mathcal{H}$ and it is directed with respect to the inclusion.

Define

$$P(G) = \sup_{H \in \Lambda_G} \sigma(H) \quad (5)$$

for every open set G , and for $M \in \Sigma$ set

$$\gamma(M) = \inf\{P(G), G \text{ open and } M \subset G\}. \quad (6)$$

As in the proof of Prokhorov's Theorem ([2]), one proves that $P \equiv \gamma$ on the open sets and that it is countably additive on Σ .

Let $\varepsilon > 0$ be fixed; for every $H \in \mathcal{H}$, there exists an open set $G_0 \supset H$ such that

$$P(G_0) < \gamma(H) + \varepsilon,$$

and so from (6) we have

$$\gamma(H) > P(G_0) - \varepsilon = \sup_{T \subset G_0} \sigma(T) - \varepsilon > \sigma(H) - \varepsilon.$$

By the arbitrariness of ε , $\gamma \geq \sigma$ on $\mathcal{H}$. Moreover, for every open set G , there exists $H_0 \subset G$, $H_0 \in \mathcal{H}$, such that

$$0 \leq P(G) - \sigma(H_0) < \frac{\varepsilon}{2}.$$

Let $H_i \in \Lambda_G$ with $H_i \supset H_0$, $i = 1, 2$. Then we have

$$\begin{aligned} 0 &\leq \sigma(H_1) + \sigma(H_2) - 2\sigma(H_1 \cap H_2) \leq \\ &\leq \gamma(H_1) + \gamma(H_2) - 2\sigma(H_0) \leq 2[P(G) - \sigma(H_0)] < \varepsilon. \end{aligned}$$

Since H_1, H_2 and $H_1 \cap H_2 \in \mathcal{H}'$, from (3) it follows that there exists j^* such that for $j > j^*$

$$\sigma(H_1) + \sigma(H_2) - 2\sigma(H_1 \cap H_2) - |m_{k_j}|(H_1 \triangle H_2) < \varepsilon.$$

Therefore $|m_{k_j}|(H_1 \triangle H_2) < 2\varepsilon$, for $j > j^*$.

So, for every $x^* \in X_1^*$, it is

$$|x^*m_{k_j}|(H_1 \triangle H_2) < 2\varepsilon,$$

for every $j > j^*$.

Then we have

$$\begin{aligned} |x^*\alpha(H_1) - x^*\alpha(H_2)| &= \lim_j |x^*m_{k_j}(H_1) - x^*m_{k_j}(H_2)| \leq \\ &\leq \limsup_j |x^*m_{k_j}|(H_1 \triangle H_2) \leq 2\varepsilon. \end{aligned}$$

Hence the net $\{\alpha(H), H \in \mathcal{H}, H \subset G\}$ is strongly Cauchy. $\square$

Now for every open set G we can define

$$\beta(G) = \lim_{H \in \Lambda_G} \alpha(H) \tag{7}$$

Claim 2) *The set function defined in (7) is finitely additive.*

Proof: We observe first that α is finitely additive on $\mathcal{H}$. In fact if $H_1, H_2 \in \mathcal{H}$, for every $x^* \in X_1^*$ it is

$$x^*\alpha(H_1 \cup H_2) = \lim_j x^*m_{k_j}(H_1 \cup H_2) =$$

Weak convergence

$$\begin{aligned}
&= \lim_j x^* m_{k_j}(H_1) + \lim_j x^* m_{k_j}(H_2) - \lim_j x^* m_{k_j}(H_1 \cap H_2) = \\
&= x^*[\alpha(H_1) + \alpha(H_2) - \alpha(H_1 \cap H_2)].
\end{aligned}$$

By the arbitrariness of x^* we obtain

$$\|\alpha(H_1 \cup H_2) - [\alpha(H_1) + \alpha(H_2) - \alpha(H_1 \cap H_2)]\| = 0,$$

namely

$$\alpha(H_1 \cup H_2) = \alpha(H_1) + \alpha(H_2) - \alpha(H_1 \cap H_2). \quad (8)$$

Note that, if $H_1 \cap H_2 = \emptyset$, then from (8)

$$\alpha(H_1 \cup H_2) = \alpha(H_1) + \alpha(H_2).$$

Now we prove that β is finitely additive.

By definition of β we have $\beta(\emptyset) = 0$.

Let G_1, G_2 be two disjoint open sets. We define the family

$$\mathcal{H}_{G_1 \cup G_2} = \{H_1 \cup H_2 : H_i \in \mathcal{H}, H_i \subset G_i, i = 1, 2\}.$$

We want to prove that $\{\alpha(H_1 \cup H_2) : H_i \in \mathcal{H}, H_i \subset G_i, i = 1, 2\}$ is a subnet of the net $\{\alpha(H) : H \in \mathcal{H}, H \subset G_1 \cup G_2\}$. Obviously $\mathcal{H}_{G_1 \cup G_2}$ is directed with respect to the inclusion.

Let $\varphi : \mathcal{H}_{G_1 \cup G_2} \rightarrow \{H \in \mathcal{H} : H \subset G_1 \cup G_2\}$ be the identity function. φ is clearly monotone with respect to the inclusion. It remains to prove that φ is cofinal. Let $H \in \Lambda_{G_1 \cup G_2}$. Then H can be written as

$$H = \bigcup_{j=1}^p (\overline{B}(x_j, \rho_j) \cap K_{u_j}) \subset G_1 \cup G_2,$$

where $B(x_j, \rho_j) \in S$, the countable basis fixed in the beginning.

Define

$$F_1 = \{\omega \in H : d(\omega, G_1^c) \geq d(\omega, G_2^c)\}$$

$$F_2 = \{\omega \in H : d(\omega, G_2^c) \geq d(\omega, G_1^c)\}$$

As in the proof of Prokhorov's Theorem [2], it follows that $F_1 \subset G_1$ and $F_2 \subset G_2$. Moreover $H = F_1 \cup F_2$. Let $j \in \{1, \dots, p\}$ be fixed. The set $F_1 \cap \overline{B}(x_j, \rho_j) \cap K_{u_j}$ is compact in G_1 , and so there exist $y_1, \dots, y_n \in D$ and $r_1, \dots, r_n \in \mathbf{Q}^+$ such that

$$F_1 \cap \overline{B}(x_j, \rho_j) \cap K_{u_j} \subset \bigcup_{i=1}^n \overline{B}(y_i, r_i) \cap K_{u_j} \subset G_1.$$

Set $H_{1,j} = \bigcup_{i=1}^n \overline{B}(x_i, \rho_i) \cap K_{u_j}$. Obviously $H_{1,j} \in \Lambda_{G_1}$. Iterating this construction for $j \in \{1, \dots, p\}$, we have

$$F_1 = F_1 \cap H \subset \bigcup_{j=1}^p H_{1,j}.$$

Since $\mathcal{H}$ is closed under finite unions and every $H_{1,j} \subset G_1$, setting $H_1 = \bigcup_{j=1}^p H_{1,j}$, we obtain $F_1 \subset H_1 \subset G_1$, where $H_1 \in \mathcal{H}$. So we have proved that for every $H \in \Lambda_{G_1 \cup G_2}$ there exist $H_1 \in \Lambda_{G_1}, H_2 \in \Lambda_{G_2}$ such that $H \subset H_1 \cup H_2$, namely φ is cofinal. (Note that φ is cofinal even if G_1, G_2 are not disjoint sets). Then, by e) pag. 75 of [21], it is

$$\lim_{H_1 \cup H_2 \in \Lambda_{G_1 \cup G_2}} \alpha(H_1 \cup H_2) = \lim_{H \in \Lambda_{G_1 \cup G_2}, H \in \mathcal{H}} \alpha(H),$$

and so

$$\beta(G_1 \cup G_2) = \lim_{H_1 \cup H_2 \in \Lambda_{G_1 \cup G_2}} (\alpha(H_1) + \alpha(H_2)).$$

By definition of β it follows that for every $\varepsilon > 0$ there exist $\tilde{H}_1$ and $\tilde{H}_2$ such that for every $H'_i \supset \tilde{H}_i$, with $H'_i \in \Lambda_{G_i}$, $i = 1, 2$,

$$\|\beta(G_1 \cup G_2) - \alpha(H'_1 \cup H'_2)\| \leq \frac{\varepsilon}{3};$$

$$\|\beta(G_1) - \alpha(H'_1)\| \leq \frac{\varepsilon}{3};$$

$$\|\beta(G_2) - \alpha(H'_2)\| \leq \frac{\varepsilon}{3}.$$

Weak convergence

Thus β is finitely additive on open sets. $\square$

By (7) it follows that for every open set A in Ω there exists $H_A^\varepsilon \in \mathcal{H}$ such that $\|\beta(A) - \alpha(H)\| < \frac{\varepsilon}{4}$, for every $H \in \mathcal{H}$ with $H_A^\varepsilon \subset H \subset A$.

Define $\Lambda'_G = \{H' \in \mathcal{H}' : H' \subset G\}$. We shall show that

Claim 3) β can be equivalently defined as

$$\beta(G) = \lim_{H' \in \Lambda'_G} \alpha(H').$$

Proof: Let the open set G and $\varepsilon > 0$ be fixed. It is enough to show that there exists $H_G^\varepsilon \in \Lambda_G$ such that for every $H \in \mathcal{H}$ and $H' \in \mathcal{H}'$ with $H_G^\varepsilon \subset H \subset G$ and $H_G^\varepsilon \subset H' \subset G$ it is $\|\alpha(H') - \alpha(H)\| < \varepsilon$. Let G be an open set of Ω , and let $\varepsilon > 0$ be fixed. From (5) there exists $H_G^\varepsilon \subset G$, $H_G^\varepsilon \in \mathcal{H}$, such that

$$0 \leq P(G) - \sigma(H_G^\varepsilon) < \frac{\varepsilon}{4}.$$

Let $H' \in \mathcal{H}'$ be a subset of G . We will prove that there exists $H_0 \in \mathcal{H}$ such that $H' \subset H_0 \subset G$.

Since $H' \in \mathcal{H}'$, there exist $H_1, H_2 \in \mathcal{H}$ such that $H' = H_1 \cap H_2$. So H' is the finite union of sets of the form $\overline{A}_l \cap \overline{B}_l \cap K_{u_l} \cap K_{v_l}$, where $A_l, B_l \in S$, $u_l, v_l \geq 1$ and $1 \leq l \leq t$, for some $t \in \mathbb{N}$. Since $H' \subset G$ and G is open, for every compact set $\overline{A}_l \cap \overline{B}_l \cap K_{u_l} \cap K_{v_l}$ there exists a finite family $\{D_i^l\}_{1 \leq i \leq n_l} \subset S$ such that

$$\overline{A}_l \cap \overline{B}_l \cap K_{u_l} \cap K_{v_l} \subset \bigcup_{1 \leq i \leq n_l} D_i^l \subset \bigcup_{1 \leq i \leq n_l} \overline{D}_i^l \subset G,$$

and therefore

$$\overline{A}_l \cap \overline{B}_l \cap K_{u_l} \cap K_{v_l} \subset \bigcup_{1 \leq i \leq n_l} \overline{D}_i^l \cap K_{v_l} \subset G.$$

It follows that

$$H' = \bigcup_{1 \leq l \leq t} \overline{A}_l \cap \overline{B}_l \cap K_{u_l} \cap K_{v_l} \subset \bigcup_{1 \leq l \leq t} \left(\bigcup_{1 \leq i \leq n_l} \overline{D}_i^l \cap K_{v_l} \right).$$

Then, setting $H_0 = \bigcup_{1 \leq l \leq t, 1 \leq i \leq n_l} (\overline{D}_i^l \cap K_{v_l})$, we find that

$$H_0 \in \mathcal{H}, \quad H' \subset H_0 \subset G. \quad (9)$$

If $H \in \mathcal{H}$ and $H' \in \mathcal{H}'$ are such that $H_G^\varepsilon \subset H \cap H'$, and $H_0 \in \mathcal{H}$ is obtained from (9), we

have that

$$\sigma(H_G^\varepsilon) \leq \sigma(H') \leq \sigma(H_0) \leq P(G);$$

$$\sigma(H_G^\varepsilon) \leq \sigma(H) \leq P(G);$$

$$\sigma(H_G^\varepsilon) \leq \sigma(H \cap H') \leq \sigma(H) \leq P(G);$$

and for every $x^* \in X_1^*$,

$$\begin{aligned} & |x^* \alpha(H) - x^* \alpha(H')| = \left| \lim_j x^* m_j(H) - \lim_j x^* m_j(H') \right| = \\ & = \left| \lim_j (x^* m_j(H) - x^* m_j(H')) \right| \leq \limsup_j [|x^* m_j| (H - H \cap H') + \\ & + |x^* m_j| (H' - H \cap H')] = \limsup_j |x^* m_j| (H \triangle H') = \\ & = \limsup_j [|x^* m_j| (H) + |x^* m_j| (H') - 2 |x^* m_j| (H \cap H')] = \\ & = \lim_j [|x^* m_j| (H) + |x^* m_j| (H') - 2 |x^* m_j| (H \cap H')] = \\ & = \sigma(H) + \sigma(H') - 2\sigma(H \cap H') \leq \\ & \leq |\sigma(H) - P(G)| + |\sigma(H') - P(G)| + 2|P(G) - \sigma(H \cap H')| = \\ & = (P(G) - \sigma(H)) + (P(G) - \sigma(H')) + 2(P(G) - \sigma(H \cap H')) \leq \\ & \leq 4|P(G) - \sigma(H_G^\varepsilon)| < \varepsilon. \end{aligned}$$

Weak convergence

It follows that

$$\|\alpha(H) - \alpha(H')\| < \varepsilon.$$

□

Claim 4) *Let $\mathcal{G}$ be the family of the open sets of Ω . Then the set function β defined on $\mathcal{G}$ is strongly additive and it has a unique countably additive extension to Σ .*

Proof: The family $\mathcal{G}$ is a lattice of sets because it is closed under finite joints and meets and $\emptyset \in \mathcal{G}$. Obviously $\beta(\emptyset) = 0$.

We will prove that

$$\beta(A) + \beta(B) = \beta(A \cup B) + \beta(A \cap B), \quad (10)$$

for every $A, B \in \mathcal{G}$.

Let $A, B \in \mathcal{G}$ be fixed. We consider the collection

$$\mathcal{H}_A \oplus \mathcal{H}_B = \{H_1 \cup H_2 : H_1 \subset A, H_2 \subset B, H_i \in \mathcal{H}, i = 1, 2\}.$$

Obviously $\mathcal{H}_A \oplus \mathcal{H}_B$ is directed by inclusion.

Let $\varphi : \mathcal{H}_A \oplus \mathcal{H}_B \rightarrow \{H \in \mathcal{H} : H \subset A \cup B\}$ be the identity. Clearly φ is monotone. Moreover as in the proof of **Claim 2)** φ is cofinal.

Let $\varepsilon > 0$ be fixed. From the definition of β and from **Claim 3)** there exist $H_{A \cup B}^\varepsilon \in \Lambda_{A \cup B}$, $H_{A \cap B}^\varepsilon \in \Lambda_{A \cap B}$, $H_A^\varepsilon \in \Lambda_A$, $H_B^\varepsilon \in \Lambda_B$ such that

$$\|\beta(A \cup B) - \alpha(H)\| < \varepsilon,$$

if $H \in \mathcal{H}$ and $H_{A \cup B}^\varepsilon \subset H \subset A \cup B$;

$$\|\beta(A \cap B) - \alpha(H)\| < \varepsilon,$$

if $H \in \mathcal{H}$ and $H_{A \cap B}^\varepsilon \subset H \subset A \cap B$;

$$\|\beta(A) - \alpha(H)\| < \varepsilon,$$

if $H \in \mathcal{H}$ and $H_A^\varepsilon \subset H \subset A$;

$$\|\beta(B) - \alpha(H_B^\varepsilon)\| < \varepsilon,$$

and

$$\|\alpha(H') - \alpha(H_B^\varepsilon)\| < \varepsilon,$$

if $H' \in \mathcal{H}'$ and $H_B^\varepsilon \subset H' \subset B$.

Since φ is cofinal there exist $H_1, H_2 \in \mathcal{H}$ with $H_1 \subset A, H_2 \subset B$ such that $H_{A \cup B}^\varepsilon \subset H_1 \cup H_2 \subset A \cup B$. We set $H' = H_A^\varepsilon \cup H_1 \subset A$ and $H'' = H_B^\varepsilon \cup H_2 \subset B$; so it follows that $H_{A \cup B}^\varepsilon \subset H' \cup H'' \subset A \cup B$. Set $H_1^* = H_{A \cap B}^\varepsilon \cup H'$ and $H_2^* = H_{A \cap B}^\varepsilon \cup H''$. Then it follows that $H_1^* \supset H_A^\varepsilon$, $H_2^* \supset H_B^\varepsilon$, $H_1^* \cup H_2^* \supset H_{A \cup B}^\varepsilon$, $A \cap B \supset H_1^* \cap H_2^* \supset H_{A \cap B}^\varepsilon$. Remind that, since $H_i \in \mathcal{H}, i = 1, 2$, $H_1^* \cap H_2^* \in \mathcal{H}'$. Therefore we have

$$\|\alpha(H_{A \cap B}^\varepsilon) - \alpha(H_1^* \cap H_2^*)\| < \varepsilon,$$

hence, clearly

$$\|\beta(A \cap B) - \alpha(H_1^* \cap H_2^*)\| < 2\varepsilon.$$

Moreover, observing that $\mathcal{H} \subset \mathcal{H}'$, from (8)

$$\|\alpha(H_1^* \cup H_2^*) + \alpha(H_1^* \cap H_2^*) - \alpha(H_1^*) - \alpha(H_2^*)\| = 0.$$

Therefore we easily find

$$\|\beta(A \cup B) + \beta(A \cap B) - \beta(A) - \beta(B)\| < 5\varepsilon.$$

By the arbitrariness of ε , (10) is proved.

By **Theorem 2.17** β has a unique finitely additive extension defined on the algebra $\mathcal{R}$

Weak convergence

generated by $\mathcal{G}$.

Moreover, by (4) we have, for every $G \in \mathcal{G}$

$$\|\beta(G)\| = \lim_{H \in \Lambda_G} \|\alpha(H)\| \leq \lim_{H \in \Lambda_G} \sigma(H) = \sup_{H \in \Lambda_G} \sigma(H) \leq \gamma(G); \quad (11)$$

since Ω is a metrizable space, the result is true also for every $G \in \mathcal{R}$ (see [2]), and so β is countably additive and strongly bounded on $\mathcal{R}$, since γ is (see [2]). By **Theorem 2.15**, since strong boundedness implies that for every monotone sequence $(A_n)_n$ $\lim_n \beta(A_n)$ exists, there exists a unique countably additive extension of β to Σ . Without loss of generality we will continue to denote with β this extension. $\square$

Step 2) We will prove that m_{k_j} converges weakly to β .

Claim 1) *For every open set G of γ -continuity it is*

$$\lim_{H \in \Lambda_G} m_{k_j}(H) = m_{k_j}(G) \quad (12)$$

uniformly in k_j .

Proof: Let $(m_{k_j})_j$ be the subsequence chosen in the proof of **Step 1** (at the beginning).

Fix k_j ; since m_{k_j} is regular (by **Lemma 2.7**), for every $\varepsilon > 0$ there exists a compact set

$F_{\varepsilon, k_j} \subset G$ such that

$$\|m_{k_j}\|(G \setminus F_{\varepsilon, k_j}) < \varepsilon,$$

being $F_{\varepsilon, k_j} = F \cap K_u$, where $F \subset G$ is closed, $u \in \mathcal{N}$ and $\|m_{k_j}\|(\Omega \setminus K_u) < \frac{1}{2\varepsilon}$, $\|m_{k_j}\|(G \setminus F) < \frac{1}{2\varepsilon}$, and so there are $B(x_1^{(\varepsilon, k_j)}, r_1^{(\varepsilon, k_j)}), \dots, B(x_n^{(\varepsilon, k_j)}, r_n^{(\varepsilon, k_j)}) \in \mathcal{S}$ such that

$$F_{\varepsilon, k_j} \subset \bigcup_{i=1}^n \overline{B}(x_i^{(\varepsilon, k_j)}, r_i^{(\varepsilon, k_j)}) \subset G.$$

Let $\overline{H} = \bigcup_{i=1}^n \overline{B}(x_i^{(\varepsilon, k_j)}, r_i^{(\varepsilon, k_j)}) \cap K_u \in \Lambda_G$. By inclusion, $\|m_{k_j}\|(G \setminus \overline{H}) < \varepsilon$. Hence for every $H \in \Lambda_G, H \supset \overline{H}$ it is $\|m_{k_j}\|(G \setminus H) < \varepsilon$, and a fortiori $\|m_{k_j}(G \setminus H)\| < \varepsilon$.

This proves that

$$\lim_{H \in \Lambda_G} m_{k_j}(H) = m_{k_j}(G). \quad (13)$$

We want to show now that the limit in (13) is uniform with respect to j .

Let γ be the measure defined in (6) (see [2]) and let G be an open set such that $\gamma(\partial G) = 0$.

Then by Portmanteau's Theorem [2] it is

$$P(G) = \gamma(\overline{G}) \geq \limsup_j |m_{k_j}|(\overline{G}).$$

By definition of P , for every $\varepsilon > 0$ there exists $H_0 \in \Lambda_G$ such that

$$P(G) - \sigma(H_0) < \varepsilon.$$

So it follows that

$$\begin{aligned} \limsup_j \{|m_{k_j}|(G) - |m_{k_j}|(H_0)\} &\leq \\ &\leq \limsup_j |m_{k_j}|(\overline{G}) - \lim_j |m_{k_j}|(H_0) \leq \\ &\leq P(G) - \sigma(H_0) < \varepsilon. \end{aligned}$$

Thus there exists j^* such that for $H \supset H_0, H \subset G$,

$$|m_{k_j}|(G) - |m_{k_j}|(H) < \varepsilon,$$

for every $j > j^*$. Thus

$$\|m_{k_j}(G \setminus H)\| \leq |m_{k_j}|(G \setminus H) = |m_{k_j}|(G) - |m_{k_j}|(H) < \varepsilon,$$

for every $j > j^*$. Since $\lim_{H \in \Lambda_G} m_{k_j}(H) = m_{k_j}(G)$ for each k_j , we can determine $H_1^\varepsilon, \dots, H_{j^*}^\varepsilon \in \Lambda_G$ such that

$$\|m_{k_p}(G) - m_{k_p}(H)\| < \varepsilon,$$

Weak convergence

for every $H \supset H_p^\varepsilon$ and for every $p = 1, \dots, j^*$.

We set $\tilde{H} = H_0 \cup H_1^\varepsilon \cup \dots \cup H_{j^*}^\varepsilon$ with $\tilde{H} \in \Lambda_G$. Let $H \supset \tilde{H}$. Then we have

$$\|m_{k_j}(G) - m_{k_j}(H)\| < \varepsilon,$$

for every $j \in \mathbb{N}$. □

Now we can apply Lemma I.7.6 of [12] and we obtain

$$\begin{aligned} \beta(G) &= \lim_{H \in \Lambda_G} \alpha(H) = \lim_{H \in \Lambda_G} [(w) - \lim_j m_{k_j}(H)] = \\ &= (w) - \lim_j [\lim_{H \in \Lambda_G} m_{k_j}(H)] = (w) - \lim_j m_{k_j}(G) \end{aligned}$$

Therefore, for every open set G of γ -continuity, we have

$$\beta(G) = (w) - \lim_j m_{k_j}(G) \tag{14}$$

By complementation (14) is true also for every closed set of γ -continuity.

Let $\mathcal{A}$ be the algebra of γ -continuity sets: we want to prove that the equality (14) is verified

for every $M \in \mathcal{A}$.

Let $x^* \in X_1^*$ and $\varepsilon > 0$ be fixed. Then we want to find a $\bar{j} \in \mathbb{N}$ such that for each $j > \bar{j}$

$$|x^*[\beta(M) - m_{k_j}(M)]| < \varepsilon. \tag{15}$$

Observe that

$$\begin{aligned} &|x^*m_{k_j}(M) - x^*\beta(M)| \leq \\ &\leq |x^*m_{k_j}(M) - x^*m_{k_j}(M^0)| + |x^*m_{k_j}(M^0) - x^*\beta(M^0)| + \\ &+ |x^*\beta(M^0) - x^*\beta(M)| \leq |x^*m_{k_j}(M \setminus M^0)| + \\ &+ |x^*[m_{k_j}(M^0) - \beta(M^0)]| + \|\beta(M \setminus M^0)\|. \end{aligned}$$

By (11) it is $\|\beta(M \setminus M^0)\| \leq \gamma(M \setminus M^0) = 0$.

Moreover, since M^0 is a γ -continuity open set, from (14) there exists $j^* \in \mathbb{N}$ such that for every $j > j^*$ it is

$$|x^* m_{k_j}(M^0) - x^* \beta(M^0)| \leq \frac{\varepsilon}{2}.$$

Since $|m_{k_j}| \rightarrow \gamma$ (see [2]), from the Portmanteau's Theorem it follows that there exists $j^{**} \in \mathbb{N}$ such that for every $j > j^{**}$

$$|\gamma(\overline{M}) - |m_{k_j}|(\overline{M})| < \frac{\varepsilon}{4},$$

$$|P(M^0) - |m_{k_j}|(M^0)| < \frac{\varepsilon}{4}.$$

Since $M \in \mathcal{A}$, we have $\gamma(\overline{M}) = P(M^0)$. Therefore it is

$$\begin{aligned} |x^* m_{k_j}(M \setminus M^0)| &\leq \|m_{k_j}\|(M \setminus M^0) \leq \\ &\leq |m_{k_j}|(M \setminus M^0) \leq |m_{k_j}|(\overline{M} \setminus M^0) < \frac{\varepsilon}{2}. \end{aligned}$$

Hence (15) is proved.

The next step will be finally to prove weak convergence.

Let $A \in \mathcal{A}$ and let $f \in C_b(\Omega)$. We set $H = \{t : \gamma(f = t) > 0\}$. Since γ is bounded, H is at most countable. We observe that for every $t \notin H$ it is $(f = t) \in \mathcal{A}$. Since $\|\beta(\bullet)\| \leq \gamma(\bullet)$, if $t \notin H$ then the set $(f > t)$ is a β -continuity set. Let $x^* \in X^*$ be fixed and assume f to be positive. Since f is bounded, then

$$\begin{aligned} &\left| x^* \widehat{\int_{\Omega}} f dm_{k_j} - x^* \widehat{\int_{\Omega}} f d\beta \right| = \left| \int_{\Omega} f d(x^* m_{k_j}) - \int_{\Omega} f d(x^* \beta) \right| = \\ &= \left| \int_0^{+\infty} x^* m_{k_j}(f > t) dt - \int_0^{+\infty} x^* \beta(f > t) dt \right| = \\ &= \left| \int_{[0, \sup f] \setminus H} [x^* m_{k_j}(f > t) - x^* \beta(f > t)] dt \right| \leq \\ &\leq \int_{[0, \sup f] \setminus H} |x^* m_{k_j}(f > t) - x^* \beta(f > t)| dt. \end{aligned}$$

Weak convergence

By (15)

$$|x^*\beta(f > t) - x^*m_{k_j}(f > t)| \rightarrow 0$$

in $[0, \sup f] \setminus H$. Moreover it is

$$\begin{aligned} |x^*m_{k_j}(f > t) - x^*\beta(f > t)| &\leq |x^*\beta(f > t)| + |x^*m_{k_j}(f > t)| \leq \\ &\leq \|x^*\| \cdot [\|m_{k_j}\|(\Omega) + \|\beta\|(\Omega)] < +\infty. \end{aligned}$$

By the Dominated Convergence Theorem the assertion follows.

Remark 3.5 For an $m \in ca(\Omega, \Sigma, X)$, and $A \in \Sigma$ we denote by m^A the induced measure $m^A(E) = m(A \cap E)$. Then from Prokhorov's Theorem it follows that for every $A \in \mathcal{A}$, $m_{k_j}^A \rightharpoonup \beta^A$. It suffices to observe that, for $A \in \mathcal{A}$ and $B \in \Sigma$ fixed, B is a β^A -continuity set if $B \cap A$ is a β -continuity set.

References

- [1] P. BILLINGSLEY "Convergence of Probability Measures", John Wiley and Sons, Inc., New York, (1968).
- [2] P. BILLINGSLEY "Weak Convergence of Measures: Applications in Probability", Regional conference series in applied mathematics, Siam, (1977).
- [3] J. K. BROOKS "On the existence of a control measure for strongly bounded vector measures", Bull. Amer. Math. Soc., **77**, (1971), 999-1001.
- [4] J. K. BROOKS-A. MARTELLOTTI "On the De Giorgi-Letta integral in infinite dimensions", Atti Sem. Mat. Fis. Univ. Modena, **4**, (1992), 285-302.

- [5] M. J. MUNOZ BOUZO "A Prokhorov's theorem for Banach lattice valued measures", R. Acad. Ci. Madrid, **89**, (I-II), (1995), 111-127.
- [6] D. CANDELORO "Uniforme esaustività e assoluta continuità", Boll. Umi., **4-B (6)**, (1985), 709-724.
- [7] E. D'ANIELLO–R. M. SHORTT "Vector-valued capacities", Rend. Circ. Mat. Palermo, **46**, (1997).
- [8] M. DEKIERT "Kompaktheit, Fortsetzbarkeit und Konvergenz von Vectormassen", Dissertation, University of Essen, (1991).
- [9] E. DE GIORGI–G. LETTA "Une notion générale de convergence faible des fonctions croissantes d'ensemble", Ann. Scuola Norm. Sup. Pisa, **33**, (1977), 61-99 .
- [10] N. DINCULEANU "Vector Measures", Pergamon Press, (1967).
- [11] J. DIESTEL–J.J. UHL "Vector measures", Amer. Math. Soc. Surveys, **15**, Providence, (1968).
- [12] N. DUNFORD–J.T. SCHWARTZ "Linear Operators I; General Theory ", Interscienze, New York, (1958).
- [13] F. J. FERNÁNDEZ–P. JIMÉNEZ GUERRA "On the Prokhorov's Theorem for Vector Measures", Comm. Math. Univ. S. Pauli, **39**, (1990), 129-141.
- [14] F. J. FERNÁNDEZ–P. JIMÉNEZ GUERRA "Projective Limits of Vector Measures", Revista Matematica de la Univ. Complutense de Madrid, **3**, (1990), 125-142.
- [15] G. H. GRECO "Integrale monotono", Rend. Sem. Mat. Univ. Padova, **57**, (1977), 149-166.

Weak convergence

- [16] M. P. KATS "On the extension of vector measures", Sib. Math. Zhurn., **13** (**5**), (1972), 1158-1168.
- [17] Z. LIPECKI "On strongly additive set functions", Coll. Math., **22**, (1971), 255-256.
- [18] M. MÄRZ–R. M. SHORTT "Weak convergence of vector measures" Publ. Math. Debrecen, **45**, (1994), 71-92.
- [19] W. RUDIN "Analisi reale e complessa", Boringhieri, (1974).
- [20] M. SION "A theory of semigroup valued measures", Lecture Notes in Math., Springer-Verlag, New York, (1973).
- [21] S. WILLARD "General Topology", Addison-Wesley Publishing Company (1970).

Fitted mesh B-spline collocation method for solving singularly perturbed reaction-diffusion problems

Mohan K. Kadalbajoo¹ and Vivek K. Aggarwal²

^{1,2}*Department of Mathematics and Statistics, Indian Institute of Technology, Kanpur, India*
email: kadal@iitk.ac.in, vivekkumar.ag@gmail.com

Abstract.

In this paper we develop B-spline collocation method for solving a class of singularly perturbed reaction diffusion equations given by

$$(0.1) \quad \begin{aligned} -\epsilon u'' + a(x)u(x) &= f(x), \quad a(x) \geq a^* > 0 \\ u(0) &= \alpha, \quad u(1) = \beta \end{aligned}$$

Here we use the fitted mesh technique to generate piecewise uniform mesh. Our scheme leads to a tridiagonal linear system. The convergence analysis is given and the method is shown to have uniform convergence of second order. Numerical illustrations are given in the end to demonstrate the efficiency of our method.

Key words: Singularly perturbed boundary value problems, B-spline collocation method, fitted mesh method, boundary layers, uniform convergence.

1 Introduction

We consider the following class of singularly perturbed two-point boundary value problems

$$(1.1) \quad Lu \equiv -\epsilon u''(x) + a(x)u(x) = f(x) \quad \text{where } 0 \leq x \leq 1$$

subject to

$$(1.2) \quad u(0) = \alpha \quad u(1) = \beta, \quad \alpha, \beta \in \mathbb{R}$$

where ϵ is a small positive parameter and $a(x)$ and $f(x)$ are bounded continuous functions. For $a(x) > 0$, Most the equations of type, (1.1) possesses boundary layers [1], i.e regions of rapid change in the solution near the end-points, at the both of the end points. Due to this unique characteristic, these type of problems are very important in application point of view. These class of problems arises in various fields of science and engineering, for instance, fluid mechanics, quantum mechanics, optimal control, chemical-reactor theory, aerodynamics, reaction-diffusion process, geophysics etc.

There are three principal approaches to solve numerically these type of problems, namely, the Finite Difference methods [2] and [3], the Finite Element methods [4] and the Spline Approximation methods. First two methods have been used by numerous researchers. Enright [5] reported on an ongoing investigation into the performance of numerical methods for two-point boundary value problems. He outlined how methods based on multiple shooting, collocation

and other discretizations can share a common structure. Surla and Jerkovic [6] considered the singularly perturbed boundary value problem using spline collocation method. Sakai and Usmani [7] gave a new concept of B-spline in term of hyperbolic and trigonometric splines which are different from earlier ones. It is proved that the hyperbolic and trigonometric B-spline are characterized by a convolution of some special exponential functions and a characteristic function on the interval $[0,1]$.

In this paper we use the third approach, namely, B-Spline collocation method [8], [9] and [10], to solve the problems of above type. In section 2 we have given the derivation for the B-spline method and mesh strategy. Uniform convergence for the method have been discussed in section 3. In sections 4 and 5 we have solved some problems using the method and the results and graphs have also been shown.

2 B-spline collocation method

We subdivide the interval $[0, 1]$, and we choose piecewise uniform mesh points represented by $\pi = \{x_0, x_1, x_2, \dots, x_N\}$, such that $x_0 = 0$ and $x_N = 1$ and $\tilde{h}$ (Discussed in the next section) is the piecewise uniform spacing. We define $L_2[0, 1]$ is a vector space of all the square integrable function on $[0,1]$, and X be the linear subspace of $L_2[0, 1]$. Now we define for $i = 0, 1, 2, \dots, N$.

$$(2.1) \quad B_i(x) = \frac{1}{\tilde{h}^3} \begin{cases} (x - x_{i-2})^3, & \text{if } x \in [x_{i-2}, x_{i-1}] \\ \tilde{h}^3 + 3\tilde{h}^2(x - x_{i-1}) + 3\tilde{h}(x - x_{i-1})^2 - 3(x - x_{i-1})^3, & \text{if } x \in [x_{i-1}, x_i] \\ \tilde{h}^3 + 3\tilde{h}^2(x_{i+1} - x) + 3\tilde{h}(x_{i+1} - x)^2 - 3(x_{i+1} - x)^3, & \text{if } x \in [x_i, x_{i+1}] \\ (x_{i+2} - x)^3, & \text{if } x \in [x_{i+1}, x_{i+2}] \\ 0, & \text{otherwise.} \end{cases}$$

We introduce four additional knots as $x_{-2} < x_{-1} < x_0$ and $x_{N+2} > x_{N+1} > x_N$. From the above equation (2.1) we can simply check that each of the functions $B_i(x)$ is twice continuously differentiable on the entire real line, Also

$$(2.2) \quad B_i(x_j) = \begin{cases} 4, & \text{if } i = j \\ 1, & \text{if } i - j = \pm 1 \\ 0, & \text{if } i - j = \pm 2. \end{cases}$$

and that $B_i(x) = 0$ for $x \geq x_{i+2}$ and $x \leq x_{i-2}$. Similarly we can show that

$$(2.3) \quad B'_i(x_j) = \begin{cases} 0, & \text{if } i = j \\ \pm \frac{3}{\tilde{h}}, & \text{if } i - j = \pm 1 \\ 0, & \text{if } i - j = \pm 2. \end{cases}$$

and

$$(2.4) \quad B''_i(x_j) = \begin{cases} \frac{-12}{\tilde{h}^2}, & \text{if } i = j \\ \frac{6}{\tilde{h}^2}, & \text{if } i - j = \pm 1 \\ 0, & \text{if } i - j = \pm 2. \end{cases}$$

Each $B_i(x)$ is also a piece-wise cubic with knots at π , and $B_i(x) \in X$. Let $\Omega = \{B_{-1}, B_0, B_1, \dots, B_{N+1}\}$ and let $\Phi_3(\pi) = \text{span } \Omega$. The functions Ω are linearly independent on $[0, 1]$, thus $\Phi_3(\pi)$ is $(N+3)$ -dimensional. Even one can show that $\Phi_3(\pi) \subseteq_{sp} X$. Let L be a linear operator whose domain is X and whose range is also in X . Now we define

$$(2.5) \quad S(x) = c_{-1}B_{-1}(x) + c_0B_0(x) + c_1B_1(x) + \dots + c_NB_N(x) + c_{N+1}B_{N+1}(x)$$

Then force $S(x)$ to satisfy the collocation equations plus the boundary conditions. We have

$$(2.6) \quad LS(x_i) = f(x_i) \quad 0 \leq x_i \leq N$$

and

$$(2.7) \quad S(0) = \alpha, \quad S(1) = \beta$$

On solving the equation (2.6) we get

$$(2.8) \quad \begin{aligned} c_{i-1}(-\epsilon B''_{i-1}(x_i) + a_i B_{i-1}(x_i)) + c_i(-\epsilon B''_i(x_i) + a_i B_i(x_i)) \\ + c_{i+1}(-\epsilon B''_{i+1}(x_i) + a_i B_{i+1}(x_i)) = f_i \quad \forall i = 0, 1, 2, \dots, N \end{aligned}$$

where $a(x_i) = a_i$ and $f(x_i) = f_i$. Solving equation (2.8) we get

$$(2.9) \quad (-6\epsilon + a_i \tilde{h}^2)c_{i-1} + (12\epsilon + 4a_i \tilde{h}^2)c_i + (-6\epsilon + a_i \tilde{h}^2)c_{i+1} = \tilde{h}^2 f_i, \quad \forall i = 0, 1, \dots, N$$

The given boundary conditions (2.7) becomes

$$(2.10) \quad c_{-1} + 4c_0 + c_1 = \alpha$$

and

$$(2.11) \quad c_{N-1} + 4c_N + c_{N+1} = \beta$$

Equations (2.9), (2.10) and (2.11) lead to a $(N+3) \times (N+3)$ tridiagonal system with $(N+3)$ unknowns $C_N = (c_{-1}, c_0, \dots, c_{N+1})^t$ (where t stands for transpose). Now eliminating c_{-1} from first equation of (2.9) and (2.10): we find

$$(2.12) \quad 36\epsilon c_0 = f_0 \tilde{h}^2 - \alpha(-6\epsilon + a_0 \tilde{h}^2)$$

Similarly, eliminating c_{N+1} from the last equation of (2.9) and from (2.11), we find

$$(2.13) \quad 36\epsilon c_N = f_N \tilde{h}^2 - \beta(-6\epsilon + a_N \tilde{h}^2)$$

Coupling equations (2.12) and (2.13) with the second through $(N-1)$ st equations of (2.9), we are lead to the system of $(N+1)$ linear equations $Tx_N = d_N$ in the $(N+1)$ unknowns $x_N = (c_0, \dots, c_N)^t$ with right-hand side $d_N = (f_0 \tilde{h}^2 - \alpha(-6\epsilon +$

$a_0\tilde{h}^2), \tilde{h}^2 f_1, \tilde{h}^2 f_2, \dots, \tilde{h}^2 f_{N-1}, f_N \tilde{h}^2 - \beta(-6\epsilon + a_N \tilde{h}^2))$ and the coefficient matrix is

$$(2.14) \quad \begin{bmatrix} 36\epsilon & 0 & 0 & 0 & \dots & 0 \\ -6\epsilon + a_1 \tilde{h}^2 & 12\epsilon + 4a_1 \tilde{h}^2 & -6\epsilon + a_1 \tilde{h}^2 & 0 & \dots & 0 \\ 0 & \vdots & \vdots & \vdots & 0 & \dots \\ 0 & 0 & -6\epsilon + a_i \tilde{h}^2 & 12\epsilon + 4a_i \tilde{h}^2 & -6\epsilon + a_i \tilde{h}^2 & \dots \\ \vdots & \vdots & \vdots & \vdots & \vdots & \vdots \\ 0 & \dots & 0 & -6\epsilon + a_{N-1} \tilde{h}^2 & 12\epsilon + 4a_{N-1} \tilde{h}^2 & -6\epsilon + a_{N-1} \tilde{h}^2 \\ 0 & \dots & 0 & 0 & 0 & 36\epsilon \end{bmatrix}$$

Since $a(x) > 0$, it is easily seen that the matrix T is strictly diagonally dominant and hence nonsingular. Since T is nonsingular, we can solve the system $Tx_N = d_N$ for $c_0, c_1, \dots, c_N$ and substitute into the boundary equations (2.10) and (2.11) to obtain c_{-1} and c_{N+1} . Hence this method of collocation applied to (1.1) using a basis of cubic B-spline has a unique solution $S(x)$.

2.1 Mesh selection strategy

We form the piecewise -uniform grid in such a way that more pts are generated in the boundary layers region than outside of it.

we divide the interval $[0,1]$ into three sub-intervals $(0, \kappa)$, $(\kappa, 1-\kappa)$ and $(1-\kappa, 1)$, where

$$(2.15) \quad \kappa = \min\{1/4, c\sqrt{\epsilon} \ln N\}; \quad c \text{ is a constant.}$$

Assuming $N = 2^r$ with $r \geq 3$ be the total no. of subintervals in the partitions of $[0,1]$.

we divide the intervals $(0, \kappa)$ and $(1-\kappa, 1)$ each into $N/4$ equal mesh elements, while the interval $(\kappa, 1-\kappa)$ is divided into $N/2$ equal mesh elements. The resulting piecewise mesh depends upon just one parameter κ . Obviously Now we consider

$$(2.16) \quad \begin{aligned} \tilde{h} &= h_1 \quad \text{in the interval } [0, \kappa] \\ h_1 &= \frac{4\kappa}{N} \\ x_i &= x_{i-1} + h_1 \quad \text{for } i = 1, 2, 3, \dots, N/4. \end{aligned}$$

where $x_0 = 0$. Also

$$(2.17) \quad \begin{aligned} \tilde{h} &= h_2 \quad \text{in the interval } [\kappa, 1-\kappa] \\ h_2 &= \frac{2(1-2\kappa)}{N} \\ x_i &= x_{i-1} + h_2 \quad \text{for } i = N/4 + 1, \dots, 3N/4 \end{aligned}$$

Similarly

$$\begin{aligned}
 \tilde{h} &= h_3 \quad \text{in the interval } [1 - \kappa, 1] \\
 (2.18) \quad h_3 &= \frac{4\kappa}{N} \\
 x_i &= x_{i-1} + h_3 \quad \text{for } i = 3N/4 + 1, \dots, N.
 \end{aligned}$$

3 Derivation for Uniform Convergence

First we prove the following lemma

LEMMA 3.1. *The B-splines $\{B_{-1}, B_0, \dots, B_{N+1}\}$ defined in equation (2.1), satisfy the inequality*

$$\sum_{i=-1}^{N+1} |B_i(x)| \leq 10, \quad 0 \leq x \leq 1$$

PROOF. We know that

$$\left| \sum_{i=-1}^{N+1} B_i(x) \right| \leq \sum_{i=-1}^{N+1} |B_i(x)|$$

At any i th nodal point x_i , we have

$$\sum_{i=-1}^{N+1} |B_i| = |B_{i-1}| + |B_i| + |B_{i+1}| = 6 < 10$$

Also we have

$$|B_i(x)| \leq 4 \text{ and } |B_{i-1}(x)| \leq 4 \text{ for } x \in [x_{i-1}, x_i]$$

similarly

$$|B_{i-2}(x)| \leq 1 \text{ and } |B_{i+1}(x)| \leq 1 \text{ for } x \in [x_{i-1}, x_i]$$

Now for any point $x \in [x_{i-1}, x_i]$ we have

$$\sum_{i=-1}^{N+1} |B_i(x)| = |B_{i-2}| + |B_{i-1}| + |B_i| + |B_{i+1}| \leq 10$$

Hence this proves the lemma. $\square$

Now to estimate the error $\|u(x) - S(x)\|_\infty$, let Y_n be the unique spline interpolate from $\Phi_3(\pi)$ to the solution $u(x)$ of our boundary value problem (1.1)-(1.2). If $f(x) \in C^2[0, 1]$, and $u(x) \in C^4[0, 1]$, and it follows from the De Boor-Hall error estimates that

$$(3.1) \quad \|D^j(u(x) - Y_n)\|_\infty \leq \gamma_j h_c^{4-j}, \quad j = 0, 1, 2.$$

where $h_c = \max\{h_1, h_2, h_3\}$ and γ_j 's are independent of h_c and N . Let

$$(3.2) \quad Y_n(x) = \sum_{i=-1}^{N+1} b_i B_i(x)$$

Where

$$S(x) = \sum_{i=-1}^{N+1} c_i B_i(x)$$

is our collocation solution. It is clear that from equation (3.1) that

$$(3.3) \quad |LS(x_i) - LY_n(x_i)| = |f(x_i) - LY_n(x_i) + Lu(x_i) - Lu(x_i)| \leq \lambda h_c^2$$

where $\lambda = [\epsilon \gamma_2 + \gamma_0 \|a(x)\|_\infty h_c^2]$. Also $LS(x_i) = Lu(x_i) = f(x_i)$. Let $LY_n(x_i) = \hat{f}_n(x_i) \quad \forall i$, and $\hat{f}^n = (\hat{f}_n(x_0), \hat{f}_n(x_1), \dots, \hat{f}_n(x_N))^t$. Now from the system $Tx_N = d_N$ and equation (3.3) that the i th coordinate $[T(x_N - y^n)]_i$ of $T(x_N - y^n)$, $y^n = (b_0, b_1, \dots, b_N)^t$, satisfy the inequality

$$(3.4) \quad |[T(x_N - y^n)]_i| = h_c^2 |f_i - \hat{f}_i| \leq \lambda h_c^4$$

since $(Tx_N)_i = h_c^2 f_i$ and $(Ty^n)_i = h_c^2 \hat{f}_i(x_i)$ for $i = 1, 2, 3, \dots, N-1$. Also

$$(3.5) \quad (Tx_N)_0 = h_c^2 f_0 - \alpha(-6\epsilon + a_0 \tilde{h}_c^2), \text{ and } (Ty^n)_0 = h_c^2 \hat{f}_n(x_0) - \alpha(-6\epsilon + a_0 \tilde{h}_c^2)$$

similarly

$$(3.6) \quad (Tx_N)_N = h_c^2 f_N - \beta(-6\epsilon + a_N \tilde{h}_c^2), \text{ and } (Ty^n)_N = h_c^2 \hat{f}_n(x_N) - \beta(-6\epsilon + a_N \tilde{h}_c^2)$$

But the i th coordinate of $[T(x_N - y^n)]$ is the i th equation

$$(3.7) \quad (-6\epsilon + a_i h_c^2) \delta_{i-1} + (12\epsilon + 4a_i h_c^2) \delta_i + (-6\epsilon + a_i h_c^2) \delta_{i+1} = \xi_i \quad \forall i = 1, 2, \dots, N-1.$$

where $\delta_i = b_i - c_i$, $-1 \leq i \leq N+1$, and

$$(3.8) \quad \xi_i = h_c^2 [f(x_i) - \hat{f}_n(x_i)] \quad \forall i = 1, 2, 3, \dots, N-1.$$

Obviously from equation (3.3)

$$(3.9) \quad |\xi_i| \leq \lambda h_c^4$$

Let $\xi = \max_{1 \leq i \leq N-1} |\xi_i|$. Also consider $\delta = (\delta_{-1}, \delta_0, \dots, \delta_{N+1})^t$, then we define $e_i = |\delta_i|$, and $\tilde{e} = \max_{1 \leq i \leq N-1} |e_i|$. Now equation (3.7) becomes

$$(3.10) \quad (12\epsilon + 4a_i h_c^2) \delta_i = \xi_i + (6\epsilon - a_i h_c^2) (\delta_{i-1} + \delta_{i+1}) \quad \forall i = 1, 2, \dots, N-1.$$

Taking absolute values with sufficiently small h_c . We have

$$(3.11) \quad (12\epsilon + 4a_i h_c^2) e_i \leq \xi + 2\tilde{e}(6\epsilon - a_i h_c^2)$$

Also $0 < a^* \leq a(x)$ for all x . We get

$$(3.12) \quad (12\epsilon + 4a^*h_c^2) e_i \leq \xi + 2\tilde{e}(6\epsilon - a^*h_c^2) \leq \xi + 2\tilde{e}(6\epsilon + a^*h_c^2)$$

In particular

$$(3.13) \quad (12\epsilon + 4a^*h_c^2) \tilde{e} \leq \xi + 2\tilde{e}(6\epsilon + a^*h_c^2)$$

which gives

$$(3.14) \quad 2a^*h_c^2\tilde{e} \leq \xi \leq \lambda h_c^4 \quad \text{implies } \tilde{e} \leq \frac{\lambda h_c^2}{2a^*}$$

To estimate e_{-1}, e_0, e_N and e_{N+1} , first we observe that the first equation of the system $T(x_N - y^n) = h_c^2[F_N - \hat{f}^n]$ (where $F_N = (f_0, f_1, \dots, f_N)$) gives

$$(3.15) \quad 36\epsilon \delta_0 = h_c^2[f_0 - \hat{f}_0]$$

which gives

$$(3.16) \quad e_0 \leq \frac{\lambda h_c^4}{36\epsilon}$$

similarly we can calculate

$$(3.17) \quad e_N \leq \frac{\lambda h_c^4}{36\epsilon}$$

Now e_{-1} and e_{N+1} can be calculated using equations (2.10) and (2.11) as $\delta_{-1} = (-4\delta_0 - \delta_1)$ and $\delta_N + 1 = (-4\delta_N - \delta_{N-1})$

$$(3.18) \quad e_{-1} \leq \frac{\lambda h_c^4}{9\epsilon} + \frac{\lambda h_c^2}{2a^*}$$

also

$$(3.19) \quad e_{N+1} \leq \frac{\lambda h_c^4}{9\epsilon} + \frac{\lambda h_c^2}{2a^*}$$

using value $\lambda = [\epsilon \gamma_2 + \gamma_0 \|a(x)\|_\infty h_c^2]$, we get

$$(3.20) \quad e = \max_{-1 \leq i \leq N+1} \{e_i\} \Rightarrow e \leq \omega h_c^2$$

where $\omega = \frac{\gamma_2 h_c^2}{9} + \frac{\lambda}{2a^*}$ where $(\frac{h_c^6}{\epsilon} \sim 0)$.

The above inequality enables us to estimate $\|S(x) - Y_n(x)\|_\infty$, and hence $\|u(x) - S(x)\|_\infty$. In particular

$$(3.21) \quad S(x) - Y_n(x) = \sum_{i=-1}^{N+1} (c_i - b_i) B_i(x).$$

thus

$$(3.22) \quad |S(x) - Y_n(x)| \leq \max |c_i - b_i| \sum_{i=-1}^{N+1} |B_i(x)|.$$

But

$$(3.23) \quad \sum_{i=-1}^{N+1} |B_i(x)| \leq 10, \quad 0 \leq x \leq 1 \quad (\text{using lemma 1}).$$

Combining equations (3.20) (3.22) and (3.23), we see that

$$(3.24) \quad \|S - Y_n\|_\infty \leq 10\omega h_c^2.$$

But $\|u - Y_n\|_\infty \leq \gamma_0 h_c^4$. Since $\|u - S\|_\infty \leq \|u - Y_n\|_\infty + \|Y_n - S\|_\infty$, we see that

$$(3.25) \quad \|u - S\|_\infty \leq M h_c^2$$

Where $M = 10\omega + \gamma_0 h_c^2$. Combining the results we have proved.

THEOREM 3.2. *The collocation approximation $S(x)$ from the space $\Phi_3(\pi)$ to the solution $u(x)$ of the boundary value problem (1.1) with (1.2) exists. If $f \in C^2[0, 1]$, then*

$$\|u(x) - S(x)\|_\infty \leq M h_c^2$$

for h_c^2 sufficiently small and M is a positive constant (independent of ϵ).

4 Numerical Results

In this section the numerical results of some model problems are presented.

Problem 1.: We take

$$(4.1) \quad -\epsilon u'' + u = 1 + 2\sqrt{\epsilon}[\exp(-x/\sqrt{\epsilon}) + \exp((x-1)/\sqrt{\epsilon})],$$

subject to

$$(4.2) \quad u(0) = 0 \text{ and } u(1) = 0$$

which has the exact solution

$$u(x) = 1 - (1-x)\exp(-x/\sqrt{\epsilon}) - x\exp[(x-1)/\sqrt{\epsilon}].$$

Problem 2. : We solve [11]

$$(4.3) \quad -\epsilon u'' + u = -\cos^2(\pi x) - 2\epsilon\pi^2 \cos(2\pi x),$$

subject to

$$(4.4) \quad u(0) = 0 \text{ and } u(1) = 0$$

which has the exact solution

$$u(x) = [\exp(-1(1-x)/\sqrt{\epsilon}) + \exp(-x/\sqrt{\epsilon})]/[1 + \exp(-1/\sqrt{\epsilon})] - \cos^2(\pi x).$$

Problem 3.: We consider

$$(4.5) \quad -\epsilon u'' + u = x,$$

subject to

$$(4.6) \quad u(0) = 1 \text{ and } u(1) = 1 + \exp(-1/\sqrt{\epsilon}).$$

which has the exact solution

$$u(x) = \exp(-x/\sqrt{\epsilon}) + x.$$

Problem 4.: We solve [12]

$$(4.7) \quad -\epsilon u'' + (1+x)u = -40[x(x^2-1) - 2\epsilon],$$

subject to

$$(4.8) \quad u(0) = 0 \text{ and } u(1) = 0$$

which has the exact solution

$$u(x) = 40x(1-x).$$

The maximum error for problem 1 for different values of ϵ (uniform mesh) and grid points N is presented in Table 1.

TABLE 1
Maximum error for Problem 1
Uniform Mesh

$\epsilon = 2^{-k}$	N=16	N=32	N=64	N=128	N=256	N=512	N=1024
k =2	9.424E-04	2.352E-04	5.879E-05	1.469E-05	3.674E-06	2.175E-06	2.296E-07
4	2.240E-03	5.576E-04	1.392E-04	3.480E-05	8.700E-06	5.203E-06	5.437E-07
6	5.462E-03	1.339E-03	3.333E-04	8.325E-05	2.081E-05	1.778E-05	1.300E-06
8	2.015E-02	4.663E-03	1.146E-03	2.848E-04	7.113E-05	6.551E-05	4.445E-06
10	6.921E-02	1.851E-02	4.292E-03	1.054E-03	2.623E-04	2.510E-04	1.637E-05
12	1.582E-01	6.637E-02	1.769E-02	4.107E-03	1.008E-03	9.861E-04	6.270E-05
14	2.318E-01	1.552E-01	6.495E-02	1.728E-02	4.014E-03	3.968E-03	2.454E-04
16	2.593E-01	2.299E-01	1.537E-01	6.424E-02	1.708E-02	1.698E-02	9.748E-04
18	2.662E-01	2.583E-01	2.289E-01	1.530E-01	6.389E-02	6.371E-02	3.941E-03
20	2.677E-01	2.657E-01	2.578E-01	2.284E-01	1.526E-01	2.471E-01	1.693E-02
25	2.680E-01	2.679E-01	2.677E-01	2.666E-01	2.626E-01	2.471E-01	1.961E-01
30	2.679E-01	2.679E-01	2.679E-01	2.679E-01	2.677E-01	2.672E-01	2.652E-01

The maximum error for problem 1 for different values of ϵ (piecewise uniform mesh) and grid points N is presented in Table 2.

TABLE 2
Maximum error for Problem 1
Using Fitting Mesh

$\epsilon = 2^{-k}$	N=16	N=32	N=64	N=128	N=256	N=512	N=1024
k=2	9.424E-04	2.352E-04	5.879E-05	1.469E-05	3.674E-06	9.043E-07	2.296E-07
4	2.240E-03	5.576E-04	1.392E-04	3.480E-05	8.700E-06	2.141E-06	5.437E-07
6	5.462E-03	1.339E-03	3.333E-04	8.325E-05	2.081E-05	5.123E-06	1.300E-06
8	2.866E-02	5.755E-03	1.144E-03	2.848E-04	7.113E-05	1.750E-05	4.445E-06
10	5.226E-02	2.152E-02	8.340E-03	3.114E-03	1.122E-03	3.737E-04	1.094E-04
12	5.962E-02	2.811E-02	1.283E-02	5.743E-03	2.559E-03	1.136E-03	5.159E-04
14	6.162E-02	3.028E-02	1.469E-02	7.019E-03	3.322E-03	1.560E-03	7.484E-04
16	6.219E-02	3.092E-02	1.531E-02	7.540E-03	3.681E-03	1.775E-03	8.706E-04
18	6.237E-02	3.112E-02	1.551E-02	7.720E-03	3.828E-03	1.875E-03	9.312E-04
20	6.244E-02	3.119E-02	1.558E-02	7.778E-03	3.879E-03	1.915E-03	9.590E-04
25	6.249E-02	3.124E-02	1.561E-02	7.808E-03	3.903E-03	1.936E-03	9.752E-04
30	6.249E-02	3.124E-02	1.562E-02	7.811E-03	3.903E-03	1.937E-03	9.763E-04

The maximum error for problem 2 for different values of ϵ and grid points N is presented in Table 3.

TABLE 3
Maximum error for Problem 2
Using Fitting Mesh

$\epsilon = 2^{-k}$	N=16	N=32	N=64	N=128	N=256	N=512	N=1024
k=4	7.098E-03	1.779E-03	4.450E-04	1.112E-04	2.782E-05	6.955E-06	1.738E-06
6	4.073E-03	1.016E-03	2.541E-04	6.353E-05	1.588E-05	3.970E-06	9.926E-07
8	6.305E-02	2.441E-03	9.336E-04	2.323E-04	5.803E-05	1.450E-05	3.626E-06
10	7.752E-02	5.022E-02	3.928E-03	9.612E-04	2.392E-04	5.028E-05	1.276E-05
12	6.341E-02	3.066E-02	2.021E-02	3.803E-03	9.633E-04	2.768E-04	5.988E-05
14	6.237E-02	3.174E-02	1.576E-02	6.303E-03	5.367E-03	9.917E-04	2.398E-04
16	6.251E-02	3.119E-02	1.581E-02	7.909E-03	3.468E-03	9.731E-04	9.635E-04
18	6.250E-02	3.124E-02	1.560E-02	7.871E-03	3.940E-03	1.826E-03	6.840E-04
20	6.250E-02	3.125E-02	1.562E-02	7.804E-03	3.921E-03	1.963E-03	9.404E-04
25	6.250E-02	3.125E-02	1.562E-02	7.812E-03	3.906E-03	1.952E-03	9.759E-04
30	6.250E-02	3.125E-02	1.562E-02	7.812E-03	3.906E-03	1.953E-03	9.765E-04

The maximum error for problem 3 for different values of ϵ and grid points N is presented in Table 4.

TABLE 4
Maximum error for Problem 3
Using Fitting Mesh

$\epsilon = 2^{-k}$	N=16	N=32	N=64	N=128	N=256	N=512	N=1024
k =2	1.910E-04	4.771E-05	1.193E-05	2.983E-06	7.458E-07	1.864E-07	4.661E-08
4	9.553E-04	2.337E-04	5.938E-05	1.484E-05	3.710E-06	9.276E-07	2.319E-07
6	3.921E-03	9.635E-04	2.398E-04	5.989E-05	1.497E-05	3.742E-06	9.355E-07
8	1.465E-02	7.796E-03	9.635E-04	2.398E-04	5.989E-05	1.497E-05	3.742E-06
10	2.208E-02	1.507E-02	3.921E-03	9.635E-04	2.398E-04	2.989E-05	1.256E-05
12	2.761E-02	1.581E-02	4.842E-03	7.921E-03	9.635E-04	7.118E-04	5.269E-05
14	3.133E-02	1.670E-02	9.525E-03	6.947E-03	6.252E-03	9.074E-04	5.977E-04
16	3.353E-02	1.774E-02	9.125E-03	5.266E-03	3.898E-03	3.522E-03	3.432E-03
18	3.472E-02	1.848E-02	9.279E-03	4.747E-03	2.783E-03	2.080E-03	1.879E-03
20	3.534E-02	1.890E-02	9.476E-03	4.683E-03	2.416E-03	1.435E-03	1.078E-03
25	3.586E-02	1.928E-02	9.703E-03	4.763E-03	2.334E-03	1.156E-03	1.374E-03
30	3.595E-02	1.935E-02	9.749E-03	4.790E-03	2.346E-03	1.154E-03	1.381E-03

The comparison Table for problem 4 with the existing methods for different values of ϵ and grid points N is presented in Tables 5 and 6.

TABLE 5
Maximum error for Problem 4
Fitted Mesh: N=16

ϵ	Miller [13]	Nijima [14]	Nijima [12]	Our method (Unif. Mesh)
0.1E-03	0.25E-01	0.26E-01	0.65E-04	1.776E-15
0.1E-04	0.21E-01	0.24E-01	0.36E-04	1.776E-15
0.1E-05	0.70E-02	0.17E-01	0.33E-04	1.776E-15
0.1E-06	0.75E-03	0.69E-02	0.26E-04	1.776E-15
0.1E-07	0.74E-04	0.23E-02	0.20E-04	1.783E-15
0.1E-08	0.67E-05	0.76E-03	0.20E-04	1.783E-15
0.1E-09	0.00E+00	0.24E-03	0.11E-04	1.784E-15

TABLE 6
Maximum error for Problem 4
Fitted Mesh : N=32

ϵ	Miller [13]	Nijima [14]	Nijima [12]	Our method (Unif. Mesh)
0.1E-03	0.64E-02	0.65E-02	0.59E-04	1.776E-15
0.1E-04	0.61E-02	0.64E-02	0.21E-04	1.776E-15
0.1E-05	0.41E-02	0.56E-02	0.35E-04	1.776E-15
0.1E-06	0.77E-03	0.31E-02	0.39E-04	1.776E-15
0.1E-07	0.76E-04	0.12E-02	0.21E-04	1.776E-15
0.1E-08	0.67E-05	0.38E-03	0.21E-04	1.788E-15
0.1E-09	0.00E+00	0.13E-03	0.14E-04	1.789E-15

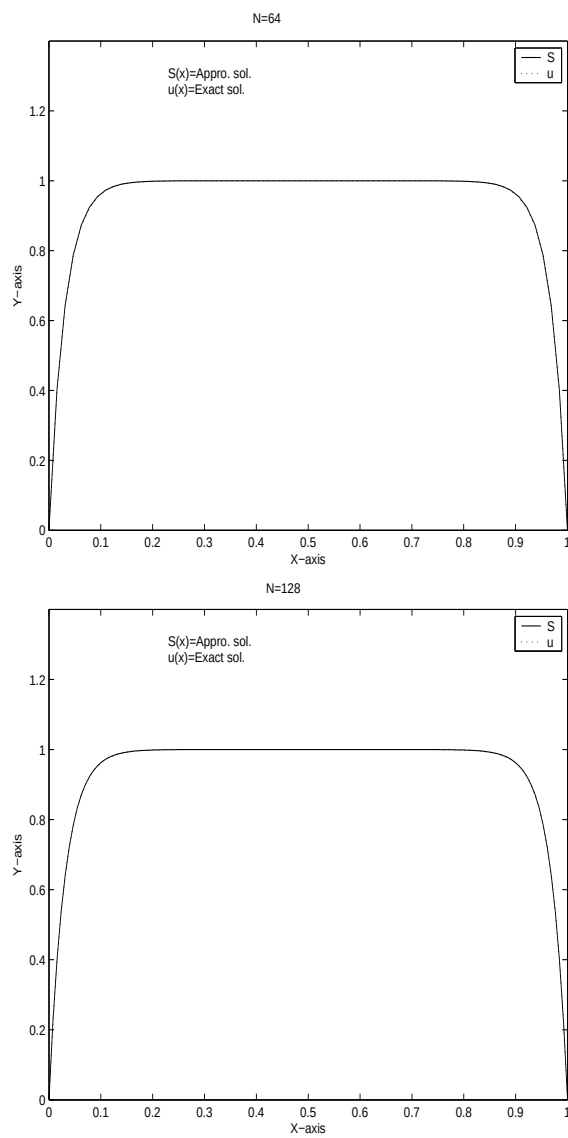
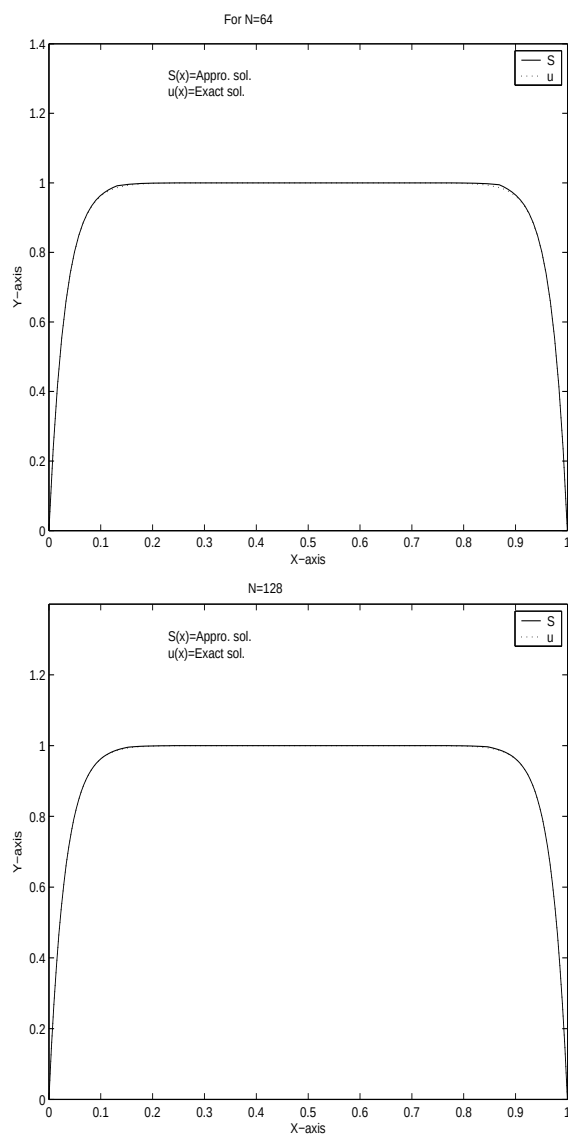


Figure 4.1: Problem 1. with uniform mesh for $\epsilon = 10^{-3}$

Figure 4.2: Problem 1. with Fitted mesh for $\epsilon = 10^{-3}$

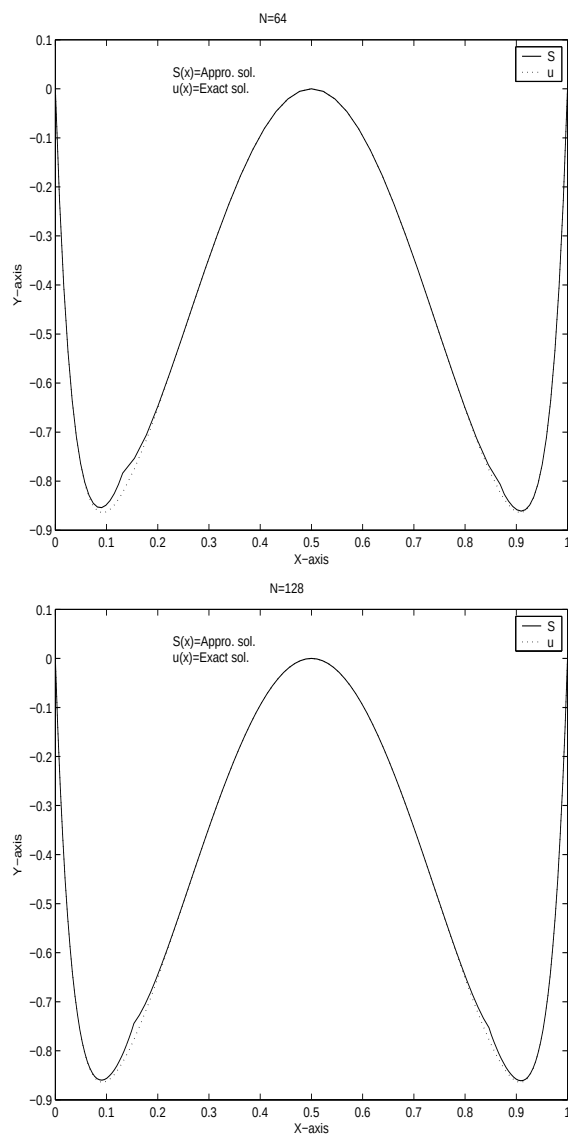
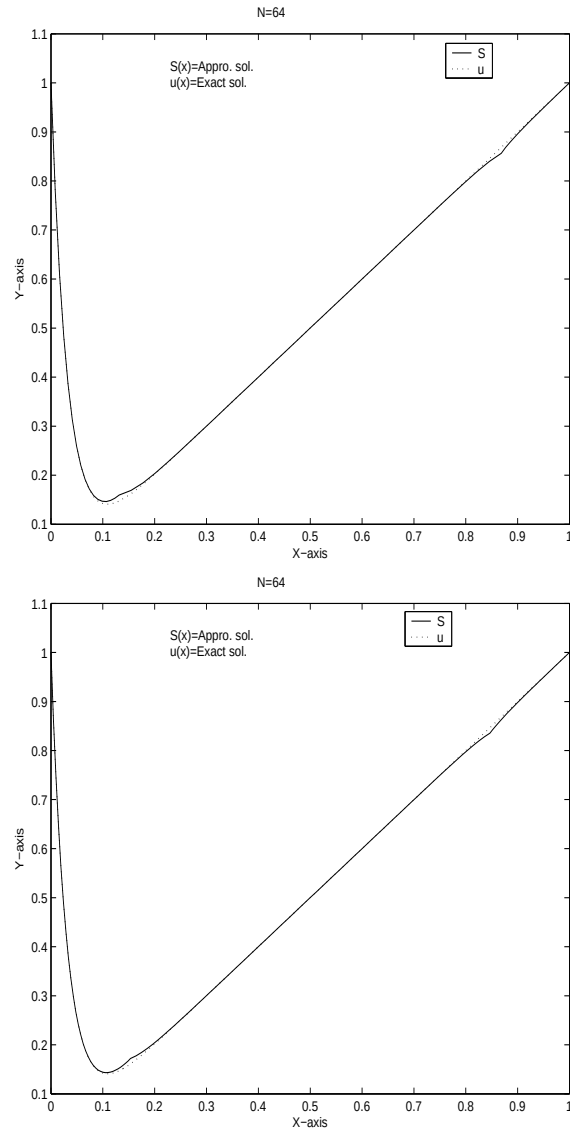


Figure 4.3: Problem 2. with Fitted mesh for $\epsilon = 10^{-3}$

Figure 4.4: Problem 3. with Fitted mesh for $\epsilon = 10^{-3}$

5 Discussion

We have described a B-spline collocation method for solving singularly perturbed problems using fitted mesh technique. We have applied this method for problems having second derivative term and function term. These type of problems are very important physically and difficult to solve because of two boundary layers at both of the end points. In Table 1, we have shown the maximum error for the problem 1 using B-spline collocation method with uniform mesh. It is very clear from the Table that for large value of ϵ we have better results but as we decrease the value of ϵ , the accuracy decreases rapidly. So we conclude that uniform mesh is not a good technique for such kind of problems, because number of mesh points in the boundary layer region should be much more than that from the outer region. In Tables 2,3 and 4 we have shown that maximum error for problems 2,3 and 4 using fitted mesh technique (piecewise -uniform mesh) respectively. For

$$N \geq \exp \frac{1}{4\sqrt{\epsilon}}$$

our mesh behaves as uniform mesh. It can be seen from the respective Tables that if we take more points in the boundary layer region then we obtain better results which implies that the use of piecewise-uniform mesh is quite advantageous. In Tables 5 and 6, we have given the comparison with the different methods applied to the same kind of problems.

To further corroborate the applicability of the proposed method, graphs have been plotted for problems for $x \in [0, 1]$ versus the computed solution obtained at different values of x for a fixed ϵ .

6 Conclusion

As is evident from the numerical results, this method gives $O(h^2)$ accuracy. The results obtained using this method are better than using the stated existing methods with the same no. of knots. Also this method produce a spline function which may be used to obtain the solution at any point in the range, whereas the finite-difference methods [12]-[14], gives the solution only at the chosen knots.

REFERENCES

1. P.W.Hemker, J.J.H.Miller(Eds.), "Numerical analysis of singular perturbation problems", *Academic Press*, New York, 1979.
2. P.A.Markowich, "A finite difference method for the basic stationary semiconductor device equations,in progress in Science and Computations", Vol. 5, Numerical boundary value ODEs, Birkhäuser, Boston, MA, 1985, pp 285-301.
3. D.Herceg, "Uniform fourth order difference scheme for a singular perturbation problem", *Numer.Math.*, Vol. 56, (7) pp 675-693, 1990.
4. E.O'Riordan, M.Stynas, "A uniformly accurate finite element method for a singularly perturbed one-dimensional reaction-diffusion problem", *Math. Comput.*, 47 (176), pp 555-570, 1986.
5. W.H.Enright, "Improving the performance of numerical methods for two-point boundary value problems,in progress in Science and Computations",

- Vol. 5, Numerical boundary value ODEs, Birkhäuser, Boston, MA, 1985, pp 107-119.”,
6. K.Surla, V.Jerkovic, "Some possibilities of applying spline collocations to singular perturbation problems", *Numerical methods and Approximation Theory*, Vol. II, Noviad, pp 19-25, 1985.
 7. M.Sakai, R.A.Usmani, "On exponential splines", *J.Appro. Theory*, Vol. 47, pp 122-131, 1986.
 8. L.L.Schumaker, "Spline functions: Basic Theory", *Krieger Publishing Company*, Florida , 1981.
 9. J.M.Ahlberg, E.N.Nilson and J.L.Walsh, "The Theory of splines and their applications", *Academic Press*, New York, 1967.
 10. G.Micula, "Handbook of Splines", *Kluwer Academic Publishers*, Dordrecht, The Netherlands 1999.
 11. E.P. Doolan, J.J.H. Miller, W.H.A. Schilders, "Uniform Numerical methods for problems with initial and boundary layers", *Boole Press*, Dublin , 1980.
 12. K. Nijima, "On a three-point difference scheme for a singular perturbation problem without a first derivative term II", *Mem. Numer. Math.*, Vol. 7, pp 11-27, 1980.
 13. J.J.H. Miller, "On the convergence, uniformly in ϵ , of difference schemes for a two-point boundary singular perturbation problem, Numerical Analysis of Singular Perturbation Problems", *Academic Press*, New York , pp 467-474, 1979.
 14. K. Nijima, "On a three-point difference scheme for a singular perturbation problem without a first derivative term I", *Mem. Numer. Math.*, Vol. 7, pp 1-10, 1980.

A Procedure for Improving Convexity Preserving Quasi-Interpolation

Costanza Conti and Rossana Morandi

*Dip. di Energetica "Sergio Stecco", Università di Firenze
Via Lombroso 6/17, 50134 Firenze, Italy.*

costanza.conti@unifi.it, rossana.morandi@unifi.it

Abstract

In this paper a procedure is proposed to improve the shape preserving quasi-interpolant scheme defined by positive quadratic splines with local support proposed in [3]. In fact, starting from that quasi-interpolant, a new shape preserving quasi-interpolant is defined by quasi-interpolating the weighted error obtained at the previous step. This strategy provides an error reduction at the data points and guarantees the $\mathbb{P}_1$ -reproduction property as well as the end point interpolation condition.

Keywords: Quasi-interpolation, Convexity preservation.

1 Introduction

The problem of constructing a quasi-interpolant shape preserving function fitting a data set $\{(x_i, f_i)\}_{i=1}^n$ is known to be of great interest in the application to graphical problems and in the construction of functions and curves following the shape suggested by the data (see for instance [2], [4] and references therein). In particular, a shape preserving quasi-interpolant in its simplest form, i.e. expressed as $\sigma(x) = \sum_{i=0}^{n+1} f_i C_i(x)$ for $x \in [x_1, x_n]$, was defined in [1], [5] and in [3] by using multiquadric functions and natural quadratic splines with local support, respectively. Here, dealing with convexity preservation, an improving procedure based on a more general quasi-interpolation scheme is proposed. The procedure defines a new quasi-interpolant function, namely $\tilde{\sigma}$, of the more general type

$$\tilde{\sigma}(x) := \sum_{i=0}^{n+1} \Lambda_i(f) C_i(x), \quad x \in [x_1, x_n]$$

where $\{\Lambda\}_{i=0}^{n+1}$ are suitable functionals and $\{C_i\}_{i=0}^{n+1}$ are positive quadratic splines with local support. The quasi-interpolant is obtained by quasi-interpolating the error function, that is

$$\tilde{\sigma}(x) := \sigma(x) + \sum_{i=0}^{n+1} w_i e(x_i) C_i(x), \quad x \in [x_1, x_n]$$

with $\{w_i\}_{i=0}^{n+1}$ suitably chosen weights, $x_0, x_{n+1}, f_0, f_{n+1}$ additional data and $e(x_i) := f_i - \sigma(x_i)$, $i = 0, \dots, n+1$.

This strategy provides the convexity preservation, the error reduction at the data points, the $\mathbb{P}_1$ -reproduction and the end point interpolation.

The outline of the paper is as follows. In section 2 some background information is given and in section 3 the strategy for improving the quasi-interpolant is presented. Also the reduction of the error at the data points is investigated together with the shape preservation properties. Finally, in section 4 some pictures are given to illustrate the good performance of the method.

2 Preliminaries

As we refer to the quasi-interpolant scheme given in [3], we start by recalling it shortly. We also include in this section two new theorems useful for further discussions. We begin by recalling the following

Definition 2.1 *A set of real values $\{(x_i, f_i)\}_{i=1, \dots, n}$ with $x_1 < \dots < x_n$, is called monotone increasing (decreasing) if $f_i \leq f_{i+1}$, $i = 1, \dots, n-1$ ($f_i > f_{i+1}$, $i = 1, \dots, n-1$) and convex (concave) if $\Delta_{i-1} \leq \Delta_i \leq \Delta_{i+1}$, $i = 1, \dots, n-2$ ($\Delta_{i-1} \geq \Delta_i \geq \Delta_{i+1}$, $i = 1, \dots, n-2$), where $\Delta_i = \frac{f_{i+1} - f_i}{x_{i+1} - x_i}$, $i = 1, \dots, n-1$.*

Given the data set $\{f_i, x_i\}_{i=1}^n$ we define $h_i = x_{i+1} - x_i$, $i = 1, \dots, n-1$ and the extra knots $x_{1-i} = x_1 - i \cdot h_1$, $i = 1, 2, 3$ and $x_{n+i} = x_{n-1} + i \cdot h_{n-1}$, $i = 1, 2, 3$. Then, we approximate the first derivatives of any function g at the point x_i by $D_{2,i}^1 g(x)|_{x=x_i} = \sum_{k=-1}^1 \gamma_k^i g(x_{i+k})$, where the coefficients $\{\gamma_k^i\}_{k=-1,0,1}$ are $\gamma_{-1}^i = \frac{-h_i}{h_{i-1}(h_{i-1}+h_i)}$, $\gamma_0^i = \frac{h_i - h_{i-1}}{h_{i-1}h_i}$, $\gamma_1^i = \frac{h_{i-1}}{h_i(h_{i-1}+h_i)}$. It is easy to show that $D_{2,i}^1 g(x)|_{x=x_i}$ is a $\mathbb{P}_2$ -exact approximation of the first derivative of g at x_i ,

i.e. it is exact for every polynomial of degree two. Next, we define the functions

$$C_i(x) = \frac{1}{h_{i-1}} D_{2,i-1}^1 v(x-y)|_{y=x_{i-1}} - \frac{(h_i+h_{i-1})}{h_i h_{i-1}} D_{2,i}^1 v(x-y)|_{y=x_i} + \frac{1}{h_i} D_{2,i+1}^1 v(x-y)|_{y=x_{i+1}}, \quad i = 0, \dots, n+1, \quad (1)$$

where the discrete derivatives of $v(x-y) = -\frac{1}{4}|x-y|(x-y)$ are done with respect to the y variable at the specified points. To be more specific, the functions C_i , $i = 0, \dots, n+1$ are

$$C_i(x) = \frac{-1}{h_{i-2}(h_{i-2}+h_{i-1})} v(x-x_{i-2}) + \frac{1}{h_{i-2}h_{i-1}} v(x-x_{i-1}) + \left(\frac{1}{h_{i-1}(h_{i-2}+h_{i-1})} + \frac{1}{h_i(h_{i-1}+h_i)} \right) v(x-x_i) + \frac{-1}{h_i h_{i+1}} v(x-x_{i+1}) + \frac{1}{h_{i+1}(h_i+h_{i+1})} v(x-x_{i+2}), \quad (2)$$

that is quadratic splines. These functions are used to define the quasi-interpolant $\sigma(x) := \sum_{i=0}^{n+1} f_i C_i(x)$ where the extra values are $f_0 := 2f_1 - f_2$ and $f_{n+1} := 2f_n - f_{n-1}$. The properties of σ are summarized in the following theorem whose proof is given in [3].

Theorem 2.1 *The functions $\{C_i\}_{i \in \mathbb{Z}}$ defined in (2) are piecewise quadratic positive functions with local support on $[x_{i-2}, x_{i+2}]$, for $i \in \mathbb{Z}$. The quasi-interpolant $\sigma(x) = \sum_{i=0}^{n+1} f_i C_i(x)$ is P_1 -reproducing and interpolating the end points (x_1, f_1) , (x_n, f_n) . Furthermore, if the data set is monotone and/or convex (concave), so is the quasi-interpolant σ .*

The second theorem is about derivative interpolation at the end points. This can be useful when dealing with locally monotone and locally convex/concave data. The proof easily follows from the C_i definition by differentiation.

Theorem 2.2 *The quasi-interpolant $\sigma(x) = \sum_{i=0}^{n+1} f_i C_i(x)$ is such that $\sigma'(x_1) = \Delta_1$ and $\sigma'(x_n) = \Delta_{n-1}$.*

We conclude this section with a theorem discussing the error behaviour at the data points i.e. $e(x_i) = f_i - \sigma(x_i)$, $1 \leq i \leq n$.

Theorem 2.3 *The sign of the error at the data points depends on the convexity/concavity of the data.*

Proof. For $1 \leq i \leq n$, we have $C_{i-1}(x_i) = \frac{h_i}{2(h_{i-1}+h_i)}$, $C_i(x_i) = \frac{1}{2}$ and $C_{i+1}(x_i) = \frac{h_{i-1}}{2(h_{i-1}+h_i)}$. Thus, the error at the knots can be written as

$$\begin{aligned} e(x_i) &:= f_i - f_{i-1}C_{i-1}(x_i) - f_iC_i(x_i) - f_{i+1}C_{i+1}(x_i) \\ &= \frac{1}{2}f_i - f_{i-1}\frac{h_i}{2(h_{i-1}+h_i)} - f_{i+1}\frac{h_{i-1}}{2(h_{i-1}+h_i)} \\ &= \frac{h_{i-1}h_i}{2(h_{i-1}+h_i)}(\Delta_{i-1} - \Delta_i), \quad 1 \leq i \leq n. \quad \blacksquare \end{aligned} \tag{3}$$

3 Improving the convexity preserving quasi-interpolation

In this section we present a strategy to define a new convexity preserving quasi-interpolant method. Without loss of generality, convex data is assumed. The concave case can be treated similarly.

Based on the properties of σ , the idea is to define a new quasi-interpolant $\tilde{\sigma}$ by quasi-interpolating the error at the data points, i.e. by considering

$$\tilde{\sigma}(x) := \sigma(x) + \sum_{i=0}^{n+1} w_i e(x_i) C_i(x), \tag{4}$$

with $e(x_i) := f_i - \sigma(x_i)$, $0 \leq i \leq n$ and the weights w_i , $0 \leq i \leq n$ positive real values. The extra values are defined as $e(x_0) := -e(x_2)$, $e(x_{n+1}) := -e(x_{n-1})$. Obviously, the quasi-interpolant $\tilde{\sigma}$ can be also written as $\sum_{i=0}^{n+1} \tilde{f}_i C_i(x)$, having set $\{\tilde{f}_i := f_i + w_i e(x_i)\}_{i=0}^{n+1}$. A lemma states the properties of $\tilde{\sigma}$.

Lemma 3.1 *The quasi-interpolant $\tilde{\sigma}$ defined in (4) is $\mathbb{P}_1$ -reproducing. Furthermore, setting $w_0 := w_2$ and $w_{n+1} := w_{n-1}$ the end point interpolation is kept. Last, defining the new error at the data points as $\tilde{e}(x_i) := f(x_i) - \tilde{\sigma}(x_i)$, $i = 1, \dots, n$, then the following holds for $i = 1, \dots, n$,*

$$\tilde{e}(x_i) = e(x_i) - w_{i-1} \frac{h_i e(x_{i-1})}{2(h_{i-1} + h_i)} - w_i \frac{e(x_i)}{2} - w_{i+1} \frac{e(x_{i+1}) h_{i-1}}{2(h_{i-1} + h_i)}.$$

Proof. The first statement obviously follows from the $\mathbb{P}_1$ -reproduction of σ , while the second one is just a matter of computation. The last one is the

consequence of the local support of C_i and their values at the knots being $C_{i-1}(x_i) = \frac{h_i}{2(h_{i-1}+h_i)}$, $C_i(x_i) = \frac{1}{2}$ and $C_{i+1}(x_i) = \frac{h_{i-1}}{2(h_{i-1}+h_i)}$. ■

3.1 Error reduction

Let us now investigate the error behaviour of the proposed procedure and the influence of the weight choice. First, we give the following theorem whose proof is trivial.

Theorem 3.1 *Let $\{w_i\}_{i=0}^{n+1}$ be equal to a given positive constant less or equal to 2. Then $|\tilde{e}(x_i)| \leq \max_{j=1,\dots,n} |e(x_j)|$, for $1 \leq i \leq n$.*

Figure 3.1. shows the behaviour of the quasi-interpolants and of piecewise linear functions interpolating the error at the knots. The graphs are obtained by iterating three times the procedure (4) assuming $w_i = 1$, $i = 0, \dots, n+1$.

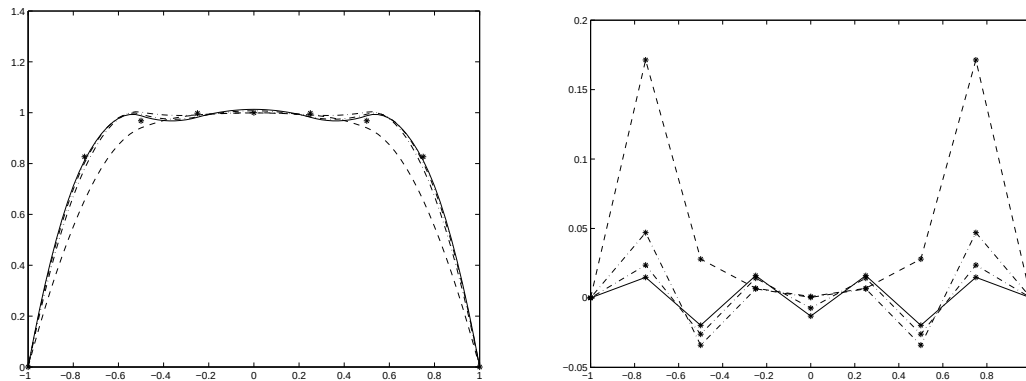


Fig. 3.1. *Quasi-interpolants (left) and errors (right)*

Next, we discuss the case of non-uniform weights. In particular, Theorem 3.2. analyzes the influence of the weights determined by the algorithm below.

Algorithm 3.1

1. Set $w_i := 0$ for $i = 0, \dots, n+1$
2. If $e(x_2) \neq 0 \& e(x_3) \neq 0$, set $w_2 := \mathcal{C} \cdot \min\{1, \frac{(h_2+h_3)e(x_3)}{h_3e(x_2)}\}$
3. For $i = 3, \dots, n-2$
 - 3.1. If $e(x_{i-1}) \neq 0 \& e(x_i) \neq 0 \& e(x_{i+1}) \neq 0$
set $w_i := \min\{a_i, b_i, c_i\}$
with $a_i := \mathcal{C} - \frac{w_{i-1}h_i e(x_{i-1})}{(h_{i-1}+h_i)e(x_i)}$, $b_i := \frac{\mathcal{C}(h_i+h_{i+1})e(x_{i+1})}{h_{i+1}e(x_i)}$ and
 $c_i := \frac{-(h_{i-2}+h_i)}{h_{i-2}e(x_i)}[(w_{i-1} - \mathcal{C})e(x_{i-1}) + \frac{w_{i-2}h_{i-1}e(x_{i-2})}{(h_{i-2}+h_{i-1})}]$
4. If $e(x_{n-2}) \neq 0 \& e(x_{n-1}) \neq 0$
set $w_{n-1} := \min\{a_{n-1}, b_{n-1}\}$
with $a_{n-1} := \mathcal{C} - \frac{w_{n-2}h_{n-1}e(x_{n-2})}{(h_{n-2}+h_{n-1})e(x_{n-2})}$,
and $b_{n-1} := \frac{-(h_{n-3}+h_{n-2})}{h_{n-3}e(x_{n-1})}[(w_{n-2} - \mathcal{C})e(x_{n-2}) + \frac{w_{n-3}h_{n-2}e(x_{n-3})}{(h_{n-3}+h_{n-2})}]$
5. set $w_0 := w_2$
6. set $w_{n+1} := w_{n-1}$

Theorem 3.2 Let $\{w_i\}_{i=0}^{n+1}$ non negative real values defined by means of the Algorithm 3.1, assuming $\mathcal{C} \leq 4$. If $\{f_i\}_{i=1}^n$ is a convex data set, then $|\tilde{e}(x_i)| \leq |e(x_i)|$ for $1 \leq i \leq n$. Furthermore, if $\mathcal{C} \leq 2$ it also follows that $\tilde{e}(x_i) \leq 0$ for $1 \leq i \leq n$.

Proof. Assuming $|\tilde{e}(x_i)| \leq |e(x_i)|$ for $i = 2, \dots, n-2$ leads to

$$2e(x_i) \leq w_{i-1} \frac{h_i e(x_{i-1})}{2(h_{i-1} + h_i)} + w_i \frac{e(x_i)}{2} + w_{i+1} \frac{e(x_{i+1})h_{i-1}}{2(h_{i-1} + h_i)} \leq 0.$$

The inequality on the right hand side is always satisfied. For the inequality on the left hand side the following six cases are possible:

If $e(x_i) = 0$, then $w_{i-1} := 0$, $w_{i+1} := 0$.

If $e(x_i) \neq 0$, then

if $e(x_{i-1}) = 0 \& e(x_{i+1}) = 0$, then $w_i \leq \mathcal{C}$,

if $e(x_{i-1}) = 0$ & $e(x_{i+1}) \neq 0$, then $w_i \leq \mathcal{C}$, $w_{i+1} \leq \frac{(h_{i-1}+h_i)[(w_i-\mathcal{C})e(x_i)]}{-h_{i-1}e(x_{i+1})}$,

if $e(x_{i-1}) \neq 0$ & $e(x_{i+1}) = 0$, then

$$w_{i-1} \leq \frac{\mathcal{C}(h_{i-1}+h_i)e(x_i)}{h_i e(x_{i-1})}, \quad w_i \leq \mathcal{C} - \frac{w_{i-1}h_i e(x_{i-1})}{(h_{i-1}+h_i)e(x_i)},$$

if $e(x_{i-1}) \neq 0$ & $e(x_{i+1}) \neq 0$, then

$$w_{i-1} \leq \frac{\mathcal{C}(h_{i-1}+h_i)e(x_i)}{h_i e(x_{i-1})}, \quad w_i \leq \mathcal{C} - \frac{w_{i-1}h_i e(x_{i-1})}{(h_{i-1}+h_i)e(x_i)}$$

$$w_{i+1} \leq -\frac{(h_{i-1}+h_i)[(w_i-\mathcal{C})e(x_i) + \frac{w_{i-1}h_i e(x_{i-1})}{h_{i-1}+h_i}]}{h_{i-1}e(x_{i+1})}.$$

Combining all the previous conditions on w_i , we conclude that the choice $\mathcal{C} \leq 4$ implies $|\tilde{e}(x_i)| \leq |e(x_i)|$, $i = 1, \dots, n$.

Moreover, if we also require the non-positivity of $\tilde{e}(x_i)$, then

$$w_{i-1} \frac{h_i e(x_{i-1})}{2(h_{i-1}+h_i)} + w_i \frac{e(x_i)}{2} + w_{i+1} \frac{e(x_{i+1})h_{i-1}}{2(h_{i-1}+h_i)} \geq e(x_i).$$

This is true whenever $\{w_i\}_{i=0}^{n+1}$ are set as in Algorithm 3.1 with $\mathcal{C} \leq 2$. ■

3.2 Convexity preservation

To keep convexity, it is convenient to set the weights w_i in such a way that the new data set $\{\tilde{f}_i = f_i + w_i e(x_i)\}_{i=0}^{n+1}$ is convex as well. Thus, let us consider the quantities $\tilde{\Delta}_i = \frac{\tilde{f}_{i+1} - \tilde{f}_i}{h_i}$, $i = 0, \dots, n$. It is easy to see that $\tilde{\Delta}_1 - \tilde{\Delta}_0 = \Delta_1 - \Delta_0$, $\tilde{\Delta}_n - \tilde{\Delta}_{n-1} = \Delta_n - \Delta_{n-1}$ and

$$\begin{aligned} \tilde{\Delta}_2 - \tilde{\Delta}_1 &= (\Delta_2 - \Delta_1) + \frac{w_3 e(x_3) - w_2 e(x_2)}{h_2} - \frac{w_2 e(x_2)}{h_1}, \\ \tilde{\Delta}_{i+1} - \tilde{\Delta}_i &= (\Delta_{i+1} - \Delta_i) + \frac{w_{i+2} e(x_{i+2}) - w_{i+1} e(x_{i+1})}{h_{i+1}} \\ &\quad - \frac{w_{i+1} e(x_{i+1}) - w_i e(x_i)}{h_i}, \quad i = 2, \dots, n-3, \\ \tilde{\Delta}_{n-1} - \tilde{\Delta}_{n-2} &= (\Delta_{n-1} - \Delta_{n-2}) - \frac{w_{n-1} e(x_{n-1})}{h_{n-1}} \\ &\quad - \frac{w_{n-1} e(x_{n-1}) - w_{n-2} e(x_{n-2})}{h_{n-2}}. \end{aligned} \tag{5}$$

The following algorithm can be used for choosing the weights in order to preserve convexity.

Algorithm 3.2

1. Set $w_i := 0$, $i = 0, \dots, n+1$
2. If $e(x_2) \neq 0$, set

$$w_2 := \frac{h_2(\Delta_2 - \Delta_3)}{e(x_2)}$$
3. For $i = 3, \dots, n-2$
 3.1. If $e(x_i) \neq 0$ set

$$w_i := \min\left\{\frac{h_i(\Delta_i - \Delta_{i+1})}{e(x_i)}, \frac{h_{i-1}(\Delta_{i-2} - \Delta_{i-1} + w_{i-1}e(x_{i-1})\frac{h_{i-1}+h_i}{h_{i-1}h_i} - w_{i-2}\frac{e(x_{i-2})}{h_{i-2}})}{e(x_i)}\right\}$$
4. If $e(x_{n-1}) \neq 0$, set

$$w_{n-1} := \frac{h_{n-2}(\Delta_{n-3} - \Delta_{n-2} + w_{n-2}e(x_{n-2})\frac{h_{n-2}+h_{n-3}}{h_{n-2}h_{n-3}} - w_{n-3}\frac{e(x_{n-3})}{h_{n-3}})}{e(x_{n-1})}$$
5. set $w_0 := w_2$
6. set $w_{n+1} := w_{n-1}$

Theorem 3.3 Let $\{w_i\}_{i=0}^{n+1}$ non negative real values defined via Algorithm 3.2 and $\{f_i\}_{i=1}^n$ be a convex data set. Then the new data set $\{\tilde{f}_i\}_{i=1}^n$ is also convex.

The proof of this result based on (5) is along the same lines as the proof of Theorem 3.2. ■

3.3 Monotonicity preservation

Now, assuming the data set is convex and monotone, the combination of the following algorithms with the previous one guarantees the convexity and the monotonicity preservation. The first algorithm is related to the monotone increasing case, while the second one to the monotone decreasing case.

Algorithm 3.3

1. For $i = 2, \dots, n-1$, set $w_i := \mathcal{C}$
2. For $i = 2, \dots, n-1$
if $e(x_{i+1}) \neq 0$ set $w_{i+1} := \frac{-h_i \Delta_i + w_i e(x_i)}{e(x_{i+1})}$
3. set $w_0 := w_2$
4. set $w_{n+1} := w_{n-1}$

Algorithm 3.4

1. For $i = 2, \dots, n-1$, set $w_i := \mathcal{C}$
2. For $i = n-1, \dots, 2$
if $e(x_i) \neq 0$, set $w_i := \frac{h_i \Delta_i + w_{i+1} e(x_{i+1})}{e(x_i)}$
3. set $w_0 := w_2$
4. set $w_{n+1} := w_{n-1}$

Theorem 3.4 Let $\{w_i\}_{i=0}^{n+1}$ non negative real values defined via the previous algorithms where $\mathcal{C}$ is any non-negative real value. Let $\{f_i\}_{i=1}^n$ be convex and monotone data. Then, the new data set $\{\tilde{f}_i\}_{i=0}^{n+1}$ is monotone as well.

Now, as the reduction of the error is obviously desirable, in case convex data or convex and monotone data are given, the combination of the above algorithms (3.1, 3.2, 3.3 or 3.4) yields an improved shape preserving quasi-interpolant. Here, combination means that the final set of weights is defined by taking, elementwise, the minimum among the set produced by the algorithms.

Remarks. It is important to note that, due to the non symmetry of the proposed algorithms, even in case of symmetric data the improved quasi-interpolant $\tilde{\sigma}$ given in (4) turns out to be non symmetric. A trivial solution of this drawback is the weight “symmetrization” by taking the minimum between the corresponding weights. In addition, we could performe algorithms analogous to the one discussed above also in a converse way (i.e. checking the inequality from right to left). After symmetrization of this second sweep the highest symmetric weights (between the two sweeps) can be used in (4) to reduce the error as much as possible.

According to the results given in Theorem 3.2, the choice $\mathcal{C} \leq 2$ in Algorithm 3.1 allows us to iterate the procedure until an index i , such that both $w_i \neq 0$ and $e(x_i) \neq 0$, is found. Furthermore, locally convex and monotone data can be handled taking into consideration the results of Theorem 2.2.

4 Examples

We conclude this paper by showing some applications of the described procedure. Three examples are given to put in evidence its effectiveness. For each of them two pictures are presented. The ones on the left display the shape preserving quasi-interpolants $\sigma(-\cdot)$ and $\tilde{\sigma}(-)$ and the ones on the right the piecewise linear functions interpolating the values of the error at the knots, i.e. $e(x_i), i = 1 \dots, n$ ($-\cdot$) and $\tilde{e}(x_i), i = 1 \dots, n$ ($-$), respectively. In the pictures on the left the data sets $\{f_i\}_{i=1}^n$ and $\{\tilde{f}_i\}_{i=1}^n$ are denoted by the symbols ‘*’ and ‘+’, respectively. In all the examples the constant is set as $\mathcal{C} = 2$.

The first example is related to data coming from the function $f(x) = -\log(x)$ when evaluated at the abscissae $\{0.01, 0.7, 3.6, 8, 13.4, 17\}$. This set of data is convex and monotone decreasing. Thus, the weights are obtained combining algorithms 3.1, 3.2, 3.4.

The second one concerns the function $f(x) = x - \sqrt{x^2/2 - x - 4}$ and the points $\{4.1, 4.8, 1 + 3\sqrt{2}, 5.9, 6.4, 7, 7.5\}$. Here we combine algorithms 3.1, 3.2 as we deal with convex data.

The last example deals with concave data taken by evaluating the function $f(x) = \sqrt{1 - x^2}$ at the points $\{-1, -0.75, -0.35, 0, .5, .75, 1\}$. The concave quasi-interpolant is performed via algorithm 3.1 and the analogous to algorithm 3.2 for concave data.

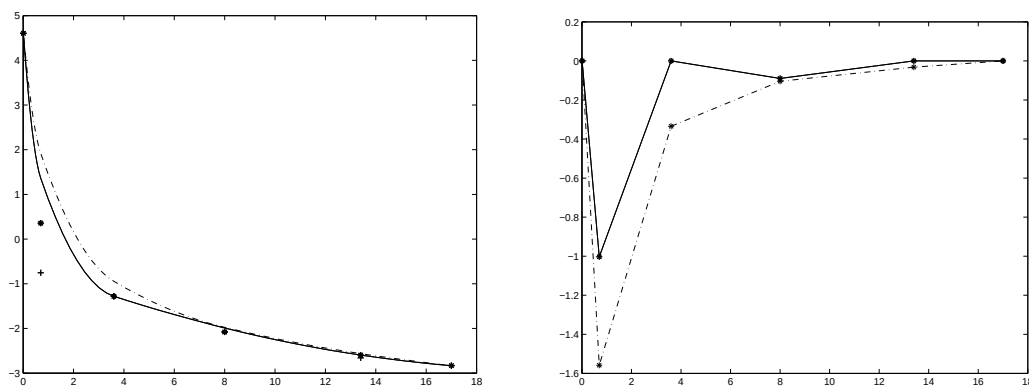


Fig. 4.1. *Quasi-interpolants (left) and errors (right)*

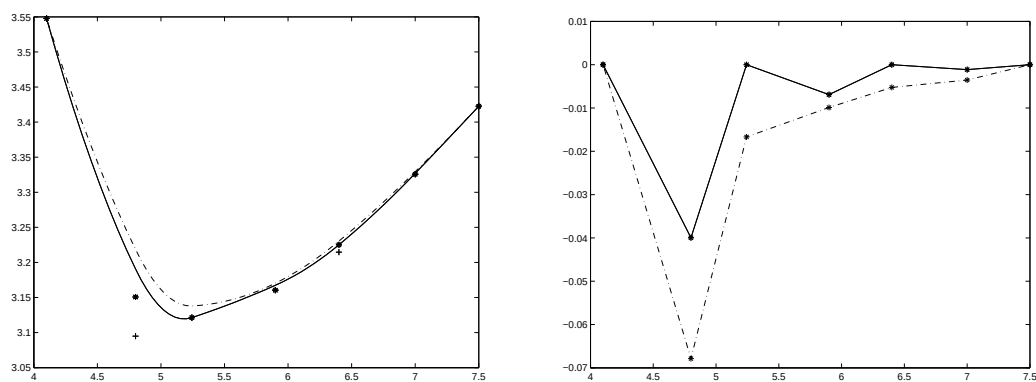


Fig. 4.2. *Quasi-interpolants (left) and errors (right)*

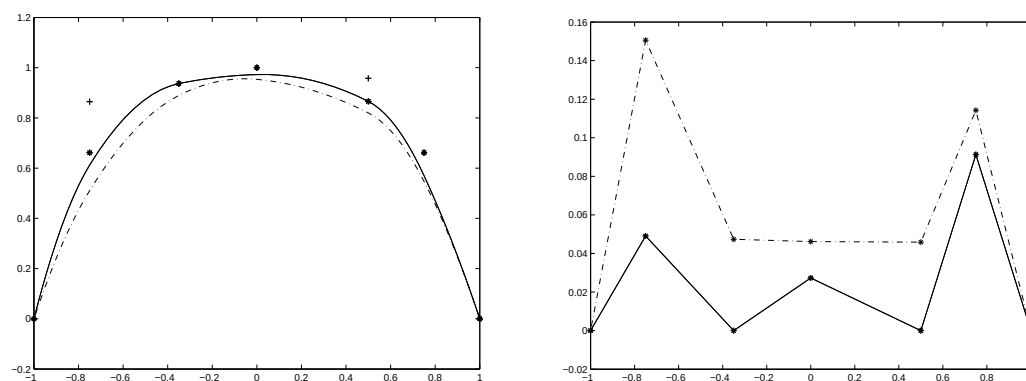


Fig. 4.3. *Quasi-interpolants (left) and errors (right)*

References

- [1] R. K. Beatson and M. J. D. Powell, Univariate multiquadric approximation, *Constructive Approximation*, 8, 275–288 (1992).
- [2] C. Conti and R. Morandi, Piecewise C^1 Shape-Preserving Hermite Interpolation, *Computing*, 56, 323–341 (1996).
- [3] C. Conti, R. Morandi and C. Rabut, $\mathbb{P}_1$ -reproducing, Shape-preserving Quasi-Interpolant Using Positive Quadratic Splines with Local Support, *Journal of Computational and Applied Mathematics*, 104, 111–122 (1999).

- [4] P. Costantini, An Algorithm for Computing Shape-Preserving Interpolating Spline of Arbitrary Degree, *Journal of Computational and Applied Mathematics*, 22, 89–136 (1988).
- [5] Z. Wu and R. Schaback, Shape preserving properties and convergence of univariate multiquadric quasi-interpolation, *Acta Mathematicae Sinica, Study Bulletin*, 10, No. 4, 441–446 (1994).

Instructions to Contributors
Journal of Concrete and Applicable Mathematics

A quarterly international publication of Eudoxus Press, LLC, of TN.

Editor in Chief: George Anastassiou

Department of Mathematical Sciences
 University of Memphis
 Memphis, TN 38152-3240, U.S.A.

1. Manuscripts hard copies in triplicate, and in English, should be submitted to the Editor-in-Chief:

Prof. George A. Anastassiou
 Department of Mathematical Sciences
 The University of Memphis
 Memphis, TN 38152, USA.
 Tel. 001. 901.678.3144
 e-mail: ganastss@memphis.edu.

Authors may want to recommend an associate editor the most related to the submission to possibly handle it.

Also authors may want to submit a list of six possible referees, to be used in case we cannot find related referees by ourselves.

2. Manuscripts should be typed using any of TEX, LaTeX, AMS-TEX, or AMS-LaTeX and according to EUDOXUS PRESS, LLC. LATEX STYLE FILE. (VISIT www.msci.memphis.edu/~ganastss/jcaam/ to save a copy of the style file.) They should be carefully prepared in all respects. Submitted copies should be brightly printed (not dot-matrix), double spaced, in ten point type size, on one side high quality paper 8(1/2)x11 inch. Manuscripts should have generous margins on all sides and should not exceed 24 pages.

3. Submission is a representation that the manuscript has not been published previously in this or any other similar form and is not currently under consideration for publication elsewhere. A statement transferring from the authors (or their employers, if they hold the copyright) to Eudoxus Press, LLC, will be required before the manuscript can be accepted for publication. The Editor-in-Chief will supply the necessary forms for this transfer. Such a written transfer of copyright, which previously was assumed to be implicit in the act of submitting a manuscript, is necessary under the U.S. Copyright Law in order for the publisher to carry through the dissemination of research results and reviews as widely and effectively as possible.

4. The paper starts with the title of the article, author's name(s) (no titles or degrees), author's affiliation(s) and e-mail addresses. The affiliation should comprise the department, institution (usually university or company), city, state (and/or nation) and mail code.

The following items, 5 and 6, should be on page no. 1 of the paper.

5. An abstract is to be provided, preferably no longer than 150 words.
 6. A list of 5 key words is to be provided directly below the abstract. Key words should express the precise content of the manuscript, as they are used for indexing purposes. The main body of the paper should begin on page no. 1, if possible.

7. All sections should be numbered with Arabic numerals (such as: 1. INTRODUCTION) .

Subsections should be identified with section and subsection numbers (such as 6.1. Second-Value Subheading).

If applicable, an independent single-number system (one for each category) should be used to label all theorems, lemmas, propositions, corollaries, definitions, remarks, examples, etc. The label (such as Lemma 7) should be typed with paragraph indentation, followed by a period and the lemma itself.

8. Mathematical notation must be typeset. Equations should be numbered consecutively with Arabic numerals in parentheses placed flush right, and should be thusly referred to in the text [such as Eqs.(2) and (5)]. The running title must be placed at the top of even numbered pages and the first author's name, et al., must be placed at the top of the odd numbered pages.

9. Illustrations (photographs, drawings, diagrams, and charts) are to be numbered in one consecutive series of Arabic numerals. The captions for illustrations should be typed double space. All illustrations, charts, tables, etc., must be embedded in the body of the manuscript in proper, final, print position. In particular, manuscript, source, and PDF file version must be at camera ready stage for publication or they cannot be considered.

Tables are to be numbered (with Roman numerals) and referred to by number in the text. Center the title above the table, and type explanatory footnotes (indicated by superscript lowercase letters) below the table.

10. List references alphabetically at the end of the paper and number them consecutively. Each must be cited in the text by the appropriate Arabic numeral in square brackets on the baseline.

References should include (in the following order):

initials of first and middle name, last name of author(s)

title of article,

name of publication, volume number, inclusive pages, and year of publication.

Authors should follow these examples:

Journal Article

1. H.H.Gonska, Degree of simultaneous approximation of bivariate functions by Gordon operators, (journal name in italics) *J. Approx. Theory*, 62,170-191(1990).

Book

2. G.G.Lorentz, (title of book in italics) *Bernstein Polynomials* (2nd ed.), Chelsea, New York, 1986.

Contribution to a Book

3. M.K.Khan, Approximation properties of beta operators, in (title of book in italics) *Progress in Approximation Theory* (P.Nevai and A.Pinkus, eds.), Academic Press, New York, 1991, pp.483-495.

11. All acknowledgements (including those for a grant and financial support) should occur in one paragraph that directly precedes the References section.

12. Footnotes should be avoided. When their use is absolutely necessary, footnotes should be numbered consecutively using Arabic numerals and should be typed at the bottom of the page to which they refer. Place a line above the footnote, so that it is set

off from the text. Use the appropriate superscript numeral for citation in the text.

13. After each revision is made please again submit three hard copies of the revised manuscript, including in the final one. And after a manuscript has been accepted for publication and with all revisions incorporated, manuscripts, including the TEX/LaTeX source file and the PDF file, are to be submitted to the Editor's Office on a personal-computer disk, 3.5 inch size. Label the disk with clearly written identifying information and properly ship, such as:

Your name, title of article, kind of computer used, kind of software and version number, disk format and files names of article, as well as abbreviated journal name.

Package the disk in a disk mailer or protective cardboard. Make sure contents of disks are identical with the ones of final hard copies submitted!

Note: The Editor's Office cannot accept the disk without the accompanying matching hard copies of manuscript. No e-mail final submissions are allowed! The disk submission must be used.

14. Effective 1 Nov. 2005 the journal's page charges are \$8.00 per PDF file page. Upon acceptance of the paper an invoice will be sent to the contact author. The fee payment will be due one month from the invoice date. The article will proceed to publication only after the fee is paid. The charges are to be sent, by money order or certified check, in US dollars, payable to Eudoxus Press, LLC, to the address shown in the Scope and Prices section.

No galleys will be sent and the contact author will receive one(1) electronic copy of the journal issue in which the article appears.

15. This journal will consider for publication only papers that contain proofs for their listed results.

TABLE OF CONTENTS, JOURNAL OF CONCRETE AND APPLICABLE
MATHEMATICS, VOL.4, NO.3, 2006

ON THE UNIFORM EXPONENTIAL STABILITY OF EVOLUTION FAMILIES IN TERMS OF THE ADMISSIBILITY OF AN ORLICZ SEQUENCE SPACE, C.PREDA.S.DRAGOMIR,C.CHILARESCU,.....	253
GENERALIZED QUASILINEARIZATION TECHNIQUE FOR THE SECOND ORDER DIFFERENTIAL EQUATIONS WITH GENERAL MIXED BOUNDARY CONDITIONS,B.AHMAD,R.KHAN,.....	267
EXISTENCE AND CONSTRUCTION OF FINITE TIGHT FRAMES, P.CASAZZA,M.LEON,.....	277
SEMIGROUPS AND GENETIC REPRESSION,G.FRAGNELLI,.....	291
THE IRREGULAR NESTING PROBLEM INVOLVING TRIANGLES AND RECTANGLES,L.FAINA,.....	307
WEAK CONVERGENCE AND PROKHOROV'S THEOREM FOR MEASURES IN A BANACH SPACE,M.ISIDORI,.....	325
FITTED MESH B-SPLINE COLLOCATION METHOD FOR SOLVING SINGULARLY PERTURBED REACTION-DIFFUSION PROBLEMS, M.KADALBAJOO,V.AGGARWAL,.....	349
A PROCEDURE FOR IMPROVING CONVEXITY PRESERVING QUASI- INTERPOLATION,C.CONTI,R.MORANDI,.....	367

VOLUME 4,NUMBER 4 OCTOBER 2006

ISSN:1548-5390 PRINT,1559-176X ONLINE
Special Issue on “Wavelets and Applications”



**JOURNAL
OF CONCRETE
AND APPLICABLE
MATHEMATICS**

EUDOXUS PRESS,LLC

SCOPE AND PRICES OF THE JOURNAL
Journal of Concrete and Applicable Mathematics

A quartely international publication of **Eudoxus Press, LLC**

Editor in Chief: George Anastassiou

Department of Mathematical Sciences,
University of Memphis
Memphis, TN 38152, U.S.A.
ganastss@memphis.edu

The main purpose of the "Journal of Concrete and Applicable Mathematics" is to publish high quality original research articles from all subareas of Non-Pure and/or Applicable Mathematics and its many real life applications, as well connections to other areas of Mathematical Sciences, as long as they are presented in a Concrete way. It welcomes also related research survey articles and book reviews. A sample list of connected mathematical areas with this publication includes and is not restricted to: Applied Analysis, Applied Functional Analysis, Probability theory, Stochastic Processes, Approximation Theory, O.D.E, P.D.E, Wavelet, Neural Networks, Difference Equations, Summability, Fractals, Special Functions, Splines, Asymptotic Analysis, Fractional Analysis, Inequalities, Moment Theory, Numerical Functional Analysis, Tomography, Asymptotic Expansions, Fourier Analysis, Applied Harmonic Analysis, Integral Equations, Signal Analysis, Numerical Analysis, Optimization, Operations Research, Linear Programming, Fuzzyness, Mathematical Finance, Stochastic Analysis, Game Theory, Math. Physics aspects, Applied Real and Complex Analysis, Computational Number Theory, Graph Theory, Combinatorics, Computer Science Math. related topics, combinations of the above, etc. In general any kind of Concretely presented Mathematics which is Applicable fits to the scope of this journal. Working Concretely and in Applicable Mathematics has become a main trend in many recent years, so we can understand better and deeper and solve the important problems of our real and scientific world. "Journal of Concrete and Applicable Mathematics" is a peer-reviewed International Quarterly Journal. We are calling for papers for possible publication. The contributor should send three copies of the contribution to the editor in-Chief typed in TEX, LATEX double spaced. [See: Instructions to Contributors]

Journal of Concrete and Applicable Mathematics(JCAAM)

ISSN:1548-5390 PRINT, 1559-176X ONLINE.

is published in January, April, July and October of each year by

EUDOXUS PRESS, LLC,

1424 Beaver Trail Drive, Cordova, TN 38016, USA,

Tel. 001-901-751-3553

anastassioug@yahoo.com

<http://www.EudoxusPress.com>.

Annual Subscription Current Prices: For USA and Canada, Institutional: Print \$250, Electronic \$220, Print and Electronic \$310. Individual: Print \$77, Electronic \$60, Print & Electronic \$110. For any other part of the world add \$25 more to the above prices for Print.

Single article PDF file for individual \$8. Single issue in PDF form for individual \$25.
The journal carries page charges \$8 per page of the pdf file of an article, payable upon acceptance of the article within one month and before publication.
No credit card payments. Only certified check, money order or international check in US dollars are acceptable.
Combination orders of any two from JoCAAA, JCAAM, JAFA receive 25% discount, all three receive 30% discount.

Copyright©2004 by Eudoxus Press, LLC all rights reserved. JCAAM is printed in USA.

JCAAM is reviewed and abstracted by AMS Mathematical Reviews, MATHSCI, and Zentralblatt MATH.

It is strictly prohibited the reproduction and transmission of any part of JCAAM and in any form and by any means without the written permission of the publisher. It is only allowed to educators to Xerox articles for educational purposes. The publisher assumes no responsibility for the content of published papers.

JCAAM IS A JOURNAL OF RAPID PUBLICATION

Editorial Board Associate Editors

Editor in -Chief:
George Anastassiou
Department of Mathematical Sciences
The University Of Memphis
Memphis, TN 38152, USA
tel. 901-678-3144, fax 901-678-2480
e-mail ganastss@memphis.edu
www.msci.memphis.edu/~ganastss/jcaam
Areas: Approximation Theory,
Probability, Moments, Wavelet,
Neural Networks, Inequalities, Fuzzyness.

Associate Editors:

1) Ravi Agarwal
Florida Institute of Technology
Applied Mathematics Program
150 W. University Blvd.
Melbourne, FL 32901, USA
agarwal@fit.edu
Differential Equations, Difference
Equations, inequalities

2) Shair Ahmad
University of Texas at San Antonio
Division of Math. & Stat.
San Antonio, TX 78249-0664, USA
SAHMAD@utsa.edu
Differential Equations, Mathematical
Biology

3) Drumi D. Bainov
Medical University of Sofia
P.O. Box 45, 1504 Sofia, Bulgaria
drumibainov@yahoo.com
Differential Equations, Optimal Control,
Numerical Analysis, Approximation Theory

4) Carlo Bardaro
Dipartimento di Matematica & Informatica
Universita' di Perugia
Via Vanvitelli 1
06123 Perugia, ITALY
tel. +390755855034, +390755853822,
fax +390755855024
bardaro@unipg.it ,
bardaro@dipmat.unipg.it
Functional Analysis and Approximation Th.,

19) Rupert Lasser
Institut fur Biomathematik & Biomertie, GSF
-National Research Center for environment and
health
Ingolstaedter landstr.1
D-85764 Neuherberg, Germany
lasser@gsf.de
Orthogonal Polynomials, Fourier Analysis,
Mathematical Biology

20) Alexandru Lupas
University of Sibiu
Faculty of Sciences
Department of Mathematics
Str. I. Ratiu nr. 7
2400-Sibiu, Romania
lupas@ulbsibiu.ro
Classical Analysis, Inequalities,
Special Functions, Umbral Calculus,
Approximation Th., Numerical Analysis
and Methods

21) Ram N. Mohapatra
Department of Mathematics
University of Central Florida
Orlando, FL 32816-1364
tel. 407-823-5080
ramm@mail.ucf.edu
Real and Complex analysis, Approximation Th.,
Fourier Analysis, Fuzzy Sets and Systems

22) Rainer Nagel
Arbeitsbereich Funktionalanalysis
Mathematisches Institut
Auf der Morgenstelle 10
D-72076 Tuebingen
Germany
tel. 49-7071-2973242
fax 49-7071-294322
rana@fa.uni-tuebingen.de
evolution equations, semigroups, spectral th.,
positivity

23) Panos M. Pardalos
Center for Appl. Optimization
University of Florida
303 Weil Hall
P.O. Box 116595
Gainesville, FL 32611-6595

Summability, Signal Analysis, Integral
Equations,
Measure Th., Real Analysis

5) Francoise Bastin
Institute of Mathematics
University of Liege
4000 Liege
BELGIUM
f.bastin@ulg.ac.be
Functional Analysis, Wavelets

6) Paul L. Butzer
RWTH Aachen
Lehrstuhl A für Mathematik
D-52056 Aachen
Germany
tel. 0049/241/80-94627 office,
0049/241/72833 home,
fax 0049/241/80-92212
Butzer@rwth-aachen.de
Approximation Th., Sampling Th., Signals,
Semigroups of Operators, Fourier Analysis

7) Yeol Je Cho
Department of Mathematics Education
College of Education
Gyeongsang National University
Chinju 660-701
KOREA
tel. 055-751-5673 Office,
055-755-3644 home,
fax 055-751-6117
yjcho@nongae.gsnu.ac.kr
Nonlinear operator Th., Inequalities,
Geometry of Banach Spaces

8) Sever S. Dragomir
School of Communications and Informatics
Victoria University of Technology
PO Box 14428
Melbourne City M.C
Victoria 8001, Australia
tel 61 3 9688 4437, fax 61 3 9688 4050
sever.dragomir@vu.edu.au,
sever@sci.vu.edu.au
Math. Analysis, Inequalities, Approximation
Th.,
Numerical Analysis, Geometry of Banach
Spaces,
Information Th. and Coding

9) A.M. Fink
Department of Mathematics
Iowa State University
Ames, IA 50011-0001, USA

tel. 352-392-9011
pardalos@ufl.edu
Optimization, Operations Research

24) Svetlozar T. Rachev
Dept. of Statistics and Applied Probability
Program
University of California, Santa Barbara
CA 93106-3110, USA
tel. 805-893-4869
rachev@pstat.ucsb.edu
AND
Chair of Econometrics and Statistics
School of Economics and Business Engineering
University of Karlsruhe
Kollegium am Schloss, Bau II, 20.12, R210
Postfach 6980, D-76128, Karlsruhe, Germany
tel. 011-49-721-608-7535
rachev@lsoe.uni-karlsruhe.de
Mathematical and Empirical Finance,
Applied Probability, Statistics and Econometrics

25) Paolo Emilio Ricci
Universita' degli Studi di Roma "La Sapienza"
Dipartimento di Matematica-Istituto
"G. Castelnuovo"
P.le A. Moro, 2-00185 Roma, ITALY
tel. ++39 0649913201, fax ++39 0644701007
riccip@uniroma1.it, Paoloemilio.Ricci@uniroma1.it
Orthogonal Polynomials and Special functions,
Numerical Analysis, Transforms, Operational
Calculus,
Differential and Difference equations

26) Cecil C. Rousseau
Department of Mathematical Sciences
The University of Memphis
Memphis, TN 38152, USA
tel. 901-678-2490, fax 901-678-2480
ccrousse@memphis.edu
Combinatorics, Graph Th.,
Asymptotic Approximations,
Applications to Physics

27) Tomasz Rychlik
Institute of Mathematics
Polish Academy of Sciences
Chopina 12, 87100 Torun, Poland
T.Rychlik@impan.gov.pl
Mathematical Statistics, Probabilistic
Inequalities

28) Bl. Sendov
Institute of Mathematics and Informatics
Bulgarian Academy of Sciences
Sofia 1090, Bulgaria

tel.515-294-8150
fink@math.iastate.edu
Inequalities, Ordinary Differential
Equations

10) Sorin Gal
Department of Mathematics
University of Oradea
Str. Armatei Romane 5
3700 Oradea, Romania
galso@uoradea.ro
Approximation Th., Fuzzyness, Complex
Analysis

11) Jerome A. Goldstein
Department of Mathematical Sciences
The University of Memphis,
Memphis, TN 38152, USA
tel. 901-678-2484
jgoldste@memphis.edu
Partial Differential Equations,
Semigroups of Operators

12) Heiner H. Gonska
Department of Mathematics
University of Duisburg
Duisburg, D-47048
Germany
tel. 0049-203-379-3542 office
gonska@informatik.uni-duisburg.de
Approximation Th., Computer Aided
Geometric Design

13) Dmitry Khavinson
Department of Mathematical Sciences
University of Arkansas
Fayetteville, AR 72701, USA
tel. (479) 575-6331, fax (479) 575-8630
dmitry@uark.edu
Potential Th., Complex Analysis, Holomorphic
PDE, Approximation Th., Function Th.

14) Virginia S. Kiryakova
Institute of Mathematics and Informatics
Bulgarian Academy of Sciences
Sofia 1090, Bulgaria
virginia@diogenes.bg
Special Functions, Integral Transforms,
Fractional Calculus

15) Hans-Bernd Knoop
Institute of Mathematics
Gerhard Mercator University
D-47048 Duisburg
Germany
tel. 0049-203-379-2676

bSENDov@bas.bg
Approximation Th., Geometry of Polynomials,
Image Compression

29) Igor Shevchuk
Faculty of Mathematics and Mechanics
National Taras Shevchenko
University of Kyiv
252017 Kyiv
UKRAINE
shevchuk@univ.kiev.ua
Approximation Theory

30) H.M. Srivastava
Department of Mathematics and Statistics
University of Victoria
Victoria, British Columbia V8W 3P4
Canada
tel. 250-721-7455 office, 250-477-6960 home,
fax 250-721-8962
harimsri@math.uvic.ca
Real and Complex Analysis, Fractional Calculus
and Appl.,
Integral Equations and Transforms, Higher
Transcendental
Functions and Appl., q-Series and q-Polynomials,
Analytic Number Th.

31) Ferenc Szidarovszky
Dept. Systems and Industrial Engineering
The University of Arizona
Engineering Building, 111
PO Box 210020
Tucson, AZ 85721-0020, USA
szidar@sie.arizona.edu
Numerical Methods, Game Th., Dynamic Systems,
Multicriteria Decision making,
Conflict Resolution, Applications
in Economics and Natural Resources
Management

32) Gancho Tachev
Dept. of Mathematics
Univ. of Architecture, Civil Eng. and Geodesy
1 Hr. Smirnenski blvd
BG-1421 Sofia, Bulgaria
Approximation Theory

33) Manfred Tasche
Department of Mathematics
University of Rostock
D-18051 Rostock
Germany
manfred.tasche@mathematik.uni-rostock.de
Approximation Th., Wavelet, Fourier Analysis,
Numerical Methods, Signal Processing,

knoop@math.uni-duisburg.de
Approximation Theory, Interpolation

16) Jerry Koliha
Dept. of Mathematics & Statistics
University of Melbourne
VIC 3010, Melbourne
Australia
koliha@unimelb.edu.au
Inequalities, Operator Theory,
Matrix Analysis, Generalized Inverses

17) Mustafa Kulenovic
Department of Mathematics
University of Rhode Island
Kingston, RI 02881, USA
kulenm@math.uri.edu
Differential and Difference Equations

18) Gerassimos Ladas
Department of Mathematics
University of Rhode Island
Kingston, RI 02881, USA
gladas@math.uri.edu
Differential and Difference Equations

Image Processing, Harmonic Analysis

34) Chris P. Tsokos
Department of Mathematics
University of South Florida
4202 E. Fowler Ave., PHY 114
Tampa, FL 33620-5700, USA
profcpt@math.usf.edu, profcpt@chuma1.cas.usf.edu
Stochastic Systems, Biomathematics,
Environmental Systems, Reliability Th.

35) Lutz Volkmann
Lehrstuhl II fuer Mathematik
RWTH-Aachen
Templergraben 55
D-52062 Aachen
Germany
volkm@math2.rwth-aachen.de
Complex Analysis, Combinatorics, Graph Theory

Preface

This special issue contains invited research papers on “Wavelets and Applications.” Results in wavelets, frames, subdivision schemes, spline wavelets, and wavelet applications in statistics/bio-statistics were among the areas discussed.

The paper by *Bruce Kessler*, following a macroelement approach, deals with construction of an orthogonal scaling vector of differentiable fractal interpolation functions, with symmetry properties, that generates a space including C^1 cubic splines.

In his paper, *Qun Mo* shows that normalized tight multiwavelet frames with highest possible order of vanishing moments can be derived from refinable functions under the multiwavelet settings.

The paper by *Bin Han and Thomas P.-Y. Yu*, entitled *Face-Based Hermite Subdivision Schemes*, studies the symmetry and structure of the vectors for the sum rule orders of the refinement mask for a type of vector Hermite subdivision scheme, called “face-based Hermite subdivision scheme” by the authors. Examples of the refinement masks with the quincunx dilation matrix and $2I_2$ are constructed.

In the paper by *Don Hong and Qingbo Xue*, piecewise linear prewavelets over a bounded domain with regular triangulation are constructed by investigating the orthogonal conditions directly.

The contribution by *Lutai Guan and Tan Yang* gives results on spline-wavelets on refined grids and applications. So-called addition spline-wavelets are constructed using natural spline interpolation skills. Applications in image compression are also included.

The following two papers are on wavelet applications in statistics/bio-statistics. In the paper by *Jianzhong Wang*, the author first gives a brief review of wavelet shrinkage and then applies the variational method for wavelet regression and shows the relation between the variational method and the wavelet shrinkage.

In the paper by *Don Hong and Yu Shyr*, many applications of wavelets in cancer study are reviewed, in particular, in medical imaging such as in detection of microcalcifications in digital mammography and functional magnetic resonance imaging, and molecular biology area such as in genome sequence analysis, protein structure investigation, and gene expression data analysis. Wavelet-based method for MALDI-TOF mass spectrometry data analysis is also reviewed.

We are grateful to all the authors who contributed to the special issue. We appreciate all the referees who reviewed the submitted papers. Their efforts and suggestions made it possible to publish scientific papers of high quality. We would like to express our special thanks to *George Anastassiou* for his patience and support and to *Benedict Bobga, Shuo Chen, Yuannan Diao, Brad Dyer, Qingtang Jiang, David Roach, Bashar Shakhtour, Wasin So, and Qiyu Sun* for their assistance in the editing of this special issue. Last but not least we also thank the support of the Department of Biostatistics and Ingram Cancer Center of Vanderbilt University, and the Department of Mathematics, Middle Tennessee State University.

Don Hong
Middle Tennessee State University
Murfreesboro, TN, USA

Yu Shyr
Vanderbilt University
Nashville, TN, USA

An Orthogonal Scaling Vector Generating a Space of C^1 Cubic Splines Using Macroelements

Bruce Kessler

Department of Mathematics
Western Kentucky University
Bowling Green, KY, USA
bruce.kessler@wku.edu

Abstract

The main result of this paper is the creation of an orthogonal scaling vector of four differentiable functions, two supported on $[-1, 1]$ and two supported on $[0, 1]$, that generates a space containing the classical spline space $\mathcal{S}_3^1(\mathbf{Z})$ of piecewise cubic polynomials on integer knots with one derivative at each knot. The author uses a macroelement approach to the construction, using differentiable fractal function elements defined on $[0, 1]$ to construct the scaling vector. An application of this new basis in an image compression example is provided.

AMS Subject Classification Numbers: 42C40, 65D15

Keywords: orthogonal bases, scaling vectors, spline spaces, macroelements, wavelets, multi-wavelets

1 Introduction

Orthogonal scaling vectors have numerous applications in signal and image processing, including image compression, denoising, and edge detection. Geronimo, Hardin, and Massopust were among the first to use more than one scaling function, in their construction of the GHM orthogonal scaling vector in [7]. The space generated by the GHM functions and their integer translates had approximation order 2, like the space generated by the single compactly-supported D4 orthogonal scaling function constructed by Daubechies in [3], but also contained the space of continuous piecewise linear polynomials on integer knots. Also, by using more than one function, they were able to build scaling functions and wavelets with symmetry properties. The only single compactly-supported scaling function with symmetry properties is the characteristic function on $[0, 1]$, known as the Haar basis. Other types of orthogonal scaling vectors have been constructed in [4], [8], [11], and [13] that have compactly-supported elements that lack continuous derivatives. Han and Jiang constructed a length-4 scaling vector in [9] that has approximation order 4 and two continuous derivatives.

Hardin and Kessler put the constructions of Geronimo and Hardin from [7] (with Massopust) and [4] (with Donovan) into a macroelement framework in [10]. The construction of scaling vectors from a single macroelement supported on $[0, 1]$ is important, since these types of scaling functions will be orthogonal interval-by-interval, and the intervals over which the basis is defined need not be of uniform length. This adaptability to an arbitrary partition of $\mathbf{R}$ promises to give greater flexibility to the use of these types of bases in applications. While building bases on variable length intervals is outside the scope of this paper, we should keep in mind that most of the results here can be easily adapted to intervals of arbitrary length.

The main result in this paper is the creation of a orthogonal scaling vector of length 4 of compactly-supported, differentiable piecewise fractal interpolation functions, with symmetry properties, that generates a space including the classical spline space $\mathcal{S}_3^1(\mathbf{Z})$ of differentiable cubic polynomials on integer knots. This scaling vector has a symmetric/antisymmetric pair supported on $[-1, 1]$ and a symmetric/antisymmetric pair supported on $[0, 1]$. The construction will follow a macroelement approach; that is, we construct a macroelement with the desirable properties using fractal functions and then construct a scaling vector from the macroelement. Fractal functions have been used previously in [4], [6], [7], and [11] to construct orthogonal scaling vectors. We show that a shorter-length scaling vector can not be found with the same properties, and then show this new basis in use in an image compression example. We believe this to be the first scaling vector to be constructed with all of the above-mentioned properties.

1.1 Scaling Vectors

A vector $\Phi = (\phi_1, \phi_2, \dots, \phi_r)^T$ of functions defined on $\mathbf{R}^k$ is said to be *refinable* if

$$\Phi = N^{\frac{k}{2}} \sum g_i \Phi(N \cdot -i) \quad (1)$$

for some integer dilation $N > 1$, $i \in \mathbf{Z}^k$, and for some sequence of $r \times r$ matrices g_i . (The normalization factor $N^{\frac{k}{2}}$ can be dropped, but is convenient for applications.) A *scaling vector* is a refinable vector Φ of square-integrable functions where the set of the components of Φ and their integer translates are linearly independent. An *orthogonal* scaling vector Φ is a scaling vector where the functions $\phi_1, \dots, \phi_r$ are compactly supported and satisfy

$$\langle \phi_i, \phi_j(\cdot - n) \rangle = \delta_{i,j} \delta_{0,n}, \quad i, j \in \{1, \dots, r\}, \quad n \in \mathbf{Z}^k,$$

where the inner product is the standard $L^2(\mathbf{R}^k)$ integral inner product

$$\langle f, g \rangle = \int_{\mathbf{R}^k} f(x)g(x)dx$$

and δ is Kronecker's delta (1 if indices are equal, 0 otherwise.) A scaling vector Φ is said to *generate* a closed linear space denoted by

$$S(\Phi) = \text{clos}_{L^2} \text{span} \{ \phi_i(\cdot - j) : i = 1, \dots, r, j \in \mathbf{Z} \}.$$

Two scaling vectors Φ and $\tilde{\Phi}$ are *equivalent* if $S(\Phi) = S(\tilde{\Phi})$. The scaling vector $\tilde{\Phi}$ is said to *extend* Φ , or be an *extension* of Φ , if $S(\Phi) \subset S(\tilde{\Phi})$.

Scaling vectors are important because they provide a framework for analyzing functions in $L^2(\mathbf{R}^k)$. A *multiresolution analysis* (MRA) of $L^2(\mathbf{R}^k)$ of multiplicity r is a set of closed linear spaces (V_p) such that

1. $\dots \supset V_{-2} \supset V_{-1} \supset V_0 \supset V_1 \supset V_2 \dots$,
2. $\overline{\bigcup_{p \in \mathbf{Z}} V_p} = L^2(\mathbf{R}^k)$,
3. $\bigcap_{p \in \mathbf{Z}} V_p = \{0\}$,

4. $f \in V_0$ iff $f(N^{-j}\cdot) \in V_j$, and
5. there exists a set of functions $\phi_1, \dots, \phi_r$ whose integer translates form a Riesz basis of V_0 .

From the above definitions, it is clear that scaling vectors can be used to generate MRA's, with $V_0 = S(\Phi)$. Jia and Shen proved in [15] that if the components of a scaling vector Φ are compactly-supported, then Φ will always generate an MRA. All the scaling vectors discussed in this paper will consist of compactly-supported functions, and therefore, will generate MRA's. A function vector $\Psi = (\psi_1, \dots, \psi_{r(N^k-1)})^T$, such that $\psi_i \in V_{-1}$ for $i = 1, \dots, r(N^k - 1)$ (see [14]) and such that $S(\Psi) = V_{-1} \ominus V_0$, is called a *multiwavelet*, and the individual ψ_i are called *wavelets*.

1.2 Macroelements on $[0, 1]$

We will use the notation $f^{(j)}(x)$ to denote the j^{th} derivative of $f(x)$, with the convention $f^{(0)}(x) = f(x)$. As a convenience, we will use the notation $f^{(j)}(0)$ and $f^{(j)}(1)$ to denote $\lim_{x \rightarrow 0^+} f^{(j)}(x)$ and $\lim_{x \rightarrow 1^-} f^{(j)}(x)$, respectively, although the notation may not always be mathematically rigorous.

A C^p *macroelement* defined on $[0, 1]$ is a vector of the form $(l_1, \dots, l_k, r_1, \dots, r_k, m_1, \dots, m_n)^T$ where the set of elements are linearly independent functions supported on $[0, 1]$ with p continuous derivatives such that

1. $l_i^{(j)}(0) = r_i^{(j)}(1) = 0$ for $i = 1, \dots, k$,
2. $m_i^{(j)}(0) = m_i^{(j)}(1) = 0$ for $i = 1, \dots, n$, and
3. $l_i^{(j)}(1) = r_i^{(j)}(0)$ for $i = 1, \dots, k$

for $j = 0, \dots, p$. A macroelement is *orthonormal* if $\langle l_i, r_j \rangle = 0$ for $i, j \in \{1, \dots, k\}$, $\langle l_i, m_j \rangle = \langle r_i, m_j \rangle = 0$ for $i \in \{1, \dots, k\}$ and $j \in \{1, \dots, n\}$, and each element is normalized.

A macroelement Λ is *refinable* if there are $(2k+n) \times (2k+n)$ matrices $p_0, \dots, p_{N-1}$ such that

$$\Lambda(x) = \sqrt{N} p_i \Lambda(Nx - i) \text{ for } x \in \left[\frac{i}{N}, \frac{i+1}{N} \right], \quad i = 0, \dots, N-1. \quad (2)$$

Because of the linear independence of the components of Λ , the matrix coefficients will be unique if they exist. Note that a refinable C^0 macroelement by necessity has $l_i(1) = r_i(0) \neq 0$ for some i , since if all elements were 0 at $x = 0, 1$, then each element would be 0 at N^j -adic points on $[0, 1]$ as $j \rightarrow \infty$. Hence, each element would be 0, a contradiction. Likewise, note that a refinable C^1 macroelement by necessity has $l'_i(1) = r'_i(0) \neq 0$ for some i , since if all elements had a zero derivative at $x = 0, 1$, then each element would be a constant function, also a contradiction.

Lemma 1. *A refinable C^p macroelement $\Lambda = (l_1, \dots, l_k, r_1, \dots, r_k, m_1, \dots, m_n)^T$ defined on $[0, 1]$ has an associated scaling vector Φ of length $k+n$ and support $[-1, 1]$. If the macroelement Λ is orthonormal, then the scaling vector Φ is equivalent to an orthogonal scaling vector.*

Proof. Let $l = (l_1, \dots, l_k)^T$, $r = (r_1, \dots, r_k)^T$, and $m = (m_1, \dots, m_n)^T$. Let a_i^l , a_i^r , and a_i^m be the $k \times k$, $k \times k$, and $k \times n$ matrices, respectively, such that

$$\begin{bmatrix} a_i^l & a_i^r & a_i^m \end{bmatrix} \Lambda(2x - i) = l(x) \text{ if } x \in \left[\frac{i}{N}, \frac{i+1}{N} \right], \quad i = 0, \dots, N-1.$$

Likewise, let b_i^l , b_i^r , and b_i^m be the $k \times k$, $k \times k$, and $k \times n$ matrices, respectively, such that

$$\begin{bmatrix} b_i^l & b_i^r & b_i^m \end{bmatrix} \Lambda(2x - i) = r(x) \text{ if } x \in \left[\frac{i}{N}, \frac{i+1}{N} \right], \quad i = 0, \dots, N-1$$

and let c_i^l , c_i^r , and c_i^m be the $n \times k$, $n \times k$, and $n \times n$ matrices, respectively, such that

$$\begin{bmatrix} c_i^l & c_i^r & c_i^m \end{bmatrix} \Lambda(2x - i) = m(x) \text{ if } x \in \left[\frac{i}{N}, \frac{i+1}{N} \right], \quad i = 0, \dots, N-1.$$

Then the matrix coefficients in (2) can be written in block-matrix form

$$p_i = \begin{bmatrix} a_i^l & a_i^r & a_i^m \\ b_i^l & b_i^r & b_i^m \\ c_i^l & c_i^r & c_i^m \end{bmatrix}.$$

Note that many of the block matrices are redundant: $a_{i-1}^l = a_i^r$, $b_{i-1}^l = b_i^r$, and $c_{i-1}^l = c_i^r$ for $i = 0, \dots, N-1$ since the macroelement components are continuous, and $a_{N-1}^l = b_0^r$ due to the endpoint conditions of the macroelement. Also, many of the block matrices are zero matrices: $a_0^r = b_{N-1}^l = 0_{k \times k}$ and $c_0^r = c_{N-1}^l = 0_{n \times k}$ due to the endpoint conditions of the macroelement.

Define

$$\phi_i(x) = \frac{1}{\sqrt{2}} \begin{cases} l_i(x+1) & \text{for } x \in [-1, 0] \\ r_i(x) & \text{for } x \in [0, 1] \end{cases}, \quad i = 1, \dots, k, \text{ and} \quad (3)$$

$$\phi_{k+i}(x) = m_i(x) \text{ for } x \in [0, 1], \quad i = 1, \dots, n. \quad (4)$$

Then the function vector $\Phi = (\phi_1, \dots, \phi_{k+n})^T$ satisfies (1), with

$$g_i = \begin{bmatrix} a_{N+i}^r & a_{N+i}^m \\ 0_{n \times k} & 0_{n \times n} \end{bmatrix} \quad i = -N, \dots, -1 \text{ and } g_i = \begin{bmatrix} b_i^r & b_i^m \\ c_i^r & c_i^m \end{bmatrix} \quad i = 0, \dots, N-1.$$

Hence, Φ is refinable, and supported completely in $[-1, 1]$.

If Λ is orthonormal, then by definition, Φ meets the criteria of an orthogonal scaling vector, except that possibly $\langle \phi_i, \phi_j \rangle \neq 0$ for $i, j \in \{1, \dots, k\}$, $i \neq j$, and for $i, j \in \{k+1, \dots, k+n\}$, $i \neq j$. However, we may replace $\{\phi_1, \dots, \phi_k\}$ with an orthonormal set $\{\tilde{\phi}_1, \dots, \tilde{\phi}_k\}$ and $\{\phi_{k+1}, \dots, \phi_{k+n}\}$ with an orthonormal set $\{\tilde{\phi}_{k+1}, \dots, \tilde{\phi}_{k+n}\}$, so that $\{\tilde{\phi}_1, \dots, \tilde{\phi}_k, \tilde{\phi}_{k+1}, \dots, \tilde{\phi}_{k+n}\}$ is an orthogonal scaling vector. $\square$

Let $\text{span } \Lambda$ refer to the span of the elements of Λ . Two macroelements Λ and $\tilde{\Lambda}$ are *equivalent* if $\text{span } \Lambda = \text{span } \tilde{\Lambda}$. The macroelement $\tilde{\Lambda}$ is said to *extend* Λ , or be an *extension* of Λ , if $\text{span } \Lambda \subset \text{span } \tilde{\Lambda}$. In this paper, we will extend macroelements for the purpose of extending scaling vectors, using the following lemma. We use the notation $\chi_{[a,b]}$ to be the characteristic function defined by

$$\chi_{[a,b]} = \begin{cases} 1 & \text{for } x \in [a, b], \\ 0 & \text{otherwise.} \end{cases}$$

Lemma 2. Let Λ be a refinable C^p macroelement defined on $[0, 1]$, and let Φ be the associated scaling vector as defined in (3) and (4), for $p = 0, 1$. If $\tilde{\Lambda}$ is a C^p macroelement extension of Λ , then the associated scaling vector $\tilde{\Phi}$ as defined in (3) and (4) is an extension of Φ .

Proof. Let $\Lambda = \{l_1, \dots, l_k, r_1, \dots, r_k, m_1, \dots, m_n\}$ and $\tilde{\Lambda} = \{\tilde{l}_1, \dots, \tilde{l}_{k'}, \tilde{r}_1, \dots, \tilde{r}_{k'}, \tilde{m}_1, \dots, \tilde{m}_{n'}\}$, where $\tilde{\Lambda}$ is an extension of Λ , so $k \leq k'$ and $n \leq n'$. Consider a basis element $\phi \in \{\phi_i(\cdot - j) : \phi_i \in \Phi, i \in \{1, \dots, k+n\}, j \in \mathbf{Z}\}$. If $\text{supp } \phi \subset [j, j+1]$ for some $j \in \mathbf{Z}$, then $\phi(x+j) \in \text{span}\{m_1, \dots, m_n\} \subset \text{span}\{\tilde{m}_1, \dots, \tilde{m}_{n'}\}$. From the definition of $\tilde{\Phi}$ in (4), then $\phi \in S(\tilde{\Phi})$.

Otherwise, $\text{supp } \phi \subset [j, j+2]$ for some $j \in \mathbf{Z}$. Let $l = \phi\chi_{[j,j+1]}$ and $r = \phi\chi_{[j+1,j+2]}$. Then $l(x+j) \in \text{span}\{l_1, \dots, l_k, m_1, \dots, m_n\} \subset \text{span}\{\tilde{l}_1, \dots, \tilde{l}_{k'}, \tilde{m}_1, \dots, \tilde{m}_{n'}\}$ and $r(x+j+1) \in \text{span}\{r_1, \dots, r_k, m_1, \dots, m_n\} \subset \text{span}\{\tilde{r}_1, \dots, \tilde{r}_{k'}, \tilde{m}_1, \dots, \tilde{m}_{n'}\}$. From the match-up conditions in the definition of the macroelement and the definition of $\tilde{\Phi}$ in (3), then $\phi \in S(\tilde{\Phi})$. Thus, $S(\Phi) \subset S(\tilde{\Phi})$ and $\tilde{\Phi}$ is an extension of Φ . $\square$

In the following example, we show a simple way to extend a macroelement, and hence, a scaling vector.

Example 1. Let $l_1 = x\chi_{[0,1]}$ and $r_1 = (1-x)\chi_{[0,1]}$. Then $\Lambda = (l_1, r_1)^T$ is a C^0 macroelement. It is also refinable, since

$$\Lambda(x) = \sqrt{2} \begin{cases} \begin{bmatrix} \frac{1}{2\sqrt{2}} & 0 \\ \frac{1}{2\sqrt{2}} & \frac{1}{\sqrt{2}} \end{bmatrix} \Lambda(2x) & \text{if } x \in [0, \frac{1}{2}] \\ \begin{bmatrix} \frac{1}{\sqrt{2}} & \frac{1}{2\sqrt{2}} \\ 0 & \frac{1}{2\sqrt{2}} \end{bmatrix} \Lambda(2x-1) & \text{if } x \in [\frac{1}{2}, 1] \end{cases}.$$

Therefore, we have the scaling vector $\Phi = (\phi_1)$ given in (3) and (4), with

$$\Phi(x) = \sqrt{2} \left[\frac{1}{2\sqrt{2}} \Phi(2x+1) + \frac{1}{\sqrt{2}} \Phi(2x) + \frac{1}{2\sqrt{2}} \Phi(2x-1) \right].$$

Then Φ is the linear B-spline, and generates $S(\Phi) = \mathcal{S}_1^0(\mathbf{Z})$, the spline space of continuous piecewise linear polynomials on integer knots.

Both Λ and Φ can be extended by the addition of the function $m_1(x) = \phi_2 = 4x(1-x)\chi_{[0,1]}$. Then $\tilde{\Lambda} = (l_1, r_1, m_1)^T$ is refinable, since

$$\tilde{\Lambda}(x) = \sqrt{2} \begin{cases} \begin{bmatrix} \frac{1}{2\sqrt{2}} & 0 & 0 \\ \frac{1}{2\sqrt{2}} & \frac{1}{\sqrt{2}} & 0 \\ \frac{1}{\sqrt{2}} & 0 & \frac{1}{4\sqrt{2}} \end{bmatrix} \tilde{\Lambda}(2x) & \text{if } x \in [0, \frac{1}{2}] \\ \begin{bmatrix} \frac{1}{\sqrt{2}} & \frac{1}{2\sqrt{2}} & 0 \\ 0 & \frac{1}{2\sqrt{2}} & 0 \\ 0 & \frac{1}{\sqrt{2}} & \frac{1}{4\sqrt{2}} \end{bmatrix} \tilde{\Lambda}(2x-1) & \text{if } x \in [\frac{1}{2}, 1] \end{cases},$$

and, by Lemma 1, we have the scaling vector $\tilde{\Phi} = (\phi_1, \phi_2)^T$, with

$$\tilde{\Phi}(x) = \sqrt{2} \left(\begin{bmatrix} \frac{1}{2\sqrt{2}} & 0 \\ 0 & 0 \end{bmatrix} \tilde{\Phi}(2x+1) + \begin{bmatrix} \frac{1}{\sqrt{2}} & 0 \\ 0 & \frac{1}{4\sqrt{2}} \end{bmatrix} \tilde{\Phi}(2x) + \begin{bmatrix} \frac{1}{2\sqrt{2}} & 0 \\ \frac{1}{\sqrt{2}} & \frac{1}{4\sqrt{2}} \end{bmatrix} \tilde{\Phi}(2x-1) \right).$$

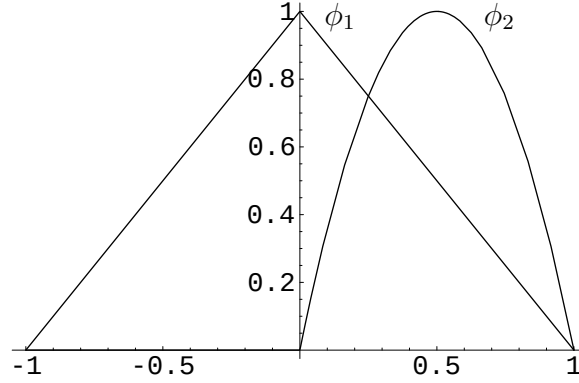


Figure 1: The scaling functions ϕ_1 and ϕ_2 from Example 1.

By Lemma 2, $\tilde{\Phi}$ is an extension of Φ . In fact, $S(\tilde{\Phi}) = \mathcal{S}_2^0(\mathbf{Z}) \supset \mathcal{S}_1^0(\mathbf{Z})$. Both functions are illustrated in Figure 1.

1.3 Fractal Interpolation Functions

Let $C_0([0, 1])$ denote the space of continuous functions defined over $[0, 1]$ that are 0 at $x = 0, 1$. Let $C_1([0, 1])$ denote the subspace of differentiable functions in $C_0([0, 1])$ whose derivatives are 0 at $x = 0, 1$. Let Λ be a refinable macroelement of length n , and let Π be a function vector of length k defined by

$$\Pi(x) = p_i \Lambda(Nx - i) \text{ for } x \in \left[\frac{i}{N}, \frac{i+1}{N} \right], \quad i = 0, \dots, N-1,$$

for some $k \times n$ matrices p_i such that $\Pi(x) \in C_0([0, 1])^k$. Then a vector Γ of the form

$$\Gamma(x) = \Pi(x) + \sum_{i=0}^{N-1} s_i \Gamma(Nx - i) \in C_0([0, 1])^k, \quad (5)$$

where each s_i is a $k \times k$ matrix and $\max_i \|s_i\|_\infty < 1$, is a vector of *fractal interpolation functions* (FIF's). (See [1] and [2] for a more detailed introduction to FIF's.) By definition, the vector $\tilde{\Lambda} = (\Lambda^T, \Gamma^T)^T$ is a refinable C^0 macroelement that extends Λ .

Lemma 3. *Let Γ be a FIF satisfying (5). If $\Pi(x) \in C_1([0, 1])^k$ and $\max_{i=1, \dots, N-1} \|s_i\|_\infty < \frac{1}{N}$, then $\Gamma \in C_1([0, 1])^k$.*

Proof. Since Γ satisfies (5), then

$$\Gamma'(x) = \Pi'(x) + \sum_{i=0}^{N-1} N s_i \Gamma'(Nx - i).$$

Since $\Pi'(x) \in C_0([0, 1])^k$ and $\max_{i=1, \dots, N-1} \|N s_i\|_\infty = N \max_{i=1, \dots, N-1} \|s_i\|_\infty < N \frac{1}{N} = 1$, then $\Gamma'(x)$ is by definition a FIF, and $\Gamma'(x) \in C_0([0, 1])^k$. Therefore, $\Gamma(x) \in C_1([0, 1])^k$. $\square$

Consider a C^0 or C^1 macroelement $\Lambda = (l_1, \dots, l_k, r_1, \dots, r_k, m_1, \dots, m_n)^T$ defined on $[0, 1]$ that is not orthogonal. We can not simply apply the Gram-Schmidt process to the components of Λ to obtain an orthonormal macroelement, since the resulting functions will not satisfy the endpoint criteria. In fact, we can not apply the process to any subset of elements that includes a l_i and r_j and still have the same type of macroelement. (We could go from a C^1 macroelement to a C^0 macroelement, but the associated scaling vector would not be an extension of the original.) However, we can apply the Gram-Schmidt process to the set of functions $M = \{m_1, \dots, m_n\}$ to get $\tilde{M} = \{\tilde{m}_1, \dots, \tilde{m}_n\}$, and then subtract P_M , the orthogonal projection onto the space spanned by M , from each of the other elements, giving the equivalent macroelement

$$\tilde{\Lambda} = ((I - P_M)l_1, \dots, (I - P_M)l_k, (I - P_M)r_1, \dots, (I - P_M)r_k, \tilde{m}_1, \dots, \tilde{m}_n)^T.$$

If

$$\langle (I - P_M)l_i, (I - P_M)r_j \rangle = 0 \text{ for } i, j = 1, \dots, k, \quad (6)$$

(and each element is normalized), then $\tilde{\Lambda}$ is an orthonormal macroelement. This is the fractal function approach for extending a macroelement: add FIF's to the set M , hence the macroelement, so that (6) is satisfied. (See [5] for a broader discussion on constructing intertwined MRA's.)

Example 2. The scaling vector shown in this example was originally constructed by Geronimo, Hardin, and Massopust in [7], although not in the macroelement context, and is reconstructed by Hardin and Kessler in detail using macroelements in [10]. It is widely known as the GHM scaling vector.

Consider from Example 1 the C^0 macroelement Λ and the scaling vector $\Phi = (\phi_1)$ that generates $\mathcal{S}_1^0(\mathbf{Z})$. In order to extend Λ to an orthonormal C^0 macroelement, we construct a FIF satisfying

$$u(x) = \phi_1(2x - 1) + s_0 u(2x) + s_1 u(2x - 1), \quad \max_{i=0,1} |s_i| < 1,$$

such that $\langle (I - P_u)l_1, (I - P_u)r_1 \rangle = 0$. It was shown in [7] and [10] that the orthogonality condition is satisfied by $s_0 = s_1 = -\frac{1}{5}$. By letting

$$\check{l}_1 = \frac{(I - P_u)l_1}{\|(I - P_u)l_1\|}, \quad \check{r}_1 = \frac{(I - P_u)r_1}{\|(I - P_u)r_1\|}, \quad \text{and} \quad \check{m}_1 = \frac{u}{\|u\|},$$

we have the orthonormal C^0 macroelement $\check{\Lambda} = (\check{l}_1, \check{r}_1, \check{m}_1)^T$, equivalent to $(l_1, r_1, u)^T$ and an extension of Λ . The associated scaling vector $\check{\Phi} = (\phi_1, \phi_2)^T$ defined in (3) and (4) is the orthogonal GHM scaling vector, and is illustrated in Figure 2.

Example 3. The scaling vector shown in this example was originally constructed by Donovan, Geronimo, and Hardin in [4], although not in the macroelement context, and again by Hardin and Kessler in detail in [10] using a macroelement approach.

Consider from Example 1 the C^0 macroelement $\tilde{\Lambda}$ and the scaling vector $\tilde{\Phi} = (\phi_1, \phi_2)^T$ that generates $\mathcal{S}_2^0(\mathbf{Z})$. In order to extend $\tilde{\Lambda}$ to an orthonormal C^0 macroelement, we construct a FIF satisfying

$$u(x) = \phi_2(2x) - \phi_2(2x - 1) + su(2x) + su(2x - 1) \text{ for } |s| < 1,$$

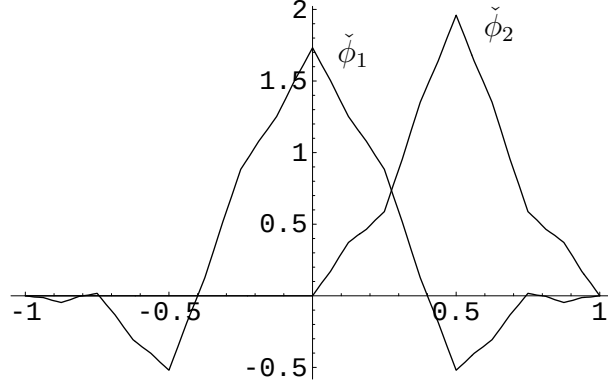


Figure 2: The orthogonal GHM scaling vector $\check{\Phi}$ from Example 2.

such that $\langle (I - P_M)l_1, (I - P_M)r_1 \rangle = 0$, where $M = \{\phi_2, u\}$. (Note that $\langle \phi_2, u \rangle = 0$, since ϕ_2 and u are symmetric and antisymmetric, respectively, about $x = \frac{1}{2}$.) It was shown in [4] and [10] that the orthogonality condition is satisfied by $s = \frac{2-\sqrt{10}}{6} \approx -0.1937$. By letting

$$\hat{l}_1 = \frac{(I - P_M)l_1}{\|(I - P_M)l_1\|}, \quad \hat{r}_1 = \frac{(I - P_M)r_1}{\|(I - P_M)r_1\|}, \quad \hat{m}_1 = \frac{m_1}{\|m_1\|}, \quad \text{and} \quad \hat{m}_2 = \frac{u}{\|u\|},$$

we have the orthonormal C^0 macroelement $\hat{\Lambda} = (\hat{l}_1, \hat{r}_1, \hat{m}_1, \hat{m}_2)^T$, equivalent to $(l_1, r_1, m_1, m_2)^T$ and an extension of $\tilde{\Lambda}$. The associated scaling vector $\hat{\Phi} = (\hat{\phi}_1, \hat{\phi}_2, \hat{\phi}_3)^T$ defined in (3) and (4) is an orthogonal scaling vector, and is illustrated in Figure 3.

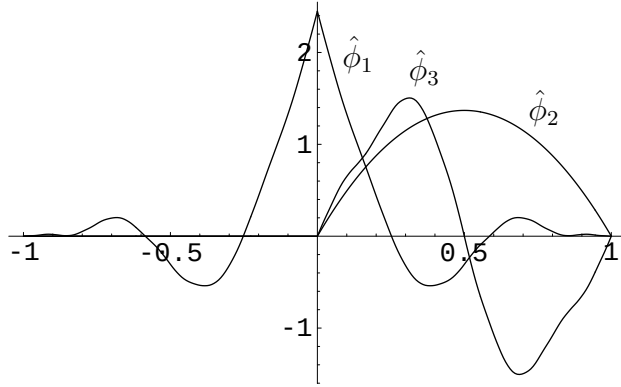


Figure 3: The orthogonal scaling vector $\hat{\Phi}$ from Example 3.

2 Main Results

In this section, we will construct a refinable C^1 macroelement on $[0, 1]$ which, in turn, provides a scaling vector Φ that generates $\mathcal{S}_3^1(\mathbf{Z})$. We will show that two additional functions are needed to extend that macroelement to an orthonormal macroelement that is refinable with dilation 2, and an explicit construction of that macroelement and the associated orthogonal scaling vector will be provided. Lastly, although general methods are available for finding a

multiwavelet associated with an orthogonal scaling vector (see [10], [16], and [17]), we will show an equivalent macroelement approach to constructing the wavelets.

2.1 Scaling Vector Generating $\mathcal{S}_3^1(\mathbf{Z})$

Consider the following basis for cubic polynomials defined on $[0, 1]$ and 0 elsewhere:

$$l_1(x) = (-2x^3 + 3x^2)\chi_{[0,1]}, \quad l_2(x) = (x^3 - x^2)\chi_{[0,1]},$$

$$r_1(x) = (2x^3 - 3x^2 + 1)\chi_{[0,1]}, \quad \text{and} \quad r_2(x) = (x^3 - 2x^2 + x)\chi_{[0,1]}.$$

One may verify that $\Lambda = (l_1, l_2, r_1, r_2)^T$ is a C^1 macroelement. It is also refinable, since

$$\Lambda = \sqrt{2} \begin{cases} \begin{bmatrix} \frac{1}{2\sqrt{2}} & \frac{3}{4\sqrt{2}} & 0 & 0 \\ -\frac{1}{8\sqrt{2}} & -\frac{1}{8\sqrt{2}} & 0 & 0 \\ \frac{1}{2\sqrt{2}} & -\frac{3}{4\sqrt{2}} & \frac{1}{\sqrt{2}} & 0 \\ \frac{1}{8\sqrt{2}} & -\frac{1}{8\sqrt{2}} & 0 & \frac{1}{2\sqrt{2}} \end{bmatrix} \Lambda(2x) \text{ if } x \in [0, \frac{1}{2}], \\ \begin{bmatrix} \frac{1}{\sqrt{2}} & 0 & \frac{1}{2\sqrt{2}} & \frac{3}{4\sqrt{2}} \\ 0 & \frac{1}{2\sqrt{2}} & -\frac{1}{8\sqrt{2}} & -\frac{1}{4\sqrt{2}} \\ 0 & 0 & \frac{1}{2\sqrt{2}} & -\frac{3}{4\sqrt{2}} \\ 0 & 0 & \frac{1}{8\sqrt{2}} & -\frac{1}{8\sqrt{2}} \end{bmatrix} \Lambda(2x-1) \text{ if } x \in [\frac{1}{2}, 1]. \end{cases} \quad (7)$$

From direct computation, we know that

$$\langle l_1, r_1 \rangle = \frac{9}{70}, \quad \langle l_1, r_2 \rangle = \frac{13}{420}, \quad \langle l_2, r_1 \rangle = -\frac{13}{420}, \quad \text{and} \quad \langle l_2, r_2 \rangle = -\frac{1}{140},$$

so Λ is not an orthonormal macroelement. As in (3), we define

$$\phi_i(x) = \frac{1}{\sqrt{2}} \begin{cases} l_i(x+1) & \text{if } x \leq 0 \\ r_i(x) & \text{if } x \geq 0 \end{cases}, \quad i = 1, 2,$$

so that we have the scaling vector $\Phi = (\phi_1, \phi_2)^T$ satisfying

$$\Phi(x) = \sqrt{2} \left(\begin{bmatrix} \frac{1}{2\sqrt{2}} & \frac{3}{4\sqrt{2}} \\ -\frac{1}{8\sqrt{2}} & -\frac{1}{8\sqrt{2}} \end{bmatrix} \Phi(2x+1) + \begin{bmatrix} \frac{1}{\sqrt{2}} & 0 \\ 0 & \frac{1}{2\sqrt{2}} \end{bmatrix} \Phi(2x) + \begin{bmatrix} \frac{1}{2\sqrt{2}} & -\frac{3}{4\sqrt{2}} \\ \frac{1}{8\sqrt{2}} & -\frac{1}{8\sqrt{2}} \end{bmatrix} \Phi(2x-1) \right),$$

where $S(\Phi) = \mathcal{S}_3^1(\mathbf{Z})$. Both functions are illustrated in Figure 4.

2.2 An Orthogonal, Refinable Extension

We may extend Φ to a refinable vector of length 3 by adding a single m_1 or by adding l_3 and r_3 functions to the macroelement Λ . However, as we show in the following theorem, neither method will produce an orthogonal extension. Finding a length-4 orthogonal extension of Φ is nontrivial. The proof of the following theorem gives a construction of one such scaling vector using the fractal function idea on the C^1 macroelement Λ defined above.

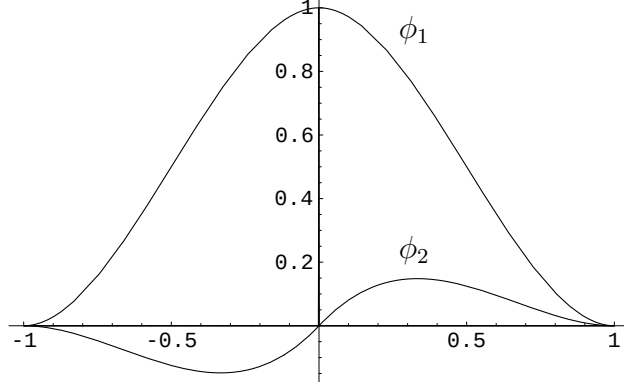


Figure 4: The scaling vector Φ generating $\mathcal{S}_3^1(\mathbf{Z})$.

Theorem 1. *The orthogonal scaling vector of least length that extends the length-2 scaling vector which generates $\mathcal{S}_3^1(\mathbf{Z})$ has length 4.*

Proof. There are two parts to the proof: we first show that a single m_1 or l_3 - r_3 pair added to the macroelement, thereby adding one function to the scaling vector, can not be found such that all of the necessary orthogonality conditions are satisfied, and then we show that a 2-vector of fractal functions (with symmetry properties, no less) can be found so that the necessary conditions are satisfied.

Suppose that the single normalized function m_1 added to the macroelement Λ gives a C^1 extension of the macroelement. Since we are assuming no symmetry properties for m_1 , we have four orthogonality conditions from (6), which reduce to the four equations

$$\langle l_i, r_j \rangle = \langle l_i, m_1 \rangle \langle r_j, m_1 \rangle, \quad i, j \in \{1, 2\}.$$

The system of four equations with the unknowns $\langle l_1, m_1 \rangle$, $\langle l_2, m_1 \rangle$, $\langle r_1, m_1 \rangle$, and $\langle r_2, m_1 \rangle$ is inconsistent, so no single function m_1 can satisfy all of the necessary conditions.

Now suppose that two orthonormal functions l_3 and r_3 added to the macroelement Λ give a C^1 extension of the macroelement. To be an orthogonal extension, l_3 and r_3 will need to satisfy the four orthogonality conditions

$$\langle l_i - \langle l_i, l_3 \rangle l_3, r_j - \langle r_j, r_3 \rangle r_3 \rangle = 0 \quad i, j \in \{1, 2\},$$

while satisfying the endpoint conditions

$$l_i(1) - \langle l_i, l_3 \rangle l_3(1) = r_i(0) - \langle r_i, r_3 \rangle r_3(0) \quad i = 1, 2.$$

The endpoint conditions force $\langle l_1, l_3 \rangle = \langle r_1, r_3 \rangle$ and $\langle l_2, l_3 \rangle = \langle r_2, r_3 \rangle$, the values of which we will denote α and β , respectively. Then the four orthogonality conditions become

$$\begin{cases} \alpha(\langle l_1, r_3 \rangle + \langle r_1, l_3 \rangle) &= \frac{9}{70} \\ \alpha\langle r_2, l_3 \rangle + \beta\langle l_1, r_3 \rangle &= \frac{13}{420} \\ \alpha\langle l_2, r_3 \rangle + \beta\langle r_1, l_3 \rangle &= -\frac{13}{420} \\ \beta(\langle l_2, r_3 \rangle + \langle r_2, l_3 \rangle) &= -\frac{1}{140} \end{cases}$$

with unknowns $\langle l_1, r_3 \rangle$, $\langle l_2, r_3 \rangle$, $\langle r_1, l_3 \rangle$, and $\langle r_2, l_3 \rangle$. Again, the system is inconsistent, so no pair of functions l_3 and r_3 can satisfy all of the necessary conditions.

To show that a length-4 orthogonal scaling vector is possible, we will actually construct one by adding two functions m_1 and m_2 to the macroelement Λ . Consider two FIF m_1 and m_2 and function vector $\Gamma = (m_1, m_2)^T$ satisfying

$$\Gamma(x) = \begin{bmatrix} \alpha & 0 \\ 0 & \beta \end{bmatrix} \Phi(2x-1) + \begin{bmatrix} 0 & q \\ s & t \end{bmatrix} \Gamma(2x) + \begin{bmatrix} 0 & -q \\ -s & t \end{bmatrix} \Gamma(2x-1)$$

where the maximum ∞ -norm of the matrix coefficients of $\Gamma(2x)$ and $\Gamma(2x-1)$ is 1 and $\alpha, \beta \neq 0$ are chosen so that $\|m_1\| = \|m_2\| = 1$. The extended macroelement $\hat{\Lambda} = (l_1, l_2, r_1, r_2, m_1, m_2)^T$ is refinable, with

$$\hat{\Lambda}(x) = \sqrt{2} \begin{cases} \begin{bmatrix} \frac{1}{2\sqrt{2}} & \frac{3}{4\sqrt{2}} & 0 & 0 & 0 & 0 \\ -\frac{1}{8\sqrt{2}} & -\frac{3}{8\sqrt{2}} & 0 & 0 & 0 & 0 \\ \frac{1}{2\sqrt{2}} & \frac{3}{4\sqrt{2}} & \frac{1}{\sqrt{2}} & 0 & 0 & 0 \\ \frac{1}{8\sqrt{2}} & -\frac{3}{8\sqrt{2}} & 0 & \frac{1}{2\sqrt{2}} & 0 & 0 \\ \frac{\alpha}{\sqrt{2}} & 0 & 0 & 0 & 0 & \frac{q}{\sqrt{2}} \\ 0 & \frac{\beta}{\sqrt{2}} & 0 & 0 & \frac{s}{\sqrt{2}} & \frac{t}{\sqrt{2}} \end{bmatrix} \hat{\Lambda}(2x) \text{ if } x \in [0, \frac{1}{2}], \\ \begin{bmatrix} \frac{1}{\sqrt{2}} & 0 & \frac{1}{2\sqrt{2}} & \frac{3}{4\sqrt{2}} & 0 & 0 \\ 0 & \frac{1}{2\sqrt{2}} & -\frac{1}{8\sqrt{2}} & -\frac{3}{8\sqrt{2}} & 0 & 0 \\ 0 & 0 & \frac{1}{2\sqrt{2}} & \frac{3}{4\sqrt{2}} & 0 & 0 \\ 0 & 0 & \frac{1}{8\sqrt{2}} & -\frac{3}{8\sqrt{2}} & 0 & 0 \\ 0 & 0 & \frac{\alpha}{\sqrt{2}} & 0 & 0 & -\frac{q}{\sqrt{2}} \\ 0 & 0 & 0 & \frac{\beta}{\sqrt{2}} & -\frac{s}{\sqrt{2}} & \frac{t}{\sqrt{2}} \end{bmatrix} \hat{\Lambda}(2x-1) \text{ if } x \in [\frac{1}{2}, 1]. \end{cases} \quad (8)$$

Let $M = \{m_1, m_2\}$. In order to construct an orthonormal macroelement equivalent to $\hat{\Lambda}$, the elements of M must satisfy the following conditions:

$$\langle (I - P_M)l_i, (I - P_M)r_j \rangle = 0, \quad i, j \in \{1, 2\}.$$

One can verify that m_1 and m_2 are symmetric and antisymmetric, respectively, about $x = \frac{1}{2}$, so that $\langle m_1, m_2 \rangle = 0$. Also, due to symmetry,

$$\langle r_1, m_1 \rangle = \langle l_1, m_1 \rangle, \quad \langle r_1, m_2 \rangle = -\langle l_1, m_2 \rangle, \quad \langle r_2, m_1 \rangle = -\langle l_2, m_1 \rangle, \quad \text{and} \quad \langle r_2, m_2 \rangle = \langle l_2, m_2 \rangle.$$

Thus, the four orthogonality conditions reduce to the three equations

$$\begin{cases} \frac{9}{70} - \langle r_1, m_1 \rangle^2 + \langle r_1, m_2 \rangle^2 = 0 \\ \frac{1}{140} - \langle r_2, m_1 \rangle^2 + \langle r_2, m_2 \rangle^2 = 0 \\ \frac{13}{420} - \langle r_1, m_1 \rangle \langle r_2, m_1 \rangle + \langle r_1, m_2 \rangle \langle r_2, m_2 \rangle = 0. \end{cases} \quad (9)$$

Using direct computation, we know that

$$\langle r_1, r_1 \rangle = \langle l_1, l_1 \rangle = \frac{13}{35}, \quad \langle r_2, r_2 \rangle = \langle l_2, l_2 \rangle = \frac{1}{105}, \quad \text{and} \quad \langle r_1, r_2 \rangle = -\langle l_1, l_2 \rangle = \frac{11}{210}.$$

Using the coefficients in (8), we may expand and solve for $\langle r_1, m_1 \rangle$, $\langle r_1, m_2 \rangle$, $\langle r_2, m_1 \rangle$, and $\langle r_2, m_2 \rangle$ in the following system:

$$\begin{cases} 2\langle r_1, m_1 \rangle = \frac{\alpha}{2} \\ 2\langle r_2, m_1 \rangle = -\frac{1}{4}q(\langle r_1, m_2 \rangle - 2\langle r_2, m_2 \rangle) + \frac{13\alpha}{120} \\ 2\langle r_1, m_2 \rangle = s\langle r_1, m_1 \rangle + \frac{3s}{2}\langle r_2, m_1 \rangle + t\langle r_1, m_2 \rangle - \frac{3t}{2}\langle r_2, m_2 \rangle - \frac{19\beta}{420} \\ 2\langle r_2, m_2 \rangle = \frac{3s}{4}\langle r_2, m_1 \rangle + \frac{t}{4}\langle r_2, m_2 \rangle - \frac{\beta}{168}. \end{cases} \quad (10)$$

The systems (9) and (10) have the solutions

$$\begin{aligned} \langle r_1, m_1 \rangle &= \frac{\alpha}{4}, \quad \langle r_1, m_2 \rangle = -\frac{1}{4}\sqrt{\alpha^2 - \frac{72}{35}}, \quad \langle r_2, m_1 \rangle = \frac{1}{216} \left(13\alpha - \sqrt{7\alpha^2 - \frac{72}{5}} \right), \\ \langle r_2, m_2 \rangle &= \frac{1}{216} \left(\sqrt{7}\alpha - 13\sqrt{\alpha^2 - \frac{72}{35}} \right), \quad q = \frac{26\alpha - 20\sqrt{7\alpha^2 - \frac{72}{5}}}{5\sqrt{7}\alpha + 70\sqrt{\alpha^2 - \frac{72}{35}}}, \\ s &= \frac{13104 - 6125\alpha^2 + 35\alpha\sqrt{7\alpha^2 - \frac{72}{5}} + 51\sqrt{7}\alpha\beta + 147\beta\sqrt{\alpha^2 - \frac{72}{35}}}{\sqrt{7}(504 + 140\alpha^2 + 29\sqrt{5}\alpha\sqrt{35\alpha^2 - 72})}, \text{ and} \\ t &= \frac{35280 - 7\sqrt{5}\sqrt{35\alpha^2 - 72}(1465\alpha - 9\sqrt{7}\beta) + 5\alpha(8575\alpha - 9\sqrt{7}\beta)}{35(504 + 140\alpha^2 + 29\sqrt{5}\alpha\sqrt{35\alpha^2 - 72})}. \end{aligned}$$

Again using the coefficients in (8), we may expand $\langle m_1, m_1 \rangle$ and $\langle m_2, m_2 \rangle$ and numerically solve for α and β . After substituting the above results into the initial expansions

$$q^2 - 2q\alpha\langle r_1, m_2 \rangle + \frac{13}{35}\alpha^2 = 1 \text{ and } s^2 + t^2 - 2s\beta\langle r_2, m_1 \rangle + 2t\beta\langle r_2, m_2 \rangle + \frac{\beta^2}{105} = 1,$$

we find the approximate solutions $\alpha \approx 1.63240645$ and $\beta \approx 14.19575451$, so that

$$\langle r_1, m_1 \rangle \approx 0.40810161, \quad \langle r_1, m_2 \rangle \approx -0.19487303, \quad \langle r_2, m_1 \rangle \approx 0.08869880,$$

$$\langle r_2, m_2 \rangle \approx -0.02691878, \quad q \approx 0.01570030, \quad s \approx 0.45783086, \text{ and } t \approx -0.03034240.$$

Since

$$\left\| \begin{bmatrix} 0 & q \\ s & t \end{bmatrix} \right\|_\infty = \left\| \begin{bmatrix} 0 & -q \\ -s & t \end{bmatrix} \right\|_\infty \approx 0.48817326 < \frac{1}{2},$$

then by Lemma 3, $m_1, m_2 \in C^1([0, 1])$. By setting

$$\tilde{l}_1 = \frac{(I - P_M)l_1}{\|(I - P_M)l_1\|}, \quad \tilde{l}_2 = \frac{(I - P_M)l_2}{\|(I - P_M)l_2\|}, \quad \tilde{r}_1 = \frac{(I - P_M)r_1}{\|(I - P_M)r_1\|}, \text{ and } \tilde{r}_2 = \frac{(I - P_M)r_2}{\|(I - P_M)r_2\|},$$

we have the orthonormal, refinable C^1 macroelement $\tilde{\Lambda} = (\tilde{l}_1, \tilde{l}_2, \tilde{r}_1, \tilde{r}_2, m_1, m_2)^T$ that is equivalent to $\hat{\Lambda}$ and an extension of Λ . By Lemma 2, we have the scaling vector $\tilde{\Phi}$ of length 4 as defined in (3) and (4) that is an extension of Φ . Since $\langle \tilde{\phi}_1, \tilde{\phi}_2 \rangle = 0$ due to their symmetry-antisymmetry about $x = 0$, $\tilde{\Phi}$ is an orthogonal scaling vector. $\square$

The elements of the orthogonal scaling vector constructed in the above proof are illustrated in Figure 5.

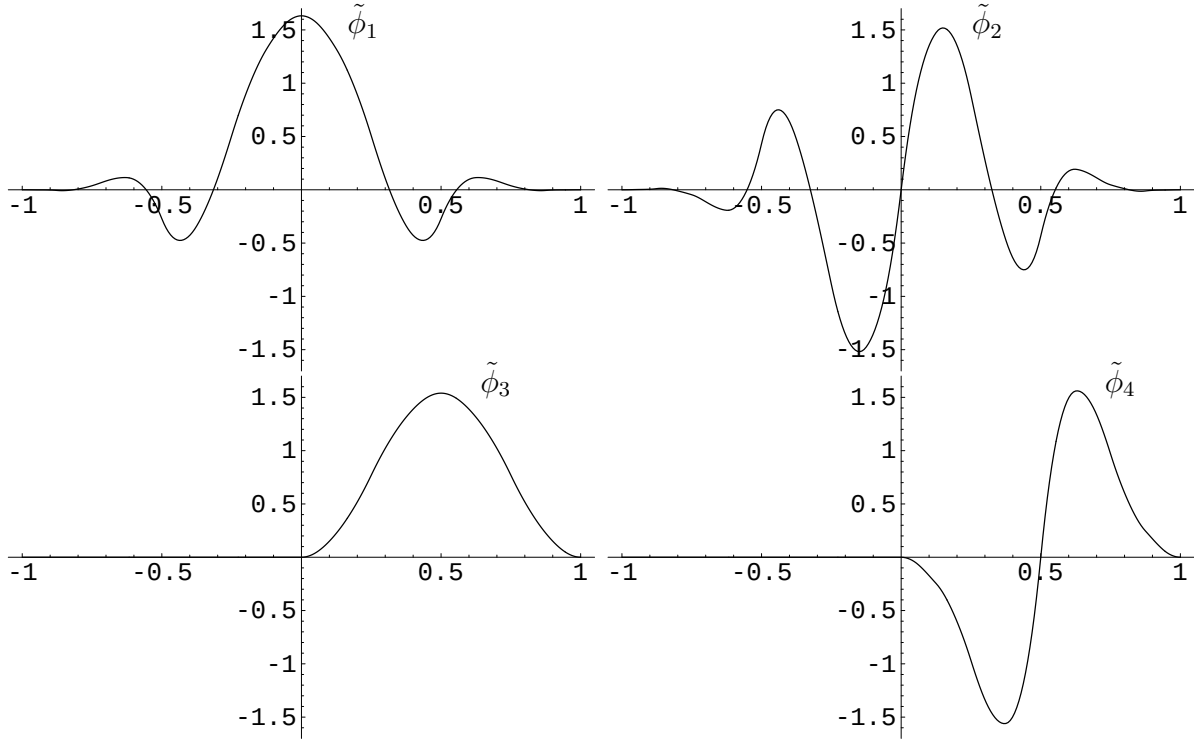


Figure 5: The orthogonal scaling vector $\tilde{\Phi} = (\tilde{\phi}_1, \tilde{\phi}_2, \tilde{\phi}_3, \tilde{\phi}_4)^T$, where $S(\tilde{\Phi}) \supset \mathcal{S}_3^1(\mathbf{Z})$.

2.3 The Associated Multiwavelets

A general construction for a multiwavelet associated with a scaling vector defined on $\mathbf{R}$ was found in [16] (see also [10]), and that technique could be used here. However, I will use an equivalent approach that will produce wavelets with symmetry properties. From Jia in [14], we know that, since $\tilde{\Phi}$ was of length 4, then $\tilde{\Psi}$ will also have length $4(2 - 1) = 4$.

Define the 6-dimensional space $V = \{f : f \in V_{-1}, \text{ supp } f \subseteq [0, 1]\}$ and let P_V denote the orthogonal projection onto V . Note that, in this construction, $V \cap V_0^\perp = \emptyset$, since there are 6 independent orthogonality conditions affecting elements of V . Recall the original macroelement $\tilde{\Lambda} = (\tilde{l}_1, \tilde{l}_2, \tilde{r}_1, \tilde{r}_2, m_1, m_2)^T$. Let

$$\bar{l}_i(x) = P_V \tilde{l}_i(x) + a_i \tilde{\phi}_1(Nx - N) + b_i \tilde{\phi}_2(Nx - N) \text{ for } i = 1, 2,$$

where a_i and b_i are chosen so that $\langle \bar{l}_i(x), \phi_1(x - 1) \rangle = 0$ and $\langle \bar{l}_i(x), \phi_2(x - 1) \rangle = 0$. Likewise, let

$$\bar{r}_i(x) = P_V \tilde{r}_i(x) + c_i \tilde{\phi}_1(Nx) + d_i \tilde{\phi}_2(Nx) \text{ for } i = 1, 2,$$

where c_i and d_i are chosen so that $\langle \bar{r}_i(x), \phi_j(x) \rangle = 0$. Note that the properties $\langle \bar{l}_i, \bar{r}_j \rangle = 0$, $\langle \bar{l}_i, m_j \rangle = 0$, and $\langle \bar{r}_j, m_j \rangle = 0$ for $i, j \in \{1, 2\}$ are maintained from the original macroelement, and that $a_1 = c_1$, $b_1 = -d_1$, $a_2 = -c_2$, and $b_2 = d_2$ due to symmetry properties.

We may construct functions

$$f_1(x) = \bar{l}_1(x + 1) + \bar{r}_1(x) \quad \text{and} \quad g_1(x) = -\bar{l}_2(x + 1) + \bar{r}_2(x)$$

that are symmetric with respect to $x = 0$, and

$$f_3(x) = -\bar{l}_1(x + 1) + \bar{r}_1(x) \quad \text{and} \quad g_3(x) = \bar{l}_2(x + 1) + \bar{r}_2(x)$$

that are antisymmetric with respect to $x = 0$, such that all are orthogonal to V_0 . Set

$$f_2 = g_1 - \frac{\langle g_1, f_1 \rangle}{\langle f_1, f_1 \rangle} f_1 \quad \text{and} \quad f_4 = g_3 - \frac{\langle g_3, f_3 \rangle}{\langle f_3, f_3 \rangle} f_3$$

to handle the last remaining orthogonalities, and set $\tilde{\psi}_i = \frac{f_i}{\|f_i\|}$ for $i = 1, \dots, 4$. Then $\tilde{\Psi} = (\tilde{\psi}_1, \tilde{\psi}_2, \tilde{\psi}_3, \tilde{\psi}_4)^T$ is a multiwavelet that generates W_0 , and is illustrated in Figure 6.

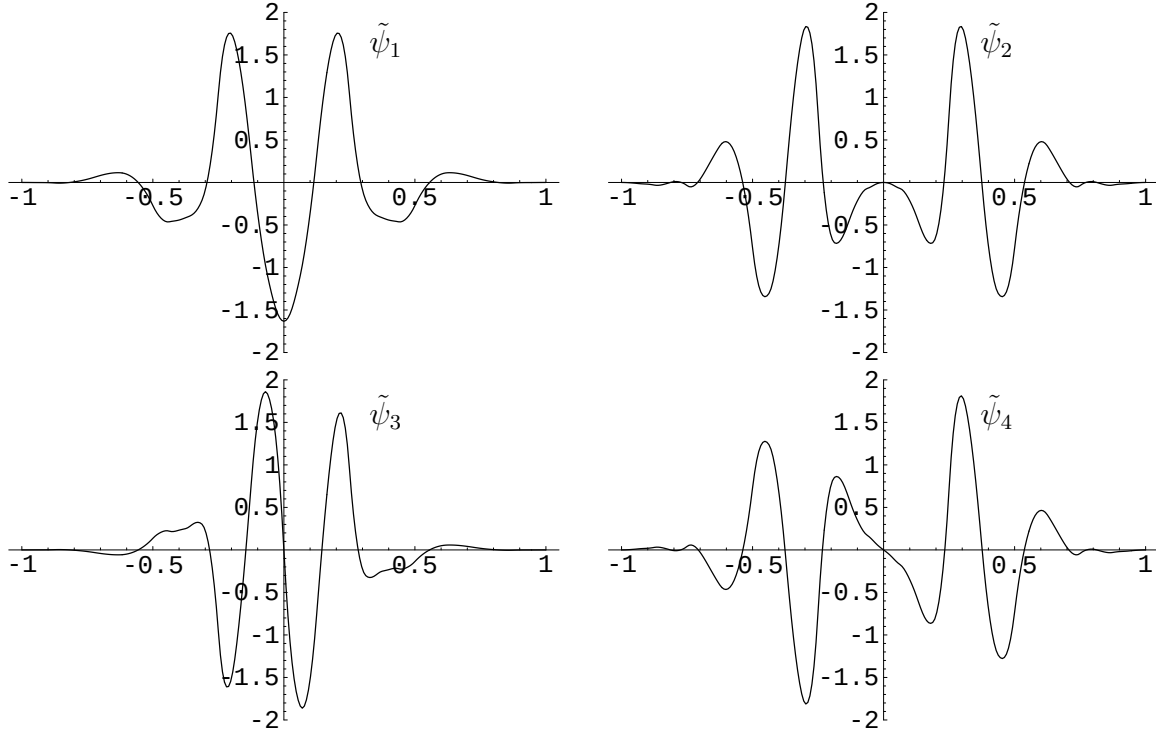


Figure 6: The multiwavelet $\tilde{\Psi} = (\tilde{\psi}_1, \tilde{\psi}_2, \tilde{\psi}_3, \tilde{\psi}_4)^T$, where $S(\tilde{\Psi}) = W_0$.

3 An Application of the Bases

One possible use for this smoother scaling vector is in image compression applications. The original version of JPEG used the discrete cosine transform, a non-wavelet technique, on disjoint 8×8 blocks of pixels to decompose the image data. This caused noticeable distortions, or “artifacts,” to appear in the reconstructed image at higher compression ratios. The discontinuous Haar wavelet basis can also be used to decompose the image data, but with similar results at higher compression ratios. Even continuous bases like Daubechies’s D4 scaling function or the GHM scaling vector constructed in Example 2, which are not differentiable everywhere, leave sharp distortions in the image at higher compression ratios. We would expect that by using a smoother, differentiable basis like we have constructed in this paper, we should be able to produce a smoother compressed image with no sharp artifacts.

Let c_i denote the sequence of scaling function coefficients in V_i , and let d_i denote the sequence of wavelet coefficients in W_i . The wavelet approach to producing a compressed

image is to take a signal (or function) in V_0 and find its best approximation in the smoother, nested function space V_1 , keeping the error in the wavelet space W_1 . If the function is close to being in V_1 , then the wavelet coefficients will be very small (close to zero). We may then repeat this process with the function in V_1 , and then V_2 , etc. The process is illustrated in Figure 7 for a decomposition to the fourth level. The image may be reconstructed exactly from c_n , where n is the last level of decomposition, and from the wavelet coefficients $d_1, \dots, d_n$. However, wavelet schemes typically trade accuracy of the image for more efficient storage

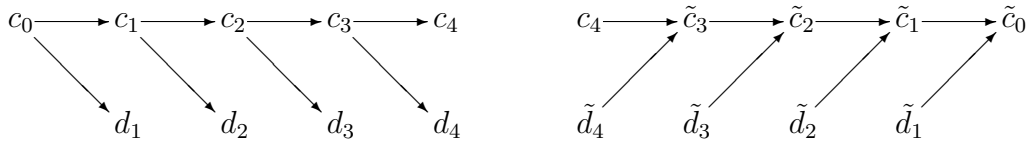


Figure 7: The decomposition of the signal c_0 and the reconstruction of $\tilde{c}_0$.

and/or transmission of the image. In order to achieve compression of the signal, we *quantize* the wavelet coefficients, replacing coefficients with values in a certain range with a central value. This creates a signal with lower *entropy*, a measure of the smallest average bits per character needed to store a signal without losing information. This measure is quantifiable, using the formula

$$E = - \sum_{i=1}^N p(i) \log_2 p(i),$$

where N is the number of distinct characters in the signal, and $p(i)$ is the relative frequency of the i^{th} character. The altered signal is then stored near this bit-rate using an entropy encoder. This is called *lossless* compression. A less accurate version of the original image can then be reconstructed from the quantized wavelet coefficients, as illustrated in Figure 7. This is called *lossy* compression.

3.1 Prefiltering

The process of turning discrete data into a function in V_0 is called *prefiltering*. Ideally, a prefilter (usually a set of matrices) should be orthogonal (norm-preserving) and send data sampled from a polynomial of degree n to a multiple of the same polynomial in V_0 , up to the approximation order of V_0 , using as few matrices in the prefilter as possible. Prefiltering is not an issue when using a single scaling function, as using the raw data as the basis coefficients (the identity prefilter) accomplishes each of these goals. However, prefiltering becomes vitally important when using multiple scaling functions: the best basis will perform poorly in applications if the data is not prefiltered efficiently. See [12] for a more thorough discussion on prefiltering.

Multiple-matrix prefilters that accomplish all of the above goals can be difficult to find. We may always find an orthogonal single-matrix prefilter that preserves constant data by using the following method. Consider a scaling vector Φ of length r that generates the function space V_0 of approximation order $n \leq r$. Let a_k be the r -vector of function values sampled uniformly over $[0, 1]$ from $f_k(x) = x^k$, $k = 0, \dots, n-1$, and let a_k fill out the space of r -vectors for $k = n, \dots, r-1$. Likewise, let α_k be the r -vector of basis coefficients of $\Phi(x)$

needed to construct $f_k(x)$, $k = 0, \dots, n-1$, and let α_k fill out the space of r -vectors for $k = n, \dots, r-1$. Then use the Gram-Schmidt process on $\{a_0, \dots, a_{r-1}\}$ and $\{\alpha_0, \dots, \alpha_{r-1}\}$ in increasing order of degrees, and normalize to get $\{\tilde{a}_0, \dots, \tilde{a}_{r-1}\}$ and $\{\tilde{\alpha}_0, \dots, \tilde{\alpha}_{r-1}\}$. We may then construct the orthogonal, single-matrix prefilter

$$Q = [\tilde{\alpha}_0 \quad \cdots \quad \tilde{\alpha}_{r-1}] \begin{bmatrix} \tilde{a}_0^T \\ \vdots \\ \tilde{a}_{r-1}^T \end{bmatrix}$$

which maps $\tilde{a}_k$ to $\tilde{\alpha}_k$. While the prefilter Q does not preserve polynomial order above constants, it gets fairly close, and has the advantages of having very short support and being orthogonal. We have observed that this type of prefilter works well in applications, and so, we will use this type of prefilter for the GHM scaling vector and the new scaling vector in the following example. We concede that better prefilters may be found for both scaling vectors. (In fact, see [12] for more elaborate prefilters for the GHM scaling vector.)

Postfiltering, turning the basis coefficients back into discrete data, is usually just an inversion of the prefiltering process. For the prefilter Q described above, we use the postfilter Q^T , since $Q^T Q = I$.

3.2 Comparison of Reconstructed Images

In this section, we consider a digital 512×512 grayscale image called “Zelda,” basically a rectangular data set ranging in value from 0 to 255, requiring 256 Kb of memory (1 byte per pixel) for storage as raw data. The original image is shown in Figure 8. We decompose the image using three different bases: the single D4 scaling function, the GHM scaling vector constructed in Example 2 from a C^0 macroelement, and the scaling vector from Theorem 1 constructed from a C^1 macroelement. We quantize the wavelet coefficients uniformly, with 0 at the center of the zero-bin, to achieve a moderate 25:1 compressed image (file size 10.24 Kb) and an extreme 50:1 compressed image (file size 5.12 Kb), based on the entropy of the quantized signal. Particularly in the 50:1 compressed images, the reader will detect distortions specific to the basis being used. The error in the images will be measured visually by the reader and with the *root mean square error (RMSE)*, defined by

$$RMSE = \sqrt{\frac{\sum_{i,j} (\text{original}_{i,j} - \text{new}_{i,j})^2}{\text{rows} \times \text{columns}}}.$$

Reconstructions from the quantized data are shown in Figure 9 for the D4 basis, in Figure 10 for the GHM scaling vector, and in Figure 11 for the new scaling vector constructed in this paper. The reconstructions for the smooth scaling vector constructed in this paper have the lowest *RMSE*’s of the three methods. Hopefully, the reader also finds the artifacts are less noticeable in the images where the new basis is used.



Figure 8: The original image 512×512 grayscale image Zelda.



$RMSE \approx 4.88688$

$RMSE \approx 6.88849$

Figure 9: A 25:1 and 50:1 compression using the D4 scaling function.

4 Appendix

4.1 Scaling Vector Coefficients

The matrix coefficients of the scaling vector constructed in Section 2.2 satisfying

$$\tilde{\Phi}(x) = \sqrt{2} \sum_{i=-2}^1 g_i \tilde{\Phi}(2x - i)$$



$RMSE \approx 4.37044$

$RMSE \approx 6.23848$

Figure 10: A 25:1 and 50:1 compression using the GHM scaling vector.



$RMSE \approx 4.13475$

$RMSE \approx 5.84525$

Figure 11: A 25:1 and 50:1 compression using the new scaling vector.

are given below.

$$g_{-2} = \begin{bmatrix} 0 & 0 & 0.02669163201881291 & 0.026188780570349884 \\ 0 & 0 & -0.03940748766209954 & -0.040769035522339014 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix}$$

$$\begin{aligned}
g_{-1} &= \begin{bmatrix} -0.1175124743509559 & -0.10652668959666897 & 0.30676124077120465 & 0.35964167838854366 \\ 0.18731879383595215 & 0.18181954032277048 & -0.4224381824999862 & -0.433225475337907 \\ 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 \end{bmatrix} \\
g_0 &= \begin{bmatrix} 0.7071067811865841 & 0 & 0.3067612407712861 & -0.3596416783882024 \\ 0 & 0.35355339059336993 & 0.42243818250022613 & -0.43322547533688205 \\ 0 & 0 & 0.47106586494422936 & 0.23604093395898632 \\ 0 & 0 & -0.5666158791523569 & -0.2916640296146426 \end{bmatrix} \\
g_1 &= \begin{bmatrix} -0.11751247435904282 & 0.10652668958949803 & 0.026691633234707347 & -0.026188779541901065 \\ -0.1873187938600181 & 0.18181954030143205 & 0.03940749128037211 & -0.040769032461869345 \\ 0.6669057455800359 & 0 & 0.47106586494422936 & -0.23604093395898632 \\ 0 & 0.4333094490577162 & 0.5666158791523569 & -0.2916640296146426 \end{bmatrix}
\end{aligned}$$

4.2 Multiwavelet Coefficients

The matrix coefficients of the multiwavelet constructed in Section 2.3 satisfying

$$\tilde{\Psi}(x) = \sqrt{2} \sum_{i=-2}^1 h_i \tilde{\Phi}(2x - i)$$

are given below.

$$\begin{aligned}
h_{-2} &= \begin{bmatrix} 0 & 0 & 0.0266916313500821 & 0.02618877991421749 \\ 0 & 0 & 0.04374365300600389 & 0.06435982959198171 \\ 0 & 0 & -0.013410540810749855 & -0.013157895717088848 \\ 0 & 0 & -0.04006995259666773 & -0.06065185761803572 \end{bmatrix} \\
h_{-1} &= \begin{bmatrix} -0.11751247140680435 & -0.10652668692775463 & 0.3067612330856239 & 0.3596416693780991 \\ -0.3334538093583893 & -0.42469284946317537 & 0.19229779050087667 & -0.40670678109361624 \\ 0.05904119433180234 & 0.053521662417019424 & -0.15412448873935367 & -0.18069293784197155 \\ 0.3166000901995356 & 0.40882389594960616 & -0.15157506646288127 & 0.4513983241700974 \end{bmatrix} \\
h_0 &= \begin{bmatrix} -0.7071067870918081 & 0 & 0.30676123308570535 & -0.3596416693777578 \\ 0 & 0 & 0.19229779050222978 & 0.40670678109948577 \\ 0 & -0.9347644627208598 & 0.15412448873939458 & -0.18069293784180007 \\ 0 & -0.034862829700229456 & 0.15157506646421717 & 0.45139832417589415 \end{bmatrix} \\
h_1 &= \begin{bmatrix} -0.11751247141489118 & 0.10652668692058369 & 0.02669163256597651 & -0.026188778885768697 \\ -0.3334538094952175 & 0.424692849341863 & 0.043743673577174125 & -0.06435981219212034 \\ -0.05904119433586537 & 0.05352166241341656 & 0.013410541421645506 & -0.013157895200370542 \\ -0.31660009033465436 & 0.40882389582980966 & 0.04006997291081938 & -0.060651840435570266 \end{bmatrix}
\end{aligned}$$

4.3 Prefilter Matrices

The prefilter used in the GHM scaling vector image compression examples in Section 3.2 is given below. This prefilter originally appeared in [12].

$$Q = \begin{bmatrix} \frac{1}{\sqrt{3}} + \frac{1}{\sqrt{6}} & -\frac{1}{\sqrt{3}} + \frac{1}{\sqrt{6}} \\ \frac{1}{\sqrt{3}} - \frac{1}{\sqrt{6}} & \frac{1}{\sqrt{3}} + \frac{1}{\sqrt{6}} \end{bmatrix}$$

The prefilter used in the new scaling vector image compression examples in Section 3.2 is given below.

$$Q = \begin{bmatrix} 0.843806675746766 & 0.48588947587427395 & -0.22260947332787062 & 0.04844309635543219 \\ -0.5359722073168798 & 0.7840831502767046 & -0.30731469556949864 & 0.05920375260967379 \\ 0.015288019449165685 & 0.26864647184429297 & 0.7701708393722542 & 0.5783011566714802 \\ -0.022144150703698316 & -0.27740602706575324 & -0.5126788258279672 & 0.8122290035974187 \end{bmatrix}$$

Acknowledgements: Research was supported by the Kentucky Science and Engineering Foundation, Grant KSEF-148-502-03-57.

References

- [1] M. Barnsley, "Fractal functions and interpolations," *Constr. Approx.* **2**, 303–329 (1986).
- [2] M. Barnsley, J. Elton, D. Hardin, and P. Massopust, "Hidden variable fractal interpolation functions," *SIAM Journ. Math. Anal.* **20:5**, 1218–1242 (1989).
- [3] I. Daubechies, "Orthonormal bases of compactly supported wavelets," *Commun. on Pure and Appl. Math.* **41**, 909–996 (1988).
- [4] G. Donovan, J. Geronimo, and D. Hardin, "Fractal functions, splines, intertwining multiresolution analysis, and wavelets," *Wavelet Application in Signal and Image Processing*, Laine and Unser, editors, SPIE Conf. Proc., San Diego, 238–243 (1994).
- [5] G. Donovan, J. Geronimo, and D. Hardin, "Intertwining multiresolution analyses and the construction of piecewise polynomial wavelets," *SIAM Journ. Math. Anal.* **27**(6), 1791–1815 (1996).
- [6] G. Donovan, J. Geronimo, D. Hardin, and P. Massopust, "Construction of orthogonal wavelets using fractal interpolation functions," *SIAM Journ. Math. Anal.* **27**, 1158–1192 (1996).
- [7] J. Geronimo, D. Hardin, and P. Massopust, "Fractal functions and wavelet expansions based on several scaling functions," *J. Approx. Theory* **78:3** 373–401 (1994).
- [8] T. N. T. Goodman, "A class of orthogonal refinable functions and wavelets," *Constr. Approx.* **19:4**, 525–540 (2003).
- [9] B. Han and Q. Jiang, "Multiwavelets on the interval," *Appl. Comp. Har. Anal.* **12:1** 100–127 (2002).
- [10] D. Hardin and B. Kessler, "Orthogonal macroelement scaling vectors and wavelets in 1-D", (invited paper for special issue on fractals and wavelets), *The Arabian Journal for Science and Engineering*, **28:1C**, 73–88 (2003).
- [11] D. Hardin, B. Kessler, and P. Massopust, "A multiresolution analysis based on fractal functions," *J. Approx. Theory* **71:1**, 104–120 (1992).

- [12] D. Hardin and D. Roach, "Multiwavelet prefilters I: orthogonal prefilters preserving approximation order $p \leq 2$," *IEEE Trans. Circuits Syst. II: Analog Digital Signal Processing* **45:8**, 1106–1112 (1998).
- [13] D. Hong and A. Wu, "Orthogonal multiwavelets of multiplicity four," *Comput. Math. Applic.* **40**, 1153–1169 (2000).
- [14] R.-Q. Jia, "Refinable shift-invariant spaces: from splines to wavelets," *Approx. Theory VIII* **2**, 179–208 (1995).
- [15] R.-Q. Jia and Z. Shen, "Multiresolution and wavelets," *Proc. Edinburgh Math. Soc.* **37**, 271–300 (1994).
- [16] G. Strang and V. Strela, "Short wavelets and matrix dilation equations," *IEEE Trans. SP* **43**, 108–115 (1995).
- [17] P. Vaidyanathan, "Multirate Systems and Filter Banks," Simon and Schuster (1993).

The existence of tight MRA multiwavelet frames

Qun Mo

Department of Mathematical and Statistical Sciences
University of Alberta
Edmonton, Alberta, Canada T6G 2G1
E-mail: mo@math.ualberta.ca

Abstract

We shall prove that under a mild assumption, we can construct tight multiwavelet frames with the highest possible order of vanishing moments from any given refinable function vector. This paper generalizes Chui, He and Stöckler's work for dilation 2 and Chui, He, Stöckler and Sun's work for any integer dilation larger than 2 on the existence of tight wavelet frames derived from any given refinable function.

Keywords: tight wavelet frames, vanishing moments, transition operator, the canonical form of a matrix mask, matrix-valued Fejér-Riesz Lemma, matrix approximation.

2000 AMS subject classification: 42C40, 42C15, 41A15, 41A25.

1 Introduction

Wavelet frames are proved to be a very useful tool in signal denoising and are being explored in other applications. Usually we derive wavelet frames from refinable functions. Thus the property of a wavelet frame derived from a refinable function is highly related to the property of the refinable function. However, some desirable properties cannot coexist very well at any refinable function. For instance, high smoothness and short support are two desirable properties of a refinable function. But the smoothness of a refinable function is restricted by its support. On the other hand, refinable function vectors, as generalizations of refinable functions, may have higher smoothness than those refinable functions with the same length limit of their support. Therefore, it is interesting to construct multiwavelet frames, especially tight multiwavelet frames, from refinable function vectors.

In literature, Daubechies, Han, Ron and Shen ([6]) and Chui, He and Stöckler ([2]) independently discovered a way to derive wavelet frames from refinable functions. There are many papers inspired by their results. For instance, Han and Mo have generalized their results on the construction of pairs of dual wavelet frames from the classical wavelet setting to the multiwavelet setting ([12]), have

given a general construction of tight wavelet frames with three (anti)symmetric generators derived from B -splines ([13]), and have given a characterization of tight wavelet frames with two (anti)symmetric generators derived from any given refinable function ([14]). Also, Chui, He, Stöckler and Sun ([3]) have generalized Chui, He and Stöckler's result ([2]) on the existence of tight wavelet frames from dilation 2 to any integer dilation larger than 2. We can see there is an interesting problem left unknown: *under a mild assumption, can we always derive tight multiwavelet frames with the highest possible vanishing moments from any given refinable function vector?* In this paper, we shall give a positive answer to the above question.

Before proceeding further, let us give some definitions first. From now on, by d we denote a dilation factor which is an integer such that $|d| > 1$. For a function $f \in L_2(\mathbb{R})$, we define

$$f_{j,k} := |d|^{j/2} f(d^j \cdot -k), \quad j, k \in \mathbb{Z}.$$

Define

$$\langle f, g \rangle := \int_{\mathbb{R}} f(t) \overline{g(t)} dt \quad \forall f, g \in L_2(\mathbb{R})$$

and $\|f\|^2 := \langle f, f \rangle$. Let $\{\psi^1, \dots, \psi^L\}$ be a finite set of functions in $L_2(\mathbb{R})$, we say that $\{\psi^1, \dots, \psi^L\}$ generates a **wavelet frame** $\{\psi_{j,k}^\ell\}_{j,k \in \mathbb{Z}; \ell=1, \dots, L}$ in $L_2(\mathbb{R})$ if there exist positive constants C_1 and C_2 such that

$$C_1 \|f\|^2 \leq \sum_{\ell=1}^L \sum_{j \in \mathbb{Z}} \sum_{k \in \mathbb{Z}} |\langle f, \psi_{j,k}^\ell \rangle|^2 \leq C_2 \|f\|^2 \quad \forall f \in L_2(\mathbb{R}). \quad (1.1)$$

In particular, when $C_1 = C_2 (= 1)$ in (1.1), we say that $\{\psi^1, \dots, \psi^L\}$ generates a (normalized) **tight wavelet frame** in $L_2(\mathbb{R})$. We shall assume that every tight wavelet frame is normalized in this paper.

A tight wavelet frame $\{\psi_{j,k}^\ell\}_{j,k \in \mathbb{Z}; \ell=1, \dots, L}$ can represent a function $f \in L_2(\mathbb{R})$ as

$$f = \sum_{\ell=1}^L \sum_{j,k \in \mathbb{Z}} \langle f, \psi_{j,k}^\ell \rangle \psi_{j,k}^\ell.$$

Furthermore, if a tight wavelet frame $\{\psi_{j,k}^\ell\}_{j,k \in \mathbb{Z}; \ell=1, \dots, L}$ is linearly independent, then it is an orthonormal wavelet basis. In wavelet applications, wavelet frames have already been proved to be very useful for signal denoising and currently are being explored for image compression. It is of interest in both theory and applications to construct wavelet frames with certain properties.

An important property of a wavelet frame is its order of vanishing moments. A function f with enough decay is said to have **vanishing moments** of order m if

$$\int_{\mathbb{R}} t^k f(t) dt = 0 \quad \forall k = 0, \dots, m-1.$$

A wavelet frame $\{\psi_{j,k}^\ell\}_{j,k \in \mathbb{Z}; \ell=1, \dots, L}$ has **vanishing moments** of order m if all its generators $\psi^1, \dots, \psi^L$ have vanishing moments of order m . The order of

vanishing moments of a wavelet frame has a great impact on how efficient the wavelet frame is in representing a function.

We can construct the generators of a wavelet frame from refinable function vectors. In this case, the wavelet frame is called an “**MRA multiwavelet frame**”. An $r \times 1$ function vector $\phi = (\phi_1, \dots, \phi_r)^T \in (L_2(\mathbb{R}))^r$ is a **refinable function vector**, where r is called the “multiplicity”, if ϕ satisfies the following **refinement equation**,

$$\phi = |d| \sum_{k \in \mathbb{Z}} a_k \phi(d \cdot -k), \quad (1.2)$$

where a is a finitely supported sequence of $r \times r$ complex-valued matrices on $\mathbb{Z}$, called the “**matrix mask**” for ϕ . Under an appropriate mild condition (see [1, 5, 16, 20, 18, 22, 23]) on the matrix mask a , there exists a unique normalized distributional solution of the refinement equation (1.2). The refinement equation as well as the various properties of its refinable function vector has been well studied in the literature. To name a few, please see [1, 5, 7, 8, 10, 16, 18, 20, 21, 22, 23, 26, 27, 28] and the references therein. A refinable function vector with the multiplicity $r = 1$ is a scalar refinable function.

The remaining part of this paper is as follows. In Section 2, we shall give more definitions and discuss some properties of the transition operator T_a . In Section 3, we shall build some lemmas for matrix inequalities. In Section 4, we shall prove the existence of tight MRA multiwavelet frames with the maximal possible vanishing moments. In Section 5, we shall show one example of tight multiwavelet frames for the case $r = 2$ and $d = 2$.

2 The transition operator T_a

Let $\phi = (\phi_1, \dots, \phi_r)^T \in (L_2(\mathbb{R}))^r$ be a refinable function vector and let a be the matrix mask of ϕ . We say that a mask a with multiplicity r satisfies the **sum rules** of order m with respect to the lattice $d\mathbb{Z}$ if there exists a $1 \times r$ row vector $y(\xi)$ of 2π -periodic trigonometric polynomials such that $y(0) \neq 0$ (we assume $y(0)\hat{\phi}(0) \neq 0$ in this paper) and for $k = 0, \dots, |d| - 1$,

$$[y(d \cdot) a(\cdot)]^{(j)}(2\pi k/d) = \delta_k y^{(j)}(0), \quad j = 0, \dots, m-1, \quad (2.1)$$

where δ is the **Dirac sequence** such that $\delta_0 = 1$ and $\delta_k = 0$ for all $k \in \mathbb{Z} \setminus \{0\}$, and $y^{(j)}(\xi)$ denotes the j th derivative of $y(\xi)$ for all $j \in \mathbb{Z}$. As discussed in [12, Theorem 2.3], if a matrix mask a satisfies the summer rule of order m , we can assume that a takes the canonical form

$$a(\xi) = \begin{bmatrix} (1 + e^{-i\xi} + \dots + e^{-i(|d|-1)\xi})^m P_{1,1}(\xi) & (1 - e^{-i|d|\xi})^m P_{1,2}(\xi) \\ (1 - e^{-i\xi})^m P_{2,1}(\xi) & P_{2,2}(\xi) \end{bmatrix}, \quad (2.2)$$

where $P_{1,1}, P_{1,2}, P_{2,1}$ and $P_{2,2}$ are some $1 \times 1, 1 \times (r-1), (r-1) \times 1$ and $(r-1) \times (r-1)$ matrices of 2π -periodic trigonometric polynomials respectively and $P_{1,1}(0) = |d|^{-m}$. We shall keep this assumption from now on.

Define $\mathbb{T} := \mathbb{R}/2\pi\mathbb{R}$. Denote $\mathcal{P}(\mathbb{T})^{r \times r}$ the space of $r \times r$ matrices of 2π -periodic trigonometric polynomials, $C^\infty(\mathbb{T})^{r \times r}$ the space of $r \times r$ matrices of 2π -periodic C^∞ -functions and $C(\mathbb{T})^{r \times r}$ the space of $r \times r$ matrices of 2π -periodic continuous functions. Given a matrix A , we denote A^T the transpose of A , A^* the transpose of the complex conjugate of A , A^{adj} the adjoint matrix of A , and $A_{i,j}$ the (i, j) -entry of A . For any $A \in C(\mathbb{T})^{r \times r}$, we say $A > 0$ (or $A \geq 0$) if for all $\xi \in \mathbb{T}$, $A(\xi) > 0$ which means $A(\xi)$ is positive definite (or $A(\xi) \geq 0$ which means $A(\xi)$ is positive semi-definite).

The **transition operator** T_a is defined by

$$(T_a F)(d\xi) := \sum_{j=0}^{|d|-1} a(\xi + 2\pi j/d) F(\xi + 2\pi j/d) a(\xi + 2\pi j/d)^* \quad \forall F \in C(\mathbb{T})^{r \times r}. \quad (2.3)$$

By the definition, we know that T_a is well defined and T_a maps $C(\mathbb{T})^{r \times r}$ and $\mathcal{P}(\mathbb{T})^{r \times r}$ into $C(\mathbb{T})^{r \times r}$ and $\mathcal{P}(\mathbb{T})^{r \times r}$, respectively. A special eigenfunction of T_a is Φ , the bracket product of $\widehat{\phi}$ and $\widehat{\phi}$. For $f, g \in (L_2(\mathbb{R}))^r$, define the **bracket product** (see [19]) as

$$[f, g](\xi) := \sum_{j \in \mathbb{Z}} f(\xi + 2\pi j) g(\xi + 2\pi j)^* \quad \forall \xi \in \mathbb{T}.$$

Since $f, g \in (L_2(\mathbb{R}))^r$, $[f, g]$ is an $r \times r$ matrix of functions in $L_1(\mathbb{T})$. Define

$$\Phi := [\widehat{\phi}, \widehat{\phi}].$$

Then we can verify that Φ is an eigenfunction of T_a as follows.

$$\begin{aligned} T_a \Phi(d\xi) &= \sum_{j=0}^{|d|-1} a(\xi + 2\pi j/d) \Phi(\xi + 2\pi j/d) a(\xi + 2\pi j/d)^* \\ &= \sum_{j=0}^{|d|-1} \sum_{k \in \mathbb{Z}} a\left(\xi + \frac{2\pi j}{d}\right) \widehat{\phi}\left(\xi + \frac{2\pi j}{d} + 2\pi k\right) \widehat{\phi}\left(\xi + \frac{2\pi j}{d} + 2\pi k\right)^* a\left(\xi + \frac{2\pi j}{d}\right)^* \\ &= \sum_{j=0}^{|d|-1} \sum_{k \in \mathbb{Z}} \widehat{\phi}(d\xi + 2\pi j + 2\pi dk) \widehat{\phi}(d\xi + 2\pi j + 2\pi dk)^* \\ &= \Phi(d\xi). \end{aligned}$$

Thus we have $T_a \Phi = \Phi$.

Moreover, we have the following lemma.

Lemma 2.1 *Let $\phi = (\phi_1, \dots, \phi_r)^T \in (L_2(\mathbb{R}))^r$ be a refinable function vector with matrix mask a . Define $\Phi := [\widehat{\phi}, \widehat{\phi}]$. It is evident that $\Phi^* = \Phi$. If the matrix mask a takes the canonical form (2.2), then we have*

$$(1 - e^{-i\xi})^m \mid \Phi(\xi)_{1,j} = \overline{\Phi(\xi)_{(j,1)}}, \quad j = 2, \dots, r$$

and

$$(1 - e^{-i\xi})^m \mid \Phi^{adj}(\xi)_{1,j} = \overline{\Phi^{adj}(\xi)_{j,1}}, \quad j = 2, \dots, r.$$

Moreover, if $\Phi > 0$, then we have

$$(1 - e^{-i\xi})^m \mid \Phi^{-1}(\xi)_{1,j} = \overline{\Phi^{-1}(\xi)_{j,1}}, \quad j = 2, \dots, r$$

and

$$(1 - e^{-i\xi})^{2m} \mid [\Phi^{-1}(\xi) - a(\xi)^* \Phi^{-1}(d\xi) a(\xi)]_{1,1}.$$

Proof: By assumption, we have

$$a(\xi) = \begin{bmatrix} (1 - e^{-id\xi})^m / (1 - e^{-i\xi})^m * & (1 - e^{-id\xi})^m * \\ (1 - e^{-i\xi})^m * & * \end{bmatrix}, \quad (2.4)$$

where $*$ denotes some 1×1 , $1 \times (r-1)$, $(r-1) \times 1$ and $(r-1) \times (r-1)$ matrices of trigonometric polynomials. Since ϕ is refinable, we have

$$\widehat{\phi}(d\xi) = a(\xi) \widehat{\phi}(\xi).$$

Taking the 0th, $\dots$, $(m-1)$ th derivatives at $\xi = 0$ from both sides of the equation above, since $a(\xi)$ has the canonical form (2.4), we have

$$\widehat{\phi}_j^{(\ell)}(0) = 0, \quad j = 2, \dots, r; \ell = 0, 1, \dots, m-1.$$

Notice that ϕ_1 is a superfunction. We have for $\ell = 0, 1, \dots, m-1$,

$$\sum_{k \in \mathbb{Z}} k^\ell \phi_1(x+k) = p_\ell(x),$$

where p_ℓ is a polynomial with degree at most ℓ . Notice

$$\Phi(\xi)_{1,j} = \sum_{k \in \mathbb{Z}} \left[\int \phi_1(x+k) \overline{\phi_j(x)} dx \right] e^{-ik\xi} =: \sum_{k \in \mathbb{Z}} b_{k,j} e^{-ik\xi}$$

and for $\ell = 0, \dots, m-1$, $j = 2, \dots, r$,

$$\begin{aligned} \sum_{k \in \mathbb{Z}} \overline{k^\ell b_{k,j}} &= \int \sum_{k \in \mathbb{Z}} \overline{[k^\ell \phi_1(x+k)]} \phi_j(x) dx \\ &= \int \overline{p_\ell(x)} \phi_j(x) dx \\ &= \left(\overline{p_\ell(-iD)} \widehat{\phi_j} \right)(0) = 0. \end{aligned}$$

Therefore we have

$$(1 - e^{-i\xi})^m \mid \Phi(\xi)_{1,j} = \overline{\Phi(\xi)_{j,1}}, \quad j = 2, \dots, r. \quad (2.5)$$

By the definition of adjoint matrices, $(-1)^{j+1}\Phi_{1,j}^{\text{adj}}$, $j = 2, \dots, r$, is the determinant of a sub-matrix of Φ . Since the first row of this sub-matrix has a common factor $(1 - e^{-i\xi})^m$, we have

$$(1 - e^{-i\xi})^m \mid \Phi^{\text{adj}}(\xi)_{1,j} = \overline{\Phi^{\text{adj}}(\xi)_{j,1}}, \quad j = 2, \dots, r. \quad (2.6)$$

Moreover, if $\Phi > 0$, then $\det\Phi(\xi) > 0$ for all $\xi \in \mathbb{T}$. Therefore, $\det\Phi(0) \neq 0$. By the matrix identity $A^{-1} = A^{\text{adj}}/\det A$ and (2.6), we have

$$(1 - e^{-i\xi})^m \mid \Phi^{-1}(\xi)_{1,j} = \overline{\Phi^{-1}(\xi)_{j,1}}, \quad j = 2, \dots, r. \quad (2.7)$$

As proved before, $T_a\Phi = \Phi$, i.e.,

$$\Phi(d\xi) = \sum_{j=0}^{|d|-1} a(\xi + 2\pi j/d)\Phi(\xi + 2\pi j/d)a(\xi + 2\pi j/d)^*.$$

By (2.4), we have

$$(1 - e^{-i\xi})^{2m} \mid [a(\xi + 2\pi j/d)\Phi(\xi + 2\pi j/d)a(\xi + 2\pi j/d)^*]_{1,1}, \quad j = 2, \dots, r.$$

Therefore,

$$\Phi(d\xi)_{1,1} = [a(\xi)\Phi(\xi)a(\xi)^*]_{1,1} + O(|\xi|^{2m}) \text{ as } \xi \rightarrow 0.$$

Combining the equality above with (2.4) and (2.5), we have

$$\Phi(d\xi)_{1,1} = |a(\xi)_{1,1}|^2\Phi(\xi)_{1,1} + O(|\xi|^{2m}) \text{ as } \xi \rightarrow 0.$$

Hence

$$(\Phi(\xi)_{1,1})^{-1} = |a(\xi)_{1,1}|^2(\Phi(\xi)_{1,1})^{-1} + O(|\xi|^{2m}) \text{ as } \xi \rightarrow 0. \quad (2.8)$$

By (2.5), (2.7) and $\Phi(\xi)\Phi^{-1}(\xi) = I_r$, we have

$$1 = \Phi(\xi)_{1,1}\Phi^{-1}(\xi)_{1,1} + O(|\xi|^{2m}) \text{ as } \xi \rightarrow 0.$$

Therefore,

$$(\Phi(\xi)_{1,1})^{-1} = \Phi^{-1}(\xi)_{1,1} + O(|\xi|^{2m}) \text{ as } \xi \rightarrow 0. \quad (2.9)$$

By (2.4), (2.7), (2.8) and (2.9), we have

$$\begin{aligned} \Phi^{-1}(\xi)_{1,1} &= (\Phi(\xi)_{1,1})^{-1} + O(|\xi|^{2m}) \\ &= |a(\xi)_{1,1}|^2(\Phi(d\xi)_{1,1})^{-1} + O(|\xi|^{2m}) \\ &= |a(\xi)_{1,1}|^2\Phi^{-1}(d\xi)_{1,1} + O(|\xi|^{2m}) \\ &= [a(\xi)^*\Phi^{-1}(d\xi)a(\xi)]_{1,1} + O(|\xi|^{2m}) \text{ as } \xi \rightarrow 0. \end{aligned}$$

Therefore

$$(1 - e^{-i\xi})^{2m} \mid [\Phi^{-1}(\xi) - a(\xi)^*\Phi^{-1}(d\xi)a(\xi)]_{1,1}.$$

■

Now let us consider the eigenvectors of T_a . As been proved, we have $T_a\Phi = \Phi$. Moreover, if the integer shifts of ϕ are stable, then the cascade algorithm associated with the mask a converges in L_2 . Hence, we have that T_a has a simple eigenvalue 1 and all other eigenvalues of T_a are less than 1 in modulus (see [26]). Thus under the condition the integer shifts of ϕ are stable, we have that Φ is the unique 1-eigenfunction of T_a and for all $\xi \in \mathbb{T}$, $\Phi(\xi) > 0$. Define $\tilde{a}$ as

$$\tilde{a}(\xi) := \begin{bmatrix} (1 - e^{-id\xi})^m & 0 \\ 0 & I_{r-1} \end{bmatrix}^{-1} a(\xi) \begin{bmatrix} (1 - e^{-i\xi})^m & 0 \\ 0 & I_{r-1} \end{bmatrix}.$$

By our assumption, a is of the form (2.4). Therefore $\tilde{a}$ is an $r \times r$ matrix of trigonometric polynomials. Define a new operator $T_{\tilde{a}}$ mapping $C(\mathbb{T})^{r \times r}$ into $C(\mathbb{T})^{r \times r}$ by

$$T_{\tilde{a}}F(d\xi) := \sum_{j=0}^{|d|-1} \tilde{a}(\xi + 2j\pi/d)F(\xi + 2j\pi/d)\tilde{a}(\xi + 2j\pi/d)^*. \quad (2.10)$$

Since $\tilde{a}$ is an $r \times r$ matrix of trigonometric polynomials, we denote $\mathcal{P}_{\tilde{a}}(\mathbb{T})^{r \times r}$ a subspace containing all $r \times r$ matrices of trigonometric polynomials up to a finite degree determined by $\tilde{a}$ such that it is an invariant subspace of $T_{\tilde{a}}$ and $\rho(T_{\tilde{a}}) = \rho(T_{\tilde{a}}|_{\mathcal{P}_{\tilde{a}}(\mathbb{T})^{r \times r}})$ (see [11]). Similarly, since a is an $r \times r$ matrix of trigonometric polynomials, denote $\mathcal{P}_a(\mathbb{T})^{r \times r}$ a subspace containing all $r \times r$ matrices of trigonometric polynomials up to a finite degree determined by a such that it is an invariant subspace of T_a and $\rho(T_a) = \rho(T_a|_{\mathcal{P}_a(\mathbb{T})^{r \times r}})$. It is obvious that

$$T_{\tilde{a}}\tilde{F} = \lambda\tilde{F} \implies T_aF = \lambda F,$$

where F is derived from $\tilde{F}$ by

$$F(\xi) = \begin{bmatrix} (1 - e^{-i\xi})^m & 0 \\ 0 & I_{r-1} \end{bmatrix} \tilde{F}(\xi) \begin{bmatrix} (1 - e^{-i\xi})^m & 0 \\ 0 & I_{r-1} \end{bmatrix}^* \quad \forall \xi \in \mathbb{T}. \quad (2.11)$$

Therefore, we have

$$\rho(T_{\tilde{a}}) \leq \rho(T_a). \quad (2.12)$$

Now we are in the position to prove the following theorem.

Theorem 2.2 *Let ϕ and a be defined as in Lemma 2.1. If ϕ has stable integer shifts, then there exist a positive number $\rho < 1$ and some $\tilde{F} \in \mathcal{P}(\mathbb{T})^{r \times r}$ such that $\tilde{F} > 0$ and $T_aF \leq \rho F$ where F is derived from $\tilde{F}$ by (2.11).*

Proof: Since $T_{\tilde{a}}$ is a linear operator acting on $\mathcal{P}_{\tilde{a}}(\mathbb{T})^{r \times r}$ which is a finite dimensional space, by the definition of spectrum, we know that there exists $0 \neq \tilde{G}(\xi) \in \mathcal{P}_{\tilde{a}}(\mathbb{T})^{r \times r}$ such that $T_{\tilde{a}}\tilde{G} = \lambda_0\tilde{G}$ and $|\lambda_0| = \rho(T_{\tilde{a}}|_{\mathcal{P}_{\tilde{a}}(\mathbb{T})^{r \times r}})$. Define

$$G(\xi) = \begin{bmatrix} (1 - e^{-i\xi})^m & 0 \\ 0 & I_{r-1} \end{bmatrix} \tilde{G}(\xi) \begin{bmatrix} (1 - e^{-i\xi})^m & 0 \\ 0 & I_{r-1} \end{bmatrix}^*,$$

then it is evident to see that we have $T_a G = \lambda_0 G$.

Since $\Phi(0)_{1,1} \geq |\widehat{\phi}_1(0)|^2 \neq 0$, $\Phi(0)_{1,1} \neq 0$ and $G(0)_{1,1} = 0$ by the definition of G . Thus for all $\lambda \in \mathbb{C}$, we have $\Phi \neq \lambda G$. Notice that Φ is the unique 1-eigenfunction of T_a , we have $|\lambda_0| = \rho(T_a) \neq 1$. Since ϕ has stable integer shifts, T_a has a simple eigenvalue 1 and all other eigenvalues of T_a are less than 1 in modulus (see [26]). By (2.12), we have $\rho(T_a) \leq 1$. Combining the fact that $\rho(T_a) \neq 1$, we have $|\lambda_0| = \rho(T_a) < 1$.

Choose $\rho_1 := [1 + \rho(T_a)]/2$. Then we have $\rho(T_a) < \rho_1 < 1$. Borrowing the idea from the proof of [24, Theorem 3], since $\rho(T_a/\rho_1) < 1$, $(Id - T_a/\rho_1)^{-1}$ is a well defined operator acting on $\mathcal{P}_a(\mathbb{T})^{r \times r}$ and

$$(Id - T_a/\rho_1)^{-1} = Id + T_a/\rho_1 + T_a^2/\rho_1^2 + \cdots,$$

where Id denotes the identity operator. Define $\tilde{F} = (Id - T_a/\rho_1)^{-1}I_r$, then $\tilde{F} = I_r + T_a I_r/\rho_1 + \cdots$. Thus $\tilde{F} \in \mathcal{P}_a(\mathbb{T})^{r \times r}$ and $\tilde{F} > 0$. By the definition of $\tilde{F}$, we have $(Id - T_a/\rho_1)\tilde{F} > 0$. Therefore, $T_a \tilde{F} < \rho_1 \tilde{F}$. Let F be derived from $\tilde{F}$ by (2.11), then we have $T_a F \leq \rho_1 F$. ■

After Theorem 2.2, our next major step will be proving that there exist a positive number ϵ and some $\Theta \in \mathcal{P}(\mathbb{T})^{r \times r}$ such that for all $\xi \in \mathbb{T}$,

$$(\Phi(\xi) + \epsilon F(\xi))^{-1} \leq \Theta(\xi) \leq (\Phi(\xi) + \epsilon \rho F(\xi))^{-1}.$$

Before going to this step, we need some lemmas to build up matrix approximation first.

3 Matrix approximation

Lemma 3.1 *If $A, B \in C(\mathbb{T})^{r \times r}$ such that $A^* = A$, $B^* = B$ and $A < B$. Then there exists some $P \in \mathcal{P}(\mathbb{T})^{r \times r}$ such that $A < P < B$.*

Proof: For every $\xi \in \mathbb{T}$, suppose $\lambda_1(\xi), \dots, \lambda_r(\xi)$ are all the eigenvalues of $(B - A)(\xi)$ and we have $\lambda_1(\xi) \geq \cdots \geq \lambda_r(\xi)$. Thus $B(\xi) - A(\xi) \geq \lambda_r(\xi)I_r$. Since $B(\xi) - A(\xi) > 0$, we have

$$\lambda_1(\xi) \geq \cdots \geq \lambda_r(\xi) > 0.$$

Hence we have, for all $\xi \in \mathbb{T}$,

$$\lambda_r(\xi) = \frac{\det(B(\xi) - A(\xi))}{\lambda_1(\xi) \cdots \lambda_{r-1}(\xi)} \geq \frac{\det(B(\xi) - A(\xi))}{[\text{trace}(B(\xi) - A(\xi))]^{r-1}} \geq c_1,$$

where c_1 is a suitable positive constant number we choose to satisfy the above inequality. Therefore, for all $u \in \mathbb{C}^r$, $u^*(B - A)u \geq c_1 u^*u$. Then we can get a trigonometric polynomial $P_1 \in \mathcal{P}(\mathbb{T})^{r \times r}$ such that for $1 \leq i, j \leq r$ and for all $\xi \in \mathbb{T}$,

$$|[P_1 - (A + B)/2](\xi)_{i,j}| < c_1/(4r^2).$$

Define $P := (P_1 + P_1^*)/2$, then we have $P^* = P$ and for all $u \in \mathbb{C}^r$,

$$\begin{aligned} u^*(P - A)u &= u^*[(A + B)/2 - A]u + u^*[P_1 - (A + B)/2]u/2 \\ &\quad + u^*[P_1^* - (A + B)^*/2]u/2 \\ &\geq c_1 u^*u/2 - c_1 u^*u/4/2 - c_1 u^*u/4/2 \\ &= c_1 u^*u/4. \end{aligned}$$

Thus $P - A \geq \frac{c_1}{4} I_r > 0$. Similarly, we have $B - P \geq \frac{c_1}{4} I_r > 0$. Thus $B < P < A$. $\blacksquare$

Lemma 3.2 *Let $A, B \in C^\infty(\mathbb{T})^{r \times r}$ be such that $A^* = A$, $B^* = B$ and $B - A = P_1 F P_1^*$, where $P_1 \in \mathcal{P}(\mathbb{T})^{r \times r}$, $\det P_1 \neq 0$ and $F \in C(\mathbb{T})^{r \times r}$, $F > 0$. Then there exists some $P \in \mathcal{P}(\mathbb{T})^{r \times r}$ such that $A \leq P \leq B$.*

Proof: For all $\xi \in \mathbb{T}$, define $p_1(\xi) := \det P_1(\xi)$. Since $P_1 \in \mathcal{P}(\mathbb{T})^{r \times r}$ and $\det P_1 \neq 0$, p_1 is a non-zero 2π -periodic trigonometric polynomial. Therefore, $|p_1|^2$ has finitely many roots in $\mathbb{T}$. Suppose all the roots of $|p_1|^2$ in $\mathbb{T}$ are

$$-\pi \leq \xi_1 < \dots < \xi_N < \pi$$

with multiplicities $\alpha_1, \dots, \alpha_N$, respectively. Then we have $|p_1|^2 = p_2 p_3$, where p_2 and p_3 are two 2π -periodic trigonometric polynomials such that $p_2(\xi) \neq 0$ for all $\xi \in \mathbb{T}$ and

$$p_3(\xi) = \prod_{k=1}^N (e^{-i\xi} - e^{-i\xi_k})^{\alpha_k}.$$

By the Lagrange Interpolation Theorem, for $1 \leq i, j \leq r$, there exist unique $p_{i,j} \in \mathcal{P}(\mathbb{T})^{1 \times 1}$ such that $\deg p_{i,j} < \deg p_3$ and

$$p_{i,j}^{(\ell)}(\xi_k) = A_{i,j}^{(\ell)}(\xi_k), \quad \ell = 0, \dots, \alpha_k - 1; \quad k = 1, \dots, N.$$

Define $\tilde{A}_{i,j} := |p_1|^{-2}(A_{i,j} - p_{i,j})$, $P_0 := [p_{i,j}]_{1 \leq i,j \leq r}$ and $\tilde{A} = [\tilde{A}_{i,j}]_{1 \leq i,j \leq r}$. It is obvious that $P_0 \in \mathcal{P}(\mathbb{T})^{r \times r}$ and $A = P_0 + p_1 \tilde{A} p_1^*$. Also, by the definition of $p_{i,j}$ and the fact that $A_{i,j}$ is a C^∞ -function, it is evident to see that $\tilde{A}_{i,j} = p_2^{-1} \frac{A_{i,j} - p_{i,j}}{p_3}$ is a continuous function on $\mathbb{T}$. Since $A^* = A$ and P_0 is uniquely determined by A , we have $P_0^* = P_0$ and $\tilde{A}^* = \tilde{A}$. By $F > 0$, $\tilde{A} \in C(\mathbb{T})^{r \times r}$ and Lemma 3.1, there exists some $P_2 \in \mathcal{P}(\mathbb{T})^{r \times r}$ such that

$$P_1^{\text{adj}} \tilde{A} (P_1^{\text{adj}})^* < P_2 < P_1^{\text{adj}} \tilde{A} (P_1^{\text{adj}})^* + F.$$

Let $P := P_0 + P_1 P_2 P_1^*$. We have

$$P - A = P_1 [P_2 - P_1^{\text{adj}} \tilde{A} (P_1^{\text{adj}})^*] P_1^* \geq 0.$$

Similarly $B - P \geq 0$. Thus $A \leq P \leq B$. $\blacksquare$

Recall $\Phi := [\hat{\phi}, \hat{\phi}]$, now we can prove the following proposition.

Proposition 3.3 Suppose $0 < \rho < 1$, $\Phi > 0$, $\tilde{F} \in C(\mathbb{T})^{r \times r}$, $\tilde{F} > 0$, F is derived from $\tilde{F}$ by (2.11), then there exist $\epsilon > 0$ and $\Theta \in \mathcal{P}(\mathbb{T})^{r \times r}$ such that

$$(\Phi + \epsilon F)^{-1} \leq \Theta \leq (\Phi + \epsilon \rho F)^{-1}.$$

Proof: For a constant $r \times r$ matrix $A \geq 0$, we have

$$(I + A)^{-1} = I - A + A^2(I + A)^{-1} = I - A + A(I + A)^{-1}A.$$

Hence we have

$$(I + A)^{-1} \geq I - A \quad \text{when } A \geq 0.$$

Moreover for a given positive number λ , if $0 \leq A \leq \lambda I_r$, we have

$$(I + A)^{-1} = I - A + A(I + A)^{-1}A \leq I - A + A^2 \leq I - A + \lambda A.$$

Hence

$$(I + A)^{-1} \leq I - A + \lambda A \quad \text{when } 0 \leq A \leq \lambda I_r. \quad (3.1)$$

Since $\Phi > 0$, we have $\Phi_1 := \Phi^{1/2} > 0$. Thus we have

$$\begin{aligned} (\Phi + \epsilon \rho F)^{-1} &= \Phi_1^{-1} [I + \epsilon \rho \Phi_1^{-1} F \Phi_1^{-1}]^{-1} \Phi_1^{-1} \\ &\geq \Phi_1^{-1} [I - \epsilon \rho \Phi_1^{-1} F \Phi_1^{-1}] \Phi_1^{-1} = \Phi^{-1} - \epsilon \rho \Phi^{-1} F \Phi^{-1}, \end{aligned}$$

i.e.,

$$(\Phi + \epsilon \rho F)^{-1} \geq \Phi^{-1} - \epsilon \rho \Phi^{-1} F \Phi^{-1}. \quad (3.2)$$

Choose $\epsilon > 0$ small enough such that $\epsilon F \leq (1 - \rho)\Phi/2$ since $\Phi \geq cI_r$ for some positive number c . By inequality (3.1), choosing $\lambda = (1 - \rho)/2$, similarly to the proof of inequality (3.2), we have

$$(\Phi + \epsilon F)^{-1} \leq \Phi^{-1} - \epsilon(1 + \rho)\Phi^{-1} F \Phi^{-1}/2. \quad (3.3)$$

Define

$$A_0(\xi) := \text{diag}[(1 - e^{-i\xi}), 1, \dots, 1] \quad (3.4)$$

and

$$\tilde{\Phi} := A_0^{-m} \Phi^{-1} A_0^m.$$

By Lemma 2.1 we know that $\tilde{\Phi} \in C^\infty(\mathbb{T})^{r \times r}$. By $\Phi > 0$ we know that the determinant of $\tilde{\Phi}$ is positive. By the definition of $\tilde{\Phi}$, we have

$$\begin{aligned} &[\Phi^{-1} - \epsilon \rho \Phi^{-1} F \Phi^{-1}] - [\Phi^{-1} - \epsilon(1 + \rho)\Phi^{-1} F \Phi^{-1}/2] \\ &= \epsilon(1 - \rho)\Phi^{-1} F \Phi^{-1}/2 \\ &= \epsilon(1 - \rho)A_0^m [\tilde{\Phi} A_0^{-m} F (A_0^{-m})^* \tilde{\Phi}^*] (A_0^m)^*/2 \\ &= \epsilon(1 - \rho)A_0^m [\tilde{\Phi} \tilde{F} \tilde{\Phi}^*] (A_0^m)^*/2. \end{aligned} \quad (3.5)$$

It is obvious that $\tilde{\Phi}\tilde{F}\tilde{\Phi}^* > 0$. By Lemma 3.2, there exists some $P \in \mathcal{P}(\mathbb{T})^{r \times r}$ such that

$$\Phi^{-1} - \epsilon(1 + \rho)\Phi^{-1}F\Phi^{-1}/2 \leq P \leq \Phi^{-1} - \epsilon\rho\Phi^{-1}F\Phi^{-1}. \quad (3.6)$$

Applying inequalities (3.2) and (3.3), we have

$$(\Phi + \epsilon F)^{-1} \leq P \leq (\Phi + \epsilon\rho F)^{-1}.$$

■

Now we can prove the following proposition.

Proposition 3.4 *Let ϕ and a be defined as in Theorem 2.2. If ϕ has stable integer shifts, then there exists $\Theta \in \mathcal{P}(\mathbb{T})^{r \times r}$ such that $\Theta(0)_{1,1} = 1$, $\Theta > 0$ and*

$$\Theta^{-1} - T_a(\Theta^{-1}) \geq 0. \quad (3.7)$$

Proof: By Theorem 2.2, there exists some positive number $\rho < 1$ and $\tilde{F} > 0$ such that $T_a F \leq \rho F$ where F is defined as in (2.11). As we discussed before, Φ satisfies $T_a \Phi = \Phi > 0$. Hence by Proposition 3.3, there exist $\epsilon > 0$ and some $P \in \mathcal{P}(\mathbb{T})^{r \times r}$ such that

$$0 < (\Phi + \epsilon F)^{-1} \leq P \leq (\Phi + \epsilon\rho F)^{-1},$$

i.e.,

$$\Phi + \epsilon\rho F \leq P^{-1} \leq \Phi + \epsilon F. \quad (3.8)$$

Let $\Theta := P > 0$, then by inequality (3.8) we have $\Theta(0)_{1,1} = 1$ and

$$\Theta^{-1} - T_a(\Theta^{-1}) \geq (\Phi + \epsilon\rho F) - T_a(\Theta^{-1}) \geq (\Phi + \epsilon\rho F) - T_a(\Phi + \epsilon F) \geq 0.$$

■

4 The existence of tight multiwavelet frames derived from refinable function vectors

Our final major step will be proving the main theorem in this paper. We shall go through the following well-known lemmas first.

Lemma 4.1 *Suppose A is an $m \times n$ matrix and B is an $n \times m$ matrix. Then we have*

$$\lambda^n \det(\lambda I_m - AB) = \lambda^m \det(\lambda I_n - BA).$$

Proof: We have the following identities:

$$\begin{bmatrix} I_m & A \\ B & \lambda I_n \end{bmatrix} \begin{bmatrix} \lambda I_m & 0 \\ -B & I_n \end{bmatrix} = \begin{bmatrix} \lambda I_m - AB & A \\ 0 & \lambda I_n \end{bmatrix}$$

and

$$\begin{bmatrix} I_m & A \\ B & \lambda I_n \end{bmatrix} \begin{bmatrix} \lambda I_m & -A \\ 0 & I_n \end{bmatrix} = \begin{bmatrix} \lambda I_m & 0 \\ \lambda B & \lambda I_n - BA \end{bmatrix}.$$

Taking the determinants of the matrices in the above two equations, we have

$$\lambda^n \det(\lambda I_m - AB) = \lambda^m \det(\lambda I_n - BA).$$

■

Lemma 4.2 *Let A be an $m \times n$ matrix and B be an $n \times m$ matrix. If $(AB)^* = AB$ and $(BA)^* = BA$, then we have*

$$I_m - AB \geq 0 \iff I_n - BA \geq 0.$$

Proof: Suppose $I_m - AB \geq 0$. Then all the eigenvalues of $(I_m - AB)$ are nonnegative. Hence

$$\det[\lambda I_m - (I_m - AB)] = \prod_{j=1}^m (\lambda - \lambda_j),$$

where for $j = 1, \dots, m$, λ_j is nonnegative. By Lemma 4.1, we have

$$\begin{aligned} (\lambda - 1)^{m-n} \det[\lambda I_n - (I_n - BA)] &= (\lambda - 1)^{m-n} \det[(\lambda - 1)I_n + BA] \\ &= \det[(\lambda - 1)I_m + AB] = \prod_{j=1}^m (\lambda - \lambda_j). \end{aligned}$$

Thus all the eigenvalues of $(I_n - BA)$ are nonnegative. Combining that $(BA)^* = BA$, we know that $I_n - BA \geq 0$. Therefore, $I_m - AB \geq 0$ implies $I_n - BA \geq 0$. Similarly, $I_n - BA \geq 0$ implies $I_m - AB \geq 0$. ■

Finally, we are in the position to state the main theorem in this paper.

Theorem 4.3 *Let $\phi \in (L_2(\mathbb{R}))^r$ be a refinable function vector. If the integer shifts of ϕ are stable and the matrix mask a of ϕ satisfies sum rules of order m , then there exists a tight multiwavelet frame derived from ϕ and it has vanishing moments of order m .*

Proof: By assumption, the matrix mask a takes the canonical form (2.4). By Proposition 3.4, there exist $\Theta \in \mathcal{P}(\mathbb{T})^{r \times r}$ such that $\Theta(0)_{1,1} = 1$, $\Theta > 0$ and $\Theta^{-1} - T_a(\Theta^{-1}) \geq 0$.

First we want to prove that $M \geq 0$, where M is defined by

$$M(\xi) := \begin{bmatrix} \Theta(\xi) & & \\ & \ddots & \\ & & \Theta\left(\xi + \frac{2\pi(|d|-1)}{d}\right) \end{bmatrix} - \begin{bmatrix} a(\xi)^* & & \\ a\left(\xi + \frac{2\pi}{d}\right)^* & & \\ \vdots & & \\ a\left(\xi + \frac{2\pi(|d|-1)}{d}\right)^* & & \end{bmatrix} \Theta(d\xi) \begin{bmatrix} a(\xi)^* & & \\ a\left(\xi + \frac{2\pi}{d}\right)^* & & \\ \vdots & & \\ a\left(\xi + \frac{2\pi(|d|-1)}{d}\right)^* & & \end{bmatrix}^*. \quad (4.1)$$

By the fact $\Theta > 0$ we have $\Theta_1 := \Theta^{1/2} > 0$. Therefore, by the definition of M in (4.1), we have

$$\begin{aligned} & \begin{bmatrix} \Theta_1(\xi) & & \\ & \ddots & \\ & & \Theta_1\left(\xi + \frac{2\pi(|d|-1)}{d}\right) \end{bmatrix}^{-1} M(\xi) \begin{bmatrix} \Theta_1(\xi) & & \\ & \ddots & \\ & & \Theta_1\left(\xi + \frac{2\pi(|d|-1)}{d}\right) \end{bmatrix}^{-1} \\ &= I_{|d|r} - \begin{bmatrix} (a\Theta_1^{-1})^*(\xi)\Theta_1(d\xi) \\ (a\Theta_1^{-1})^*\left(\xi + \frac{2\pi}{d}\right)\Theta_1(d\xi) \\ \vdots \\ (a\Theta_1^{-1})^*\left(\xi + \frac{2\pi(|d|-1)}{d}\right)\Theta_1(d\xi) \end{bmatrix} \begin{bmatrix} (a\Theta_1^{-1})^*(\xi)\Theta_1^*(d\xi) \\ (a\Theta_1^{-1})^*\left(\xi + \frac{2\pi}{d}\right)\Theta_1^*(d\xi) \\ \vdots \\ (a\Theta_1^{-1})^*\left(\xi + \frac{2\pi(|d|-1)}{d}\right)\Theta_1^*(d\xi) \end{bmatrix}^*. \end{aligned}$$

Hence,

$$\begin{aligned} & M(\xi) \geq 0 \\ & \iff I_{|d|r} - \begin{bmatrix} (a\Theta_1^{-1})^*(\xi)\Theta_1(d\xi) \\ (a\Theta_1^{-1})^*\left(\xi + \frac{2\pi}{d}\right)\Theta_1(d\xi) \\ \vdots \\ (a\Theta_1^{-1})^*\left(\xi + \frac{2\pi(|d|-1)}{d}\right)\Theta_1(d\xi) \end{bmatrix} \begin{bmatrix} (a\Theta_1^{-1})^*(\xi)\Theta_1^*(d\xi) \\ (a\Theta_1^{-1})^*\left(\xi + \frac{2\pi}{d}\right)\Theta_1^*(d\xi) \\ \vdots \\ (a\Theta_1^{-1})^*\left(\xi + \frac{2\pi(|d|-1)}{d}\right)\Theta_1^*(d\xi) \end{bmatrix}^* \geq 0. \end{aligned}$$

By Lemma 4.2, for all $\xi \in \mathbb{T}$, we have that

$$\begin{aligned} & M(\xi) \geq 0 \\ & \iff I_r - \begin{bmatrix} (a\Theta_1^{-1})^*(\xi)\Theta_1^*(d\xi) \\ (a\Theta_1^{-1})^*\left(\xi + \frac{2\pi}{d}\right)\Theta_1^*(d\xi) \\ \vdots \\ (a\Theta_1^{-1})^*\left(\xi + \frac{2\pi(|d|-1)}{d}\right)\Theta_1^*(d\xi) \end{bmatrix}^* \begin{bmatrix} (a\Theta_1^{-1})^*(\xi)\Theta_1(d\xi) \\ (a\Theta_1^{-1})^*\left(\xi + \frac{2\pi}{d}\right)\Theta_1(d\xi) \\ \vdots \\ (a\Theta_1^{-1})^*\left(\xi + \frac{2\pi(|d|-1)}{d}\right)\Theta_1(d\xi) \end{bmatrix} \geq 0. \end{aligned}$$

By $\Theta = \Theta_1^2$, we have

$$M(\xi) \geq 0 \iff I_r - \Theta_1(d\xi)T_a(\Theta^{-1})(d\xi)\Theta_1(d\xi) \geq 0 \iff \Theta^{-1}(d\xi) \geq T_a(\Theta^{-1})(d\xi),$$

i.e.,

$$M \geq 0 \iff \Theta^{-1} \geq T_a(\Theta^{-1}).$$

Thus by inequality (3.7), we know $M \geq 0$.

Secondly, we want to prove that we can derive a tight wavelet frame from the given Θ . As a special case of [12, Theorem 3.1], to derive a tight wavelet frame from the given Θ , we need to find $r \times r$ matrices $a^1(\xi), \dots, a^L(\xi)$ of trigonometric polynomials such that

$$\begin{bmatrix} a^1(\xi)^* & \dots & a^L(\xi)^* \\ a^1\left(\xi + \frac{2\pi}{d}\right)^* & \dots & a^L\left(\xi + \frac{2\pi}{d}\right)^* \\ \vdots & \ddots & \vdots \\ a^1\left(\xi + \frac{2\pi(|d|-1)}{d}\right)^* & \dots & a^L\left(\xi + \frac{2\pi(|d|-1)}{d}\right)^* \end{bmatrix} \begin{bmatrix} a^1(\xi) \\ a^2(\xi) \\ \vdots \\ a^L(\xi) \end{bmatrix} = M_1(\xi), \quad (4.2)$$

where

$$M_1(\xi) := \begin{bmatrix} \Theta(\xi) - a(\xi)^* \Theta(d\xi) a(\xi) \\ -a\left(\xi + \frac{2\pi}{d}\right)^* \Theta(d\xi) a(\xi) \\ \vdots \\ -a\left(\xi + \frac{2\pi(|d|-1)}{d}\right)^* \Theta(d\xi) a(\xi) \end{bmatrix}.$$

Recall M is defined as

$$M(\xi) = \begin{bmatrix} \Theta(\xi) & & \\ & \ddots & \\ & & \Theta\left(\xi + \frac{2\pi(|d|-1)}{d}\right) \end{bmatrix} - \begin{bmatrix} a(\xi)^* \\ a\left(\xi + \frac{2\pi}{d}\right)^* \\ \vdots \\ a\left(\xi + \frac{2\pi(|d|-1)}{d}\right)^* \end{bmatrix} \Theta(d\xi) \begin{bmatrix} a(\xi)^* \\ a\left(\xi + \frac{2\pi}{d}\right)^* \\ \vdots \\ a\left(\xi + \frac{2\pi(|d|-1)}{d}\right)^* \end{bmatrix}^*.$$

It is evident to see that (4.2) is equivalent to

$$A(\xi)^* A(\xi) = M(\xi), \quad (4.3)$$

where

$$A(\xi) := \begin{bmatrix} a^1(\xi) & \cdots & a^1\left(\xi + \frac{2\pi(|d|-1)}{d}\right) \\ \vdots & \ddots & \vdots \\ a^L(\xi) & \cdots & a^L\left(\xi + \frac{2\pi(|d|-1)}{d}\right) \end{bmatrix}.$$

Define

$$E(\xi) := \begin{bmatrix} I_r & e^{-i\xi} I_r & \cdots & e^{-i(|d|-1)\xi} I_r \\ I_r & e^{-i(\xi+2\pi/d)} I_r & \cdots & e^{-i(|d|-1)(\xi+2\pi/d)} I_r \\ \vdots & \vdots & \ddots & \vdots \\ I_r & e^{-i(\xi+2\pi(|d|-1)/d)} I_r & \cdots & e^{-i(|d|-1)(\xi+2\pi(|d|-1)/d)} I_r \end{bmatrix}.$$

To solve (4.3), using the polyphase decomposition, define

$$\tilde{A}(d\xi) := A(\xi)E(\xi)$$

and

$$\tilde{M}(d\xi) := E(\xi)^* M(\xi) E(\xi).$$

By direct calculation, we can verify that

$$M(\xi) \in \mathcal{P}(\mathbb{T})^{|d|r \times |d|r} \iff \tilde{M}(\xi) \in \mathcal{P}(\mathbb{T})^{|d|r \times |d|r}$$

and

$$a^1(\xi), \dots, a^L(\xi) \in \mathcal{P}(\mathbb{T})^{r \times r} \iff \tilde{A}(\xi) \in \mathcal{P}(\mathbb{T})^{Lr \times |d|r}.$$

It is evident to see that (4.3) is equivalent to

$$\tilde{A}(\xi)^* \tilde{A}(\xi) = \tilde{M}(\xi). \quad (4.4)$$

We are especially interested in the case $L = |d|$. In this case, $\tilde{A}(\xi)$ is a $|d|r \times |d|r$ matrix. Since $M(\xi) \in \mathcal{P}(\mathbb{T})^{|d|r \times |d|r}$, we know that $\tilde{M}(\xi) \in \mathcal{P}(\mathbb{T})^{|d|r \times |d|r}$. Moreover, by $M \geq 0$, we have $\tilde{M} \geq 0$. Hence by the matrix-valued Fejér-Riesz Lemma ([9], [15], [17], [25]), there exists $\tilde{A} \in \mathcal{P}(\mathbb{T})^{|d|r \times |d|r}$ such that $\tilde{A}(\xi)^* \tilde{A}(\xi) = \tilde{M}(\xi)$. Therefore we can obtain trigonometric polynomial matrices $a^1(\xi), \dots, a^{|d|}(\xi)$ by the relation $A(\xi) = \tilde{A}(d\xi)E(\xi)^{-1}$. Moreover, by the choice of $a^1, \dots, a^{|d|}$, we know that $a^1(\xi), \dots, a^{|d|}(\xi)$ are $r \times r$ matrices of trigonometric polynomials and satisfy (4.2). Hence by [12, Theorem 3.1] we can derive a tight wavelet frame with generators $\{\psi^1, \dots, \psi^{|d|}\}$ such that $\psi^1, \dots, \psi^{|d|}$ are $r \times 1$ function vectors and

$$\hat{\psi}^j(d\xi) = a^j(\xi)\hat{\phi}(\xi), \quad j = 1, \dots, |d|.$$

Finally we want to prove that $\psi^1, \dots, \psi^{|d|}$ all have vanishing moments of order m . By [12, Theorem 2.3] and a direct computation, we only need to prove that

$$(1 - e^{-i\xi})^m \mid a^j(\xi)_{k,1}, \quad k = 1, \dots, r; j = 1, \dots, |d|.$$

By (4.2), we have that

$$\begin{aligned} \sum_{j=1}^{|d|} \sum_{k=1}^r |a^j(\xi)_{k,1}|^2 &= [\Theta(\xi) - a(\xi)^* \Theta(d\xi) a(\xi)]_{1,1} \\ &= [\Phi^{-1}(\xi) + A_0^m(\xi) G(\xi) A_0^m(\xi)^* - a(\xi)^* \Phi^{-1}(d\xi) a(\xi) \\ &\quad - a(\xi)^* A_0^m(d\xi) G(d\xi) A_0^m(d\xi)^* a(\xi)]_{1,1} \\ &= [\Phi^{-1}(\xi) - a(\xi)^* \Phi^{-1}(d\xi) a(\xi)]_{1,1} + O(|\xi|^{2m}) \text{ as } \xi \rightarrow 0. \end{aligned}$$

By Lemma 2.1, we have

$$(1 - e^{-i\xi})^{2m} \mid [\Phi^{-1} - a(\xi)^* \Phi^{-1}(d\xi) a(\xi)]_{1,1}.$$

Hence

$$\sum_{j=1}^{|d|} \sum_{k=1}^r |a^j(\xi)_{k,1}|^2 = O(|\xi|^{2m}) \text{ as } \xi \rightarrow 0.$$

Therefore,

$$\sum_{j=1}^{|d|} \sum_{k=1}^r |a^j(\xi)_{k,1}/\xi^m|^2 < +\infty \text{ as } \xi \rightarrow 0.$$

Noticing the fact that the summation of nonnegative numbers is still nonnegative, we have

$$|a^j(\xi)_{k,1}/\xi^m| < +\infty \text{ as } \xi \rightarrow 0, \quad k = 1, \dots, r; j = 1, \dots, |d|.$$

Thus,

$$(1 - e^{-i\xi})^m \mid a^j(\xi)_{k,1}, \quad k = 1, \dots, r; j = 1, \dots, |d|.$$

Hence $\psi^1, \dots, \psi^{|d|}$ all have vanishing moments of order m . ■

5 An example

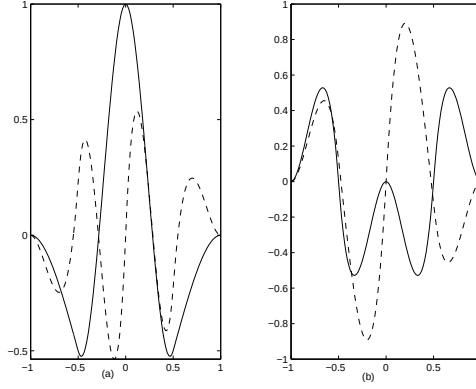


Figure 1: Generators for the tight multiwavelet frame in Example 5.1: (a) ψ^1 and ψ^2 (b) ψ^3 and ψ^4 . Functions ψ^1 , ψ^2 , ψ^3 and ψ^4 are either symmetric or antisymmetric about the origin and have vanishing moments of order 1.

Example 5.1 Define a function vector $\phi := (\phi_1, \phi_2)^T$ where

$$\phi_1 := (t+1)^2(-2t+1)\chi_{[-1,0)} + (t-1)^2(2t+1)\chi_{[0,1)}, \quad \phi_2 := t(t+1)^2\chi_{[-1,0)} + t(t-1)^2\chi_{[0,1)}.$$

The function vector ϕ is the well known Hermite cubics and it satisfies the following refinement equation

$$\phi = \begin{bmatrix} 1/2 & 3/4 \\ -1/8 & -1/8 \end{bmatrix} \phi(2 \cdot + 1) + \begin{bmatrix} 1 & 0 \\ 0 & 1/2 \end{bmatrix} \phi(2 \cdot) + \begin{bmatrix} 1/2 & -3/4 \\ 1/8 & -1/8 \end{bmatrix} \phi(2 \cdot - 1).$$

The refinable function vector ϕ has the following interpolation property: $\phi_1(0) = 1$, $\phi'_1(0) = 0$, $\phi_2(0) = 0$ and $\phi'_2(0) = 1$. For further discussion on the Hermite interpolants, please see [10] and the references therein.

The mask of the Hermite cubics ϕ is given by

$$a(\xi) := \begin{bmatrix} (e^{i\xi} + 2 + e^{-i\xi})/4 & 3(e^{i\xi} - e^{-i\xi})/8 \\ (-e^{i\xi} + e^{-i\xi})/16 & (-e^{i\xi} + 4 - e^{-i\xi})/16 \end{bmatrix}. \quad (5.1)$$

We can check out that a satisfies the sum rules of order 4 with a row vector $y(\xi) = [1, e^{i\xi}/3 + 1/2 - e^{-i\xi} + e^{-2i\xi}/6]$. Take $m = 1$. Define $U(\xi) := I_2$. Then $U(2\xi)a(\xi)U(\xi)^{-1}$ takes the form of (2.2) with $m = 1$. Define

$$\Theta(\xi) := \begin{bmatrix} 1 & 0 \\ 0 & 15 + 9\sqrt{2} \end{bmatrix}$$

and

$$a^1(\xi) := d_1 \begin{bmatrix} (2 - e^{-i\xi} - e^{i\xi})/4 & 3(e^{-i\xi} - e^{i\xi})/8 \\ (-29 + 16\sqrt{2})(e^{-i\xi} - e^{i\xi})/784 & (20 + 11e^{-i\xi} + 11e^{i\xi})/112 \end{bmatrix},$$

$$a^2(\xi) := d_2 \begin{bmatrix} 0 & (3\sqrt{6} + 4\sqrt{3})(e^{-i\xi} - e^{i\xi})/8 \\ (6 - 5\sqrt{2})(e^{-i\xi} - e^{i\xi})/196 & 3(2 - e^{-i\xi} - e^{i\xi})/28 \end{bmatrix},$$

where

$$d_1 := \text{diag} \left[1, \sqrt{105 + 63\sqrt{2}} \right], \quad d_2 := \text{diag} \left[1, \sqrt{70 + 42\sqrt{2}} \right].$$

Define functions ψ^1, ψ^2, ψ^3 and ψ^4 by

$$[\hat{\psi}^1(2\xi), \hat{\psi}^2(2\xi)]^T := a^1(\xi)\hat{\phi}(\xi), \quad [\hat{\psi}^3(2\xi), \hat{\psi}^4(2\xi)]^T := a^2(\xi)\hat{\phi}(\xi).$$

By a direct computation based on [12, Theorem 3.1], one can verify that $\{\psi^1, \psi^2, \psi^3, \psi^4\}$ generates a tight multiwavelet frame. Moreover, functions ψ^1, ψ^2, ψ^3 and ψ^4 are real-valued and symmetric, and all of them have vanishing moments of order 1. For their graphs, see Figure 1.

Remark: Recently, Chui and Stöckler ([4]) have constructed another tight multiwavelet frame derived from the Hermite cubics. Their example has five (anti)symmetric generators with vanishing moments of order 4.

Acknowledgements: The autor would like to thank the anonymous referees for their reports to improve the presentation of this paper.

References

- [1] C. Cabrelli, C. Heil, and U. Molter, Accuracy of lattice translates of several multidimensional refinable functions, *J. Approx. Theory* 95, 5–52, (1998).
- [2] C. K. Chui, W. He, and J. Stöckler, Compactly supported tight and sibling frames with maximum vanishing moments, *Appl. Comput. Harmon. Anal.* 13, 224–262, (2002).
- [3] C. K. Chui, W. He, J. Stöckler, and Q. Y. Sun, Compactly supported tight affine frames with integer dilations and maximum vanishing moments, *Adv. Comput. Math.* 18, 159–187, (2003).
- [4] C. K. Chui and J. Stöckler, Recent development of spline wavelet frames with compact support, preprint, 2002.
- [5] W. Dahmen and C. A. Micchelli, Biorthogonal wavelet expansions, *Constr. Approx.* 13, 293–328, (1997).
- [6] I. Daubechies, B. Han, A. Ron, and Z. W. Shen, Framelets: MRA-based constructions of wavelet frames, *Appl. Comput. Harmon. Anal.* 14, 1–46, (2003).

- [7] G. C. Donovan, J. S. Geronimo, and D. P. Hardin, Intertwining multiresolution analyses and the construction of piecewise-polynomial wavelets, *SIAM J. Math. Anal.* 27, 1791–1815, (1996).
- [8] G. C. Donovan, J. S. Geronimo, D. P. Hardin, and P. Massopust, Construction of orthogonal wavelets using fractal interpolation function, *SIAM J. Math. Anal.* 27, 1158–1192, (1996).
- [9] I. Gohberg, S. Goldberg, and M. A. Kaashoek, *Classes of Linear Operators*, Birkhauser Verlag, Boston, 1993.
- [10] B. Han, Approximation properties and construction of Hermite interpolants and biorthogonal multiwavelets, *J. Approx. Theory* 110, 18–53, (2001).
- [11] B. Han and R. Q. Jia, Multivariate refinement equations and convergence of subdivision schemes, *SIAM J. Math. Anal.* 29, 1177–1199, (1998).
- [12] B. Han and Q. Mo, Multiwavelet frames from refinable function vectors, *Adv. Comput. Math.* 18, 211–245, (2003).
- [13] B. Han and Q. Mo, Tight wavelet frames generated by three symmetric B-spline functions with high vanishing moments, *Proc. Amer. Math. Soc.* 132, 77–86, (2004).
- [14] B. Han and Q. Mo, Splitting a matrix of Laurent polynomials with symmetry and its application to symmetric framelet filter banks, *SIAM, J. Matrix Anal. Appl.*, to appear.
- [15] D. P. Hardin, T. A. Hogan, and Q. Y. Sun, The matrix-valued Riesz lemma and local orthonormal bases in shift-invariant spaces, preprint.
- [16] C. Heil, G. Strang, and V. Strela, Approximation by translates of refinable functions, *Numer. Math.* 73, 75–94, (1996).
- [17] H. Helson and D. Lowdenslager, Prediction theory and Fourier series in several variables, *Acta Math.* 99, 165–201, (1958).
- [18] R. Q. Jia and Q. Jiang, Approximation power of refinable vectors of functions, *Wavelet analysis and applications (Guangzhou, 1999)*, AMS/IP Stud. Adv. Math., 25, Amer. Math. Soc., Providence, RI, 2002, pp. 155–178.
- [19] R. Q. Jia and C. A. Micchelli, Using the refinement equations for the construction of pre-wavelets. II. Powers of two, *Curves and surfaces (Chamonix-Mont-Blanc, 1990)*, Academic Press, Boston, MA, 1991, pp. 209–246.

- [20] R. Q. Jia, S. Riemenschneider, and D. X. Zhou, Approximation by multiple refinable functions and multiple functions, *Canadian J. Math.* 49, 944–962, (1997).
- [21] Q. Jiang and Z. Shen, On the existence and weak stability of matrix refinable functions, *Constr. Approx.* 15, 337–353, (1999).
- [22] G. Plonka, Approximation order provided by refinable function vectors, *Constr. Approx.* 13, 221–244, (1997).
- [23] G. Plonka and A. Ron, A new factorization technique of the matrix mask of univariate refinable functions, *Numer. Math.* 87, 555–595, (2001).
- [24] W. C. Rheinboldt and J. S. Vandergraft, A Simple approach to the Perron-Frobenius theory for positive operators on general partially-ordered finite-dimensional linear spaces, *Math. Compu.* 121(27), 139–145, (1973).
- [25] M. Rosenblatt, A multi-dimensional prediction problem, *Arkiv Math.* 3, 407–424, (1958).
- [26] Z. Shen, Refinable function vectors, *SIAM J. Math. Anal.* 29, 235–250, (1998).
- [27] Q. Sun, Compactly supported distributional solutions of nonstationary nonhomogeneous refinement equation, *Acta Math. Sinica* 17, 1–14, (2001).
- [28] D. X. Zhou, Existence of multiple refinable distribution, *Michigan Math. J.* 44, 317–329, (1997).

Face-based Hermite Subdivision Schemes

Bin Han

Department of Mathematical and Statistical Sciences

University of Alberta

Edmonton, Alberta, Canada T6G 2G1

and

Thomas P.-Y. Yu

Department of Mathematical Science

Rensselaer Polytechnic Institute

Troy, New York 12180-3590

Abstract

Interpolatory and non-interpolatory multivariate Hermite type subdivision schemes are introduced in [8, 7]. In their applications in free-form surfaces, *symmetry* properties play a fundamental role: one can essentially argue that a subdivision scheme without a symmetry property simply *cannot* be used for the purpose of modelling free-form surfaces. The symmetry properties defined in the article [8] are formulated based on an underlying conception that Hermite data produced by the subdivision process is attached exactly to the *vertices* of the subsequently refined tessellations of the Euclidean space. As such, certain interesting possibilities of symmetric Hermite subdivision schemes are disallowed under our vertex-based symmetry definition. In this article, we formulate new symmetry conditions based on the conception that Hermite data produced in the subdivision process is attached to the *faces* instead of *vertices* of the subsequently refined tessellations. New examples of symmetric faced-based schemes are then constructed.

Similar to our earlier work in vertex-based interpolatory and non-interpolatory Hermite subdivision schemes, a key step in our analysis is that we make use of the *strong convergence theory* of refinement equation to convert a *prescribed geometric condition* on the subdivision scheme – namely, the subdivision scheme is of Hermite type – to an *algebraic condition* on the subdivision mask. Our quest for face-based schemes in this article leads also to a refined result in this direction.

Mathematics Subject Classification. 41A05, 41A15, 41A63, 42C40, 65T60, 65F15

Keywords. Refinable Function, Vector Refinability, Subdivision Scheme, Hermite Subdivision Scheme, Shift Invariant Subspace, Symmetry, Subdivision Surface, Spline

1 Introduction

Subdivision schemes are used in both curve and surface design as well as wavelet construction. Because of the two different applications, there comes also a certain degree of confusion in what subdivision is supposed to mean even in the most standard *linear, stationary and shift-invariant* case:

- **[Geo]** To a number of geometric modelling and computer graphics scientists, subdivision, in their so-called *regular* setting, typically consists of the following 4 components:

- [T1] a choice of isohedral tiling of $\mathbb{R}^s$;
- [T2] a *refinement rule* that is used to transform the tiling to a *similar* tiling but with a smaller tile size; this refinement rule is used repeatedly in the subdivision process to create finer and finer tilings of $\mathbb{R}^s$;
- [G1] a specification of *how* data (measuring geometric positions) is attached to the individual tiles; in 2-D this comes with the choice of attaching positional data to vertex and/or edge and/or face;
- [G2] a fixed linear rule for how data attached to any specific tile is determined from the data attached to the coarser scale tiles.

[T1]-[T2] are typically referred to as the “topological” part of the subdivision scheme. So far the main application of subdivision is in surface modelling, in which $s = 2$ and the isohedral tiling is usually based on equilateral triangles or squares.

(In the setting of free-form subdivision surfaces, the above is the so-called regular part of a subdivision scheme, one needs also extraordinary and boundary subdivision rules for producing a free-form surface. We do not discuss the latter in this article.)

- **[Wav]** To a number of mathematicians working in wavelet analysis, subdivision is defined as a linear operator $S := S_{\mathbf{a}, M}$ of the form:

$$Sv(\alpha) = \sum_{\beta \in \mathbb{Z}^s} v(\beta) \mathbf{a}(\alpha - M\beta), \quad (1.1)$$

where $\mathbf{a}$ is the mask and M is a so-called isotropic dilation matrix. See [8, 7] and the references therein. This operator has a very close connection to the refinement equation

$$\phi(x) = \sum_{\beta} \mathbf{a}(\beta) \phi(Mx - \beta), \quad (1.2)$$

which is directly related to wavelet construction under the MRA framework of Mallet and Meyer. Iteratively applying S to a sequence of initial data v produces the data $S^n v$, $n = 0, 1, 2, \dots$, with

$$(S^n v)(\alpha) \approx f(M^{-n}\alpha), \quad n \text{ large}, \quad (1.3)$$

where f is the limit function – exists for a “good” mask $\mathbf{a}$ – of the subdivision process.

The geometric entities that underly (1.3) are, of course, the successively refined *lattices*

$$\mathbb{Z}^s \subset M^{-1}\mathbb{Z}^s \subset M^{-2}\mathbb{Z}^s \subset \dots \subset \bigcup_{n=0}^{\infty} M^{-n}\mathbb{Z}^s \stackrel{\text{dense}}{\subset} \mathbb{R}^s. \quad (1.4)$$

Most subdivision schemes we are aware of can be described under both settings. However, it seems unclear whether the two settings are exactly the same in general. In (1.4) only discrete *points* are involved, these points are not connected and hence there is no edge or face, and consequently no concept of tiling/graph/tessellation is directly involved in setting **[Wav]**, so in this sense one may think that **[Wav]** is more general than **[Geo]**. On the other hand, the lattices in (1.4) are all isomorphic to $\mathbb{Z}^s$, whereas the isohedral tiling **[T1]** involved in **[Geo]** may have no structural similarity whatsoever with $\mathbb{Z}^s$ – so from this point of view **[Geo]** is perhaps more general than **[Wav]**.

At the analysis level, the setting **[Wav]** is very well-understood: there is by now an extensive collection of analytical and computational tools available for the analysis of (1.1)-(1.2).

In this paper, we propose face-based Hermite subdivision schemes which we first describe under the setting **[Geo]**, we then translate them back to the form (1.1) and analyze them using the general theory available. To this end, we reuse the main ideas *mutatis mutandis* developed in [8]; see Section 2. We follow closely the notations and vocabularies in [8].

We recall here a couple of notations from [8] that will be used very often in the rest of the paper: Λ_r is the set of s -tuples of non-negative integers with sum no greater than r , ordered by the lexicographic ordering. $\mathcal{S}(E, \Lambda_r)$ is the $\#\Lambda_r \times \#\Lambda_r$ matrix that measures how Hermite data of a function changes upon a linear change of variable by $E \in \mathbb{R}^{s \times s}$: let $\partial^{\leq r} f := [D^\nu f]_{\nu \in \Lambda_r} \in \mathbb{R}^{1 \times \#\Lambda_r}$, then

$$f \in C^r(\mathbb{R}^s), \quad g = f(E \cdot) \implies \partial^{\leq r} g = \partial^{\leq r} f(E \cdot) \mathcal{S}(E, \Lambda_r). \quad (1.5)$$

Why face-based Hermite schemes? Besides theoretical motivations, *some vertex-based schemes are simply unnecessarily smooth for our purposes*. We have been considering applications of symmetric Hermite subdivision schemes in free-form surfaces [13], in which we are primarily interested in schemes that are C^2 and have very small supports. For 2-D Hermite schemes with quincunx refinement, if we insist on using vertex-based scheme then the smallest meaningful support is $[-1, 1]^2$ and the resulted order 1 Hermite scheme is C^7 [8, Section 3.5]; in fact even the scalar counterpart of this scheme is much smoother than the desired C^2 : the *scalar* quincunx scheme used in [12] (derived from the Zwart-Powell box spline) has the same $[-1, 1]^2$ support, is a vertex-based scheme and is C^4 . If one wants to use the smaller support $[0, 1]^2$ and still obtain schemes

with a meaningful symmetry property then a natural approach is to change from vertex- to face-based scheme.

At the initial phase of the research work of this article, we speculated that there existed a C^2 face-based Hermite subdivision scheme with the mask support $[0, 1]^2$ (which corresponds to the stencil in Figure 2(a)) and is based on quincunx refinement, as one could already obtain a C^1 scheme in the scalar case. See Section 3.1. To our surprise, it does not seem to be case. Despite such a “bad news”, we shall discover in Section 3.2 a C^2 Hermite face-based subdivision scheme based on quadrissection refinement (i.e. $M = 2I_2$) and the subdivision stencil depicted in Figure 2(a)). The scalar counterpart of this scheme is the regular part of what Peters and Reif called “*the simplest subdivision scheme* (for smoothing polyhedra)” [11].

From experience it is almost always possible to gain smoothness by increasing support size or by increasing multiplicity – with our experience in Section 3.1 as a rare exception. For application in free-form subdivision surfaces, the latter approach may have an advantage because of the inevitable presence of *extraordinary vertices* in free-form surfaces.

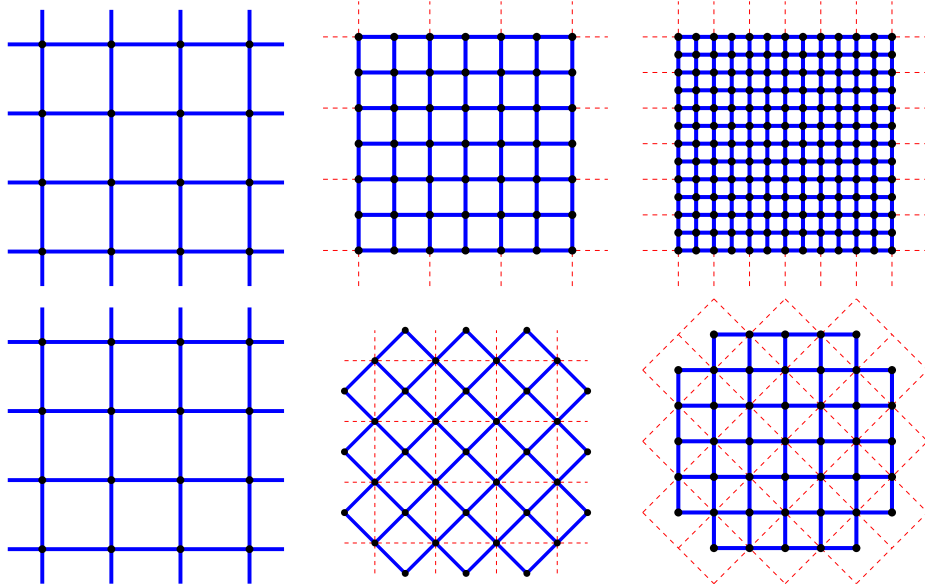


Figure 1: Two topological refinement schemes based on the square tiling

2 Face-based Hermite Subdivision Schemes

In this paper, we consider subdivision schemes based on the following specifications:

[T1] The isohedral tiling of $\mathbb{R}^s$ is chosen to be the straightforward tiling by hypercubes:

$$\mathbb{R}^s = \bigcup_{\alpha \in \mathbb{Z}^s} \alpha + [0, 1]^s.$$

[T2] Let M be an $s \times s$ integer matrix such that $M = UDU^T$ where U is unitary and D is diagonal with diagonal entries with the same modulus $\sigma > 1$. In the rest of the paper we simply call any such matrix a *dilation matrix*.¹ Then $M^{-1}[0, 1]^s$ is similar to $[0, 1]^s$. The successively refined tilings are then given by $\mathbb{R}^s = \bigcup_{\alpha \in \mathbb{Z}^s} M^{-n}(\alpha + [0, 1]^s)$. In dimension 2, two well-known examples of M are:

$$2I_2 = \begin{bmatrix} 2 & 0 \\ 0 & 2 \end{bmatrix}, \quad M_{\text{Quincunx}} = \begin{bmatrix} 1 & 1 \\ 1 & -1 \end{bmatrix}.$$

See Figure 1.

[G1] Subdivision data is a “jet” of Hermite data of order r at the **center** of each cube; conceptually, we have

$$\begin{aligned} \mathbb{R}^{1 \times \#\Lambda_r} \ni v_n(\alpha) &= \text{data attached to the tile } M^{-n}([0, 1]^s + \alpha) \\ &\approx \partial^{\leq r} f(M^{-n}(\alpha + [1/2, \dots, 1/2]^T)) \mathcal{S}(M^n, \Lambda_r). \end{aligned} \quad (2.1)$$

See [8] for the exact definition of the notations above; the vector on the right-hand side of (2.1) consists precisely of all the mixed directional derivatives of f of order up to r at the point $M^{-n}(\alpha + [\frac{1}{2}, \dots, \frac{1}{2}]^T)$ and in directions $M^{-n}e_j$, $j = 1, \dots, s$.

The informally described concept here will be made precise in definition 2.1 below.

[G2] With the choice of [T1], [T2], [G1] above, there are still many choices of subdivision rules. We are interested in those with small supports, smooth limits and *symmetry*, see the formal definitions below.

In each case above the subdivision scheme can be equivalently defined by a subdivision operator $S = S_{\mathbf{a}, M}$ of the form (1.1), i.e. there exists a subdivision operator S such that $v_n := S^n v_0$, $\forall n \geq 0$, is precisely what [T1]-[T2] & [G1]-[G2] above would produce. Notice that S is a so-called **stationary** subdivision operator — the subdivision data v_n is generated by a **fixed** subdivision mask $\mathbf{a}$ at all levels n ; we mention without a rigorous justification that if we want the data $v_n = S^n v_0$ produced by a stationary subdivision process to also enjoy a natural Hermite type convergence property (c.f. Definition 2.1 below), then the scaling by $\mathcal{S}(M^n, \Lambda_r)$ introduced in (2.1) is essentially the only sensible choice.

In 2-D, here are some subdivision stencils that we are particularly interested in:

1. When $M = M_{\text{Quincunx}}$, the Hermite data associated with any level $n + 1$ square is determined from a linear combination of the Hermite data associated with the two level n squares containing the level $n + 1$ square. See Figure 2(a). In this case, $\text{supp}(\mathbf{a}) = [0, 1] \times [-1, 0]$.

¹Notice that this is more restrictive than what is usually called an isotropic dilation matrix in the literature.

2. When $M = 2I_2$, the Hermite data associated with any level $n+1$ square is determined from a linear combination of the Hermite data associated with the 3 or 4 level n squares *closest* to the level $n+1$ square. Figure 2(b)&(c). In this case, $\text{supp}(\mathbf{a}) = [-1, 2]^2 - \{(-1, -1), (-1, 2), (2, -1), (2, 2)\}$ or $[-1, 2]^2$, respectively.

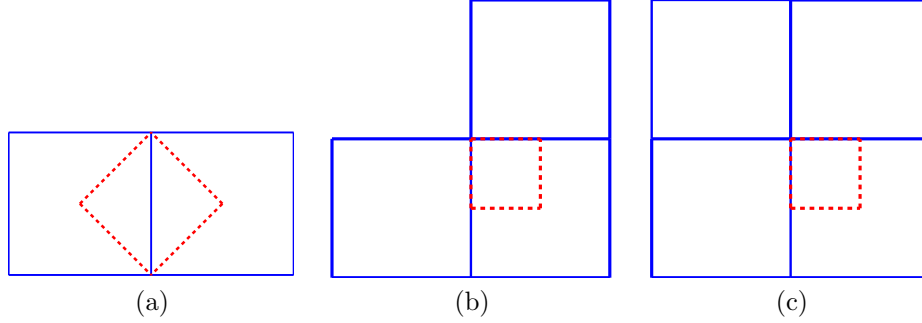


Figure 2: Subdivision Stencils for quincunx (a) and quadrissection (b)&(c) refinement

Following the geometric intuition we have so far, we are led to the following definition: Let M be a dilation matrix with the property described in [T2] above. An **order r face-based Hermite subdivision operator** $S := S_{\mathbf{a}, M}$ is a subdivision operator with multiplicity $m = \#\Lambda_r$ such that (i) for any initial sequence $v \in [l^0(\mathbb{Z}^s)]^{1 \times m}$ there exists $f_v \in C^r(\mathbb{R}^s)$ such that

$$\lim_{n \rightarrow \infty} \|[\partial^{\leq r} f_v]_{M^{-n}(\mathbb{Z}^s + [\frac{1}{2}, \dots, \frac{1}{2}]^T)} - v_n S(M^n, \Lambda_r)\|_{[l^\infty(\mathbb{Z}^s)]^{1 \times m}} = 0, \quad (2.2)$$

where $v_n = S^n v$, (ii) $f_v \neq 0$ for some $v \neq 0$. For notational convenience, we write

$$S^\infty v := f_v.$$

We now recall from [8] the following definition, proposed originally for the study of what we now call **vertex-based** Hermite subdivision scheme.

Definition 2.1 ([8, Definition 1.1]) A subdivision operator $S := S_{\mathbf{a}, M}$ is of **Hermite type** of order r if (i) $m = \#\Lambda_r$ for some $r \geq 0$, and for any initial sequence $v \in [l^0(\mathbb{Z}^s)]^{1 \times m}$ there exists $f_v \in C^r(\mathbb{R}^s)$ such that

$$\lim_{n \rightarrow \infty} \|[\partial^{\leq r} f_v]_{M^{-n}\mathbb{Z}^s} - v_n S(M^n, \Lambda_r)\|_{[l^\infty(\mathbb{Z}^s)]^{1 \times m}} = 0, \quad (2.3)$$

where $v_n = S^n v$, (ii) $f_v \neq 0$ for some $v \neq 0$.

The two definitions are clearly equivalent. Comparing (2.2) and (2.3), of course the only difference is that the latter involves sampling of Hermite data of at the lattices $M^{-n}\mathbb{Z}^s$, whereas the former involves sampling of Hermite data at the *shifted* lattices $M^{-n}(\mathbb{Z}^s + [\frac{1}{2}, \dots, \frac{1}{2}]^T)$. However, the two convergence

definitions involve n tending to infinity, and the difference induced by the shift in (2.2) becomes negligible for large n . An obvious application of triangle inequality, combined with the assumption that $f_v \in C^r$ in both (2.2) and (2.3), shows that the definition of order r face-based Hermite subdivision operator above is identical to Definition 2.1.

So what is the real difference between vertex-based and face-based subdivision schemes then? For the setup in this paper, we address only *sum-rule* and *symmetry* conditions, as these are the two conditions we use for determining a subdivision mask $\mathbf{a}$. We do not address, for instance, data-structure issues for computer implementation.

At first glance,

- Symmetry conditions must be different between vertex- and face-based schemes. As expounded in the introductory section of our earlier paper [8], the meaning of symmetry relies completely on the *geometric meaning* of subdivision data $v_n(\alpha)$; as the “geometric meaning” of $v_n(\alpha)$ has been changed from

$$v_n(\alpha) \approx \partial^{\leq r} f(M^{-n}\alpha) \mathcal{S}(M^n, \Lambda_r)$$

in vertex-based scheme to (2.1) in face-based scheme, it should be hardly surprising that symmetry properties for vertex- and face-based Hermite subdivision schemes have to be somewhat different.

As expected, this is exactly the case. See Section 2.2.

- Sum rule conditions, derived independently of symmetry conditions, are the same for vertex- and face-based Hermite subdivision schemes, because of the equivalence of (2.2) and (2.3).

As it turns out, this is *not* the case because of a technical loophole in a result in our earlier paper [8]. See Section 2.1.

2.1 Sum rules and the spectral quantity $\nu_\infty(a, M)$

Since a subdivision mask uniquely specifies a subdivision operator, **sum rules** – algebraic relations necessarily satisfied by the mask of any smooth subdivision process – are very useful for *construction* of subdivision schemes.

In the following, let us recall the definitions of sum rules in [2, 5] and the important spectral quantity $\nu_p(\mathbf{a}, M)$ in [5] in the setting of Hermite subdivision schemes.

For a given sequence $u \in (l^0(\mathbb{Z}^s))^{m \times n}$, its *Fourier series* $\hat{u}$ is defined to be

$$\hat{u}(\xi) := \sum_{\beta \in \mathbb{Z}^s} u(\beta) e^{-i\beta \cdot \xi}, \quad \xi \in \mathbb{R}^s.$$

Let $\mathbf{a}$ be a matrix mask with multiplicity m . We say that $\mathbf{a}$ satisfies the *sum rules* of order $k+1$ with respect to the dilation matrix M (see [5, Page 51]) if there exists a sequence $y \in (l^0(\mathbb{Z}^s))^{1 \times m}$ such that $\hat{y}(0) \neq 0$,

$$D^\mu[\hat{y}(M^T \cdot) \hat{\mathbf{a}}(\cdot)](0) = |\det M| D^\mu \hat{y}(0) \quad \text{and} \quad D^\mu[\hat{y}(M^T \cdot) \hat{\mathbf{a}}(\cdot)](2\pi\beta) = 0, \quad (2.4)$$

for all $|\mu| \leq k, \beta \in (M^T)^{-1}\mathbb{Z}^s \setminus \mathbb{Z}^s$.

The quantity $\nu_p(\mathbf{a}, M)$, which is defined in [5, Page 61], plays a very important role in the study of convergence of vector subdivision schemes and the characterization of smoothness of refinable function vectors. For the convenience of the reader, let us recall the definition of the quantity $\nu_p(\mathbf{a}, M)$ from [5, Page 61] as follows. The convolution of two sequences is defined to be

$$[u * v](\alpha) := \sum_{\beta \in \mathbb{Z}^s} u(\beta)v(\alpha - \beta), \quad u \in (l^0(\mathbb{Z}^s))^{m \times n}, v \in (l^0(\mathbb{Z}^s))^{n \times j}.$$

In terms of the Fourier series, we have $\widehat{u * v} = \hat{u}\hat{v}$. Let y be a sequence in $(l^0(\mathbb{Z}^s))^{1 \times m}$. We define the space $\mathcal{V}_{k,y}$ associated with the sequence y by

$$\mathcal{V}_{k,y} := \{v \in (l^0(\mathbb{Z}^s))^{m \times 1} : D^\mu[\hat{y}(\cdot)\hat{v}(\cdot)](0) = 0 \quad \forall |\mu| \leq k\}. \quad (2.5)$$

Let $1 \leq p \leq \infty$. For any $y \in (l^0(\mathbb{Z}^s))^{1 \times m}$ such that $\hat{y}(0) \neq 0$, we define

$$\rho_k(\mathbf{a}, M, p, y) := \sup \left\{ \lim_{n \rightarrow \infty} \|\mathbf{a}_n * v\|_{(\ell_p(\mathbb{Z}^s))^{m \times 1}}^{1/n} : v \in \mathcal{V}_{k,y} \right\}, \quad (2.6)$$

where $\mathbf{a}_n$ is defined to be $\hat{\mathbf{a}}_n(\xi) = \hat{\mathbf{a}}((M^T)^{n-1}\xi) \cdots \hat{\mathbf{a}}(M^T\xi)\hat{a}(\xi)$. Define

$$\rho(\mathbf{a}, M, p) := \inf \{ \rho_k(\mathbf{a}, M, p, y) : (2.4) \text{ holds for some } k \in \mathbb{N}_0$$

$$\text{and some } y \in (l^0(\mathbb{Z}^s))^{1 \times m} \text{ with } \hat{y}(0) \neq 0 \}.$$

We define the following quantity:

$$\nu_p(\mathbf{a}, M) := -\log_{\rho(M)} \left[|\det M|^{-1/p} \rho(\mathbf{a}, M, p) \right], \quad 1 \leq p \leq \infty, \quad (2.7)$$

where $\rho(M)$ denotes the spectral radius of the matrix M . The above quantity $\nu_p(\mathbf{a}, M)$ plays a key role in characterizing the convergence of a vector cascade algorithm in a Sobolev space and in characterizing the L_p smoothness of a refinable function vector. For example, it has showed in [5, Theorem 4.3] that a vector subdivision scheme associated with mask a and dilation matrix M converges in the Sobolev space $W_p^k(\mathbb{R}^s)$ if and only if $\nu_p(\mathbf{a}, M) > k$.

Since the convergence condition (2.2) for face-based Hermite subdivision operator turns out to be *exactly the same* as Definition 2.1 drawn from [8], in deriving sum rule conditions for face-based Hermite subdivision operator, one can presumably simply reuse the result from [8] pertaining to sum rule conditions, which we now recall:

Theorem 2.2 ([8, Theorem 2.2]) *Let M be an isotropic dilation matrix and $\mathbf{a}$ be a mask with multiplicity $m = \#\Lambda_r$. Suppose that $\nu_\infty(\mathbf{a}, M) > r$. Then $S_{\mathbf{a},M}$ is a subdivision operator of Hermite type of order r if $\mathbf{a}$ satisfies the sum rules of order $r+1$ with a sequence $y \in [l^0(\mathbb{Z}^s)]^{1 \times \#\Lambda_r}$ such that*

$$\frac{(-iD)^\mu}{\mu!} \hat{y}(0) = e_\mu^T, \quad \mu \in \Lambda_r. \quad (2.8)$$

At the time article [8] was written, the authors questioned whether the sufficient condition (2.8) is also necessary. If the answer of this open question were affirmative, then, in virtue of the equivalence of (2.2) and (2.3), sum rule conditions for vertex- and face-based Hermite subdivision schemes would be exactly the same. However, at the level of generality of Theorem 2.2, **(2.8) is not necessary for an order r Hermite type subdivision operator**: Theorem 2.3 below on the one hand generalizes Theorem 2.2 and on the other hand answers the above-mentioned open question *negatively*.

Recall that the cascade operator $Q_{\mathbf{a},M}$ associated with mask $\mathbf{a}$ and dilation matrix M is defined to be

$$Q_{\mathbf{a},M}f := \sum_{\beta \in \mathbb{Z}^s} \mathbf{a}(\beta) f(M \cdot -\beta).$$

A fixed point of this operator is a solution of the refinement equation (1.2). This observation is the starting point of the so-called *strong convergence theory* of refinement equation.

Theorem 2.3 *Let M be an isotropic dilation matrix and $\mathbf{a}$ be a mask with multiplicity $m = \#\Lambda_r$. Suppose that $\nu_\infty(\mathbf{a}, M) > r$. Then $S_{\mathbf{a},M}$ is a subdivision operator of Hermite type of order r if $\mathbf{a}$ satisfies the sum rules of order $r+1$ with a sequence $y \in [l^0(\mathbb{Z}^s)]^{1 \times \#\Lambda_r}$ such that*

$$\frac{(c - iD)^\mu}{\mu!} \hat{y}(0) = e_\mu^T, \quad \mu \in \Lambda_r \quad (2.9)$$

where $c \in \mathbb{R}^s$ is an arbitrary but fixed vector.

Proof: The proof follows essentially the same line of argument as in the proof of Theorem 2.2, the main new ingredient is the following observation: Let ψ be a Hermite interpolant of order r with accuracy order $r+1$, then

$$\psi_c := \psi(\cdot - c)$$

satisfies moment conditions with respect to a $y \in [l^0(\mathbb{Z}^s)]^{1 \times \#\Lambda_r}$ that satisfies (2.9).

By assumption, $\mathbf{a}$ satisfies sum rules of order $r+1$ with a y that also satisfies (2.9), and also that $\nu_\infty(\mathbf{a}, M) > r$, then the strong convergence theory of refinement equation [5, Theorem 4.3] says that

$$\lim_{n \rightarrow \infty} \|Q_{\mathbf{a},M}^n \psi_c - \phi\|_{[C^r(\mathbb{R}^s)]^{m \times 1}} = 0$$

for some $\phi \in [C^r(\mathbb{R}^s)]^{m \times 1}$; note that ϕ must be a solution of the refinement equation (1.2).

Now we show that $S_{\mathbf{a},M}$ satisfies the Hermite property in Definition 2.1. Let $v \in [l^0(\mathbb{Z}^s)]^{1 \times \#\Lambda_r}$, and $v_n = S^n v$. We can also write $v_n = \sum_{\beta} v(\beta) \mathbf{a}_n(\cdot - M^n \beta)$ where $\mathbf{a}_n = S^n(\delta I_{m \times m})$, $m = \#\Lambda_r$; on the other hand, we have

$$Q^n \psi_c = \sum_{\alpha} \mathbf{a}_n(\alpha) \psi_c(M^n \cdot -\alpha).$$

Then

$$\begin{aligned} f_n &:= \sum_{\alpha \in \mathbb{Z}^s} v_n(\alpha) \psi_c(M^n \cdot - \alpha) = \sum_{\alpha \in \mathbb{Z}^s} \sum_{\beta \in \mathbb{Z}^s} v(\beta) \mathbf{a}_n(\alpha - M^n \beta) \psi_c(M^n \cdot - \alpha) \\ &= \sum_{\beta \in \mathbb{Z}^s} v(\beta) (Q^n \psi_c)(\cdot - \beta). \end{aligned}$$

Therefore, if $f := \sum_{\alpha} v(\alpha) \phi(\cdot - \alpha)$, then $\lim_{n \rightarrow \infty} \|f_n - f\|_{C^r(\mathbb{R}^s)} = 0$ and $\lim_{n \rightarrow \infty} \|D^\mu f_n - D^\mu f\|_{L^\infty} = 0, \forall \mu \in \Lambda_r$.

If we denote by $\partial^{\leq r} \psi_c(x)$ the $\#\Lambda_r \times \#\Lambda_r$ matrix with the μ -th row equals to $\partial^{\leq r} (\psi_c)_\mu(x)$, then since ψ is a Hermite interpolant,

$$\partial^{\leq r} \psi_c(\alpha + c) = I_{\#\Lambda_r \times \#\Lambda_r} \delta(\alpha), \quad \forall \alpha \in \mathbb{Z}^s.$$

By (1.5), we have

$$\partial^{\leq r} f_n(M^{-n}(\alpha + c)) = \sum_{\beta \in \mathbb{Z}^s} v_n(\beta) (\partial^{\leq r} \psi_c)(\alpha + c - \beta) \mathcal{S}(M^n, \Lambda_r) = v_n(\alpha) \mathcal{S}(M^n, \Lambda_r).$$

But then

$$\begin{aligned} \max_{\alpha \in \mathbb{Z}^s} \|\partial^{\leq r} f(M^{-n}(\alpha - c)) - v_n(\alpha) \mathcal{S}(M^n, \Lambda_r)\|_\infty &= \max_{\alpha \in \mathbb{Z}^s} \|\partial^{\leq r} f(M^{-n}(\alpha - c)) - \partial^{\leq r} f_n(M^{-n}(\alpha - c))\|_\infty \\ &\leq \max_{\mu \in \Lambda_r} \|D^\mu f_n - D^\mu f\|_{L^\infty} \rightarrow 0, \text{ as } n \rightarrow \infty. \end{aligned}$$

So $\mathbf{a}$ satisfies condition (i) of Definition 2.1 in virtue of the fact that $f \in [C^r(\mathbb{R}^s)]^{m \times 1}$ and

$$\begin{aligned} \|\partial^{\leq r} f(M^{-n} \alpha) - v_n(\alpha) \mathcal{S}(M^n, \Lambda_r)\|_\infty &\leq \|\partial^{\leq r} f(M^{-n} \alpha) - \partial^{\leq r} f(M^{-n}(\alpha - c))\|_\infty \\ &\quad + \|\partial^{\leq r} f(M^{-n}(\alpha - c)) - v_n(\alpha) \mathcal{S}(M^n, \Lambda_r)\|_\infty \\ &\rightarrow 0 + 0, \text{ as } n \rightarrow \infty. \end{aligned}$$

The condition $\nu_\infty(\mathbf{a}, M) > r$ implies, by [5, Theorem 4.3], that $\text{span}\{\phi(\cdot - \beta) : \beta \in \mathbb{Z}^s\} \supseteq \Pi_r$, which implies $\phi \neq 0$. Thus condition (ii) of Definition 2.1 is also satisfied by $\mathbf{a}$. $\blacksquare$

Theorem 2.3 does not directly tell how to choose the shift vector c . From the proof, however, it seems the most reasonable to choose

$$c = [0, \dots, 0] \quad \text{and} \quad c = [1/2, \dots, 1/2]$$

for vertex and face-based schemes, respectively. We speculate that after one fixes the symmetry condition on mask $\mathbf{a}$ – the topic of the next section – then there corresponds a unique shift vector c such that condition (2.9) is both necessary and sufficient for $S_{\mathbf{a}, M}$ being a Hermite type subdivision operator.

2.2 Symmetry

Let G be a finite group of linear maps leaving the cube tiling of $\mathbb{R}^s$ invariant. For a technical reason, we assume also that G is compatible to dilation matrix M in the following sense ([3]):

$$MEM^{-1} \in G, \quad \forall E \in G. \quad (2.10)$$

In our bivariate examples, we use exclusively the following symmetry group:

$$D_4 = \left\{ \pm \begin{bmatrix} 1 & 0 \\ 0 & 1 \end{bmatrix}, \pm \begin{bmatrix} 0 & -1 \\ 1 & 0 \end{bmatrix}, \pm \begin{bmatrix} 1 & 0 \\ 0 & -1 \end{bmatrix}, \pm \begin{bmatrix} 0 & 1 \\ 1 & 0 \end{bmatrix} \right\}. \quad (2.11)$$

Note that D_4 is compatible to either $2I_2$ or M_{Quincunx} .

Definition 2.4 Let G be a symmetry group compatible with dilation matrix M . An order r Hermite subdivision operator has a face-based symmetry property with respect to G if the following condition is satisfied: $\forall f \in C^r(\mathbb{R}^s)$ and $E \in G$, if $v := [\partial^{\leq r} f]|_{M^{-n}(\mathbb{Z}^s + [\frac{1}{2}, \dots, \frac{1}{2}]^T)}$, $w := [\partial^{\leq r} g]|_{M^{-n}(\mathbb{Z}^s + [\frac{1}{2}, \dots, \frac{1}{2}]^T)}$ where $g := f(E \cdot)$, then $S^\infty w = (S^\infty v)(E \cdot)$.

Theorem 2.5 An order r Hermite subdivision operator $S_{\mathbf{a}, M}$ has a face-based symmetry property with respect to G if

1. 1 is a simple and dominant eigenvalue of the matrix

$$J_0 := \sum_{\beta \in \mathbb{Z}^s} a(\beta) / |\det M|$$

and the first entry of its nonzero eigenvector for the eigenvalue 1 is nonzero;

2. The following symmetry condition on the mask $\mathbf{a}$ holds

$$\mathbf{a}(E(\alpha - C_{\mathbf{a}}) + C_{\mathbf{a}}) = \mathcal{S}(M^{-1}EM, \Lambda_r) \mathbf{a}(\alpha) \mathcal{S}(E^{-1}, \Lambda_r), \quad \forall E \in G \quad (2.12)$$

where $e := [\frac{1}{2}, \dots, \frac{1}{2}]^T$ and

$$C_{\mathbf{a}} := (M - I_s) e. \quad (2.13)$$

Proof: If $\phi = [\phi_1, \dots, \phi_{\#\Lambda_r}]^T$ is the “impulse response” of $S_{\mathbf{a}, M}$, then ϕ satisfies the refinement equation (1.2) (that is, $Q_{\mathbf{a}, M}\phi = \phi$) and

$$S^\infty v = \sum_{\alpha \in \mathbb{Z}^s} v(\alpha) \phi(\cdot - \alpha).$$

We first show that the symmetry condition on $S_{\mathbf{a}, M}$ is implied by the following condition on ϕ :

$$\phi(x) = \mathcal{S}(E^{-1}, \Lambda_r) \phi(E(x - e) + e). \quad (2.14)$$

(The converse implication is also true and can be easily obtained by adapting part of the proof of [8, Proposition 2.3].)

Assume (2.14). Let f , E , g , v and w be as in Definition 2.4. Then

$$\begin{aligned} S^\infty w &= \sum_{\alpha} \partial^{\leq r} g(\alpha + e) \phi(\cdot - \alpha) \stackrel{(1.5)}{=} \sum_{\alpha} \partial^{\leq r} f(E(\alpha + e)) \mathcal{S}(E, \Lambda_r) \phi(\cdot - \alpha) \\ &= \sum_{\beta} \partial^{\leq r} f(\beta + e) \mathcal{S}(E, \Lambda_r) \phi(\cdot - (E^{-1}(\beta + e) - e)) \\ &\stackrel{(2.14)}{=} \sum_{\beta} \partial^{\leq r} f(\beta + e) \phi(E \cdot - \beta) = S^\infty v(E \cdot). \end{aligned}$$

Therefore, the order r Hermite subdivision operator $S_{\mathbf{a}, M}$ has a face-based symmetry property with respect to G .

Below we show that (2.14) follows from the assumptions of the theorem; the argument is essentially a time-domain version of the frequency-domain proof of [4, Proposition 2.1].

Let $\phi_E := S(E^{-1}, \Lambda_r) \phi(E(\cdot - e) + e)$, $E \in G$. It follows from (2.12) and the definition of $\phi_{MEM^{-1}}$ that

$$\begin{aligned} Q_{\mathbf{a}, M} \phi_{MEM^{-1}} &= \sum_{\alpha \in \mathbb{Z}^s} \mathbf{a}(\alpha) \phi_{MEM^{-1}}(M \cdot -\alpha) \\ &= S(E^{-1}, \Lambda_r) \sum_{\alpha \in \mathbb{Z}^s} S(E, \Lambda_r) \mathbf{a}(\alpha) S(ME^{-1}M^{-1}, \Lambda_r) \phi(MEM^{-1}(M \cdot -\alpha - e) + e) \\ &= S(E^{-1}, \Lambda_r) \sum_{\alpha \in \mathbb{Z}^s} \mathbf{a}(MEM^{-1}(\alpha - C_{\mathbf{a}}) + C_{\mathbf{a}}) \phi(ME \cdot -MEM^{-1}\alpha - MEM^{-1}e + e) \\ &= S(E^{-1}, \Lambda_r) \sum_{\alpha \in \mathbb{Z}^s} \mathbf{a}(\alpha) \phi(ME \cdot -\alpha + C_{\mathbf{a}} - MEM^{-1}C_{\mathbf{a}} - MEM^{-1}e + e) \\ &= S(E^{-1}, \Lambda_r) \sum_{\alpha \in \mathbb{Z}^s} \mathbf{a}(\alpha) \phi(M(E \cdot + M^{-1}C_{\mathbf{a}} - EM^{-1}C_{\mathbf{a}} - EM^{-1}e + M^{-1}e) - \alpha) \\ &= S(E^{-1}, \Lambda_r) \phi(E \cdot + M^{-1}C_{\mathbf{a}} - EM^{-1}C_{\mathbf{a}} - EM^{-1}e + M^{-1}e). \end{aligned}$$

By $C_{\mathbf{a}} = (M - I_s)e$, we deduce that

$$\begin{aligned} M^{-1}C_{\mathbf{a}} - EM^{-1}C_{\mathbf{a}} - EM^{-1}e + M^{-1}e &= (I_s - E)M^{-1}C_{\mathbf{a}} - EM^{-1}e + M^{-1}e \\ &= (I_s - E)M^{-1}(M - I_s)e - EM^{-1}e + M^{-1}e \\ &= (I_s - E)(I_s - M^{-1})e - EM^{-1}e + M^{-1}e \\ &= e - Ee - M^{-1}e + EM^{-1}e - EM^{-1}e + M^{-1}e \\ &= e - Ee. \end{aligned}$$

We conclude that

$$Q_{\mathbf{a}, M} \phi_{MEM^{-1}} = S(E^{-1}, \Lambda_r) \phi(E(\cdot - e) + e) = \phi_E \quad \forall E \in G.$$

In other words, we have $Q_{\mathbf{a}, M} \phi_E = \phi_{M^{-1}EM}$ for all $E \in G$. Therefore, $Q_{\mathbf{a}, M}^n \phi_E = \phi_{M^{-n}EM^n}$ for all $n \in \mathbb{N}$ and $E \in G$.

Since G is a finite group and $M^{-n}EM^n \in G$ for all $n \in \mathbb{N}$, there must exist a positive integer ℓ such that $M^{-\ell}EM^\ell = E$. Consequently, we have $Q_{a,M}^\ell \phi_E = \phi_{M^{-\ell}EM^\ell} = \phi_E$ for some positive integer ℓ . Since 1 is a simple dominant eigenvalue of the matrix J_0 , the same can be said to J_0^n for all $n \in \mathbb{N}$. It is known ([10] and references therein) that if J_0^n has 1 as a simple dominant eigenvalue, then ϕ is the unique solution, up to a scalar multiplicative constant, to the refinement equation $Q_{a,M}^n \phi = \phi$.

On the other hand, it follows from the refinement equation that $\hat{\phi}(0)$ and $\widehat{\phi_E}(0)$ are eigenvectors of the matrices J_0 and J_0^ℓ , respectively. Since 1 is a simple eigenvalue of J_0^ℓ , we must have $\widehat{\phi_E}(0) = c\hat{\phi}(0)$ for some complex number $c \in \mathbb{C}$ since $J_0^\ell \widehat{\phi_E}(0) = \widehat{\phi_E}(0)$ and $J_0^\ell \hat{\phi}(0) = \hat{\phi}(0)$. By the definition of ϕ_E , we have $\widehat{\phi_E}(0) = S(E^{-1}, \Lambda_r) \hat{\phi}(0)$ by $|\det E| = 1$. By our assumption, the first entry in the vector $\hat{\phi}(0)$ is nonzero; that is, $e_1^T \hat{\phi}(0) \neq 0$. Note that the first row of the matrix $S(E^{-1}, \Lambda_r)$ is e_1^T . We see that $e_1^T \widehat{\phi_E}(0) = e_1^T \hat{\phi}(0) \neq 0$. Therefore, it follows from $\widehat{\phi_E}(0) = c\hat{\phi}(0)$ that c must be 1. Hence, we conclude that we must have $\phi_E = \phi$ by $Q_{a,M}^\ell \phi_E = \phi_E$ and $\widehat{\phi_E}(0) = \hat{\phi}(0)$. In other words, (2.14) holds. ■

3 Examples

In this section, we explore examples in the three cases depicted in Figure 2. In each case, we are interested in bivariate Hermite schemes of order $r = 1$, and we reuse exactly the computational framework developed in [8, Section 3] and the associated MAPLE based solver together with the smoothness optimization code developed in [6]. Underlying this smoothness optimization code is a method by Jia and Jiang [9] which gives the critical L^2 regularity of the refinable function vector ϕ associated with a subdivision mask in the case when ϕ has stable shifts, and a lower bound for the critical L^2 regularity in the absence of stability.

The definition of sum rules, which is given in (2.4) from the frequency domain, can be equivalently rewritten in the time domain as follows: There exists a set of $m \times 1$ vectors $\{Y_\mu : \mu \in \mathbb{N}_0^s, |\mu| \leq k\}$ with $Y_0 \neq 0$ (see [2, Page 22] and [8, Equation (3.9), Section 3]) such that

$$\sum_{0 \leq \nu \leq \mu} (-1)^{|\nu|} Y_{\mu-\nu} J_\alpha^\mathbf{a}(\nu) = \sum_{\nu \in \Lambda_k} \mathcal{S}(M^{-1}, \Lambda_k)_{\mu,\nu} Y_\nu \quad \forall \mu \in \Lambda_k, \alpha \in \mathbb{Z}^s,$$

where $(\nu_1, \dots, \nu_s) \leq (\mu_1, \dots, \mu_s)$ means $\nu_j \leq \mu_j$ for all $j = 1, \dots, s$, and $J_\alpha^\mathbf{a}(\nu) := \sum_{\beta \in \mathbb{Z}^s} \mathbf{a}(\alpha + M\beta)(\beta + M^{-1}\alpha)^\nu / \nu!$ with $(\xi_1, \dots, \xi_s)^{(\nu_1, \dots, \nu_s)} := \xi_1^{\nu_1} \dots \xi_s^{\nu_s}$. The precise relation between the above definition and the definition of sum rules in (2.4) is

$$Y_\mu = \frac{(-iD)^\mu \hat{y}(0)}{\mu!}.$$

For face-based Hermite subdivision schemes of order $r = 1$ in dimension $s = 2$:

- The sum rules are now with respect to a y that satisfies (2.8) with $c = [\frac{1}{2}, \frac{1}{2}]^T$. For $s = 2$, $r = 1$ and $c = [\frac{1}{2}, \frac{1}{2}]^T$, (2.8) is equivalent to

$$Y_{(0,0)} = [1, 0, 0], \quad Y_{(1,0)} = [1/2, 1, 0], \quad Y_{(0,1)} = [1/2, 0, 1].$$

- Symmetry conditions are those in (2.12)-(2.13).

3.1 $M = M_{\text{Quincunx}}$, $\text{supp}(\mathbf{a}) = [0, 1] \times [-1, 0]$, $G = D_4$

The highest order of sum rule achievable in this case is 3, and the mask we found by our solver has one degree of freedom:

$$\mathbf{a}(0, 0) = \begin{bmatrix} \frac{1}{2} & \frac{-4t+1}{2} & \frac{4t-1}{2} \\ 0 & \frac{1}{4} & \frac{1}{4} \\ -\frac{1}{8} & t & -t \end{bmatrix}$$

and $\mathbf{a}(0, -1)$, $\mathbf{a}(1, -1)$, $\mathbf{a}(1, 0)$ are given by symmetry conditions (2.12)-(2.13). We found, first by empirical observation and then by an explicit calculation, that when $t = -1/4$, the refinable function vector has L^2 smoothness equals to 2.5 and consists of piecewise quadratic C^1 spline functions. By optimizing the L^2 smoothness over the parameter t , we found that when $t = -0.1875$, the L^2 smoothness of the subdivision scheme is 2.57793, only slightly higher than that of the spline scheme.

From a smoothness point of view, one gains almost nothing by going from scalar scheme $r = 0$ to Hermite scheme $r = 1$ – this is atypical when compared to all the examples obtained in [7, 6, 8]. When we choose $r = 0$ with the support size and symmetry group unchanged, the only mask that satisfies the first sum rule is

$$\mathbf{a}(0, 0) = \mathbf{a}(1, 0) = \mathbf{a}(0, -1) = \mathbf{a}(1, -1) = 1/2, \quad (3.1)$$

its refinable function is a box-spline by Zwart and Powell which is a piecewise quadratic C^1 spline function – same as what the above vector scheme gives when $t = -1/4$. A major difference is that the ZP element has unstable shifts whereas, from various observations, we believe that the refinable function vector of the above spline scheme has stable shifts.

Despite the bad news, we seem to be saved by the good news below.

3.2 $M = 2I_2$, $\text{supp}(\mathbf{a}) \subseteq [-1, 2]^2$, $G = D_4$

While (the regular part of) Peters and Reif's *mid-edge scheme* [11] is essentially the quincunx subdivision scheme (3.1), their scheme actually operates based on quadrissection refinement instead of quincunx refinement; this is made possible by the following observation: notice that $M_{\text{Quincunx}}^2 = 2I_2$, so if $\mathbf{b} = S_{\mathbf{a}, M_{\text{Quincunx}}}^2 \delta$, then, in our notations, we have

$$S_{\mathbf{a}, M_{\text{Quincunx}}}^2 = S_{\mathbf{b}, 2I_2}.$$

Notice also that $\text{supp}(\mathbf{b}) = [-1, 2]^2 - \{(-1, -1), (-1, 2), (2, -1), (2, 2)\}$. This scheme is C^1 but not C^2 . Here, we construct a Hermite version of $S_{\mathbf{b}, 2I_2}$ which turns out to be C^2 .

Using our solver, we found the following 3-parameter mask with the same support as $\mathbf{b}$ above which satisfies sum rules of order 4:

$$\mathbf{a}(0, 0) = \begin{bmatrix} 3/4 - 4c_3 & -c_1 + 2c_2 & -c_1 + 2c_2 \\ -3/32 + c_3 & c_1/2 + 5/16 & -c_2 - 1/16 \\ -3/32 + c_3 & -c_2 - 1/16 & c_1/2 + 5/16 \end{bmatrix},$$

$$\mathbf{a}(2, 0) = \begin{bmatrix} 1/8 + 2c_3 & c_1 & -2c_2 \\ c_3 & 1/8 + c_1/2 & -c_2 \\ -1/32 & 1/16 & 1/16 \end{bmatrix} \quad (3.2)$$

with the other entries given by symmetry conditions. By our smoothness optimization code, the L^2 smoothness of this scheme occurs to be the highest when $(c_1, c_2, c_3) \approx (-7/16, -3/32, 3/64)$, at which the L^2 smoothness is 3.5, thus the Hölder smoothness is at least 2.5, meaning that the scheme is C^2 . The associated refinable function vector ϕ is depicted in Figure 3.

support	order of sum rules	# free parameters	highest L^2 Smoothness
12 points	4 (highest possible)	3	3.5
16 points	4	7	3.5
16 points	5 (highest possible)	2	3.0

Table 1: Sum Rules and smoothness attained by some small support face-based Hermite subdivision schemes with $M = 2I_2$

The now-classical Doo-Sabin scheme [1] is based on the tensor product quadratic B-spline, which the latter can be viewed as a face-based scalar subdivision scheme with the 16 point support $[-1, 2]^2$. We have explored symmetric face-based Hermite schemes with the same support, which corresponds to the stencil in Figure 2(c). A higher order of sum rules – but not higher smoothness – can be achieved when compared to the 12 point support. See Table 1.

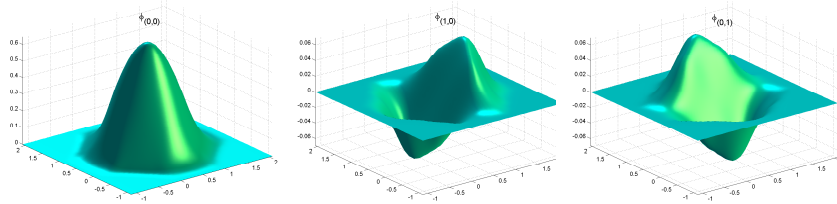


Figure 3: Refinable $\phi = [\phi_{(0,0)}, \phi_{(1,0)}, \phi_{(0,1)}]^T$ associated with the subdivision mask in (3.2)

Acknowledgements. Thomas Yu would like to thank Yonggang Xue for technical assistance and Peter Oswald for helpful discussions. Bin Han research was supported in part by NSERC Canada under Grant G121210654. Thomas Yu's research was supported in part by an NSF CAREER Award (CCR 9984501.)

References

- [1] D. Doo and M. Sabin. Analysis of the behaviour of recursive division surfaces near extraordinary points. *Comp. Aid. Geom. Des.*, 10:6 (1978), 356–360.
- [2] B. Han. Approximation properties and construction of Hermite interpolants and biorthogonal multiwavelets. *J. Approx. Theory*, 110:1 (2001), 18–53.
- [3] B. Han. Computing the smoothness exponent of a symmetric multivariate refinable function. *SIAM J. Matrix Anal. and Appl.*, 24:3 (2003), 693–714.
- [4] B. Han. Symmetric multivariate orthogonal refinable functions. *Applied and Computational Harmonic Analysis*, 17 (2003), 277–292.
- [5] B. Han. Vector cascade algorithms and refinable function vectors in Sobolev spaces. *Journal of Approximation Theory*, 124:1 (2003), 44–88.
- [6] B. Han, M. Overton, and T. P.-Y. Yu. Design of Hermite subdivision schemes aided by spectral radius optimization. *SIAM Journal on Scientific Computing*, 25:2 (2003), 643–656.
- [7] B. Han, T. P.-Y. Yu, and B. Piper. Multivariate refinable Hermite interpolants. *Mathematics of Computation*, 73 (2004), 1913–1935.
- [8] B. Han, T. P.-Y. Yu, and Y. Xue. Non-interpolatory Hermite subdivision schemes. *Mathematics of Computation* 74 (2005), 1345–1367.
- [9] R. Q. Jia and Q. T. Jiang. Spectral analysis of the transition operator and its applications to smoothness analysis of wavelets. *SIAM J. Matrix Anal. and Appl.*, 24:4 (2003), 1071–1109.
- [10] Q. Jiang and Z.W. Shen. On existence and weak stability of matrix refinable functions. *Constr. Approx.*, 15 (1999), 337–353.
- [11] J. Peters and U. Reif. The simplest subdivision scheme for smoothing polyhedra. *ACM Transactions on Graphics*, 16:4 (October 1997), 420–431.
- [12] L. Velho and D. Zorin. 4-8 subdivision. *Comp. Aid. Geom. Des.*, 18:5 (2001), 397–427.
- [13] Y. Xue, T. P.-Y. Yu, and T. Duchamp. Hermite subdivision surfaces: Applications of vector refinability to free-form surfaces. In preparation, 2003.

On the Construction of Linear Prewavelets Over a Regular Triangulation

Don Hong and Qingbo Xue
Department of Mathematical Sciences
Middle Tennessee State University
Murfreesboro, TN 37614, USA

Abstract

In this article, all the possible semi-prewavelets over uniform refinements of a regular triangulation are constructed. A corresponding theorem is given to ensure the linear independence of a set of different pre-wavelets obtained by summing pairs of these semi-prewavelets. This provides efficient multiresolution decompositions of the spaces of functions over various regular triangulation domains since the bases of the orthogonal complements of the coarse spaces can be constructed very easily.

AMS 2000 Classifications: 41A15, 41A63, 65D07, 68U05.

Key Words: Multiresolution analysis, pre-wavelets, regular triangulations, splines.

1 Introduction

Piecewise linear prewavelets with small support are useful tools in approximation theory and in the numerical solution of partial differential equations as well as in the computer graphics and practical largescale data representations. Basically speaking, a multiresolution is a decomposition of a function space into mutually orthogonal subspaces, each of which is endowed with a basis. The basis functions of each subspace are called wavelets if they are mutually orthogonal and prewavelets otherwise. The subspaces are called wavelet spaces and prewavelet spaces accordingly.

Kotyczka and Oswald [10] constructed piecewise linear prewavelets with small support in 1995. Floater and Quak [7] published their results on piecewise linear prewavelets with small support on arbitrary triangulations in 1999. Later on, they simplified the above results by introducing the idea of semi-wavelets, which can be used to construct wavelets (these semi-wavelets and wavelets are actually semi-prewavelets and prewavelets). Using this idea, Floater and Quak investigated the Type-1 triangulation in [8] and Type-2 in [7] respectively. Hong

and Mu [6] have discussed the piecewise linear prewavelets with minimal support over Type-1 triangulation. Some recent results on piecewise linear pre-wavelets and orthogonal wavelets can be found in [1] and [4]. Very recently, a hierarchical basis for C^1 cubic bivariate splines is used for surface compression in [5].

In this article, we study all the possible semi-prewavelets over uniform refinements of a regular triangulation. A corresponding theorem is given to ensure the linear independence of any set of different pre-wavelets obtained by summing pairs of these semi-prewavelets. This means that the multiresolutions of the linear function spaces over various regular triangulation domains can be done conveniently, since the bases of the orthogonal complements of the coarse spaces can be constructed very easily.

If not clearly claimed, we shall only consider the dyadic refinement. The reader will find that most symbols in this paper have commonly been used in related monographs. The paper is organized as follows. Preliminaries are introduced in Section 2. In Section 3, we construct the semi-prewavelets and obtain the main theorem on the construction of prewavelets over a regular triangulation

2 Multiresolution of Linear Spline Spaces over r-Triangulations

From a mathematical point of view, a computer graphic is nothing else but a function defined on a given region. On the other hand, a graphic on a domain Ω can be represented by functions in different level of function spaces S^j ($j = 0, 1, 2, \dots$). The difference between them is that the function from a fine space gives more detail of the original graphic than the one from a coarse space does. In an ideal situation, we can easily “switch” a function from one space into another when it is necessary. The key to choosing another function space is to use a different basis of functions. Surprisingly, in a multivariate setting, the relation between the bases of these nested spaces makes it difficult to do so. Even the case of continues piecewise linear wavelets construction is unexpectedly complicated, see [3]–[4] and [6]–[10] and the references therein.

In the following we shall discuss the multiresolution of the linear spline function space defined on any r-triangulation.

Definition 2.1 *A set of triangles $\mathcal{T} = \{T_1, \dots, T_M\}$ is called a triangulation of some subset Ω of $\mathbb{R}^2$ if $\Omega = \cup_{i=1}^M T_i$ and*

- (i) $T_i \cap T_j$ is either empty or a common vertex or a common edge, $i \neq j$,
- (ii) the number of boundary edges incident on a boundary vertex is two,
- (iii) Ω is simply connected.

A triangulation $\mathcal{T}^0 = \{T_1, \dots, T_M\}$ over some domain Ω is a regular triangulation or simply an r-triangulation if all the elements of $\mathcal{T}^0$ are equilateral triangles.

Figure 1 gives an example of an r-triangulation over a triangle shaped region. We denote by V the set of all vertices $v \in \mathbb{R}^2$ of triangles in $\mathcal{T}$ and by E the set

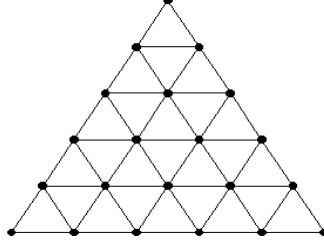


Figure 1: An r-triangulation

of all edges $e = [v, w]$ of triangles in $\mathcal{T}$. By a boundary vertex or boundary edge we mean a vertex or edge contained in the boundary of Ω . All other vertices and edges will be called interior vertices and interior edges. A boundary edge belongs to only one triangle, and an interior edge to two.

For a vertex $v \in V$, the set of neighbours of v in V is

$$V_v = \{w \in V : [v, w] \in E\}.$$

Suppose $\mathcal{T}$ is a triangulation. Given data values $f_v \in \mathbb{R}$ for $v \in V$, there is a unique function $f : \Omega \rightarrow \mathbb{R}$ which is linear on each triangle in $\mathcal{T}$ and interpolates the data: $f(v) = f_v$, $v \in V$. The function f is piecewise linear and the set of all such f constitute a linear space S with dimension $|V|$. For each $v \in V$, let $\phi_v : \Omega \rightarrow \mathbb{R}$ be the unique ‘hat’ or nodal function in S such that for all $w \in V$,

$$\phi_v(w) = \begin{cases} 1, & w = v; \\ 0, & \text{otherwise.} \end{cases}$$

The set of functions $\Phi = \{\phi_v\}_{v \in V}$ is a basis for the space S and for any function $f \in S$,

$$f(x) = \sum_{v \in V} f(v) \phi_v(x), \quad x \in \Omega. \quad (2.1)$$

The support of ϕ_v is the union of all triangles which contain v :

$$\text{supp}(\phi_v) := \cup_{v \in T \in \mathcal{T}} T.$$

For a given triangulation $\mathcal{T}^0 = \{T_1, \dots, T_M\}$, a **refined triangulation** is a triangulation $\mathcal{T}^1$ such that every triangle in $\mathcal{T}^0$ is the union of some triangles in $\mathcal{T}^1$. The result of this process is called a refinement of $\mathcal{T}^0$.

Obviously, there are various kinds of refinements. If not clearly claimed, we shall only consider the following uniform or dyadic refinement. We shall use

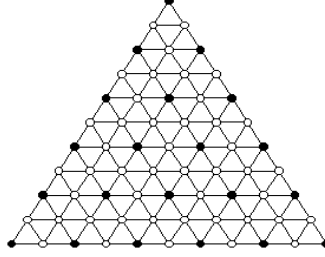


Figure 2: The first refinement of the r-triangulation in Figure 1

$[u, v]$ to denote the edge incident to two vertices u, v . A triangle with vertices u, v, w will be denoted as $[u, v, w]$.

For a given triangle $T = [x_1, x_2, x_3]$, let $y_1 = (x_2 + x_3)/2$, $y_2 = (x_1 + x_3)/2$, and $y_3 = (x_1 + x_2)/2$ denote the midpoints of its edges. Then the set of four triangles

$$\mathcal{T}_T = \{[x_1, y_2, y_3], [y_1, x_2, y_3], [y_1, y_2, x_3], [y_1, y_2, y_3]\}$$

is a triangulation and a refinement of the coarse triangle T . The set of triangles $\mathcal{T}^1 = \cup_{T \in \mathcal{T}^0} \mathcal{T}_T$ is evidently a triangulation and a refinement of $\mathcal{T}^0$. Similarly, a whole sequence of triangulations $\mathcal{T}^j, j = 0, 1, \dots$, can be generated by further refinement steps. Let V^j be the set of vertices in $\mathcal{T}^j$, and define $E^j, S^j, \phi_v^j, V_v^j$, and $\text{supp}(\phi_v^j)$ accordingly. A straightforward calculation shows that

$$\phi_v^{j-1} = \phi_v^j + \frac{1}{2} \sum_{w \in V_v^j} \phi_w^j, \quad v \in V^{j-1}, \quad (2.2)$$

and therefore we obtain a nested sequence of spaces

$$S^0 \subset S^1 \subset S^2 \subset \dots \quad (2.3)$$

Clearly, a refinement $\mathcal{T}^1$ of $\mathcal{T}^0$ is still an r-triangulation (see Figure 2). Continuing the refinement process on Ω leads us to a nested sequence of spaces of linear splines defined on the domain Ω with r-triangulations.

As usual, we use the following standard definition of the inner product of two continuous functions on $\mathcal{T}^0$,

$$\langle f, g \rangle = \sum_{T \in \mathcal{T}^0} \frac{1}{a(T)} \int_T f(x)g(x) dx, \quad f, g \in C(\Omega),$$

where $a(T)$ is the area of the triangle T .

Let c be the area of any triangle in the r -triangulation $\mathcal{T}^0$. Since all the triangles are congruent, the inner product reduces to the scaled L_2 inner product

$$\langle f, g \rangle = \frac{1}{c} \int_{\Omega} f(x)g(x) dx. \quad (2.4)$$

With this inner-product, the spaces S^j become inner-product spaces. Let W^{j-1} denote the relative orthogonal complement of the coarse space S^{j-1} in the fine space S^j , so that

$$S^j = S^{j-1} \oplus W^{j-1}. \quad (2.5)$$

We have the following decomposition:

$$S^n = S^0 \oplus W^0 \oplus W^1 \oplus \dots \oplus W^{n-1} \quad (2.6)$$

and the dimension of W^{j-1} is $|V^j| - |V^{j-1}| = |E^{j-1}|$.

In the following, we shall try to construct a basis for the unique orthogonal complement W^{j-1} of S^{j-1} in S^j . Each of these basis functions will be called a **prewavelet** and the space W^{j-1} a **prewavelets space**. By combining prewavelet bases of the spaces W^k with the nodal bases for the spaces S^k , we obtain the framework for a multiresolution analysis (MRA). Thus any function f^n in S^n can be decomposed into its $n+1$ mutually orthogonal components:

$$f^n = f^0 \oplus g^0 \oplus g^1 \oplus \dots \oplus g^n \quad (2.7)$$

where $f^0 \in S^0$ and $g^j \in W^j$ ($j = 0, 1, \dots, n-1$). We shall restrict our work for the construction of bases of W^k to the first refinement level since uniform refinement has been used.

Let b be any given non-zero real number, a_1 and a_2 be two neighboring vertices in V^0 , and denote by $u \in V^1 \setminus V^0$ their midpoint. We define the **semi-prewavelet** $\sigma_{a_1, u} \in S^1$ as the element with support contained in the support of $\phi_{a_1}^0$ and having the property that, for all $v \in V^0$,

$$\langle \phi_v^0, \sigma_{a_1, u} \rangle = \begin{cases} -b, & \text{if } v = a_1; \\ b, & \text{if } v = a_2; \\ 0, & \text{otherwise} \end{cases} \quad (2.8)$$

where

$$\sigma_{a_1, u}(x) = \sum_{v \in N_{a_1}^1} r_v \phi_v^1(x)$$

and

$$N_{a_1}^1 = \{a_1\} \cup V_{a_1}^1$$

denotes the fine neighborhood of a_1 . The only nontrivial inner products between $\sigma_{a_1, u}$ and coarse nodal functions ϕ_v^0 occur when v belongs to the coarse neighborhood of a_1 :

$$N_{a_1}^0 = \{a_1\} \cup V_{a_1}^0.$$

Thus the number of coefficients and conditions are the same and, as we will subsequently establish, the element $\sigma_{a_1,u}$ is uniquely determined for any given non-zero value of b in (2.8).

Since the dimension of W^0 is equal to the number of fine vertices in V^1 , i.e. $|V^1| - |V^0|$, it is natural to associate one prewavelet ψ_u per fine vertex $u \in V^1 \setminus V^0$. Since each u is the midpoint of some edge in E^0 connecting two coarse vertices a_1 and a_2 in V^0 , the element of S^1 ,

$$\psi_u = \sigma_{a_1,u} + \sigma_{a_2,u} \quad (2.9)$$

is a prewavelet since it is orthogonal to all nodal functions ϕ_v^0 , $v \in V^0$.

The following theorem gives a sufficient condition that all the different prewavelets obtained in this way are linearly independent and hence form a basis of W^0 .

Theorem 2.2 *The set of prewavelets $\{\psi_u\}_{u \in V^1 \setminus V^0}$ defined by (2.9) is a linearly independent set if*

$$\psi_u(u) > \sum_{w \in V^1 \setminus V^0, w \neq u} |\psi_u(w)|, \quad \forall u \in V^1 \setminus V^0$$

Proof: Let

$$|V^0| = m, \quad |V^1 \setminus V^0| = n,$$

and

$$u_1, u_2, \dots, u_m, u_{m+1}, \dots, u_{m+n}$$

be any permutation of all the vertices in V^1 such that $u_j \in V^0$ ($1 \leq j \leq m$) and $u_j \in V^1 \setminus V^0$ ($m+1 \leq j \leq m+n$). The corresponding nodal functions are

$$\phi_j(x) = \phi_{u_j}(x), \quad j = 1, 2, \dots, m+n.$$

Then the prewavelets in $\{\psi_u\}_{u \in V^1 \setminus V^0}$ can be written as

$$\psi_i(x) = \psi_{u_{m+i}}(x) = \sum_{j=1}^{m+n} r_{ij} \phi_j(x), \quad (1 \leq i \leq n)$$

where $r_{ij} = \psi_i(u_j)$ ($1 \leq i \leq n, 1 \leq j \leq m+n$). Consider the matrix

$$R = \begin{bmatrix} r_{11} & r_{12} & \cdots & r_{1,m+n} \\ r_{21} & r_{22} & \cdots & r_{2,m+n} \\ \cdots & \cdots & \cdots & \cdots \\ r_{n1} & r_{n2} & \cdots & r_{n,m+n} \end{bmatrix}$$

and its sub matrix

$$R_1 = \begin{bmatrix} r_{1,m+1} & r_{1,m+2} & \cdots & r_{1,m+n} \\ r_{2,m+1} & r_{2,m+2} & \cdots & r_{2,m+n} \\ \cdots & \cdots & \cdots & \cdots \\ r_{n,m+1} & r_{n,m+2} & \cdots & r_{n,m+n} \end{bmatrix}.$$

We claim that R_1 is diagonally dominant. Actually, keeping in mind that $r_{i,m+j} = \psi_i(u_{m+j}) = \psi_{u_{m+i}}(u_{m+j})$ ($1 \leq i, j \leq n$), we know that

$$r_{i,m+i} > \sum_{1 \leq j \leq n, j \neq i} |r_{i,m+j}|, \quad (1 \leq i \leq n)$$

is equivalent to

$$\psi_{u_{m+i}}(u_{m+i}) > \sum_{1 \leq j \leq n, j \neq i} |\psi_{u_{m+i}}(u_{m+j})|, \quad (1 \leq i \leq n)$$

or

$$\psi_u(u) > \sum_{w \in V^1 \setminus V^0, w \neq u} |\psi_u(w)|, \quad \forall u \in V^1 \setminus V^0.$$

This completes the proof of the theorem.

3 Semi-prewavelets

We are going to establish the uniqueness of the semi-prewavelets for W^0 with regard to (2.8) and to find their coefficients. To simplify our calculation, Floater and Quak's result on inner products of nodal functions [7], which we state here as a Lemma, will be used.

Lemma 3.1 *Let $t(e)$ denote the number of triangles (one or two) in $\mathcal{T}^0$ containing the edge $e \in E^0$ and $t(v)$ the number of triangles (at least one) containing the vertex $v \in V^0$. If $v \in V^0$ and $w \in V^1$ are contained in the same triangle in $\mathcal{T}^0$ then*

$$96\langle \phi_v^0, \phi_w^1 \rangle = \begin{cases} 6t(v), & \text{if } v = w; \\ 10t(e), & \text{if } w \text{ is the midpoint of } e; \\ t(e), & \text{if } e = [v, w]; \\ 4, & \text{if otherwise.} \end{cases} \quad (3.1)$$

Let $\sigma_{a_1, u}$ be a semi-prewavelet where the fine vertex u is the midpoint of a_1 and another coarse vertex a_2 . We call a_1 the **center** (vertex) of the semi-prewavelet. The **degree** of a vertex in a triangulation is the number of neighbor vertices of the vertex in that triangulation. Trivially, every coarse vertex is also a vertex in the fine triangulation and it has the same degree in both coarse and fine r-triangulations. Let $k = |V_{a_1}^0| = |V_{a_1}^1|$ be the degree of a_1 . If a_1 is an interior vertex then $k = 6$. If a_1 is a boundary vertex then the value of k could range from 2 to 6. Hence there are six possible topological structures of the support of $\sigma_{a_1, u}$, which are identical to the support of $\phi_{a_1}^0$.

For convenience, in the later steps to construct semi-prewavelets, we shall use the following permutations of the vertices in the coarse neighborhood $N_{a_1}^0$ and the fine neighborhood $N_{a_1}^1$ of a_1 :

$$N_{a_1}^0 : v_1 = a_1, v_2, \dots, v_k,$$

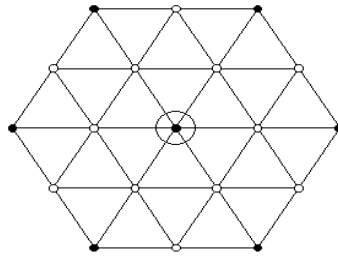


Figure 3: Interior vertex a_1 and its support

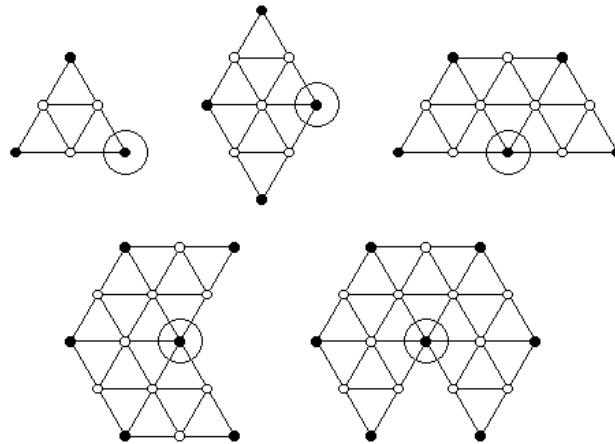


Figure 4: Boundary vertices and their supports

$$N_{a_1}^1 : u_1 = a_1, u_2, \dots, u_k,$$

where v_2 through v_k are all the neighbor vertices of a_1 in the coarse space labeled consecutively in the counterclockwise order, and u_j is the midpoint of the edge $[a_1, v_j]$, $j = 2, 3, \dots, k$.

Let us simply rewrite $\phi_{u_j}^l(x)$ as $\phi_j^l(x)$, where $j = 1, 2, \dots, k$ and $l = 0, 1$. Then,

$$\sigma_{a_1, u}(x) = \sum_{j=1}^k r_j \phi_j^1(x). \quad (3.2)$$

Let $A = (a_{ij})$ be the $k \times k$ matrix such that $a_{ij} = 96 \langle \phi_i^0, \phi_j^1 \rangle$. We have

$$\begin{aligned} \langle \phi_i^0, \sigma_{v_1, u} \rangle &= \left\langle \phi_i^0, \sum_{j=1}^k r_j \phi_j^1 \right\rangle = \sum_{j=1}^k r_j \langle \phi_i^0, \phi_j^1 \rangle \\ &= \sum_{j=1}^k \frac{1}{96} a_{ij} r_j = \frac{1}{96} [a_{i1}, a_{i2}, \dots, a_{ik}] \vec{r}, \end{aligned}$$

where

$$\vec{r} = [r_1, r_2, \dots, r_k]^T.$$

Therefore, the “semi-orthogonal condition” (2.8) is equivalent to

$$A \vec{r} = \vec{b}, \quad (3.3)$$

where $\vec{r} = [r_1, r_2, \dots, r_k]^T$,

$$A = \begin{bmatrix} a_{11} & a_{12} & \cdots & a_{1k} \\ a_{21} & a_{22} & \cdots & a_{2k} \\ \dots & \dots & \dots & \dots \\ a_{k1} & a_{k2} & \cdots & a_{kk} \end{bmatrix},$$

and

$$\vec{b} = [b_1 = -b, 0, \dots, 0, b_j = b, 0, \dots, 0]^T \text{ (here } j \text{ satisfying } u = u_j \text{)}.$$

Thus, if A is invertible then (3.3) has a unique solution of the coefficients. Figure 3 and Figure 4 give all the cases of the supports of possible semi-prewavelets up to a symmetric permutation.

Using Lemma 1, we can verify that A is invertible in every case. We choose the value of b in (2.8) as 66240 so that we can get integer coefficients r_j for all the semi-prewavelets, except the boundary one with the center vertex of degree 6.

Case 1, a_1 is an interior vertex, SPW(I6) We obtain that

$$A = \begin{bmatrix} 36 & 20 & 20 & 20 & 20 & 20 & 20 \\ 2 & 20 & 4 & 0 & 0 & 0 & 4 \\ 2 & 4 & 20 & 4 & 0 & 0 & 0 \\ 2 & 0 & 4 & 20 & 4 & 0 & 0 \\ 2 & 0 & 0 & 4 & 20 & 4 & 0 \\ 2 & 0 & 0 & 0 & 4 & 20 & 4 \\ 2 & 4 & 0 & 0 & 0 & 4 & 20 \end{bmatrix}.$$

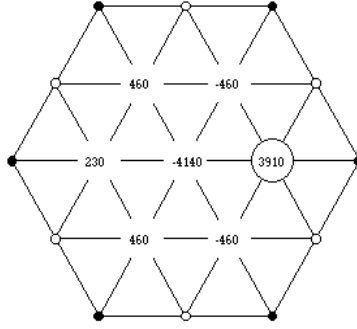


Figure 5: Semi-prewavelet SPW(I6)

A is non-singular with the inverse

$$A^{-1} = \frac{1}{1152} \begin{bmatrix} 42 & -30 & -30 & -30 & -30 & -30 & -30 \\ -3 & 65 & -11 & 5 & 1 & 5 & -11 \\ -3 & -11 & 65 & -11 & 5 & 1 & 5 \\ -3 & 5 & -11 & 65 & -11 & 5 & 1 \\ -3 & 1 & 5 & -11 & 65 & -11 & 5 \\ -3 & 5 & 1 & 5 & -11 & 65 & -11 \\ -3 & -11 & 5 & 1 & 5 & -11 & 65 \end{bmatrix}.$$

Let $\vec{b} = [-b \ b \ 0 \ 0 \ 0 \ 0 \ 0]^T$. Then

$$\vec{r} = A^{-1}\vec{b} = [-4140 \ 3910 \ -460 \ 460 \ 230 \ 460 \ -460]^T.$$

Thus, a semi-prewavelet with its center a_1 as an interior vertex has been uniquely determined as shown in Figure 5.

Simply, by turning Figure 5 around its center in the counterclockwise direction step-by-step, we can obtain all the other symmetric semi-prewavelets which share the same center a_1 . We shall see this effect in the following case from another point of view.

Case 2, a_1 is a boundary vertex with degree 2, SPW(B2)

In this case,

$$A = \begin{bmatrix} 6 & 10 & 10 \\ 1 & 10 & 4 \\ 1 & 4 & 10 \end{bmatrix}.$$

and thus,

$$A^{-1} = \frac{1}{192} \begin{bmatrix} 42 & -30 & -30 \\ -3 & 25 & -7 \\ -3 & -7 & 25 \end{bmatrix}.$$



Figure 6: Semi-prewavelet SPW(B2)

We obtain

$$\vec{r} = A^{-1}\vec{b} = \begin{bmatrix} -24840 & 9660 & -1380 \end{bmatrix}^T$$

for $\vec{b} = \begin{bmatrix} -b & b & 0 \end{bmatrix}^T$, and

$$\vec{r} = A^{-1}\vec{b} = \begin{bmatrix} -24840 & -1380 & 9660 \end{bmatrix}^T$$

for $\vec{b} = \begin{bmatrix} -b & 0 & b \end{bmatrix}^T$, respectively.

As we can see in Figure 6, the two subcases of SPW(B2) are symmetric. Note that they are the same up to symmetry. In the following cases of other boundary semi-prewavelets we shall not mention this repeatedly.

Case 3, a_1 is a boundary vertex with degree 3, SPW(B3)

We have

$$A = \begin{bmatrix} 12 & 10 & 20 & 10 \\ 1 & 10 & 4 & 0 \\ 2 & 4 & 20 & 4 \\ 1 & 0 & 4 & 10 \end{bmatrix}$$

and

$$A^{-1} = \frac{1}{1920} \begin{bmatrix} 210 & -150 & -150 & -150 \\ -15 & 221 & -35 & 29 \\ -15 & -35 & 125 & -35 \\ -15 & 29 & -35 & 221 \end{bmatrix}.$$

Therefore,

$$\vec{r} = A^{-1}\vec{b} = \begin{bmatrix} -12420 & 8142 & -690 & 1518 \end{bmatrix}^T$$

if $\vec{b} = \begin{bmatrix} -b & b & 0 & 0 \end{bmatrix}^T$, and

$$\vec{r} = A^{-1}\vec{b} = \begin{bmatrix} -12420 & -690 & 4830 & -690 \end{bmatrix}^T$$

if $\vec{b} = \begin{bmatrix} -b & 0 & b & 0 \end{bmatrix}^T$.

Case 4, a_1 is a boundary vertex with degree 4, SPW(B4)

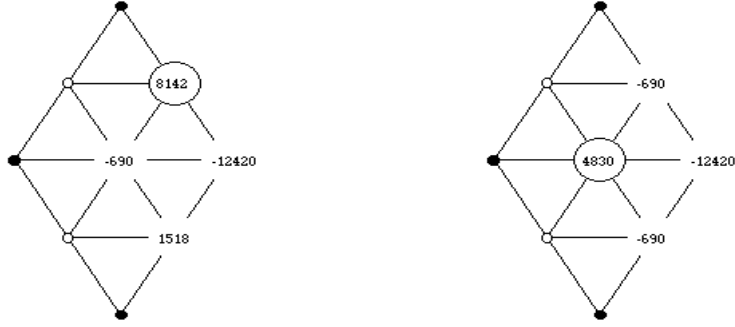


Figure 7: Semi-prewavelet SPW(B3[1]), SPW(B3[2])

In this case,

$$A = \begin{bmatrix} 18 & 10 & 20 & 20 & 10 \\ 1 & 10 & 4 & 0 & 0 \\ 2 & 4 & 20 & 4 & 0 \\ 2 & 0 & 4 & 20 & 4 \\ 1 & 0 & 0 & 4 & 10 \end{bmatrix}$$

and thus,

$$A^{-1} = \frac{1}{576} \begin{bmatrix} 42 & -30 & -30 & -30 & -30 \\ -3 & 65 & -11 & 5 & 1 \\ -3 & -11 & 35 & -5 & 5 \\ -3 & 5 & -5 & 35 & -11 \\ -3 & 1 & 5 & -11 & 65 \end{bmatrix}.$$

Therefore, $\vec{r} = [-8280 \ 7820 \ -920 \ 920 \ 460]^T$ for $\vec{b} = [-b \ b \ 0 \ 0 \ 0]^T$, and $\vec{r} = [-8280 \ -920 \ 4370 \ -230 \ 920]^T$ for $\vec{b} = [-b \ 0 \ b \ 0 \ 0]^T$, correspondingly.

Case 5, a_1 is a boundary vertex with degree 5, SPW(B5)

We have

$$A = \begin{bmatrix} 24 & 10 & 20 & 20 & 20 & 10 \\ 1 & 10 & 4 & 0 & 0 & 0 \\ 2 & 4 & 20 & 4 & 0 & 0 \\ 2 & 0 & 4 & 20 & 4 & 0 \\ 2 & 0 & 0 & 4 & 20 & 4 \\ 1 & 0 & 0 & 0 & 4 & 10 \end{bmatrix}$$

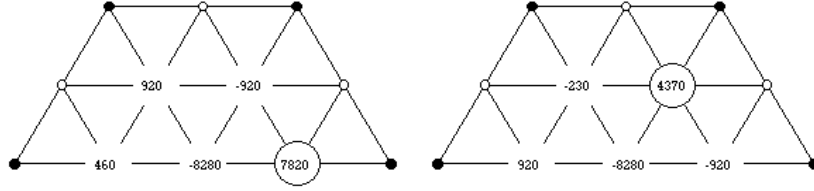


Figure 8: Semi-prewavelet SPW(B4[1]), SPW(B4[2])

and

$$A^{-1} = \frac{1}{88320} \begin{bmatrix} 4830 & -3450 & -3450 & -3450 & -3450 & -3450 \\ -345 & 9883 & -1765 & 667 & 155 & 283 \\ -345 & -1765 & 5275 & -805 & 475 & 155 \\ -345 & 667 & -805 & 5083 & -805 & 667 \\ -345 & 155 & 475 & -805 & 5275 & -1765 \\ -345 & 283 & 155 & 667 & -1765 & 9883 \end{bmatrix}.$$

Accordingly, we obtain $\vec{r} = [-6210 \ 7671 \ -1065 \ 759 \ 375 \ 471]^T$ if $\vec{b} = [-b \ b \ 0 \ 0 \ 0 \ 0]^T$, $\vec{r} = [-6210 \ 759 \ -345 \ 4071 \ -345 \ 759]^T$ if $\vec{b} = [-b \ 0 \ 0 \ b \ 0 \ 0]^T$, and $\vec{r} = [-6210 \ -1065 \ 4215 \ -345 \ 615 \ 375]^T$ if $\vec{b} = [-b \ 0 \ b \ 0 \ 0 \ 0]^T$.

Case 6, a_1 is a boundary vertex with degree 6, SPW(B6)

In the final case, we have

$$A = \begin{bmatrix} 30 & 10 & 20 & 20 & 20 & 20 & 10 \\ 1 & 10 & 4 & 0 & 0 & 0 & 0 \\ 2 & 4 & 20 & 4 & 0 & 0 & 0 \\ 2 & 0 & 4 & 20 & 4 & 0 & 0 \\ 2 & 0 & 0 & 4 & 20 & 4 & 0 \\ 2 & 0 & 0 & 0 & 4 & 20 & 4 \\ 1 & 0 & 0 & 0 & 0 & 4 & 10 \end{bmatrix}.$$

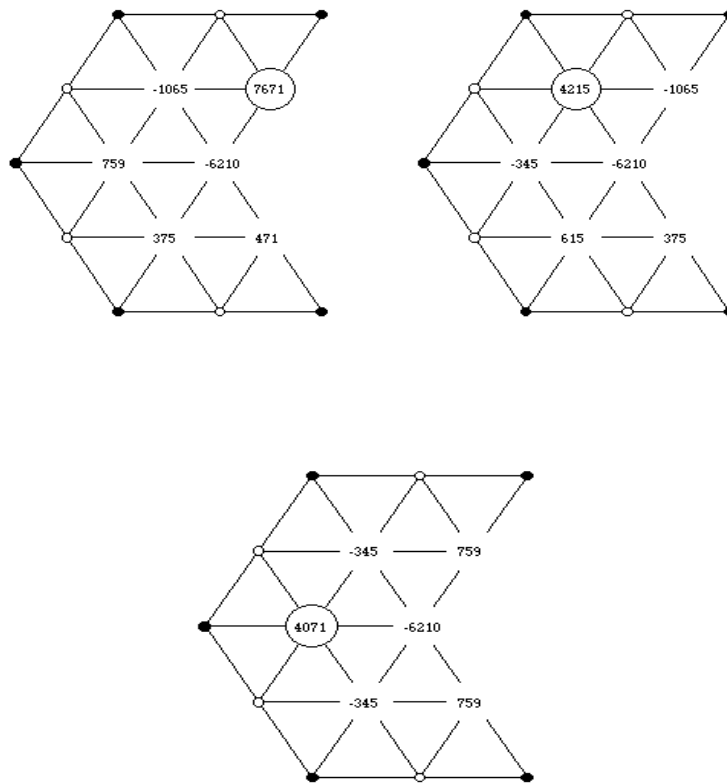


Figure 9: Semi-pretwavelet $\text{SPW}(\text{B5}[1])$, $\text{SPW}(\text{B5}[2])$, and $\text{SPW}(\text{B5}[3])$

A is non-singular with the inverse

$$A^{-1} = \frac{1}{192} \begin{bmatrix} \frac{7}{160} & \frac{-1}{32} & \frac{-1}{32} & \frac{-1}{32} & \frac{-1}{32} & \frac{-1}{32} & \frac{-1}{32} \\ \frac{-1}{320} & \frac{11779}{105792} & \frac{-2173}{105792} & \frac{739}{105792} & \frac{131}{105792} & \frac{259}{105792} & \frac{227}{105792} \\ \frac{-1}{320} & \frac{-2173}{105792} & \frac{6259}{105792} & \frac{-1021}{105792} & \frac{499}{105792} & \frac{179}{105792} & \frac{259}{105792} \\ \frac{-1}{320} & \frac{739}{105792} & \frac{-1021}{105792} & \frac{6019}{105792} & \frac{-973}{105792} & \frac{499}{105792} & \frac{131}{105792} \\ \frac{-1}{320} & \frac{131}{105792} & \frac{499}{105792} & \frac{-973}{105792} & \frac{6019}{105792} & \frac{-1021}{105792} & \frac{739}{105792} \\ \frac{-1}{320} & \frac{259}{105792} & \frac{179}{105792} & \frac{499}{105792} & \frac{-1021}{105792} & \frac{6259}{105792} & \frac{-2173}{105792} \\ \frac{-1}{320} & \frac{227}{105792} & \frac{259}{105792} & \frac{131}{105792} & \frac{739}{105792} & \frac{-2173}{105792} & \frac{11779}{105792} \end{bmatrix}.$$

We can solve for the vectors of coefficients as the following.

$$\vec{r} = A^{-1}\vec{b} = \frac{276}{551} \begin{bmatrix} -9918 & 15137 & -2303 & 1337 & 577 & 737 & 697 \end{bmatrix}^T$$

for

$$\vec{b} = \begin{bmatrix} -b & b & 0 & 0 & 0 & 0 & 0 \end{bmatrix}^T,$$

$$\vec{r} = A^{-1}\vec{b} = \frac{276}{551} \begin{bmatrix} -9918 & -2303 & 8237 & -863 & 1037 & 637 & 737 \end{bmatrix}^T$$

for

$$\vec{b} = \begin{bmatrix} -b & 0 & b & 0 & 0 & 0 & 0 \end{bmatrix}^T,$$

and

$$\vec{r} = A^{-1}\vec{b} = \frac{276}{551} \begin{bmatrix} -9918 & 1337 & -863 & 7937 & -803 & 1037 & 577 \end{bmatrix}^T$$

for

$$\vec{b} = \begin{bmatrix} -b & 0 & 0 & b & 0 & 0 & 0 \end{bmatrix}^T.$$

Now we can get any possible prewavelets of an r -triangulation, since the above semi-prewavelets included all the possible semi-prewavelets with the exception of symmetries. By (2.9), to get a prewavelet, we only need to “sum” two semi-prewavelets together in such a way that the fine vertex u (which has been circled in each figure of the semi-prewavelet) is the midpoint of the centers of these two semi-prewavelets. Some prewavelets which could be often used are given by Figure 11 through Figure 17. In these figures we denote the prewavelet obtained by summing an interior semi-prewavelet, $\text{SPW}(I6)$, and a boundary one, say $\text{SPW}(Bj)$, by $\text{PW}(I6, Bj)$. We denote the prewavelet obtained by summing two boundary semi-prewavelets, $\text{SPW}(Bi)$ and $\text{SPW}(Bj)$, by $\text{PW}(Bi, Bj)$.

With the above results on semi-prewavelets, we are now ready to state our main result on r -triangulations.

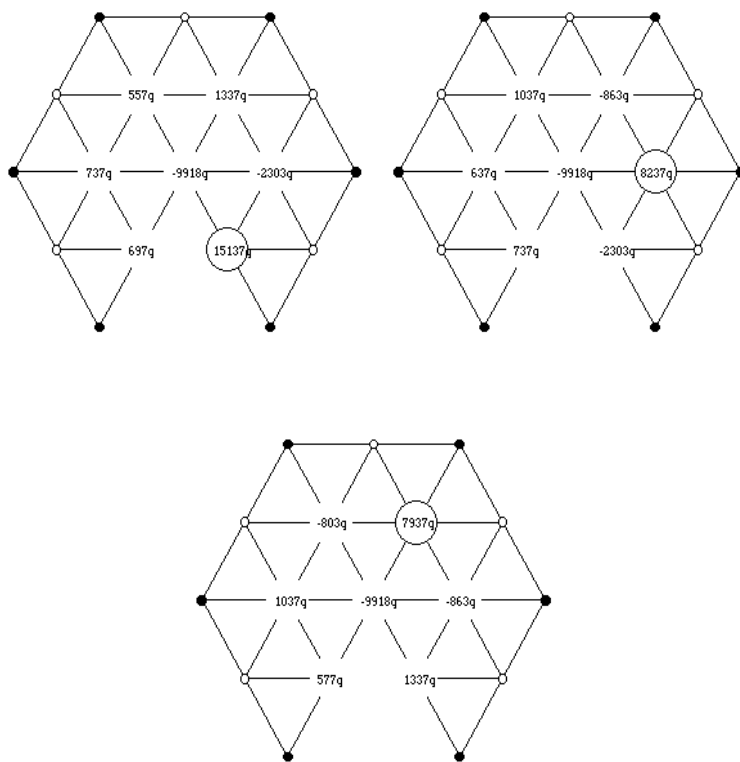


Figure 10: Semi-pretwavelet $\text{SPW}(\text{B6}[1])$, $\text{SPW}(\text{B6}[2])$ and $\text{SPW}(\text{B6}[3])$ ($q = \frac{276}{551}$)

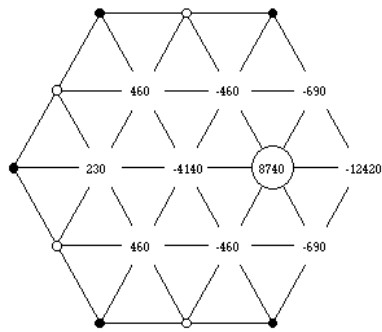


Figure 11: Pre-wavelet PW(I6,B3)

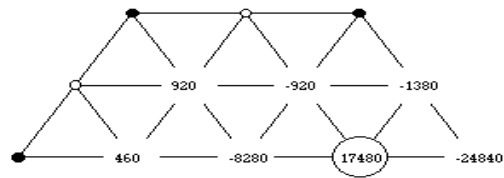


Figure 12: Pre-wavelet PW(B4,B2)

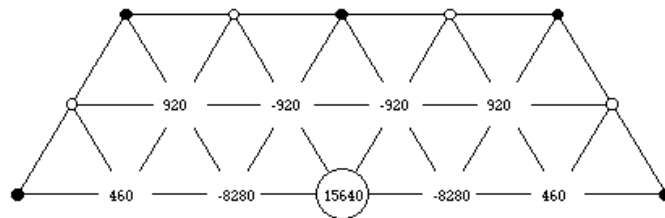


Figure 13: Pre-wavelet PW(B4,B4)

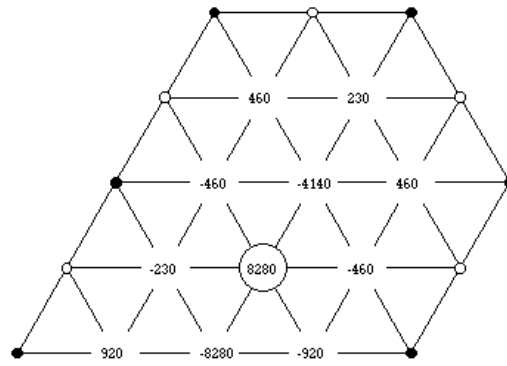


Figure 14: Pre-wavelet PW(I6,B4)

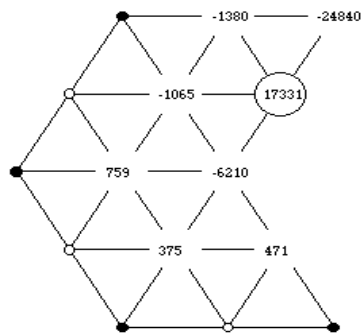


Figure 15: Pre-wavelet PW(B5,B2)

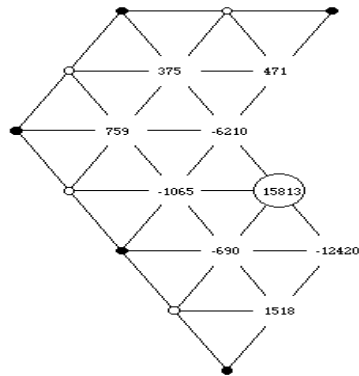


Figure 16: Pre-wavelet $PW(B_5, B_3)$

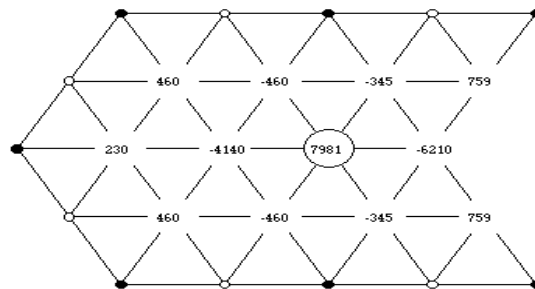


Figure 17: Pre-wavelet $PW(I_6, B_5)$

Theorem 3.2 *For any level of the refinements of any an r -triangulation, all the possible prewavelets can be constructed by simply summing up the two semi-prowavelets illustrated in Figure 5 through Figure 10.*

Proof: Each semi-prowavelet $\sigma_{a_1,u}(x)$ illustrated in Figure 5 through Figure 10, which are all the possible cases of semi-prowavelets in r -triangulations, satisfies

$$\sigma_{a_1,u}(u) > \sum_{w \in V^1 \setminus V^0, w \neq u} |\sigma_{a_1,u}(w)|, \quad \forall u \in V^1 \setminus V^0.$$

A prewavelet is the sum of two semi-prowavelets as expressed in (2.9). The sum can be done in such a way that the fine vertex u (which has been circled in each figure of semi-prowavelets) is the midpoint of the centers, a_1 and a_2 , of these two semi-prowavelets in (2.9). Since the intersection of $N_{a_1}^1$ and $N_{a_2}^1$ is the single element set $\{u\}$, the only overlapped values are the values of the two semi-prowavelets functions at vertex u . Therefore, the condition of Theorem 2.1 is satisfied and this completes the proof.

Theorem 3.2 can be used on various shaped domains and all the possible prewavelets can be found easily by simply checking if any figure in Figure 3 and Figure 4 is a subset of the given refined r -triangulation.

Acknowledgments: This work was supported in part by NSF-IGMS 0408086 and 0552377 for Hong.

References

- [1] J.S. Cao, D. Hong, and W.J. Huang, Parameterized Piecewise Linear Pre-Wavelets over Type-2 Triangulations, in *Interactions between Wavelets and Splines* (M.J. Lai ed.), Nashboro Press, Nashville, 2006, pp. 75–92.
- [2] C.K. Chui, *Multivariate Splines*, SIAM, Philadelphia, 1988.
- [3] G.C. Donovan, J.S. Geronimo, and D.P. Hardin, Compactly supported, piecewise affine scaling functions on triangulations, *Constructive Approximation* **16**, 201–219 (2000).
- [4] D.P. Hardin and D. Hong, On Piecewise Linear Wavelets and Prowavelets over Triangulations, *J. Comput. and Applied Math.*, **155**, 91–109 (2003).
- [5] D. Hong and L.L. Schumaker, Surface compression using hierarchical bases for bivariate C^1 cubic splines, *Computing*, **72** (2004), 79–92.
- [6] D. Hong and Y. Mu, On construction of minimum supported piecewise linear prewavelets over triangulations, in *Wavelets and Multiresolution Methods* (T.X. He ed.), Marcel Dekker Pub., New York, 2000, pp. 181–201.
- [7] M.S. Floater and E.G. Quak, Piecewise linear prewavelets on arbitrary triangulations, *Numer. Math.* **82**, 221–252 (1999).

- [8] M.S. Floater and E.G. Quak, A semiprewavelet approach to piecewise linear prewavelets on triangulations, in *Approximation Theory IX, Vol 2: Computational Aspects* (C.K. Chui and L.L. Schumaker eds.), Vanderbilt University Press, Nashville, 1998, pp. 63–70.
- [9] M.S. Floater and E.G. Quak, Linear independence and stability of piecewise linear prewavelets on arbitrary triangulations, *SIAM J. Numer. Anal.*, **38**, 58–79 (2000).
- [10] U. Kotyczka and P. Oswald, Piecewise linear prewavelets of small support, in *Approximation Theory VIII, Vol. 2* (C.K. Chui and L.L. Schumaker, eds.), World Scientific, Singapore, 1995, pp. 235–242.

Spline Wavelets on Refined Grids and Applications in Image Compression

Lutai Guan and Tan Yang
Dept. of Sci. Computation & Computer Application
Sun Yat-sen (Zhongshan) University
Guangzhou, Guandong 510275, P.R. China

Abstract

In this paper, one kind of new spline-wavelets on refined grids is constructed. The scale functions are natural splines on refined grids that were studied by authors for interpolation surfaces recently. Because of the advantages of these splines to describe details well in some areas, the spline-wavelets on refined grids will be applied in image compression and information processing.

2000 Mathematics Subject Classification: Primary 90C325, 41A29; Secondary 65D15, 49J52.

Keywords: *image compression, refined grids, spline-wavelets.*

1 Introduction

Wavelets play an important role in quite different fields of science and technology. There are many wavelets books available now, some of them are destined to be classics in the field, for example, [2], [4], and [9]. Image or surface compression is one of the most active fields in the applications of wavelets (see [8], [10], [11]).

In general, wavelets need a uniform grid subdivision, but for image processing, sometimes we only need detail information in some given regions and for other parts of the image we could represent them in a rough scale.

In [7], a method for compressing surfaces associated with C^1 cubic splines defined on triangulated quadrangulations is described. The refinement method for triangular regions is efficient for some applications in surface compression.

Classical natural spline interpolation can be found in [1]. Some recent research work on natural spline interpolation for scattered data has been done in [3], [5], and [6]. In [6], one kind of surface design called natural splines over refined grid points is given. The method is simple and easy to be implemented

in computers. The designed surfaces have no gaps in the connecting patches and are smoother. A change of only one de Boor point in refined grid will affect the surface in a very small region.

In this paper, based on the spline interpolation over refined grid points, we consider a new kind of wavelets called addition spline-wavelets corresponding to the splines over refined grid points, and discuss their vanishing moment properties and give algorithms for their application in image compression.

2 Compactly Supported Spline over Refined Grids

Following the notion of [6], we first define the refined grids and interpolating natural spline.

Definition 2.1 For given partitions $\pi_1 : a = x_0 < x_1 < \cdots < x_{k+1} = b$ and $\pi_2 : c = y_0 < y_1 < \cdots < y_{\ell+1} = d$, and subdivisions $\pi_3 : x_i = x_{i,0} < x_{i,1} < \cdots < x_{i,r+1} = x_{i+1}$ and $\pi_4 : y_j = y_{j,0} < y_{j,1} < \cdots < y_{j,s+1} = y_{j+1}$, the grid points:

$$\{x_p, y_q\}, \quad p = 1, \dots, k; q = 1, \dots, \ell,$$

and the sub-grid points:

$$\{x_{i,\mu}, y_{j,\nu}\}, \quad \mu = 1, \dots, r; \nu = 1, \dots, s$$

are called the refined grid points (or for short, refined grids).

Definition 2.2 The function $\sigma(x, y) \in H^{m,n}(\mathbb{R})$ is called the interpolating natural spline if it satisfies

$$\|T\sigma\|^2 = \min_{u \in X, Au=z} \{\|Tu\|^2\},$$

where

$$\|Tu\|^2 = \iint_{\mathbb{R}} (u^{(m,n)}(x, y))^2 dx dy + \sum_{\nu=0}^{n-1} \int_a^b (u^{(m,\nu)}(x, c))^2 dx + \sum_{\mu=0}^{m-1} \int_c^d (u^{(\mu,n)}(a, y))^2 dy,$$

and $(x_i, y_i, z_i), i = 1, \dots, N$ are given scattered data for $z = (z_1, \dots, z_N)$ and a linear operator $A : H^{m,n}(\mathbb{R}) \mapsto \mathbb{R}^N$ defined as

$$Au = (u(x_1, y_1), \dots, u(x_N, y_N)).$$

When the scattered data are over the refined grid points $\{x_p, y_q\}, \{x_{i,\mu}, y_{j,\nu}\}, \mu = 1, \dots, r; \nu = 1, \dots, s; p = 1, \dots, k; q = 1, \dots, \ell$, we call $\sigma(x, y)$ an interpolating spline over refined grid points. We can write the data over refined grid as $\{x_p, y_q, z_{p,q}\}, \{x_{i,\mu}, y_{j,\nu}, z_{i,j,\mu,\nu}\}, \mu = 1, \dots, r; \nu = 1, \dots, s; p = 1, \dots, k; q = 1, \dots, \ell$.

The following result on bivariate natural local basis interpolation over refined grid points is given in [6].

Lemma 2.3 *The natural spline $\sigma(x, y)$ over refined grid points has the expression:*

$$\sigma(x, y) = \sum_{\mu=1}^k \sum_{\nu=1}^{\ell} \lambda_{\mu, \nu} B_{\mu}(x) \hat{B}_{\nu}(y) + \sum_{\mu=1}^r \sum_{\nu=1}^s \gamma_{\mu, \nu} B_{i, \mu}(x) \hat{B}_{j, \nu}(y),$$

where $B_{\mu}(x)$, $\hat{B}_{\nu}(y)$, $B_{i, \mu}(x)$, and $\hat{B}_{j, \nu}(y)$ are B-splines of order $2m - 1$ for x and of order $2n - 1$ for y .

Similar to B-spline interpolation, we can add the boundary knots for the construction of natural B-splines, and the invariable natural B-spline interpolation satisfying

$$\sigma^{(\ell)}(a) = \sigma^{(\ell)}(b) = 0, \quad \ell = m, m + 1, \dots, 2m - 1.$$

3 Spline-wavelet on Refined Grids

It is not trivial at all to select the compactly supported splines over refined grids mentioned in last section as the scaling functions for constructing wavelets. For this purpose, we apply the tensor product method.

First, we notice that the basis of a spline space over refined grids can be divided into two sets: $\{B_{\mu}(x)\hat{B}_{\nu}(y)\}_{\mu=1, \nu=1}^{k, \ell}$ and $\{B_{i, \mu}(x)\hat{B}_{j, \nu}(y)\}_{\mu=1, \nu=1}^{r, s}$. Second, for equal division (uniform grid), the basis of a spline space can be grouped as $\{B(x - \mu)\hat{B}(y - \nu)\}_{\mu=1, \nu=1}^{k, \ell}$ and $\{B_i(x - \mu)\hat{B}_j(y - \nu)\}_{\mu=1, \nu=1}^{r, s}$. They are two sets of tensor product spline scaling functions.

Definition 3.1 (*Addition Wavelet Space*) If $V_0 = V_0^1 \cup V_0^2$ and there exists wavelet spaces W_1^1 and W_1^2 satisfying:

$$V_0^1 = V_1^1 \oplus W_1^1, V_0^2 = V_1^2 \oplus W_1^2,$$

then we call $W_1^1 \cup W_1^2$ an addition wavelet space, and $V_1^1 \cup V_1^2$ an addition scaling space.

Theorem 3.2 For a function $f_0 \in V_0$, we can find a decomposition $f_0 = f_1^1 + f_1^2 + g_1^1 + g_1^2$ where $f_1^i \in V_1^i$, $g_1^i \in W_1^i$, $i = 1, 2$, and f_1^1 is orthogonal to g_1^1 and f_1^2 orthogonal to g_1^2 .

Proof. Since $V_0 = V_0^1 \cup V_0^2$, $f_0 \in V_0$, there exist f_0^1 and f_0^2 satisfying $f_0 = f_0^1 + f_0^2$.

By the wavelet decomposition:

$$V_0^1 = V_1^1 \oplus W_1^1, V_0^2 = V_1^2 \oplus W_1^2$$

We can find $f_1^1 \in V_1^1$ and $g_1^1 \in W_1^1$, such that $f_0^1 = f_1^1 + g_1^1$ and f_1^1 is orthogonal to g_1^1 .

There also exist $f_1^2 \in V_1^2$ and $g_1^2 \in W_1^2$, such that $f_0^2 = f_1^2 + g_1^2$ and f_1^2 is orthogonal to g_1^2 . Hence,

$$f_0 = f_0^1 + f_0^2 = f_1^1 + f_1^2 + g_1^1 + g_1^2.$$

Definition 3.3 (*Addition Wavelet Decomposition*) To an addition wavelet space $W_1^1 \cup W_1^2$ and an addition scaling space $V_1^1 \cup V_1^2$, if for $f_0 \in V_0$, there exist $f_1^i \in V_1^i$ and $g_1^i \in W_1^i, i = 1, 2$ such that

$$f_0 = f_1^1 + f_1^2 + g_1^1 + g_1^2,$$

Then this equation is called an addition wavelet decomposition.

For an addition space $V_0 = V_0^1 \cup V_0^2$, if there exists wavelet decompositions for V_0^1 and V_0^2 , respectively, then there exists addition wavelet decomposition for V_0 and the decomposition is the sum of two decompositions.

Theorem 3.4 (*Vanishing Moments Property*) For an addition wavelet space $W_1 = W_1^1 \cup W_1^2$, let $\psi_1(x)$ be the wavelet function of W_1^1 having N_1 vanishing moments and $\psi_2(x)$ the wavelet function of W_1^2 having N_2 vanishing moments. Then $\psi(x) \in W_1$ has N vanishing moments:

$$\int_{\mathbb{R}} x^m \psi(x) dx = 0, 0 \leq m \leq N := \min\{N_1, N_2\}.$$

Proof. Assume that

$$\int_{\mathbb{R}} x^m \psi_{j,i}(x) dx = 0, 0 \leq m \leq N_j, j = 1, 2, \quad i \in \mathbb{Z},$$

where the functions $\psi_{j,i}, i \in \mathbb{Z}$ are basis functions of the wavelet spaces $W_i^j, j = 1, 2$.

By the definition of addition wavelet spaces, for any function $\psi(x) \in W_1$, there exist coefficients $\{\lambda_i\}$ and $\{\mu_i\}$ satisfying:

$$\psi(x) = \sum_i \lambda_i \psi_{1,i}(x) + \sum_i \mu_i \psi_{2,i}(x)$$

Hence, for $0 \leq m \leq \min\{N_1, N_2\}$

$$\begin{aligned} \int_{\mathbb{R}} x^m \psi(x) dx &= \int_{\mathbb{R}} x^m \left(\sum_i \lambda_i \psi_{1,i}(x) + \sum_i \mu_i \psi_{2,i}(x) \right) dx \\ &= \sum_i \lambda_i \int_{\mathbb{R}} x^m \psi_{1,i}(x) dx + \sum_i \mu_i \int_{\mathbb{R}} x^m \psi_{2,i}(x) dx \\ &= 0 \end{aligned}$$

Let scaling spaces $V_{0x}^1, V_{0y}^1, V_{0x}^2$, and V_{0y}^2 be generated by the scaling functions $\phi^1(x) = B(x)$, $\hat{\phi}^1(y) = \hat{B}(y)$, $\phi^2(x) = B_i(x)$, and $\hat{\phi}^2(y) = \hat{B}_j(y)$, respectively. For one dimension spaces, we have wavelet space decompositions:

$$V_{0x}^i = V_{1x}^i \oplus W_{1x}^i, \quad V_{0y}^i = V_{1y}^i \oplus W_{1y}^i \quad i = 1, 2.$$

The wavelet functions are: $\psi_1^i(x), \hat{\psi}_1^i(y)$ for W_{1x}^i and W_{1y}^i , scaling functions are: $\phi_1^i(x), \hat{\phi}_1^i(y)$ for V_{1x}^i and $V_{1y}^i, i = 1, 2$ correspondingly.

For two dimensional spaces, there exist tensor production wavelet decompositions:

$$\begin{aligned} V_0^i &= V_{0x}^i \otimes V_{0y}^i = (V_{1x}^i \oplus W_{1x}^i) \otimes (V_{1y}^i \oplus W_{1y}^i) \\ &= (V_{1x}^i \otimes V_{1y}^i) \oplus (V_{1x}^i \otimes W_{1y}^i) \oplus (W_{1x}^i \otimes V_{1y}^i) \oplus (W_{1x}^i \otimes W_{1y}^i) \\ &:= V_1^i \oplus W_1^i \quad i = 1, 2 \end{aligned}$$

Definition 3.5 (*Spline-wavelets on Refined Grids*) For uniform B-spline, the basis of a spline space over refined grids can be divided into two sets $\{B_\mu(x)\hat{B}_\nu(y)\}_{\mu=1, \nu=1}^k, l$ and $\{B_{i,\mu}(x)\hat{B}_{j,\nu}(y)\}_{\mu=1, \nu=1}^r, s$, let

$$\begin{aligned} V_0^1 &= \overline{\text{Span}}\{B_\mu(x)\hat{B}_\nu(y)\}_{\mu=1, \nu=1}^k, l \\ V_0^2 &= \overline{\text{Span}}\{B_{i,\mu}(x)\hat{B}_{j,\nu}(y)\}_{\mu=1, \nu=1}^r, s. \end{aligned}$$

The addition wavelet space for $V_0 = V_0^1 \cup V_0^2$ is called a spline-wavelet space on refined grids.

Theorem 3.6 (*Spline-wavelet Space Decomposition on Refined Grids*) The spline-wavelet space on refined grids is an addition wavelet space:

$$\begin{aligned} W_1 &= W_1^1 \cup W_1^2 \\ &= [(V_{1x}^1 \otimes W_{1y}^1) \oplus (W_{1x}^1 \otimes V_{1y}^1) \oplus (W_{1x}^1 \otimes W_{1y}^1)] \cup [(V_{1x}^2 \otimes W_{1y}^2) \oplus \\ &\quad (W_{1x}^2 \otimes V_{1y}^2) \oplus (W_{1x}^2 \otimes W_{1y}^2)] \end{aligned}$$

The corresponding addition scaling space is:

$$V_1 = V_1^1 \cup V_1^2 = (V_{1x}^1 \otimes V_{1y}^1) \cup (V_{1x}^2 \otimes V_{1y}^2).$$

Proof. By the definition of a spline-wavelet space on refined grids:

$$V_0 = V_0^1 \cup V_0^2$$

and using the tensor product wavelet decompositions, we have

$$\begin{aligned} V_0 &= V_1^1 \oplus W_1^1 \cup V_1^2 \oplus W_1^2 \\ &= (V_{1x}^1 \otimes V_{1y}^1) \oplus (V_{1x}^1 \otimes W_{1y}^1) \oplus (W_{1x}^1 \otimes V_{1y}^1) \oplus (W_{1x}^1 \otimes W_{1y}^1) \\ &\quad \cup (V_{1x}^2 \otimes V_{1y}^2) \oplus (V_{1x}^2 \otimes W_{1y}^2) \oplus (W_{1x}^2 \otimes V_{1y}^2) \oplus (W_{1x}^2 \otimes W_{1y}^2) \\ &= [(V_{1x}^1 \otimes V_{1y}^1) \cup (V_{1x}^2 \otimes V_{1y}^2)] \cup [(V_{1x}^1 \otimes W_{1y}^1) \oplus (W_{1x}^1 \otimes V_{1y}^1) \oplus (W_{1x}^1 \otimes W_{1y}^1)] \\ &\quad \cup [(V_{1x}^2 \otimes W_{1y}^2) \oplus (W_{1x}^2 \otimes V_{1y}^2) \oplus (W_{1x}^2 \otimes W_{1y}^2)] \\ &= (V_1^1 \cup V_1^2) \cup (W_1^1 \cup W_1^2) \\ &= V_1 \cup W_1. \end{aligned}$$

Theorem 3.7 (*Spline-wavelet Decomposition on Refined Grids*) If $\psi_1^i(x) \in W_{1x}^i$ and $\widehat{\psi}_1^i(y) \in W_{1y}^i$ are wavelet functions and $\phi_1^i(x) \in V_{1x}^i$ and $\widehat{\phi}_1^i(y) \in V_{1y}^i$ are scaling functions for $i = 1, 2$, then for $f(x, y) \in V_0$, we have spline-wavelet decomposition on refined grids:

$$f(x, y) = \sum_i \sum_{j=1}^2 \alpha_{ij} \phi_{1,i}^j(x) \widehat{\phi}_{1,i}^j(y) + \sum_{j=1}^2 \left(\sum_i \lambda_{ij} \phi_{1,i}^j(x) \widehat{\psi}_{1,i}^j(y) + \sum_i \mu_{ij} \psi_{1,i}^j(x) \widehat{\phi}_{1,i}^j(y) + \sum_i \nu_{ij} \psi_{1,i}^j(x) \widehat{\psi}_{1,i}^j(y) \right),$$

where the coefficients $\{\alpha_{ij}\}, \{\lambda_{ij}\}, \{\mu_{ij}\}, \{\nu_{ij}\}$ are determined in the tensor product wavelet decompositions.

Proof. By the definition of splines over refined grids,

$$f(x, y) = f^1(x, y) + f^2(x, y), \quad f^1(x, y) \in V_0^1, f^2(x, y) \in V_0^2.$$

From the tensor product wavelet decomposition, there exist coefficients $\{\alpha_{ij}\}, \{\lambda_{ij}\}, \{\mu_{ij}\}, \{\nu_{ij}\}$ such that

$$f^j(x, y) = \sum_i \alpha_{ij}^j \phi_{1,i}^j(x) \widehat{\phi}_{1,i}^j(y) + \sum_i \lambda_{ij} \phi_{1,i}^j(x) \widehat{\psi}_{1,i}^j(y) + \sum_i \mu_{ij} \psi_{1,i}^j(x) \widehat{\phi}_{1,i}^j(y) + \sum_i \nu_{ij} \psi_{1,i}^j(x) \widehat{\psi}_{1,i}^j(y),$$

for $j = 1, 2$. The conclusion follows from Theorem 3.6.

Corollary 3.8 For a function $f(x, y) \in V_0$, there exists a spline-wavelet decomposition on refined grids:

$$f(x, y) = \Phi(x, y) + \Psi(x, y)$$

where $\Phi(x, y)$ is the scaling approximation part of $f(x, y)$ and $\Psi(x, y)$ is the wavelet detail part of $f(x, y)$. That is, there exist coefficients $\{\alpha_{ij}\}, \{\lambda_{ij}\}, \{\mu_{ij}\}, \{\nu_{ij}\}$ such that

$$\begin{aligned} \Phi(x, y) &= \sum_i \sum_{j=1}^2 \alpha_{ij} \phi_{1,i}^j(x) \widehat{\phi}_{1,i}^j(y), \\ \Psi(x, y) &= \sum_{j=1}^2 \left(\sum_i \lambda_{ij} \phi_{1,i}^j(x) \widehat{\psi}_{1,i}^j(y) + \sum_i \mu_{ij} \psi_{1,i}^j(x) \widehat{\phi}_{1,i}^j(y) + \sum_i \nu_{ij} \psi_{1,i}^j(x) \widehat{\psi}_{1,i}^j(y) \right). \end{aligned}$$

Theorem 3.9 (*Vanishing Moments Property of Spline-wavelet on the Refined Grids*) For any function $\psi(x, y) \in W_1$, the spline-wavelet space on refined grids, we have the following vanishing moments property:

$$\iint_{\mathbb{R} \times \mathbb{R}} x^m y^n \psi(x, y) dx dy = 0, \quad 0 \leq m \leq \min\{M_1, M_1\}, 0 \leq n \leq \min\{N_1, N_2\},$$

where the wavelet functions $\psi_1^i(x) \in W_{1x}^i$ and $\widehat{\psi}_1^i(y) \in W_{1y}^i$ have vanishing moments M_i and N_i , $i = 1, 2$ respectively.

Proof. For $\psi_1^i(x) \in W_{1x}^i$ and $\widehat{\psi}_1^i(y) \in W_{1y}^i$, we have

$$\begin{aligned} \int_{\mathbb{R}} x^m \psi_{1,k}^i(x) dx &= 0, \quad 0 \leq m \leq M_i, i = 1, 2, k = 0, \pm 1, \pm 2, \dots \\ \int_{\mathbb{R}} y^n \widehat{\psi}_{1,k}^i(y) dy &= 0, \quad 0 \leq n \leq N_i, i = 1, 2, k = 0, \pm 1, \pm 2, \dots \end{aligned}$$

By Corollary 3.8 there exists coefficients $\{\lambda_{ij}\}$, $\{\mu_{ij}\}$, $\{\nu_{ij}\}$, such that

$$\psi(x, y) = \sum_{j=1}^2 \left(\sum_i \lambda_{ij} \phi_{1,i}^j(x) \widehat{\psi}_{1,i}^j(y) + \sum_i \mu_{ij} \psi_{1,i}^j(x) \widehat{\phi}_{1,i}^j(y) + \sum_i \nu_{ij} \psi_{1,i}^j(x) \widehat{\psi}_{1,i}^j(y) \right).$$

Hence

$$\begin{aligned} \iint_{\mathbb{R} \times \mathbb{R}} x^m y^n \psi(x, y) dx dy &= \iint_{\mathbb{R} \times \mathbb{R}} x^m y^n \left[\sum_{j=1}^2 \left(\sum_i \lambda_{ij} \phi_{1,i}^j(x) \widehat{\psi}_{1,i}^j(y) \right. \right. \\ &\quad \left. \left. + \sum_i \mu_{ij} \psi_{1,i}^j(x) \widehat{\phi}_{1,i}^j(y) + \sum_i \nu_{ij} \psi_{1,i}^j(x) \widehat{\psi}_{1,i}^j(y) \right) \right] dx dy \\ &= \sum_{j=1}^2 \sum_i \lambda_{ij} \iint_{\mathbb{R} \times \mathbb{R}} x^m y^n \phi_{1,i}^j(x) \widehat{\psi}_{1,i}^j(y) dx dy \\ &\quad + \sum_{j=1}^2 \sum_i \mu_{ij} \iint_{\mathbb{R} \times \mathbb{R}} x^m y^n \psi_{1,i}^j(x) \widehat{\phi}_{1,i}^j(y) dx dy \\ &\quad + \sum_{j=1}^2 \sum_i \nu_{ij} \iint_{\mathbb{R} \times \mathbb{R}} x^m y^n \psi_{1,i}^j(x) \widehat{\psi}_{1,i}^j(y) dx dy \\ &= \sum_{j=1}^2 \sum_i \lambda_{ij} \left(\int_{\mathbb{R}} x^m \phi_{1,i}^j(x) dx \right) \left(\int_{\mathbb{R}} y^n \widehat{\psi}_{1,i}^j(y) dy \right) \\ &\quad + \sum_{j=1}^2 \sum_i \mu_{ij} \left(\int_{\mathbb{R}} x^m \psi_{1,i}^j(x) dx \right) \left(\int_{\mathbb{R}} y^n \widehat{\phi}_{1,i}^j(y) dy \right) \\ &\quad + \sum_{j=1}^2 \sum_i \nu_{ij} \left(\int_{\mathbb{R}} x^m \psi_{1,i}^j(x) dx \right) \left(\int_{\mathbb{R}} y^n \widehat{\psi}_{1,i}^j(y) dy \right) \\ &= 0. \end{aligned}$$

4 Image Compression of Spline-wavelets on Refined Grids

There are many applications of wavelets in image compression. The goal of image compression is to take advantage of proper structures in the image to reduce image storages. There usually are three steps in the process: (1) the transform step, (2) the quantization step, and (3) the coding step. Symmetry, vanishing moments and size of the filters are three things to be considered with choosing a filter for the transform step.

Because of the high vanishing moments of the spline-wavelets on refined grids, we try to choose them as filters. Using spline interpolation, we can extract the part of smooth image easily. And by coding technique, we can get the non-smooth details using wavelets. Adding these two parts, we can reach the goal of compression of the images. We provide the following algorithm to show the whole process.

algorithm 4.1 1. *Input data of an image.*

2. *Design a basic grid (rough grid) for the image and select some regions in the image that need to have more details, then apply some refined grids on those regions.*

3. *To the rough grid, use the tensor product spline-wavelet decomposition to get coefficients $\{\alpha_{i1}\}, \{\lambda_{i1}\}, \{\mu_{i1}\}$ and $\{\nu_{i1}\}$, such that:*

$$\begin{aligned} f^1(x, y) &= \sum_i \alpha_{i1} \phi_{1,i}^1(x) \hat{\phi}_{1,i}^1(y) + \sum_i \lambda_{i1} \phi_{1,i}^1(x) \hat{\psi}_{1,i}^1(y) \\ &\quad + \sum_i \mu_{i1} \psi_{1,i}^1(x) \hat{\phi}_{1,i}^1(y) + \sum_i \nu_{i1} \psi_{1,i}^1(x) \hat{\psi}_{1,i}^1(y). \end{aligned}$$

4. *Let $f^2(x, y) = f(x, y) - f^1(x, y)$. On refined grids, use the tensor product spline-wavelet decomposition to get coefficients $\{\alpha_{i2}\}, \{\lambda_{i2}\}, \{\mu_{i2}\}$ and $\{\nu_{i2}\}$, such that*

$$\begin{aligned} f^2(x, y) &= \sum_i \alpha_{i2} \phi_{1,i}^2(x) \hat{\phi}_{1,i}^2(y) + \sum_i \lambda_{i2} \phi_{1,i}^2(x) \hat{\psi}_{1,i}^2(y) \\ &\quad + \sum_i \mu_{i2} \psi_{1,i}^2(x) \hat{\phi}_{1,i}^2(y) + \sum_i \nu_{i2} \psi_{1,i}^2(x) \hat{\psi}_{1,i}^2(y) \end{aligned}$$

5. *Choose a coding technique to find the wavelet detail part:*

$$g_1(x, y) = \sum_{j=1}^2 \left(\sum_i \lambda_{ij} \phi_{1,i}^j(x) \hat{\psi}_{1,i}^j(y) + \sum_i \mu_{ij} \psi_{1,i}^j(x) \hat{\phi}_{1,i}^j(y) + \sum_i \nu_{ij} \psi_{1,i}^j(x) \hat{\psi}_{1,i}^j(y) \right)$$

and the scaling approximation part of the image/sub-image:

$$f^1(x, y) = \sum_{j=1}^2 \sum_i \alpha_{ij} \phi_{1,i}^j(x) \hat{\phi}_{1,i}^j(y).$$

6 If the approximation and compression meet the criteria, stop; else, return to step 2.

5 General Cases and Image Compression Example

In the following, we show an example of image compression using spline-wavelet on refined grids, for a woman face image in high resolution (1600×1200 pixels). First, we build a rough grid and on this grid we get a low resolution image shown in Figure 1 (a) (200×150 pixels). In the “Face Recognition”, the eyes and mouth are important parts which should be described in details, specially at canthus and the corners of the mouth. We put our six refined grids on those parts (Figure 1 (b)) which uniformly refine those regions to 32×32 grids.

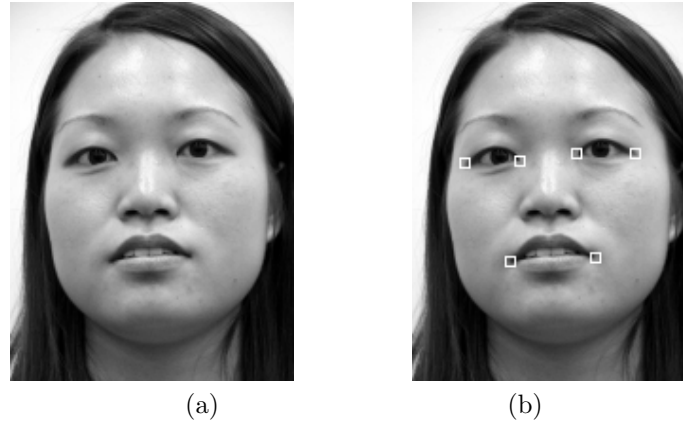


Figure 1: The woman face image with a rough grid (a) and refined grids (b).

On refined grids, we sample the original image and obtain details using spline-wavelet on refined grids. The two sets of data combined together form a new image file with compression.

In this example, there are six parts that need to be refined. We apply the algorithm and finish them one by one.

Only the necessary details have been reserved. So the image's storage space is observably reduced.

We can compare the details in rough grids (Figure 2 (a)) and in refined grids (Figure 2 (b)) for the woman's right eye's outboard canthus (marked in Figure 1 (b)). The information brought in by addition spline-wavelet improves local image quality and the two sets of data can be quantized in different resolutions, respectively.

Acknowledgments

This work is supported by Guangdong Provincial Natural Science Foundation of China (036608) and the Foundation of Zhongshan University Advanced

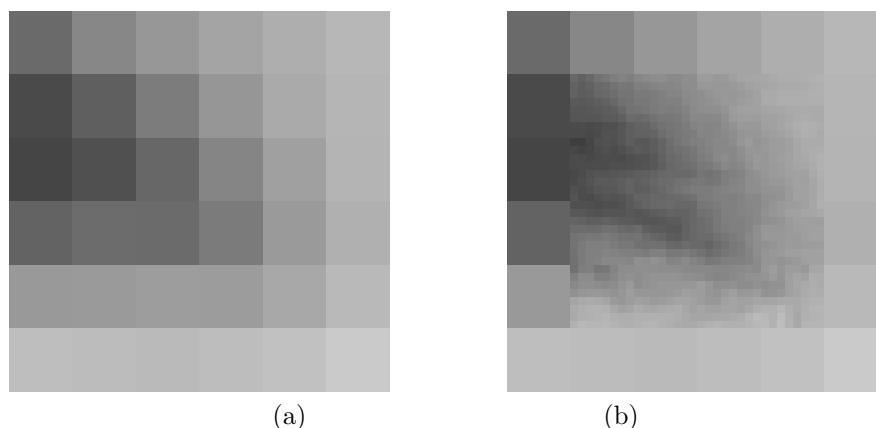


Figure 2: Details over rough grids (a) and details over refined grids (b).

Research Center, Hong Kong. The authors would like to thank the editors, professors Don Hong and Yu Shyr of this special issue, and anonymous referees for a number of comments which led to important improvements to this paper.

References

- [1] C. de Boor, *A Practical Guide in Splines(second edition)*. Springer, New York, 2002.
- [2] C.K. Chui, *An Introduction to Wavelets*, Academic Press, Boston, 1992.
- [3] C.K. Chui and Lutai Guan, Multivariate polynomial natural splines for interpolation of scattered data and other application, in *Workshop on Computational Geometry* (A.Conte et al. eds.), World Scientific, Singapore, (1993).
- [4] I. Daubechies, *Ten Lectures on Wavelets*, Philadelphia, CBMS–NSF Series in Appl.Math. (SIAM), 1991.
- [5] Lutai Guan, Bivariate polynomial natural spline interpolation algorithms with local basis for scattered data, *J. of Computational Analysis and Appl.*, **5** (2003), 77–100.
- [6] Lutai Guan and Bin Liu, On Bivariate N-spline Interpolation Local Basis for Scattered Data over Refined Grid Points, *Mathematica Numerica Sinica* **25** (2003), 375–384.
- [7] D. Hong and L.L. Schumaker, Surface Compression Using a Space of C^1 Cubic Splines with a Hierarchical Basis, *Computing*, **72** (2004), 79-92.

- [8] S. Mallat, *A Wavelet Tour of Signal Processing* (Second Edition), Academic Press, Boston, 1999.
- [9] Y. Meyer, *Wavelets and Operators*, Cambridge University Press, 1992.
- [10] N.T. PanKaj, *Wavelets Image and Compression*, Boston, Klumer Academic Publisher, 1998.
- [11] D.F. Walnut, *An Introduction to Wavelet Analysis*, Boston, Birkhäuser, 2002.

Variational method and Wavelet Regression

Jianzhong Wang

Department of Mathematics and Statistics

Sam Houston State University

Huntsville, TX 77341

email: mth_jxw@shsu.edu

Abstract

Wavelet regression is a new nonparametric regression approach. Compared with traditional methods, wavelet regression has the advantages of ideal minimax property, spatial inhomogeneous adaptivity, high dimension expansibility and fast algorithm. In this paper, we provide a brief review of wavelet shrinkage. We also apply the variational method for wavelet regression and show the relation between the variational method and wavelet shrinkage.

Keywords. Wavelets, Orthonormal wavelets, Wavelet regression, Variational Methods.

1 Introduction

Regression analysis is an important statistical technique for investigating and modeling the relationship between variables. The goal of regression is to get a model of the relationship between one variable y and one or more variables t . Let $y = \tilde{f}(t)$ be the *observation* (also called the *observed function*), which responds to the regressor t :

$$y_i(= \tilde{f}(t_i)) = f(t_i) + \varepsilon_i \quad i = 1, 2, \dots, N \quad (1)$$

where $t_i, i = 1, 2, \dots, N$, are sampling points, $f(t_i)$ are values of an unknown function $f(t)$ (called the *underlying function*), and ε_i are the statistical errors (called *noise*), which are often formulated as independent random variables. In this paper, we assume the sampling points are equal-spaced between 0 and 1. The goal of a regression is to estimate the underlying function $f(t)$ from the sampling data $(t_i, y_i)_{i=1}^N$ (called the *observed data*), i.e., to establish a method to find the estimator $\hat{f}$, which achieves the minimal risk

$$R(\hat{f}, f) = N^{-1} E \|\hat{f} - f\|_{2,N}^2.$$

The classical nonparametric regression methods mainly include orthogonal series estimators, kernel estimators and smoothing splines. Theoretical discussions about these methods can be found in [17] and [35]. Recently, a variety

of spatially adaptive methods has been developed in the statistical literature. Among them, wavelet adaptive methods have received great attention. References are given in Bock and Pliego [5], Vidakovic [36], Ogden [33], Donoho and Johnstone [13], [14], [15], [16], and their references. In wavelet regression, an important issue is to select a subset of wavelet coefficients, which well represent the underlying function f while removing most of noise.

Since 1980s, developing numerical methods for removing noise from an image while preserving edge also became an active research area. Mumford and Shah [28] proposed an energy functional for images so that the image processing such as noise removal and image segmentation can be formulated as a variational problem associated with the functional. Perona and Malik [34] proposed an anisotropic diffusion model for removing noise while enhancing edge. It is proven that the anisotropic diffusion equation in [34] was the steepest decent method for solving the variational problem associated with a certain energy functional [38]. References in this aspect can be found in the book [3].

The goal of this paper is to apply the variational method in the wavelet regression. The main idea is the following. Since the underlying function f is unknown, we cannot obtain the formula for $R(f, f)$. Hence to directly minimize $R(\hat{f}, f)$ is impossible. We shall create a substitution of the risk functional, say $GR(\hat{f}, \tilde{f})$, in which the underlying function is not involved. Thus, we can use variational method to find $\hat{f}$. The wavelet representation of data will enable us to establish the variational problems with a very simple structure. We outline the paper as follows. In Section 2, we briefly introduce orthonormal wavelet bases and adaptive wavelet regression (also called wavelet shrinkage). In Section 3, we describe the variational model for wavelet regression and establish the relation between it and wavelet shrinkages. In Section 4, we give some examples.

2 Preliminarily

2.1 Orthonormal Wavelet Bases

The regular orthonormal wavelet basis of $L^2(R)$ is constructed from a multiresolution analysis (MRA) (see [24], [25], [26], [22], and [23]).

Definition 1 A multiresolution analysis of $L^2(R)$ is a sequence $(V_j)_{j \in \mathbb{Z}}$ of closed subspaces of $L^2(R)$ such that the followings hold:

- (1) $V_j \subset V_{j+1} \quad j \in \mathbb{Z}$
- (2) $\overline{\cup_{j \in \mathbb{Z}} V_j} = L^2(R)$ and $\cap_{j \in \mathbb{Z}} V_j = \{0\}$
- (3) $f(x) \in V_j \iff f(2x) \in V_{j+1}$
- (4) $f(x) \in V_j \implies f(x - 2^{-j}k) \in V_j \quad j, k \in \mathbb{Z}$
- (5) There exists a function $\phi \in V_0$ such that $\{\phi(x - n)\}_{n \in \mathbb{Z}}$ forms an unconditional basis of V_0 , i.e., $\{\phi(x - n)\}_{n \in \mathbb{Z}}$ is a basis of V_0 and there exist two constants, $A, B > 0$ such that, $\forall (c_n) \in l^2$, the following inequality holds

$$A \sum |c_n|^2 \leq \| \sum c_n \phi(\cdot - n) \|^2 \leq B \sum |c_n|^2.$$

The function ϕ in Definition 1 is called an MRA generator. Furthermore, if $\{\phi(x - n)\}_{n \in \mathbb{Z}}$ is an orthonormal basis of V_0 , then ϕ is called an *orthonormal MRA generator* (also called an *orthonormal scaling function*). Assume ϕ satisfies the two-scale equation

$$\phi(x) = 2 \sum_{k \in \mathbb{Z}} h(k) \phi(2x - k), \quad (h(k))_{k \in \mathbb{Z}} \in l^2. \quad (2)$$

Define the wavelet function ψ by

$$\psi(x) = 2 \sum_{k \in \mathbb{Z}} (-1)^k h(1 - k) \phi(2x - k). \quad (3)$$

For a function f , we write

$$f_{j,k}(x) = 2^{\frac{j}{2}} f(2^j x - k), \quad j, k \in \mathbb{Z}. \quad (4)$$

Then $\{\phi_{j,k}\}_{k \in \mathbb{Z}}$ forms an orthonormal basis of V_j . Let $W_j \subset L^2(R)$ be the subspace spanned by $\{\psi_{j,k}\}_{k \in \mathbb{Z}}$. Then W_j is called the *wavelet subspace* at level j . It is clear that $W_j \perp V_j$ and $W_j \oplus V_j = V_{j+1}$. Hence, $\bigoplus_{j=-\infty}^{\infty} W_j = L^2(R)$ and $\{\psi_{j,k}\}_{j,k \in \mathbb{Z}}$ forms an orthonormal basis of $L^2(R)$, called a *standard orthonormal wavelet basis*. Recall that we also have

$$L^2(R) = V_{J_0} \oplus (\bigoplus_{j \geq J_0} W_j).$$

Hence, $\{\phi_{J_0,k}\}_{k \in \mathbb{Z}} \cup \{\psi_{j,k}\}_{j \geq J_0, k \in \mathbb{Z}}$ also forms an orthonormal basis of $L^2(R)$, which is called a *hybrid orthonormal wavelet basis*. Using this basis we can decompose a function $f \in L^2(R)$ into the following form

$$\begin{aligned} f &= f_{J_0} + \sum_{j \geq J_0} g_j, \quad f_{J_0} \in V_{J_0}, g_j \in W_j, \\ f_{J_0} &= \sum a_{J_0,k} \phi_{J_0,k}, \quad g_j = \sum w_{j,k} \psi_{j,k}. \end{aligned}$$

Particularly, a function $f \in V_J$ has the decomposition

$$f = f_{J_0} + \sum_{j=J_0}^{J-1} g_j, \quad f_{J_0} \in V_{J_0}, g_j \in W_j.$$

Hybrid orthonormal wavelet bases are very useful in many applications because in the decomposition above f_{J_0} provides an estimator of f while $g_j, j = J_0, \dots, J-1$, preserve the details of f at different levels (see Figures 1-3). For more discussions of wavelet bases, we refer to [12] and [24].

Daubechies ([11], [12]) created compactly supported orthonormal wavelet bases, which are very useful in wavelet regression. In regression the sample data are finite. In order to deal with finite data, we need the multiresolution analysis, scaling functions, wavelets, and wavelet bases on intervals. The details of these modifications can be found in [10]. In wavelet regression, the wavelet bases on

final intervals are used in the following way. Let $\{V_j\}_{j=0}^\infty$ be the MRA defined on the interval $[0, 1]$ and $\{W_j\}_{j=0}^\infty$ be the corresponding wavelet subspaces. For a fixed J , we assume $(\phi_{J,k})$ forms an orthonormal basis of V_J , where $\phi_{J,k}$ has been modified from its original version $2^{J/2}\phi(2^Jx - k)$ as in [10],[13]. After the modification, the subscript (J, k) indicates the spatial location of $\phi_{J,k}$. A function $f \in V_J$ is now decomposed to

$$f_J(t) = \sum a_{J,k} \phi_{J,k}(t), \quad a_{J,k} = \langle f, \phi_{J,k} \rangle. \quad (5)$$

where the sum in (5) is a finite one. Similarly, we have $V_J = V_{J_0} \oplus_{j=J_0}^{J-1} W_j$ and

$$f_J(t) = f_{J_0}(t) + \sum_{j=J_0}^{J-1} g_j(t) = \sum a_{J_0,k} \phi_{J_0,k}(t) + \sum_{j=J_0}^{J-1} \sum w_{j,k} \psi_{j,k}(t), \quad (6)$$

where $a_{J_0,k} = \langle f_J, \phi_{J_0,k} \rangle$ and $w_{j,k} = \langle f, \psi_{j,k} \rangle$.

For convenience, we write $\mathbf{a}_J = (a_{J,k})$, $\mathbf{a}_{J_0} = (a_{J_0,k})$, $\mathbf{w}_j = (w_{j,k})$, $\mathbf{v}_a = [\mathbf{a}_{J_0}, \mathbf{w}_{J_0}, \dots, \mathbf{w}_{J-1}]$, and denote the transform from $\mathbf{a}_J$ to $\mathbf{v}_a$ by $\mathbf{W}$:

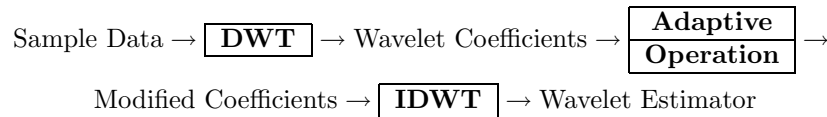
$$\mathbf{v}_a = \mathbf{W} \mathbf{a}_J.$$

We call $\mathbf{W}$ the *discrete wavelet transform* (DWT) and call its inverse $\mathbf{W}^{-1}$ the *inverse discrete wavelet transform* (IDWT). It is easy to see that when ϕ is an orthonormal MRA generator and ψ is the corresponding orthonormal wavelet, $\mathbf{W}$ is an orthonormal matrix. It follows that $\mathbf{W}^{-1} = \mathbf{W}^T$. Without loss of generality, later we shall always assume $J_0 = 0$.

Based on pyramidal algorithms, Mallat in [18] developed fast algorithms for computing DWT and IDWT, which are also called Mallat's algorithms. They have the computational complexity of $O(n \log n)$, where $n = \text{length}(\mathbf{a}_J)$.

2.2 Adaptive Wavelet Regression

We now return to the model (1). Since the sample data are finite, without loss of generality, we assume N in (1) is equal to the dimension of the space V_J for a certain level J . (If it is not, we can extend the sample data to let it be.) Then we can identify the data in (1) to a function of the space V_J . Wavelet regression adopts a wavelet estimator to estimate the underlying function. The estimator is select from a subspace of V_J and the selection is dependent on the sample data. Hence, wavelet regression usually is a nonlinear regression. For the given sample data in (1), the wavelet regression can be outlined in the following diagram.



Let $\mathbf{W}$ be the DWT. Let $\mathbf{T}_a : R^n \rightarrow R^n$ be the adaptive operator, which changes wavelet coefficients for a certain purpose. Then the algorithm that performs wavelet regression can be written as follows.

$$\begin{aligned}\mathbf{v}_y &= \mathbf{W}\mathbf{y}, \\ \hat{\mathbf{v}}_y &= \mathbf{T}_a \mathbf{v}_y, \\ \hat{\mathbf{y}} &= \mathbf{W}^{-1} \hat{\mathbf{v}}_y.\end{aligned}$$

We write

$$\mathcal{W} = \mathbf{W}^T \mathbf{T}_a \mathbf{W}.$$

Then $\mathcal{W}$ denotes the wavelet regression operator. If the adaptive operator $\mathbf{T}_a$ is a compression one, then the adaptive wavelet regression is a *wavelet shrinkage* since the number of non-vanished entries of $\hat{\mathbf{v}}_y$ is smaller than that one of $\mathbf{v}_y$.

The wavelet regression is based on the following facts.

- The orthonormal transform of white noise $\sim N(0, \sigma^2)$ is still a white noise. $\sim N(0, \sigma^2)$. (See [13].)
- If f is a noise-free function, then most of its wavelet coefficients vanish, only a very few of them have non-neglected values, which represent the details of the function.
- If a function f carries noise, then its smooth component $f_{J_0}(t)$ in (6) is not influenced by noise very much, while the wavelet components keep noise.

Figures 1 and 2 show the wavelet decompositions of noise-free functions, and Figures 3 and 4 show them of noisy functions.

By these properties, in wavelet regression, the adaptive operator $\mathbf{T}_a$ is selected so that it does not change $\mathbf{a}_0$, i.e.,

$$\mathbf{T}_a = \begin{bmatrix} \mathbf{I} & \mathbf{0} \\ \mathbf{0} & \mathbf{T} \end{bmatrix},$$

where $\mathbf{T}$ is the operator on $W := \oplus_{j=0}^{J-1} W_j$. Let $\mathbf{w} = [\mathbf{w}_0, \dots, \mathbf{w}_{J-1}]$ and assume the length of $\mathbf{w}$ is n . We simply denote $\mathbf{w} = (w_i)_{i=1}^n$. Thus,

$$\mathbf{w} = \theta + \mathbf{z}, \quad w_i = \theta_i + z_i, \quad (7)$$

where θ_i is the wavelet coefficient of the underlying function f and z_i is the wavelet transform of ϵ_i . If ϵ_i is assumed to be the white noise $\sim N(0, \sigma^2)$, so is z_i . Then the wavelet regression is essentially to find a $\mathbf{w}$ that minimizes the risk

$$R(\mathbf{w}, \theta) = \frac{1}{n} E(\|\mathbf{w} - \theta\|^2). \quad (8)$$

Since small wavelet coefficients mostly contribute to noise while large ones to signal, wavelet regression removes the small wavelet coefficients from $\mathbf{w}$.

Donaho and Johnstone in [13] used two different shrinkages to design $\mathbf{T}$. For a given threshold $\lambda > 0$, the *hard threshold* function is

$$\eta_h(x; \lambda) = \begin{cases} 0, & |x| \leq \lambda \\ x, & |x| > \lambda \end{cases}, \quad (9)$$

and the *soft threshold* function is

$$\eta_s(x; \lambda) = \begin{cases} 0, & |x| \leq \lambda \\ \text{sign}(x)(|x| - \lambda), & |x| > \lambda \end{cases}. \quad (10)$$

For a vector $\mathbf{w}$, we write $\eta_h(\mathbf{w}; \lambda) = (\eta_h(w_i; \lambda))$ and $\eta_s(\mathbf{w}; \lambda) = (\eta_s(w_i; \lambda))$ respectively. Then the adaptive operators corresponding to $\eta_h(\mathbf{w}; \lambda)$ and $\eta_s(\mathbf{w}; \lambda)$ are $\mathbf{T}^h \mathbf{w} = \eta_h(\mathbf{w}; \lambda)$ and $\mathbf{T}^s \mathbf{w} = \eta_s(\mathbf{w}; \lambda)$ respectively.

The choice of the threshold λ is a fundamental issue in wavelet shrinkage. A large threshold cuts too many coefficients and will result in an oversmoothing estimator. Conversely, a small threshold does not remove noise well and will produce a wiggly, undersmoothing estimator. The proper threshold ought to take a careful balance. A lot of work have been done in this aspect. (See [1], [2], [7], [12], [13], [14], [15], [16], [29], [30], [31], and [32].)

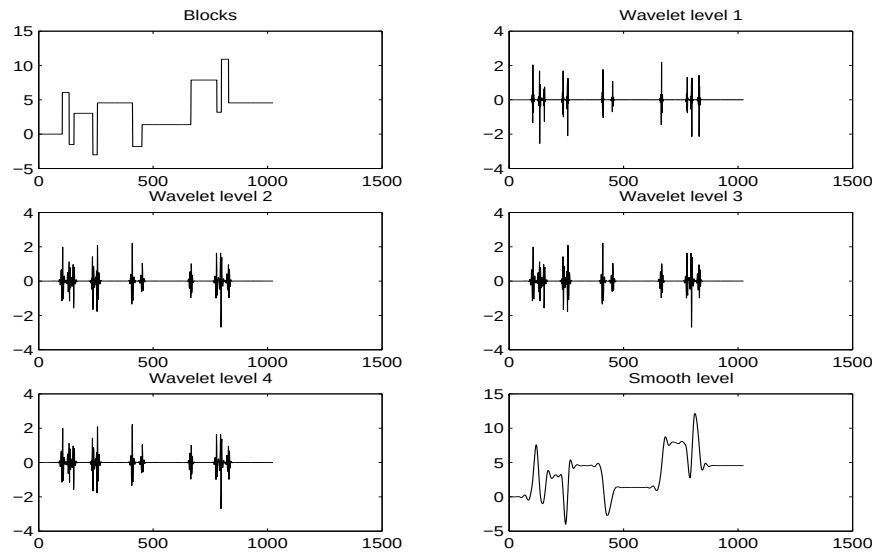


Figure 1. Wavelet decomposition of Block. Up to Level 4.

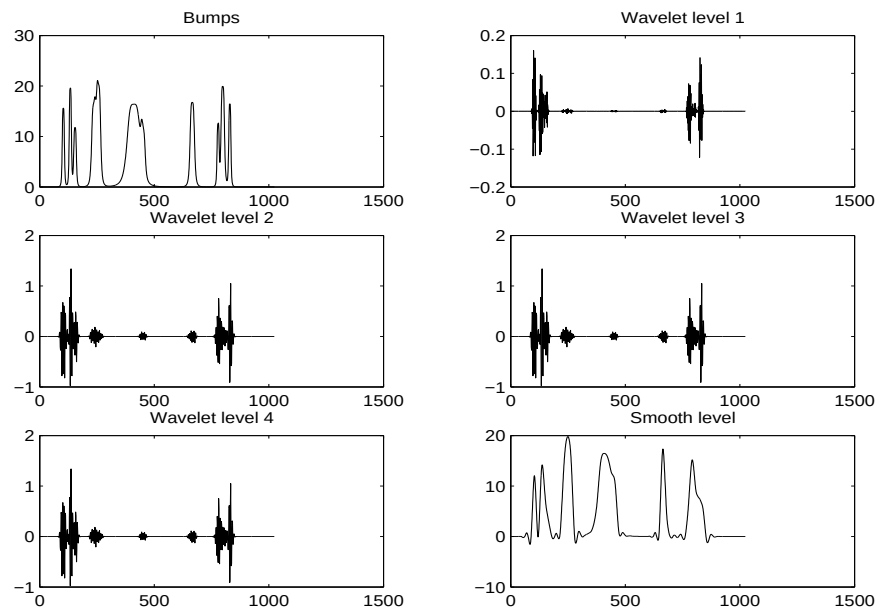


Figure 2. The wavelet decomposition of Bumps.

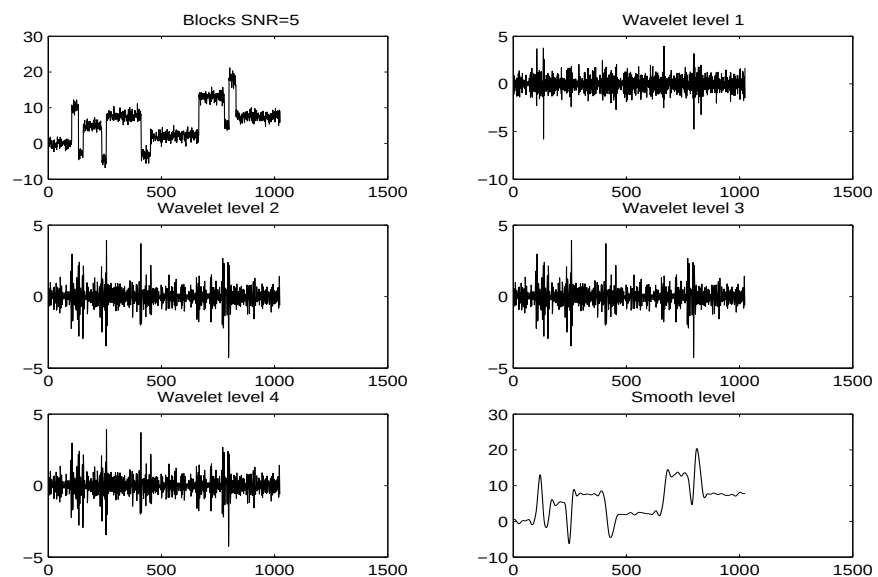


Figure 3. Wavelet decomposition of Blocks with noise.

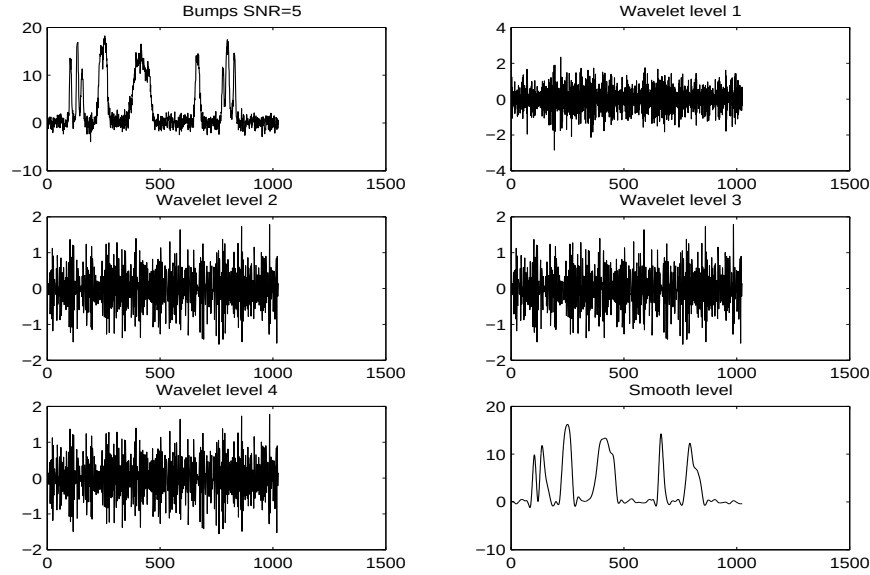


Figure 4. Wavelet decomposition of Bumps with noise.

3 Variational Method for Wavelet Regression

Wavelet shrinkage is based on the principle that each wavelet coefficient whose absolute value is smaller than the threshold contributes to noise; otherwise it contributes to the signal. The threshold in wavelet shrinkage is used to balance the oversmoothing and undersmoothing. The value of the threshold is dependent of the noise level, the size of the sampling data, and the smoothness required by the regression. The authors of [13] and [14] established several rules for determining thresholds, including Universal thresholding, MiniMax thresholding, SureShrink, Heursure and so on. (See also [15], [16], [29], [30].) These methods are very effective when noise is white and its level is known. However, in many applications, the noise type is unknown and the noise level is not uniform over all sampling areas. It is necessary to study wavelet regression in a general framework. By definition, regression is the minimization of the risk function, where the key issue is to remove noise while keeping the features of the data. Hence, we discuss the variational method for wavelet regression and reveal the relation between the variational method and the wavelet shrinkage.

Following the idea of [28], to balance the smoothness and the sharpness of an estimator, we introduce two energy functions. The first energy measures the distance between the estimator and the sampling data. It is clear that the estimator should be close to the sample data. Recall that in wavelet regression, we do not change the function $f_{J_0} \in V_{J_0}$. Hence, the distance can be measured

by the following *wavelet deviation energy*

$$A(\mathbf{u}) = \|\mathbf{u} - \mathbf{w}^0\|^2, \quad (11)$$

where $\mathbf{w}^0$ represents the vector of wavelet coefficients of the observation. By the properties of wavelets, $A(\mathbf{u})$ controls the undersmoothing. Since $\mathbf{w}^0$ carries noise, the estimator will be undersmoothing if $A(\mathbf{u})$ is close to 0. On the other hand, the wavelet components of a function represent its wiggle and noise. Therefore, we can create a weighted wavelet energy to measure the wiggle together with the noise for data $\mathbf{u}$:

$$F(\mathbf{u}) = \mathbf{u}^T G \mathbf{u}, \quad (12)$$

where $G(\mathbf{u})$ is a positive definite (or semi-definite) matrix. Under the assumption that noises are independent random variables, we may assume G is a diagonal matrix, in which each entry on the diagonal is a weight function of u_i . Recall that small wavelet coefficients contribute to noise, and large ones contribute to signal. For the purpose of regression, large weights ought to assign to small coefficients and vice versa. Hence, we have the following.

Definition 2 Let G be a nonnegative diagonal matrix defined by

$$G = \text{diag}(c_i(u_i))_{i=1}^n. \quad (13)$$

where each c_i is an even function satisfying the conditions: (1) $c_i(s) \geq 0$ and $\lim_{s \rightarrow \infty} c_i(s) = 0$; (2) $c_i(|s|)$ is decreasing on $[0, \infty)$. Then

$$F(\mathbf{u}) = \mathbf{u}^T G \mathbf{u}$$

is called the *weighted wavelet energy* and G is called the *weight matrix*.

Since G is a diagonal matrix, $F(\mathbf{u})$ can be simplified to the quadratic form $\sum_{i=1}^n s^2 c_i(s)$. Write $g_i(s) = s^2 c_i(s)$. Then $(g_i(s))$ is the *energy density*. The weighted wavelet energy controls oversmoothing. If $F(\mathbf{u}) = 0$, then all wavelet components vanish, which indicates that the estimator is in V_{J_0} , i.e., the regression will cause oversmoothing.

We now combine $G(\mathbf{u})$ and $A(\mathbf{u})$ to make an energy function

$$\Gamma(\mathbf{u}) = G(\mathbf{u}) + \lambda A(\mathbf{u}), \quad (14)$$

where λ is the parameter used to balance oversmoothing and undersmoothing. Thus, wavelet regression can be formulated to the following variational problem: Find $\mathbf{w}$ such that

$$\mathbf{w} = \arg(\min \Gamma(\mathbf{u})). \quad (15)$$

Assume each $c_i(s)$ is differentiable. Then the solution of (15) satisfies the Euler-Lagrangian equation

$$\mathbf{g}'(\mathbf{w}) + 2\lambda(\mathbf{w} - \mathbf{w}^0) = 0,$$

where $\mathbf{g}'(\mathbf{w}) = (g'_i(w_i))$. It follows that

$$2\lambda\mathbf{w} + \mathbf{g}'(\mathbf{w}) = 2\lambda\mathbf{w}^0, \quad (16)$$

i.e.,

$$2\lambda w_i + g'(w_i, w_i^0) = 2\lambda w_i^0, \quad i = 1, \dots, n. \quad (17)$$

If the matrix G is not a diagonal one, then a nonlinear matrix equation will replace (17).

We now establish the relation between the variational method and the wavelet shrinkage. A general linear wavelet shrinkage is formulated as

$$\Gamma_{DS}(w_i, \delta_i) = (\delta_i w_i)_{i=1}^n, \quad \delta_i \in [0, 1].$$

We show that if in $F(\mathbf{u}) = \sum g_i(u)$, each density g_i is linear on $[0, \infty)$, then the solution of the variational problem (15) leads to a linear shrinkage.

Theorem 3 *If a variational problem (15) satisfies*

$$g_i(s) = \mu_i |s|, \quad \mu_i \geq 0 \quad (18)$$

then it yields a linear wavelet shrinkage.

Proof. Without loss of generality, we can assume the balance parameter λ in $\Gamma(\mathbf{u})$ is $1/2$. Otherwise, we use $\lambda\mu_i/2$ to replace μ_i . We have

$$\rho'_i(s) = \mu_i \operatorname{sgn}(s),$$

which leads to the Euler-Lagrangian equation

$$s + \mu_i \operatorname{sgn}(s) = w_i^0, \quad i = 1, \dots, n. \quad (19)$$

Since $\rho'_i(0)$ does not exist, 0 is also a critical point. When $|w_i^0| < \mu_i$, the equation (19) has no solution; and when $|w_i^0| \geq \mu_i$, the equation has the unique solution

$$s = \operatorname{sgn}(w_i^0) (|w_i^0| - \mu_i). \quad (20)$$

Then it is easy to verify that the vector $\mathbf{w} = (w_i)$ with

$$w_i = \operatorname{sgn}(w_i^0) (|w_i^0| - \mu_i)_+ = \frac{(|w_i^0| - \mu_i)_+}{|w_i^0|} w_i^0, \quad i = 1, \dots, n, \quad (21)$$

minimizes $\Gamma(\mathbf{u})$, where $x_+ = \max(x, 0)$. Let $\delta_i = \frac{(|w_i^0| - \mu_i)_+}{|w_i^0|}$. Then $\delta_i \in [0, 1]$. The theorem is proven. $\square$

As applications of Theorem 3, we discuss the wavelet shrinkages by using hard threshold and soft threshold respectively.

Example 1. [Hard thresholding] Let δ be a given threshold. We assume in $\mathbf{w}^0$ each sample value whose absolute value is less than δ represents noise.

Otherwise it does not carry on noise. The assumption suggests the following weight:

$$c_i(s) = \begin{cases} \delta/|s|, & |w_i^0| \leq \delta, \\ 0, & |w_i^0| > \delta, \end{cases} \quad (22)$$

which leads

$$\rho_i(s) = \begin{cases} \delta|s|, & |w_i^0| \leq \delta, \\ 0, & |w_i^0| > \delta. \end{cases}$$

and

$$\rho'_i(s) = \begin{cases} \delta \operatorname{sgn} |s|, & |w_i^0| \leq \delta, \\ 0, & |w_i^0| > \delta. \end{cases}$$

By (20), the solution $\mathbf{w}$ of the variational problem (15) is

$$w_i = \begin{cases} (1 - \frac{1}{2\lambda})_+ w_i^0, & |w_i^0| \leq \delta, \\ w_i^0, & |w_i^0| > \delta. \end{cases} \quad (23)$$

Hence, when $\lambda \leq \frac{1}{2}$, the solution provides the wavelet shrinkage using the hard threshold δ .

Example 2. [Soft thresholding] We now assume in $\mathbf{w}^0$ each sample value carries on noise $\sim N(0, \sigma^2)$. Hence, we adopt the following weight

$$c_i(s) = \delta/|s|,$$

which leads $\rho_i(s) = \delta|s|$ and therefore

$$\rho'_i(s) = \delta \operatorname{sgn}(s).$$

By (21), the solution $\mathbf{w}$ of the variational problem (15) is

$$w_i = \operatorname{sgn}(w_i^0) \left(|w_i^0| - \frac{1}{2\lambda} \delta \right)_+,$$

which provides the wavelet shrinkage using the soft threshold $\delta/(2\lambda)$.

4 Examples

In this section, we give several examples to illustrate the function of λ in the wavelet regression. We select four functions to test the regression: “Blocks” is a step function, which has several step jumps. According to [20] and [21], Lipschitz order of a step jump is 0. “Heavy Sine” is a broken sine wave, which has two jumps. “Bumps” is a continuous function but has a lot of non-differentiable points, which are classified to be with Lipschitz orders in $(0, 1)$. “Doppler” is smooth everywhere except at the origin, where the function has infinite oscillations. The white noise is added to the sample data of them. Let $\sigma(f)$ and $\sigma(n)$

denote the standard deviation of the function and the noise respectively. In Figure 5 – Figure 8, we set $\frac{\sigma(f)}{\sigma(n)} = 5$. In the regression, we set $\delta = \sqrt{2\sigma(n) \ln N}$, where N is the size of the data. We also set $\lambda = 2, 1, 1/2, 1/4, 1/10$, which yield the thresholds $\delta/2, \delta, 2\delta, 4\delta, 10\delta$ respectively. In Figure 9, we choose $\lambda = 2$, which yield the threshold 4δ . The results show that the regression is not very sensitive to small λ , say $\lambda \leq 1/2$, but sensitive to $\lambda > 1$. To regress the functions with singularity as in “Doppler”, a larger λ should be chosen. Other tests show that the regression does not very much rely on the choices of particular wavelets. (The figures are not included in this paper.) Using different wavelets in a regression causes little differences, provided that the size of sampling data is large (say $\geq 2^{10}$ when wavelet level is ≤ 4).

When the sampling data carry on the noise with uniform distribution, the wavelet regression still works well. Figure 10 – Figure 13 show the wavelet regression of the four functions with uniform noise. Again, the estimator of “Doppler” produces a larger error.

Acknowledgment. The author thanks the referees for their valuable comments.

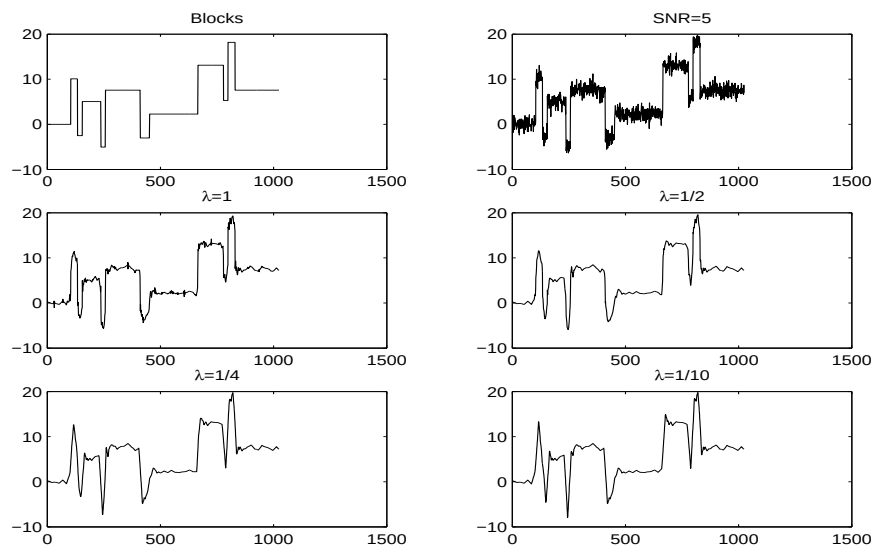


Figure 5. Wavelet regression for Blocks.

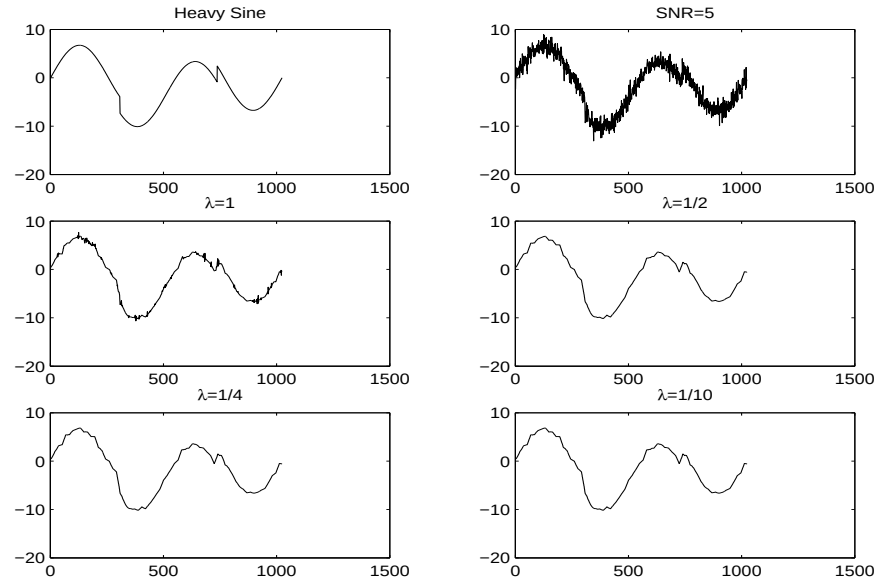


Figure 6. Wavelet regression for Heavy Sine.

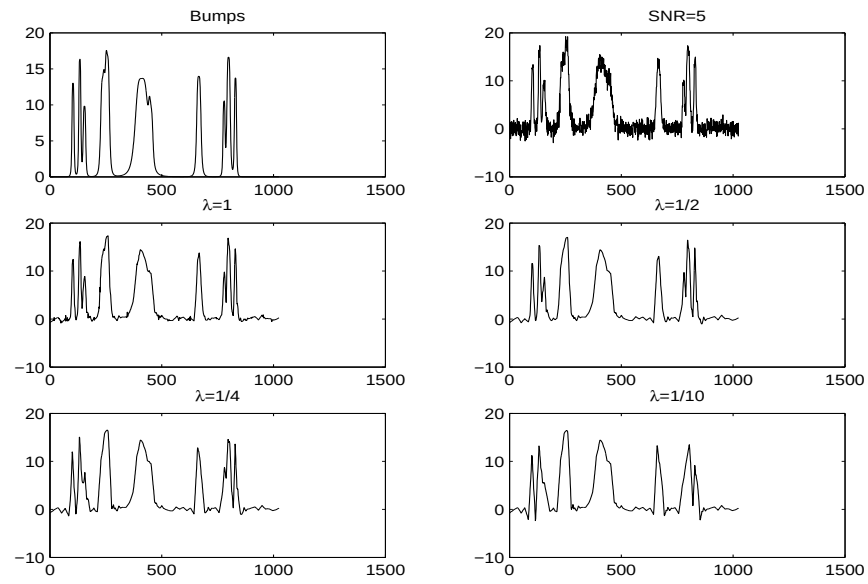


Figure 7. Wavelet regression for Bumps.

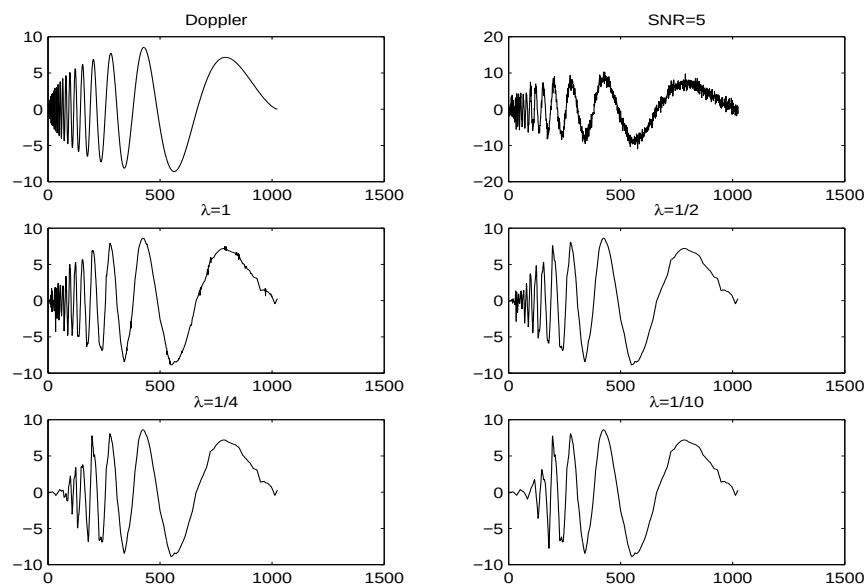
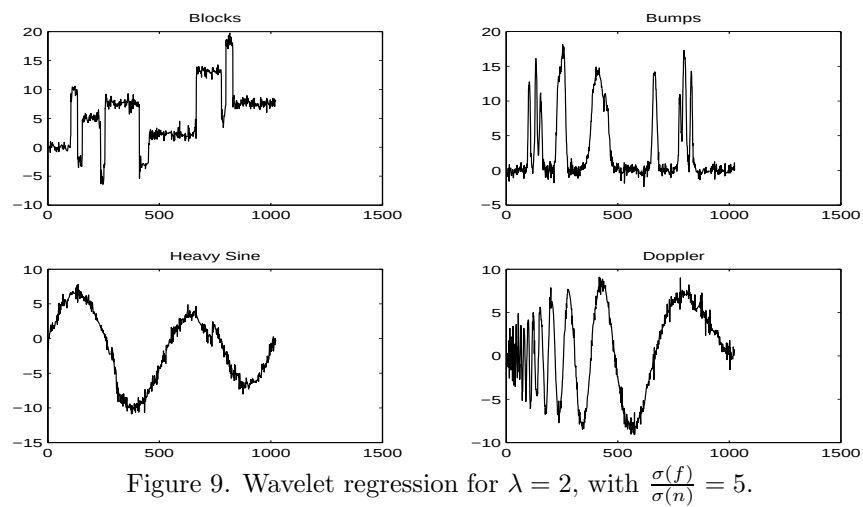


Figure 8. Wavelet regression for Doppler.

Figure 9. Wavelet regression for $\lambda = 2$, with $\frac{\sigma(f)}{\sigma(n)} = 5$.

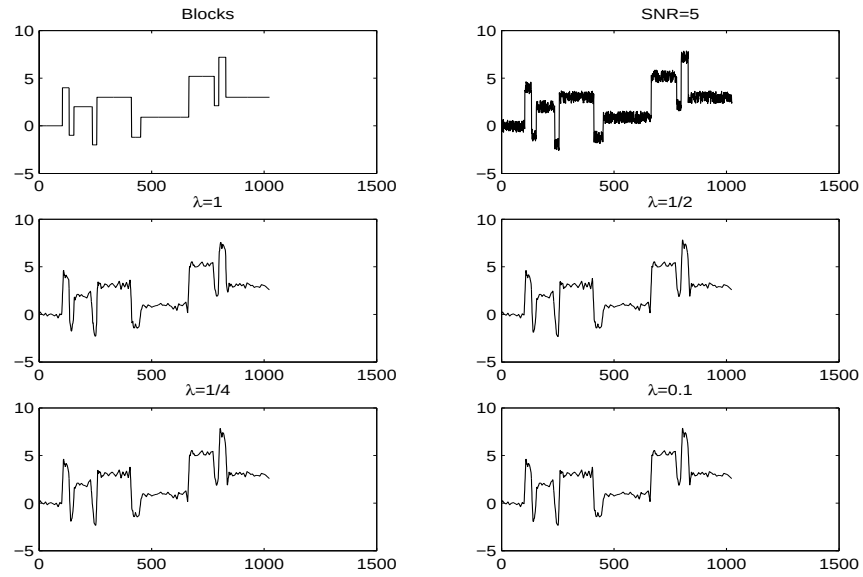


Figure 10. Wavelet regression of Block with uniform noise.

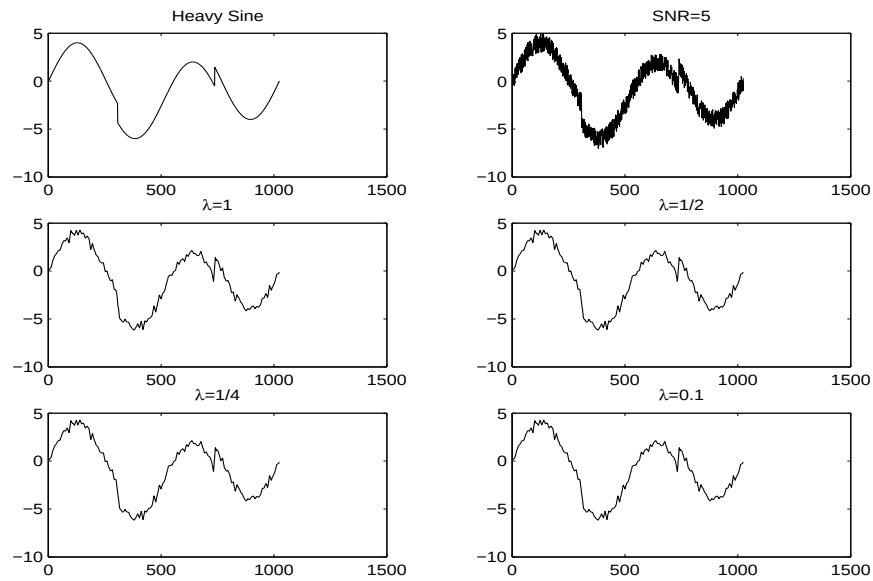


Figure 11. Wavelet regression of Heavy Sine with uniform noise.

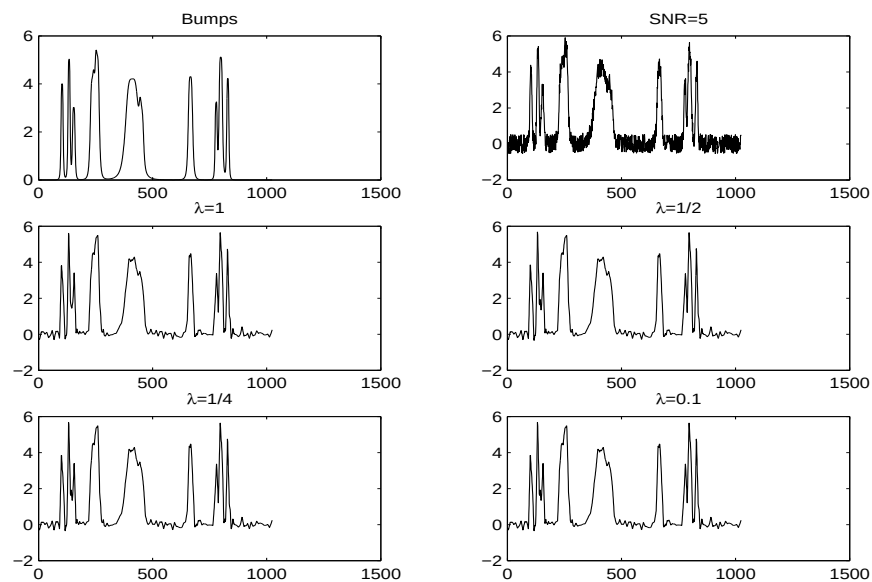


Figure 12. Wavelet regression of Bumps with uniform noise.

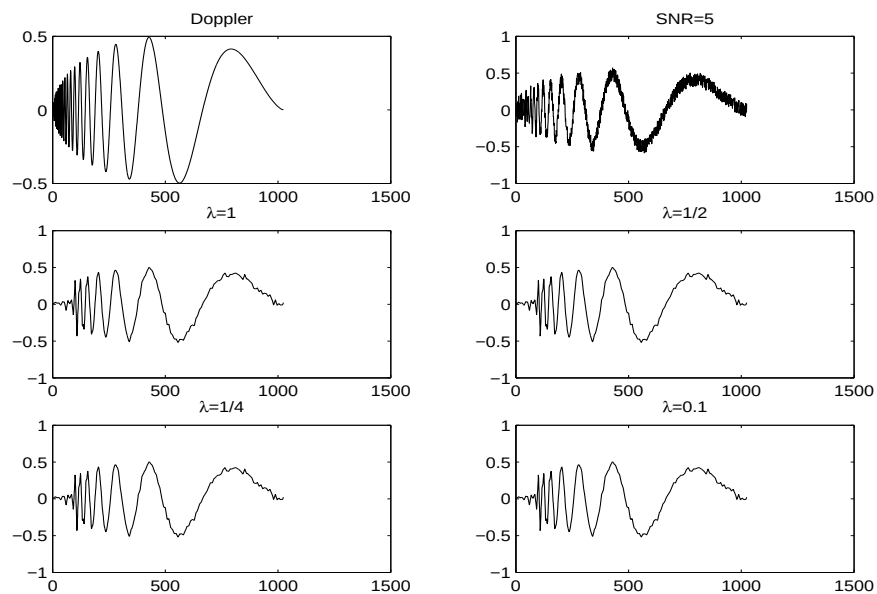


Figure 13. Wavelet regression of Doppler with uniform noise.

Acknowledgments: This work was supported in part by ER-Grant 2003 of Sam Houston State University.

References

- [1] Abramovich, F., Sapatinas, F., and Silverman, B.W, "Wavelet thresholding via a Bayesian approach", *J. Roy. Statist. Soc. B*, **60**, pp 725-749, 1998.
- [2] Abramovich, F. and Benjamini, Y., "Thresholding of wavelet coefficients as multiple hypotheses testing procedure", *Wavelets and Statistics*, (Eds. Antoniadis A. & Oppenheim G.), New York: SpringerVerlag, pp 5-14, 1995.
- [3] Aubert, G. and Kornprobst, P., "Mathematical Problems in Image Processing—Partial Differential Equations and the Calculus of Variations", Springer, 2001.
- [4] Aubert, G. and Vese, L., "A variational method in image recovery," *SIAM Journal of Numerical Analysis*, **34**(5): pp 1948-1979, 1997.
- [5] Bock, M.E. and Pliego, G., "Estimating functions with wavelets part II: Using a Daubechies wavelet in nonparametric regression", *Statistical Computing and Statistical Graphics Newsletter*, **3** (2): pp 27-34, November 1992.
- [6] Chan, A.K. and Goswami, J.C., "Fundamentals of wavelets: theory, algorithms, and application", John Wiley & Sons, Inc, 1999.
- [7] Chipman, H., Kolaczyk, E.D., and McCulloch, R.E., "Adaptive Bayesian wavelet shrinkage", *J. Amer. Statist. Ass.*, **92**, pp 1413-1421. 1997.
- [8] Chui, C.K., "An Introduction to Wavelets", Academic Press, Boston, 1992.
- [9] Chui, C.K. and Wang, J.Z., "A study of asymptotically optimal time-frequency localization by scaling functions and wavelets", *Annals of Numerical Mathematics* **4**, pp 193-216, 1997.
- [10] Cohe, A., Daubechies I., Jawerth, B., and Vial, P., "Multiresolution analysis, wavelets, and fast algorithms on an interval", *Comptes Rendus Acad Sci. Paris A* **316**, pp 417-421, 1993.
- [11] Daubechies, I., "Orthonormal bases of compactly supported wavelets", *Comms. PureAppl. Math.*, **41**, pp 909-996, 1988.
- [12] Daubechies, I., "Ten Lectures on Wavelets", *Society for Industrial and Applied Mathematics*, Philadelphia, PA, 1992.
- [13] Donoho, D.L. and Johnstone, I.M., "Ideal spatial adaptation by wavelet shrinkage", *Biometrika*, **81**, pp 425-455, 1994.
- [14] Donoho, D.L. and Johnstone, I.M., "Adapting to unknown smoothness via wavelet shrinkage", *J. Amer. Statist. Assoc.*, **90**, pp 1200-1224, 1995.
- [15] Donoho, D.L. and Johnstone, I.M., Kerkyacharian, G. and Picard, D., "Wavelet shrinkage: Asymptopia?", *Journal of the Royal Statistical Society, Series B* **57**, pp 301-369, 1995.

- [16] Donoho, D.L. and Johnstone, I.M., "Minimax estimation via wavelet shrinkage", 1998.
- [17] Eubank, R.L., "Spline Smoothing and Nonparametric Regression", Marcel Dekker, 1999.
- [18] Hall, P., McKay, I., and Turlach, B.A., "Performance of wavelet methods for functions with many discontinuities", *Annals of Statistics*, **24**(6), pp 2462-2476, 1996.
- [19] Krishnaiah, P.R. and Miao, B.Q., "Review about estimation of change-points", *Handbook of Statistics*, **7**. P.R.Krishnaiah and C.R. Rao, eds., pp 375-402, Amsterdam:Elsevier, 1988.
- [20] Mallat, S. and Hwang, W.L., "Singularity detection and processing with wavelets", *IEEE Trans. Information Theory*, **38**, pp 617-643, 1992.
- [21] Mallat, S. and Zhong, Z., "Characterization of signals from multiscale edges", *IEEE Trans. Pattern Anal. and Machine Intel.*, **14** (7), pp 710-732, 1992.
- [22] Mallat, S., "Multiresolution approximations and wavelet orthonormal bases of $L^2(R)$ ", *Transactions of the American Mathematical Society*, **315**(1), 1989.
- [23] Mallat, S., "A theory of multiresolution signal decomposition: the wavelet representation", *IEEE Trans. Pattern Anal. Machine Intell.* **11**, pp 674-693, 1989.
- [24] Meyer, Y., "Principe d'incertitude, bases hilbertiennes et algebres d'operateurs", *Bourbaki seminar*, **662**, 1985-1986.
- [25] Meyer, Y., "Ondelettes, fonctions splines et analyses graduees", *Rapport Ceremade*, **8703**, 1987.
- [26] Meyer, Y., "Wavelets: Algorithms and Applications", Philadelphia: SIAM, 1993.
- [27] Misiti, M., Misiti, Y., Oppenheim, G. and Poggi, J.M., "Wavelet toolbox: for use with MATLAB", *The MathWorks Inc.*, 1996.
- [28] Mumford, D. and Shah, J., "Optimal approximations by pieewise smooth functions and associated variational problems", *Comm. on Pure and Appl. Math.*, **XLII** no.5 (1989).
- [29] Nason, G.P., "Wavelet shrinkage using cross-validation", *J. Roy. Statist. Soc. B*, **58**, pp 463-479, 1996.
- [30] Nason, G.P. and Silverman, B.W., "The discrete wavelet transform in S", *J. Comput. Graph. Statist.*, **3**, pp 163-191, 1994.

- [31] Ogden, R.T. and Parzen, E., "Change-point approach to data analytic wavelet thresholding", *Statistics and Computing* **6**, pp 93-99, 1996.
- [32] Ogden, R.T. and Parzen, E., "Data dependent wavelet thresholding in nonparametric regression with change-point applications", *Computational Statistics and Data Analysis* **22**, pp 53-70, 1996.
- [33] Ogden, R.T., "Essential wavelets for statistical applications and data analysis", Birkhäuser, Boston, 1997.
- [34] Perona, P. and Malik, J., "Scale-space and edge detection using anisotropic diffusion", *IEEE PAML*, **12** (1990) pp. 629-639.
- [35] Schimek, M.G., "Smoothing and Regression: Approaches, Computation, and Application", Wiley, 2000.
- [36] Vidakovic, B., "Statistical modeling by wavelets", Wiley Series in Probability and Statistics, 1999.
- [37] Wang, Y.Z., "Jump and sharp cusp detection by wavelets", *Biometrika*, **82** (2), pp 385-97, 1995.
- [38] You, Y., Xu, W., Tannenbaum, A., and Kaveh, M., "Behavioral analysis of anisotropic diffusion in image processing", *IEEE Trans. Image Processing*, **5** (1996) pp 1539-1553.

Wavelet Applications in Cancer Study

Don Hong

Department of Mathematical Sciences
Middle Tennessee State University
Murfreesboro, Tennessee

and

Yu Shyr

Department of Biostatistics
Vanderbilt University
Nashville, Tennessee

Abstract

In this survey article, we review wavelet methodology in many biostatistical applications of cancer study: microcalcifications in digital mammography, functional MRI data analysis, genome sequence, protein structure and microarray data analysis, and MALDI-TOF MS data analysis.

AMS 2000 Classifications: 42C40, 65T60, 92C55.

Key Words: Splines, statistics, mass spectrometry data processing, medical image compression, wavelets.

1 Introduction

Wavelet theory has developed now into a methodology used in many disciplines: mathematics, physics, engineering, signal processing and image compression, numerical analysis, and statistics, to list a few. Wavelets are providing a rich source of useful tools for applications in time-scale types of problems. Wavelets based methods are developed in statistics in areas such as regression, density and function estimation, modeling and forecasting in time series analysis, and spatial analysis. The attention of wavelets was attracted by statisticians when Mallat (see [35] and [36]) established a connection between wavelets and signal processing and Donoho and Johnstone (see [11] – [13]) found that wavelet thresholding has desirable statistical optimality properties. Since then, wavelets have proved very useful in nonparametric statistics, and time series analysis. Wavelets based Bayesian concepts and modeling approaches have wide applications in data denoising. Wang [51] in this special issue presents a brief review of wavelet shrinkage and the application of the variational method for wavelet regression and shows the relation between the variational method and wavelet

shrinkage. In recently years, wavelets have been applied to a large variety of biomedical signals (see [1], [33] for example) and there is a growing interest in using wavelets in the analysis of sequence and functional genomics data. It is important to point out that wavelets are very useful, but of course are not, and will never be a panacea. However, in many statistical data analyses, wavelet methods can be tried and tested in practice. In the following, we first give an introduction to wavelets in a relatively accessible format, then briefly mention techniques of wavelets in medical imaging and signal processing such as in detection of microcalcifications in digital mammography and functional Magnetic Resonance Imaging (fMRI), and biostatistical applications in molecular biology areas such as in genome sequence analysis, protein structure investigation, and gene expression data analysis. Finally, we present matrix-assisted laser desorption-ionization time-of-flight mass spectrometry (MALDI-TOF MS) data analysis using wavelet-based methods.

2 Wavelets

In the following, we review some essential concepts of wavelets and describe the method of wavelet analysis for statistical data analysis. Roughly speaking, a wavelet is a little wavy function with certain useful properties. A set of wavelets, usually is constructed through a single “mother” wavelet, to form a basis of L^2 space and provide “building blocks” that can be used to describe any function in L^2 . A simple example of wavelets is the Haar function and Haar wavelets [17]:

Haar began with the function

$$h(x) = \begin{cases} 1, & x \in [0, 1/2); \\ -1, & x \in [1/2, 1); \\ 0, & \text{otherwise.} \end{cases}$$

For $n \geq 1$, let $n = 2^j + k$, $j \geq 0$, $0 \leq k < 2^j$ and define

$$h_n(x) = 2^{j/2} h(2^j x - k) = \psi_{j,k}(x).$$

It is easy to see that

$$\text{supp}(h_n) = [\frac{k}{2^j}, \frac{k+1}{2^j}].$$

Let

$$h_0(x) = \begin{cases} 1, & x \in [0, 1) \\ 0, & \text{otherwise.} \end{cases}$$

Then the Haar wavelets

$$h_0, h_1, \dots, h_n, \dots$$

form an orthonormal basis for $L^2[0, 1]$.

The uniform approximation of f by

$$S_n(f) = \langle f, h_0 \rangle h_0 + \cdots + \langle f, h_n \rangle h_n$$

is nothing but the classical approximation of a continuous function by step functions. Clearly, Haar wavelets are not continuous. However, Haar wavelets allow information to be encoded according to “levels of detail.”

In general, wavelets construction starts with a function equation:

$$\phi(\cdot) = \sum_{k=0}^{2N-1} p_k \phi(2 \cdot - k) \quad (2.1)$$

The equation (2.1) is called the two-scale equation and the function ϕ is called a scaling function (also referred to as a “father” wavelet). Under some restrictions (multiresolution analysis conditions), a “mother” wavelet ψ can be obtained by using the scaling function:

$$\psi(x) = \sum_{k=0}^{2N-1} (-1)^k p_{1-k} \phi(2x - k). \quad (2.2)$$

In the late 1980s, Daubechies constructed a family of wavelets, now called Daubechies wavelets. The easiest example is

$$\phi_{D2}(x) = \sum_{k=0}^3 p_k \psi(2x - k),$$

where $p_0 = \frac{1+\sqrt{3}}{4}$, $p_1 = \frac{3-\sqrt{3}}{4}$, $p_2 = \frac{3+\sqrt{3}}{4}$, and $p_3 = \frac{1-\sqrt{3}}{4}$. Correspondingly, the Daubechies wavelet ψ_{D2} is the solution of the equation (2.1) determined by these coefficients (see [9] for details and more orthonormal wavelets). From Haar wavelets to continuous wavelets took about 80 years!

A wavelet must be localized in time, in the sense that $\psi(t) \rightarrow 0$ quickly as $|t| \rightarrow \infty$. The wavelet should be also oscillating about zero so that $\int_{-\infty}^{\infty} \psi(t) dt = 0$ and even the first m moments are also zero: $\int_{-\infty}^{\infty} t^k \psi(t) dt = 0$ for $k = 0, 1, \dots, m-1$. The oscillation property makes the function a wave. Together the localized feature, the function ψ becomes a wavelet. Semi-orthonormal spline wavelets, often called Chui-Wang wavelets were introduced by Charles K. Chui and Jianzhong Wang (See [6] for details). These wavelets trade the translation orthogonality with an explicit expression of the wavelets in terms of B-spline functions. These wavelets are compactly supported. They are very easy to compute and to work with.

In some applications, we need to use wavelets with symmetry, orthogonality, compact support, and high approximation order, simultaneously. This is impossible to achieve according to the above construction scheme. When the coefficients in the scaling equation (2.1) and wavelet equation (2.2) are matrices and ϕ and ψ are vectors, then there are two or more scaling functions and an equal number of wavelets. We call them multi-scaling functions and multiwavelets,

respectively. These multiwavelets open new possibilities for wavelets to possess those properties simultaneously. In [24], a set of orthogonal multiwavelets of multiplicity four was constructed. More recent discussion on multiwavelets can be found in [28] and [18] and the references therein.

The continuous wavelet transform (CWT) is a linear signal transformation that uses templates $\psi_{a,b} = a^{1/2}\psi(\frac{t-b}{a})$ which are shifted (index b) and dilated (index a) of a given wavelet function using the mother wavelet function $\psi(x)$. The wavelet transform of a signal f can be written as

$$T_\psi f(a, b) = \langle f, \psi_{a,b} \rangle.$$

The CWT maps a one-dimensional signal to a two-dimensional time-scale joint representation. The resulting wavelet coefficients are highly redundant.

A basic requirement is that the transform is reversible, i.e., the signal f can be reconstructed from its wavelet coefficients. The distinction between the various types of wavelet transforms depends on the way in which the dilation and shift parameters are discretized. This is more desirable in molecular biology and genetics.

A discrete wavelet transform (DWT) decomposes a signal into several vectors of wavelet coefficients. Different coefficient vectors contain information about the signal function at different scales. Coefficients at coarse scale capture gross and global features of the signal while coefficients at fine scale contain detailed information. In practice, it is often more convenient to consider the wavelet transform for some discretized values of $a = 2^j$ (the dyadic scales) and $b = k$ (the integer shifts). The transform is reversible if and only if the wavelet coefficients define a wavelet frame $A\|f\|^2 \leq \sum_{a,b} |\langle f, \psi_{a,b} \rangle|^2 \leq B\|f\|^2$, where A and B are called bounds of wavelet frame. In particular, if $A = B = 1$, the set of templates defines an orthogonal wavelet basis. The important point is that the wavelet decomposition provides a one-to-one representation of the signal in terms of its wavelet coefficients. This makes a wavelet basis well suited for any tasks for which block transforms (Fourier, Discrete Cosine Transform, etc.) have been used traditionally.

In statistical practice, the discrete wavelet transform plays a role analogous to the fast Fourier transform in conventional frequency-domain analysis. Unlike their Fourier cousins, wavelet methods make no assumptions concerning periodicity of the data at hand. As a result, wavelets are particularly suitable for studying data exhibiting sharp changes or even discontinuities.

Recognize that the wavelet transform on a finite sequence of data points provides a linear mapping to the wavelet coefficients: $w_n = Wf_n$, where the matrix $W = W_{n \times n}$ is orthogonal and w_n and f_n are n -dimensional vectors. The wavelet approximation to a signal function f is built up over multiple scales and many localized positions. For the given family

$$\phi_{J,k}(t) = 2^{-J/2}\phi(\frac{t}{2^J} - k), \quad \psi_{j,k} = 2^{-j/2}\psi(\frac{t}{2^j} - k), \quad j = 1, 2, \dots, J.$$

The coefficients are given by the projections:

$$s_{J,k} = \int f(t)\phi_{J,k}(t)dt, \quad d_{j,k} = \int f(t)\psi_{j,k}(t)dt$$

so that

$$f(t) = \sum_k s_{J,k}\phi_{J,k}(t) + \sum_k \sum_{j=1}^J d_{j,k}\psi_{j,k}(t).$$

The large J refers to the relatively small number of coefficients for the low frequency, smooth variation of f , the small j refers to the high frequency detail coefficients.

When the sample size n , the number of observations, is divisible by 2, say $n = 2^J$, then the number of coefficients, n can be grouped as, $n/2$ coefficients $d_{1,k}$ at the finest level, $n/4$ coefficients $d_{2,k}$ at the next finest level, $\dots$, $n/2^J$ coefficients $d_{J,k}$ and $n/2^J$ coefficients $s_{J,k}$ at the coarsest level. In the following, we demonstrate this by using Haar wavelet, the simplest and crudest member of a large class of wavelet functions.

We describe a scheme for transforming large arrays of numbers into arrays that can be stored and transmitted more efficiently. For this purpose, we describe a method for transforming strings of data using Haar wavelets, called *averaging and differencing*. For instance, let us start with a row vector

$$(2854, 6334, 3226, 2277, 2030, 1323, 367, 2890).$$

Table 1 shows the results of three steps in the transform process respectively. The first row in the table is the original data string, which consists of four pairs of numbers. The first four numbers in the second row are the averages of these pairs. Similarly, the first two numbers in the third row are the averages of those four averages, taken two at a time, and the first entry in the fourth and last row is the average of the preceding two computed averages. The remaining numbers measure deviations from the various averages. For example, the first four entries in the second half of the second row are the result of subtracting the first four averages from the first elements of the pairs that generate them: $2854 - 4594 = -1740$, $3226 - 2751.5 = 474.5$, $2030 - 1676.5 = 353.5$, $367 - 1628.5 = -1261.5$. These are called *detail coefficients*, they are repeated in each subsequent row of the table. The following matrix representation illustrates this three-step process clearly and precisely.

The first step involved averaging and differencing four pairs of numbers and it can be expressed as the matrix equation

$$\begin{aligned} & (4594, 2751.5, 1676.5, 1628.5, -1740, 474.5, 353.5, -1261.5) \\ = & (2854, 6334, 3226, 2277, 2030, 1323, 367, 2890)H_1, \end{aligned}$$

2854	6334	3226	2277	2030	1323	367	2890
4594	2751.5	1676.5	1628.5	-1740	474.5	353.5	-1261.5
3672.8	1652.5	921.3	24	-1740	474.5	353.5	-1261.5
2662.6	1010.1	921.3	24	-1740	474.5	353.5	-1261.5

Table 1: Results of a vector after three steps of Haar wavelet transformation.

where H_1 is the matrix

$$\begin{bmatrix} 1/2 & 0 & 0 & 0 & 1/2 & 0 & 0 & 0 \\ 1/2 & 0 & 0 & 0 & -1/2 & 0 & 0 & 0 \\ 0 & 1/2 & 0 & 0 & 0 & 1/2 & 0 & 0 \\ 0 & 1/2 & 0 & 0 & 0 & -1/2 & 0 & 0 \\ 0 & 0 & 1/2 & 0 & 0 & 0 & 1/2 & 0 \\ 0 & 0 & 1/2 & 0 & 0 & 0 & -1/2 & 0 \\ 0 & 0 & 0 & 1/2 & 0 & 0 & 0 & 1/2 \\ 0 & 0 & 0 & 1/2 & 0 & 0 & 0 & -1/2 \end{bmatrix}.$$

The transitions of the second step and the third step are equivalent to the matrix equations

$$\begin{aligned} & (3672.8, 1652.5, 921.3, 24, -1740, 474.5, 353.5, -1261.5) \\ &= (4594, 2751.5, 1676.5, 1628.5, -1740, 474.5, 353.5, -1261.5)H_2 \end{aligned}$$

and

$$\begin{aligned} & (2662.6, 1010.1, 921.3, 24, -1740, 474.5, 353.5, -1261.5) \\ &= (3672.8, 1652.5, 921.3, 24, -1740, 474.5, 353.5, -1261.5)H_3, \end{aligned}$$

where H_2 and H_3 , respectively are the matrices

$$\begin{bmatrix} 1/2 & 0 & 1/2 & 0 & 0 & 0 & 0 & 0 \\ 1/2 & 0 & -1/2 & 0 & 0 & 0 & 0 & 0 \\ 0 & 1/2 & 0 & 1/2 & 0 & 0 & 0 & 0 \\ 0 & 1/2 & 0 & -1/2 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix}, \begin{bmatrix} 1/2 & 1/2 & 0 & 0 & 0 & 0 & 0 & 0 \\ 1/2 & -1/2 & 0 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 1 & 0 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 1 & 0 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 1 & 0 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \end{bmatrix}.$$

Thus, the combined effect of the Haar wavelet transform can be achieved with the single equation:

$$\begin{aligned} & (2662.6, 1010.1, 921.3, 24, -1740, 474.5, 353.5, -1261.5) \\ &= (2854, 6334, 3226, 2277, 2030, 1323, 367, 2890)W, \end{aligned}$$

where $W = H_1 H_2 H_3$. Noticing that the columns of H_i 's are orthogonal, each of these matrices are invertible. The inverses simply reverse the three averaging and differencing steps. In any case, we can recover the original string from the transformed version by means of the inverse Haar wavelet transform:

$$\begin{aligned} & (2854, 6334, 3226, 2277, 2030, 1323, 367, 2890) \\ = & (2662.6, 1010.1, 921.3, 24, -1740, 474.5, 353.5, -1261.5)W^{-1}, \end{aligned}$$

where $W^{-1} = H_3^{-1} H_2^{-1} H_1^{-1}$. It becomes a routine matter to construct the corresponding $2^n \times 2^n$ matrices H_i , $i = 1, \dots, n$ needed to work with strings of length 2^n . In general, there is no loss of generality in assuming that each string's length is a power of 2.

In image compression, not only the row transformations but also column transformations are applied to image matrices (see [16] and [22] for example).

Applications of wavelets in statistics can be found in [3], [39], [45], and the references therein. In the following, we only emphasize to review biostatistical wavelet applications of medical science, particularly in cancer study.

3 Wavelet-Based Detection of Microcalcifications in Digital Mammography

There are many applications of wavelets in medical image analysis. For example, quantitative neuroimaging by position emission tomography (PET) is a major research tool to investigate functional brain activity in vivo dynamic processes occurring in discrete brain region. The main difficulty with the analysis of PET data is the absence of reliable anatomical markers and the fact that the shape and size of the brain may vary substantially from one subject to another. In [49], the authors proposed an image-based statistical method for detecting differences between subject groups. The two important processing steps in their methodology are (i) the registration of the individual brain images, and (ii) the subsequent statistical analysis in order to detect differences in functional activity between subject groups. Wavelet transform plays an important role in each step, but for different reasons. In [30], wavelet compression technique is used for the low-dose chest CT data analysis. Wavelet image decomposition for the detection of myocardial viability in the early postinfarction period with the wavelet-based method for texture characterization is discussed in [40]. Ultrasonic liver tissues classification by fractal feature vector based on M-band wavelet transform can be found in [31].

Breast cancer is the main cause of death for women between the ages of 35 and 55. It has been shown that early detection and treatment of breast cancer are the most effective methods of reducing mortality. The detection of microcalcifications plays a very important role in breast cancer detection because these microcalcifications can be the only mammographic sign of nonpalpable breast disease. Screen-film mammography is widely recognized as being the only effective imaging modality for early detection of breast cancer. Though

advances in screen-film mammographical technology have resulted in significant improvements in image resolution and film contrast, images provided by screen-film mammography remain very difficult to interpret. Mammograms are of low contrast, and features in mammograms indicative of breast disease, such as the microcalcifications, are often very small. The theory of wavelets provides a common framework for numerous techniques developed independently for various signal and image processing applications. For example, multiresolution image processing, subband coding, and wavelet series expansions are the different views of this theory. In [52], an approach is presented for detecting microcalcifications in digital mammograms employing wavelet-based subband image decomposition. The authors proposed a system for microcalcification detection under the hypothesis that the microcalcifications present in mammograms can be preserved under a transform which can localize the signal characteristics in the original and the transform domain. Therefore, microcalcifications can be extracted from mammograms by suppressing the subband of the wavelets-decomposed image that carries the lowest frequencies and contains smooth background information, before the reconstruction of the image. Multiresolution-based segmentation for the detection and enhancement of the microcalcifications is also presented in [42]. Very recently, the authors in [26] present the implementation and evaluation of a novel enhancement technique for improved interpretation of digital mammography with wavelet processing. A new algorithm for enhancement of microcalcifications in mammograms is given in [20] and highly regular wavelets are used for the detection of clustered microcalcifications in mammograms in [32]. More wavelet applications of medical imaging can be found in [48].

4 Wavelet-Based Functional MRI Data Analysis

Functional Magnetic Resonance Imaging (fMRI) is a technique for determining which part of the brain are activated by different types of physical sensation or activity, such as light, sound or the movement of a subject's fingers. This "brain mapping" is achieved by setting up an advanced MRI scanner in a special way so that the increased blood flow to the activated areas of the brain shows up on functional MRI scans. Functional MRI is based on the increase in blood flow to the local vasculature that accompanies neural activity in the brain. This results in a corresponding local reduction in deoxyhemoglobin because the increase in blood flow occurs without an increase of similar magnitude in oxygen extraction. Since deoxyhemoglobin is paramagnetic, it alters the $T2^*$ weighted magnetic resonance image signal. Thus, deoxyhemoglobin is sometimes referred to as an endogenous contrast enhancing agent, and serves as the source of the signal for fMRI. Using an appropriate imaging sequence, human cortical functions can be observed without the use of exogenous contrast enhancing agents on a clinical strength (1.5 T) scanner. Functional activity of the brain determined from the magnetic resonance signal has confirmed known anatomically distinct processing areas in the visual cortex, the motor cortex, and Broca's area of speech and language-related activities. Further, a rapidly emerging body of

literature documents corresponding findings between fMRI and conventional electrophysiological techniques to localize specific functions of the human brain. Consequently, the number of medical and research centers with fMRI capabilities and investigational programs continues to escalate.

The main advantages to fMRI as a technique to image brain activity related to a specific task or sensory process include 1) the signal does not require injections of radioactive isotopes, 2) the total scan time required can be very short, i.e., on the order of 1.5 to 2.0 min per run (depending on the paradigm), and 3) the in-plane resolution of the functional image is generally about 1.5×1.5 mm although resolutions less than 1 mm are possible. To put these advantages in perspective, functional images obtained by the earlier method of positron emission tomography, PET, require injections of radioactive isotopes, multiple acquisitions, and, therefore, extended imaging times. Further, the expected resolution of PET images is much larger than the usual fMRI pixel size. Additionally, PET usually requires that multiple individual brain images are combined in order to obtain a reliable signal. Consequently, information on a single patient is compromised and limited to a finite number of imaging sessions. Although these limitations may serve many neuroscience applications, they are not optimally suitable to assist in a neurosurgical or treatment plan for a specific individual.

Wavelet transformation is a time-frequency representation which decomposes the signal according to a set of functions through translation and dilation operations. Such a time-frequency representation could provide a more efficient solution than the usual Fourier-transform or other methods presently available to process MRI data. One of the first application of wavelet transform in medical imaging was for noise reduction in MRI [53]. The approach proposed there was to compute an orthogonal wavelet decomposition of the image and apply the following soft thresholding rule on the wavelet coefficients $c_{i,k} = \langle f, \psi_{i,k} \rangle$:

$$\tilde{c}_{i,k} = \begin{cases} c_{i,k} - t_i, & c_{i,k} \geq t_i, \\ 0, & |c_{i,k}| \leq t_i, \\ c_{i,k} + t_i, & c_{i,k} \leq -t_i, \end{cases}$$

where t_i is a threshold that depends on the noise level at the i th scale. The image is then reconstructed by the inverse wavelet transform of the $\tilde{c}_{i,k}$'s.

Ruttimann et al. [41] have proposed to use the wavelet transform for the detection and localization of activation patterns in fMRI. The idea is to apply a statistical test in the wavelet domain to detect the coefficients that are significantly different from zero. In [14], the authors improved the method by replacing the original z -test by a t -test that takes into account the variability of each wavelet coefficient separately. A key issue is to find out which wavelet and which type of decomposition is best suited for the detection of a given activation pattern. Various brands of fractional spline wavelets are applied and an extensive series of tests using simulated data are performed. An interesting practical finding is that performance is strongly correlated with the number of coefficients detected in the wavelet domain, at least in the orthonormal and B-spline cases.

Therefore, it is possible to optimize the structural wavelet parameters simply by maximizing the number of wavelet counts, without any prior knowledge of the activation pattern.

A new method is proposed in [25] for activation detection in event-related fMRI. The method is based on the analysis of selected resolution levels in translation invariant wavelet transform domain. The power in different resolution levels in wavelet domain is analyzed and an optimal set of resolution levels is selected. A randomization-based statistical test is then applied in wavelet domain for activation detection. The problem of detecting significant changes in fMRI time series that are correlated to a stimulus time course is addressed in [37]. The fMRI signal is described as the sum of two effects: a smooth trend and the response to the stimulus. The trend belongs to a subspace spanned by large scale wavelets. The wavelet transform provides an approximation to the Karhunen-Loeve transform for the long memory noise. A scale space regression is developed that permits carrying out the regression in the wavelet domain while omitting the scales that are contaminated by the trend.

In [43], two iterative procedures obtained from the application of wavelet transform to noise-free magnetic resonance spectroscopy (MRS) signals composed of one and two resonances, respectively, are tested on simulated and real biomedical MRS data.

5 Wavelet-Based Molecular Biology Data Analysis

In this section, we briefly review the most interesting applications of wavelets in genome sequence analysis, protein structure investigation, and gene expression data analysis.

The Human Genome Project (HGP) was one of the great feats of exploration in history, an inward voyage of discovery rather than an outward exploration of the planet or the cosmos. An international research effort to sequence and map all of the genes - together known as the genome - of members of our species, *Homo sapiens*, the HGP was completed in April 2003. Now we can, for the first time, read nature's complete genetic blueprint for building a human being.

Integral to the HGP are similar efforts to understand the genome of various organisms commonly used in biomedical research, such as mice, fruit flies and roundworms. Such organisms are called "model organisms," because they can often serve as research models for how the human organism behaves.

In order to fully understand these genomes, the HGP also develops technologies for genomic analysis, and trains scientists who will be able to use the tools and resources created by the HGP to perform research that will improve human health. Recognizing the profound importance and seriousness of this venture, another mission of the HGP - and NHGRI - is to examine the ethical, legal and social implications of human genetic research.

Wavelets can be useful in detecting patterns in DNA sequences. In [34], it

was shown that wavelet variance decomposition of bacterial genome sequences can reveal the location of pathogenicity islands. The findings show that wavelet smoothing and scalogram are powerful tools to detect differences within and between genomes and to separate small (gene level) and large (putative pathogenicity islands) genomic regions that have different composition characteristics. An optimization procedure improving upon traditional Fourier analysis performance in distinguishing coding from noncoding regions in DNA sequences was introduced in [2]. The approach can be taken one step further by applying wavelet transforms.

Proteins are macromolecules (heteropolymers) made up from 20 different L- α -amino acids, also referred to as residues. A certain number of residues is necessary to perform a particular biochemical function, and around 40-50 residues appears to be the lower limit for a functional domain size. Protein sizes range from this lower limit to several hundred residues in multi-functional proteins. Very large aggregates can be formed from protein subunits, for example many thousand actin molecules assemble into an actin filament. Large protein complexes with RNA are found in the ribosome particles, which are in fact 'ribozymes'.

To find the similarities between two or more protein sequences is of great importance for protein sequence analysis. In [10], a comparison method based on wavelet decomposition of protein sequences and a cross-correlation study was devised that is capable of analyzing a protein sequence "hierarchically", i.e., it can examine a protein sequence at different spatial resolutions. A sequence-scale similarity vector is generated for the comparison of two sequences feasible at different spatial resolutions (scales).

In the process of protein construction, buried hydrophobic residues tend to assemble in a core of a protein. Methods used to predict these cores involve use or no use of sequential alignment. In the case of a close homology, prediction was more accurate if sequential alignment was used. If the homology was weak, prediction would be unreliable. A hydrophobicity plot involving the hydropathy index is useful for purposes of prediction. In [21], wavelet analysis is used for hydrophobicity plots to quantitatively decompose to high and low frequencies. The prediction of hydrophobic cores of proteins was made with low frequency extracted from the hydrophobicity plot.

The cosine Fourier series and discrete wavelet transforms are applied in [38] for describing replacement rate variation in genes and proteins, in which the profile of relative replacement rates along the length of a given sequence is defined as a function of the site number. The new models are applicable to testing biological hypotheses such as the statistical identity of rate variation profiles among homologous protein families.

Microarray technology allows us to analyze the expression pattern of hundreds of genes. The use of statistics is the key to extracting useful information from this technology. Recently, "False Discovery Rate" (FDR) becomes a very useful inferential approach in the analysis of microarray data [47]. The FDR is the expected proportion of rejected hypotheses that are falsely rejected. When the model is sparse, FDR-like selection yields estimators with strong large sam-

ple adaptivity properties. One natural application of FDR is threshold selection in wavelet denoising. In wavelets thresholding, FDR is the proportion of wavelet coefficients erroneously included in the reconstruction among those included. This approach might lead to other applications of wavelets in microarray data analysis.

6 Wavelet Applications in MALDI-TOF MS Data Analysis

In recent years there has been a great amount of interest in investigating the biochemical events involved in the metabolism of peptides, primarily in the brain and gut of mammals, encompassing the enzymatic breakdown of these peptides, their production from peptide and protein precursors, and the disruption of these processes by certain xenobiotics. Modern mass spectrometric techniques are used in these studies, including electrospray and matrix-assisted laser desorption-ionization mass spectrometry (MALDI MS) (for example, [5], [29], [50], and the references therein).

Matrix-assisted laser desorption-ionization, time-of-flight (MALDI-TOF) mass spectrometry (MS) is emerging as a leading technology in the proteomics revolution. Indeed, the Nobel prize in chemistry in 2002 recognized MALDI's ability to analyze intact biological macromolecules. Though MALDI-TOF MS allows direct measurement of the protein "signature" of tissue, blood, or other biological samples, and holds tremendous potential for disease diagnosis and treatment, key challenges remain in signal normalization and quantitation (for example, see [4], [7], [19], [46], and the references therein). Wavelet methods promise a principled approach for evaluating these complicated biological signals.

Mass spectrometry, especially surface enhanced laser desorption and ionization (SELDI), is increasingly being used to find disease-related proteomic patterns in complex mixtures of proteins derived from tissue samples or from easily obtained biological fluids. In [8], a wavelet-based method was proposed to the low-level processing of SELDI spectra data. This method was motivated by the idea that the total ion current is a surrogate for the total amount of protein in the sample being measured. If the baseline correction algorithm starts with the raw spectrum and ensures that the corrected signal never becomes negative, then electronic noise contributes a substantial portion of the total ion current, and one normalizes, in part, to the noise. Statistically, the low-level processing of mass spectra reduces to decomposing the observed signal into three components: true signal, baseline, and noise:

$$f(t) = B(t) + N \cdot S(t) + \epsilon(t),$$

where, $f(t)$ is the observed signal, $B(t)$ is the baseline, $S(t)$ is the true signal, N is the normalization factor, and $\epsilon(t)$ is the noise. It has been shown that the nonorthogonal discrete wavelet transform can be used to denoise MS SELDI data and the peaks in the intensity can be rapidly identified and precisely quantified.

Lung cancer kills more Americans each year than the next four leading cancer killers—cancers of the colon, breast, prostate and pancreas—combined.

What eludes medical science is a reliable method to screen those high-risk people for lung cancer so that the disease is caught early enough to make a difference in survival. Proteomics-based approaches complement the genome initiatives and may be the next step in attempts to understand the biology of cancer. It is now clear that the behavior of individual non-small cell lung cancer tumors cannot be understood through the analysis of individual or small number of genes, so cDNA microarray analysis has been employed with some success to simultaneously investigate thousands of RNA expression levels and begin to identify patterns associated with biology. However, mRNA expressions are poorly correlated with protein expression levels, and it cannot detect important post-translation modifications of proteins, all very important processes in determining protein function. Accordingly, comprehensive analysis of protein expression patterns in tissues might improve our ability to understand the molecular complexities of tumor cells.

MALDI-TOF MS can profile proteins up to 50 kDa in size in tissue. This technology can not only directly assess peptides and proteins in sections of tumor tissue, but also can be used for high resolution image of individual biomolecules present in tissue sections. The protein profiles obtained can contain thousands of data points, necessitating sophisticated data analysis algorithms.

Wavelets are providing a rich source of useful tools for applications in time-scale types of problems. In analysis of signals, the wavelet representations allow us to view a time-domain evolution in terms of scale components. The adaptivity property of wavelets certainly can help us to determine the location of peak difference(s) of MALDI-TOF MS protein expressions between cancerous and normal tissues in terms of molecular weights.

Wavelets, as building blocks of models, are well localized in both time and scale (frequency). Signals with rapid local changes (signals with discontinuities, cusps, sharp spikes, etc) can be precisely represented with just a few wavelet coefficients. In [23], we apply wavelet transform together with clustering techniques [15] and weighted flexible compound covariate method [44] to the MALDI-TOF MS data from lung cancer patients. Class-prediction models based on wavelet coefficients were found to classify lung cancer histologies and distinguish primary tumors.

Very recently, wavelets are used for neural classification of lung sounds analysis in [27].

Acknowledgments: This work was supported in part by Lung Cancer SPORC (Special Program of Research Excellence) (P50 CA90949), Breast Cancer SPORC (1P50 CA98131-01), GI (5P50 CA95103-02), and Cancer Center Support Grant (CCSG) (P30 CA68485) for Shyr and by NSF-IGMS 0408086 and 0552377 for Hong.

References

- [1] A. Aldroubi and M. Unser, Wavelets in Medicine and Biology, CRC Press, Boca Raton, FL, 1996.
- [2] D. Anastassiou, Frequency-domain analysis of biomolecular sequences, *Bioinformatics* **16** (2000), 1073–1081.
- [3] A. Antoniadis, Wavelets in Statistics: A Review, *Journal of the Italian Statistical Society*, **6** 1999, .
- [4] S. Böcker, SNP and mutation discovery using base-specific cleavage and MALDI-TOF mass spectrometry, *Bioinformatics* **19** (2003), 144–153.
- [5] P. Chaurand, M. Stoeckli, and R.M. Caprioli, Direct profiling of proteins in biological tissue sections by MALDI mass spectrometry, *Anal Chem.* **71** (1999), 5263–5270.
- [6] C.K. Chui, An Introduction to Wavelets, Academic Press, New York, NY, 1992.
- [7] K.R. Coombes, H.A. Fritsche, Jr., et al., Quality control and peak finding for proteomics data collected from nipple aspirate fluid by surface-enhanced laser desorption and ionization, *Clinical Chemistry* **49** (2003), 1615–1623.
- [8] K.R. Coombes, S. Tsavachidis, J.S. Morris, K.A. Baggerly, M.C. Hung, and H.M. Kuerer, Improved peak detection and quantification of mass spectrometry data acquired from surface-enhanced laser desorption and ionization by denoising spectra with the undecimated discrete wavelets transform, manuscript, 2003.
- [9] I. Daubechies, Ten Lectures on Wavelets, Society for Industrial and Applied Mathematics, Philadelphia, Pennsylvania, 1992.
- [10] C. H. de Trad, Q. Fang, and I. Cosic, Protein sequence comparison based on the wavelet transform approach, *Protein Engineering* **15** (2002), 193–203.
- [11] D.L. Donoho and I. M. Johnstone, Ideal spatial adaption via wavelet shrinkage, *Biometrika*, **81** (1994), 425–455.
- [12] D.L. Donoho and I. M. Johnstone, Adapting to unknown smoothness via wavelet shrinkage, *J. Amer. Statist. Assoc.*, **90** (1995), 1200–1224.
- [13] D.L. Donoho and I. M. Johnstone, Minimax estimation via wavelet shrinkage, *Ann. Statist.*, **26** (1998), 879–921.
- [14] M. Feilner, T. Blu, and M. Unser, Optimizing wavelets for the analysis of fMRI data, *Proceedings of the SPIE Conference on Mathematical Imaging: Wavelet Applications in Signal and Image Processing VIII*, San Diego CA, USA, July 31-August 4, 2000, vol. 4119, pp. 626–637.

- [15] D. Goldstein, D. Ghosh, and E. Conlon, Statistical issues in the clustering of gene expression data . *Statistica Sinica*, **12** (2002), 219 – 240.
- [16] Hao Gu, Don Hong, and Martin Barrett, On still image compression, *J. of Computational Analysis and Applications* **5**(2003), 45–75.
- [17] A. Haar, Zur Theorie der orthoganalen Funktionen-Systeme, *Annals of Mathematics*, **69** (1910), 331–371.
- [18] D. Hardin and D. Hong, On Piecewise Linear Wavelets and Prewavelets over Triangulations, *J. Computational and Applied Mathematics*, **155** (2003), 91-109.
- [19] R. Hartmer, N. Storm, et al., RNase T1 mediated base-specific cleavage and MALDI-TOF MS for high-throughput comparative sequence analysis, *Nucleic Acids Research* **31** (2003), e47, 1-10.
- [20] P. Heinlein, J. Drexler, and W. Schneider, Integrated wavelets for enhancement of microcalcifications in digit mammography, *IEEE Transactions of Medical Imaging* **22** (2003), 402–413.
- [21] H. Hirakawa, S. Muta, and S. Kuhara, The hydrophobic cores of proteins predicted by wavelet analysis, *Bioinformatics* **15** (1999), 141–148.
- [22] D. Hong, M. Barrett, and P.R. Xiao, Biorthogonal Spline Wavelets and EZW Coding for Image Compression, *Lecture Notes in Control and Information Sciences*, Vol. 299 (Tarn, Tzyh-Jong; Chen, Shan-Ben; Zhou, Changjiu (Eds.)), Springer-Verlag Berlin Heidelberg, 2004, pp.281–303.
- [23] D. Hong and Y. Shyr, Mass spectrometry data processing using wavelets, manuscript, 2004.
- [24] D. Hong and A. Wu, Orthogonal multiwavelets of multiplicity four, *Computers and Mathematics with Applications*, **40** (2000), 1153–1169
- [25] G. Hossein-Zadeh, H. Soltanian-Zadeh, and B.A. Ardekani, Multiresolution fMRI activation detection using translation invariant wavelet transform and statistical analysis based on resampling, *IEEE Trans. Medical Imaging*, **22** (2003), 302–314.
- [26] M. Kallergi, J.J. Heine, C.G. Berman, M.R. Hersh, A.P. Romilly, and R.A. Clark, Improved Interpretation of Digitized Mammography with Wavelet Processing: A Local Response Operating Characteristic Study, *American J. Roentgenology*, **182** (2004), 697 - 703.
- [27] A. Kandaswamy, C.S. Kumar, Rm. Pl. Rammanathan, S. Jayaraman, N. Malmurugan, Neural classification of lung sounds using wavelet coefficients, *Computers in Biology and Medicine*, to appear.

- [28] Bruce Kessler, An Orthogonal Scaling Vector Generating a Space of C^1 Cubic Splines Using Macroelements, *J. of Concrete and Applicable Math*, this issue, xx–xx.
- [29] Kiyoshi Yanagisawa, Yu Shyr, Baogang Xu, Pierre P Massion, Paul H Larsen, Bill C White, John R Roberts, Mary Edgerton, Adriana Gonzalez, Sorena Nadaf et al., Proteomic patterns of tumour subsets in non-small-cell lung cancer, *The Lancet*, 362 (2003), 433-439.
- [30] J.P. Ko, H. Rusinek, D.P. Naidich, G. McGuinness, A.N. Rubinowitz, B.S. Leitman, and J.M. Martino, Wavelet compression of low-dose chest CT data: effect on lung nodule detection, *Radiology* **228** (2003), 70–75.
- [31] W.-L. Lee and Y.-C. Chen, Ultrasonic liver tissues classification by fractal feature vector based on M-band wavelet transform, *IEEE Transactions on Medical Imaging*, **22** (2003), 382–392.
- [32] G. Lemaury, K. Drouiche, and J. DeConinck, Highly regular wavelets for the detection of clustered microcalcifications in mammograms, *IEEE Trans. Medical Imaging* **22** (2003), 393-401.
- [33] P. Lió, Wavelets in bioinformatics and computational biology: state of art and perspectives, *Bioinformatics* **19** (2003), 2–9.
- [34] P. Lió and M. Vannucci, Finding pathogenicity islands and gene transfer events in genome data, *Bioinformatics* **16** (2000), 932–940.
- [35] S. Mallat, A theory for multiresolution signal decomposition: the wavelets representation, *IEEE Trans. Pattern Analysis and Machine Intelligence*, **11** (1989), 674–693.
- [36] S. Mallat, *A wavelet tour of signal processing*. Academic press, San Diego, CA, 1998.
- [37] F.G. Meyer, Wavelet-based estimation of a semiparametric generalized linear model of fMRI times-series, **22** (2003), 315–322.
- [38] P. Morozov, T. Sitnikov, G. Churchill, F.J. Ayala, and A. Rzhetsky, A new method for characterizing replacement rate variation in molecular sequences: application of the Fourier and wavelet models to drosophila and mammalian protein, *Genetics* **154** (2000), 381–395.
- [39] G.P. Nason and B.W. Silverman, Wavelets for regression and other statistical problems, In: *Smoothing and Regression: Approaches, Computation and Applications* (M.G. Schimek Ed.), Wiley, New York, 159–191, 2000.
- [40] A.N. Neskovic, A. Mojsilovic, T. Jovanovic, etc., Myocardial tissue characterization after acute myocardial infection with wavelet image decomposition, *Circulation* **98** (1998), 634 - 641.

- [41] U. Ruttimann, M. Unser, R. Rawlings, D. Rio, N. Ramsey, V. Mattay, D. Hommer, J. Frank, and D. Weinberger, Statistical analysis of functional MRI data in the wavelets domain, *IEEE Trans. Medical Imaging* **17** (1998), 142–154.
- [42] S. Sentelle, C. Sentelle, and M.A. Sunton, Multiresolution-based segmentation of calcifications for the early detection of breast cancer, *Real-Time Imaging* **8** (2002), 237–252.
- [43] H. Serrai, L. Senhadji, J.D. de Certaines, and J.L. Coatrieux, Time-domain quantification of amplitude, chemical shift, apparent relaxation time T_2^* , and phase by wavelet-transform analysis: application to biomedical magnetic resonance spectroscopy, *J. Magnetic Resonance* **124** (1997), 20–34.
- [44] Y. Shyr and K.M. Kim, Weighted flexible compound covariate method for classifying microarray data, a case study of gene expression level of lung cancer, In: “A Practical Approach to Microarray Data Analysis” (Berrar, D. Ed.), pp.186–200, Klumer Academic Publishers, Norwell, MA, 2003.
- [45] B.W. Silverman, Wavelets in Statistics: Some Recent Developments, *Comp Stat: Proceedings in Computational Statistics*, Heidelberg, Physics-Verlag, 15–26, 1998.
- [46] P.R. Srinivas, M. Verma, Y.M. Zhao, and S. Srivastava, Proteomics for cancer biomarker discovery, *Clinical Chemistry* **48** (2002), 1160–1169.
- [47] J.D. Storey, The positive false discovery rate: A Bayesian interpretation and the q -value, *Annals of Statistics*, **31** (2003), 2013–2035.
- [48] M. Unser, A. Aldroubi, A. Laine, *Wavelets in Medical Imaging*, a special issue in *IEEE Transactions on Medical Imaging*, vol. 22, 2003.
- [49] M. Unser, P. Thevenaz, C. Lee, and U.E. Ruttimann, Registration and statistical analysis of PET images using the wavelet transform, *IEEE Engineering in Medicine and Biology*, **14** (1995), 603–611.
- [50] M. Verma, G.L. Wright, et al., Proteomic approaches within the NCI early detection research network for the discovery and identification of cancer biomarkers, *Annal N Y Acad. Sci.* **945** (2001), 103–115.
- [51] J.Z. Wang, Wavelet regression with an emphasis on singularity detection, *Journal of Concrete and Applicable Mathematics*, this issue, xx–xx.
- [52] T.C. Wang and N.B. Karayiannis, Detection of microcalcifications in digital mammograms using wavelets, *IEEE Transactions on Medical Imaging*, **17** (1998), 498–509.
- [53] J.B. Weaver, X. Yansun, D.M. Healy, Jr., and L.D. Cromwell, Filtering noise from images with wavelet transforms, *Magnetic Resonance in Medicine*, **21** (1991), 288–295.

Instructions to Contributors
Journal of Concrete and Applicable Mathematics

A quarterly international publication of Eudoxus Press, LLC, of TN.

Editor in Chief: George Anastassiou

Department of Mathematical Sciences
 University of Memphis
 Memphis, TN 38152-3240, U.S.A.

1. Manuscripts hard copies in triplicate, and in English, should be submitted to the Editor-in-Chief:

Prof. George A. Anastassiou
 Department of Mathematical Sciences
 The University of Memphis
 Memphis, TN 38152, USA.
 Tel. 001. 901.678.3144
 e-mail: ganastss@memphis.edu.

Authors may want to recommend an associate editor the most related to the submission to possibly handle it.

Also authors may want to submit a list of six possible referees, to be used in case we cannot find related referees by ourselves.

2. Manuscripts should be typed using any of TEX, LaTeX, AMS-TEX, or AMS-LaTeX and according to EUDOXUS PRESS, LLC. LATEX STYLE FILE. (VISIT www.msci.memphis.edu/~ganastss/jcaam/ to save a copy of the style file.) They should be carefully prepared in all respects. Submitted copies should be brightly printed (not dot-matrix), double spaced, in ten point type size, on one side high quality paper 8(1/2)x11 inch. Manuscripts should have generous margins on all sides and should not exceed 24 pages.

3. Submission is a representation that the manuscript has not been published previously in this or any other similar form and is not currently under consideration for publication elsewhere. A statement transferring from the authors (or their employers, if they hold the copyright) to Eudoxus Press, LLC, will be required before the manuscript can be accepted for publication. The Editor-in-Chief will supply the necessary forms for this transfer. Such a written transfer of copyright, which previously was assumed to be implicit in the act of submitting a manuscript, is necessary under the U.S. Copyright Law in order for the publisher to carry through the dissemination of research results and reviews as widely and effectively as possible.

4. The paper starts with the title of the article, author's name(s) (no titles or degrees), author's affiliation(s) and e-mail addresses. The affiliation should comprise the department, institution (usually university or company), city, state (and/or nation) and mail code.

The following items, 5 and 6, should be on page no. 1 of the paper.

5. An abstract is to be provided, preferably no longer than 150 words.
 6. A list of 5 key words is to be provided directly below the abstract. Key words should express the precise content of the manuscript, as they are used for indexing purposes. The main body of the paper should begin on page no. 1, if possible.

7. All sections should be numbered with Arabic numerals (such as: 1. INTRODUCTION) .

Subsections should be identified with section and subsection numbers (such as 6.1. Second-Value Subheading).

If applicable, an independent single-number system (one for each category) should be used to label all theorems, lemmas, propositions, corollaries, definitions, remarks, examples, etc. The label (such as Lemma 7) should be typed with paragraph indentation, followed by a period and the lemma itself.

8. Mathematical notation must be typeset. Equations should be numbered consecutively with Arabic numerals in parentheses placed flush right, and should be thusly referred to in the text [such as Eqs.(2) and (5)]. The running title must be placed at the top of even numbered pages and the first author's name, et al., must be placed at the top of the odd numbered pages.

9. Illustrations (photographs, drawings, diagrams, and charts) are to be numbered in one consecutive series of Arabic numerals. The captions for illustrations should be typed double space. All illustrations, charts, tables, etc., must be embedded in the body of the manuscript in proper, final, print position. In particular, manuscript, source, and PDF file version must be at camera ready stage for publication or they cannot be considered.

Tables are to be numbered (with Roman numerals) and referred to by number in the text. Center the title above the table, and type explanatory footnotes (indicated by superscript lowercase letters) below the table.

10. List references alphabetically at the end of the paper and number them consecutively. Each must be cited in the text by the appropriate Arabic numeral in square brackets on the baseline.

References should include (in the following order):

initials of first and middle name, last name of author(s)

title of article,

name of publication, volume number, inclusive pages, and year of publication.

Authors should follow these examples:

Journal Article

1. H.H.Gonska, Degree of simultaneous approximation of bivariate functions by Gordon operators, (journal name in italics) *J. Approx. Theory*, 62,170-191(1990).

Book

2. G.G.Lorentz, (title of book in italics) *Bernstein Polynomials* (2nd ed.), Chelsea, New York, 1986.

Contribution to a Book

3. M.K.Khan, Approximation properties of beta operators, in (title of book in italics) *Progress in Approximation Theory* (P.Nevai and A.Pinkus, eds.), Academic Press, New York, 1991, pp.483-495.

11. All acknowledgements (including those for a grant and financial support) should occur in one paragraph that directly precedes the References section.

12. Footnotes should be avoided. When their use is absolutely necessary, footnotes should be numbered consecutively using Arabic numerals and should be typed at the bottom of the page to which they refer. Place a line above the footnote, so that it is set

off from the text. Use the appropriate superscript numeral for citation in the text.

13. After each revision is made please again submit three hard copies of the revised manuscript, including in the final one. And after a manuscript has been accepted for publication and with all revisions incorporated, manuscripts, including the TEX/LaTeX source file and the PDF file, are to be submitted to the Editor's Office on a personal-computer disk, 3.5 inch size. Label the disk with clearly written identifying information and properly ship, such as:

Your name, title of article, kind of computer used, kind of software and version number, disk format and files names of article, as well as abbreviated journal name.

Package the disk in a disk mailer or protective cardboard. Make sure contents of disks are identical with the ones of final hard copies submitted!

Note: The Editor's Office cannot accept the disk without the accompanying matching hard copies of manuscript. No e-mail final submissions are allowed! The disk submission must be used.

14. Effective 1 Nov. 2005 the journal's page charges are \$8.00 per PDF file page. Upon acceptance of the paper an invoice will be sent to the contact author. The fee payment will be due one month from the invoice date. The article will proceed to publication only after the fee is paid. The charges are to be sent, by money order or certified check, in US dollars, payable to Eudoxus Press, LLC, to the address shown in the Scope and Prices section.

No galleys will be sent and the contact author will receive one(1) electronic copy of the journal issue in which the article appears.

15. This journal will consider for publication only papers that contain proofs for their listed results.

TABLE OF CONTENTS,JOURNAL OF CONCRETE AND APPLICABLE
MATHEMATICS,VOL.4,NO.4,2006

SPECIAL ISSUE ON “WAVELETS AND APPLICATIONS”,EDITED BY GUEST
EDITORS D.HONG,Y.SHYR

PREFACE,D.HONG,Y.SHYR,.....	391
AN ORTHOGONAL SCALING VECTOR GENERATING A SPACE C^1 CUBIC SPLINES USING MACROELEMENTS,B.KESSLER,.....	393
THE EXISTENCE OF TIGHT MRA MULTIWAVELET FRAMES,Q.MO,.....	415
FACE-BASED HERMITE SUBDIVISION SCHEMES,B.HAN,T.YU,.....	435
ON THE CONSTRUCTION OF LINEAR PREWAVELETS OVER A REGULAR TRIANGULATION,D.HONG,Q.XUE,.....	451
SPLINE WAVELETS ON REFINED GRIDS AND APPLICATIONS IN IMAGE COMPRESSION,L.GUAN,T.YANG,.....	473
VARIATIONAL METHOD AND WAVELET REGRESSION,J.WANG,.....	485
WAVELET APPLICATIONS IN CANCER STUDY,D.HONG,Y.SHYR,.....	505