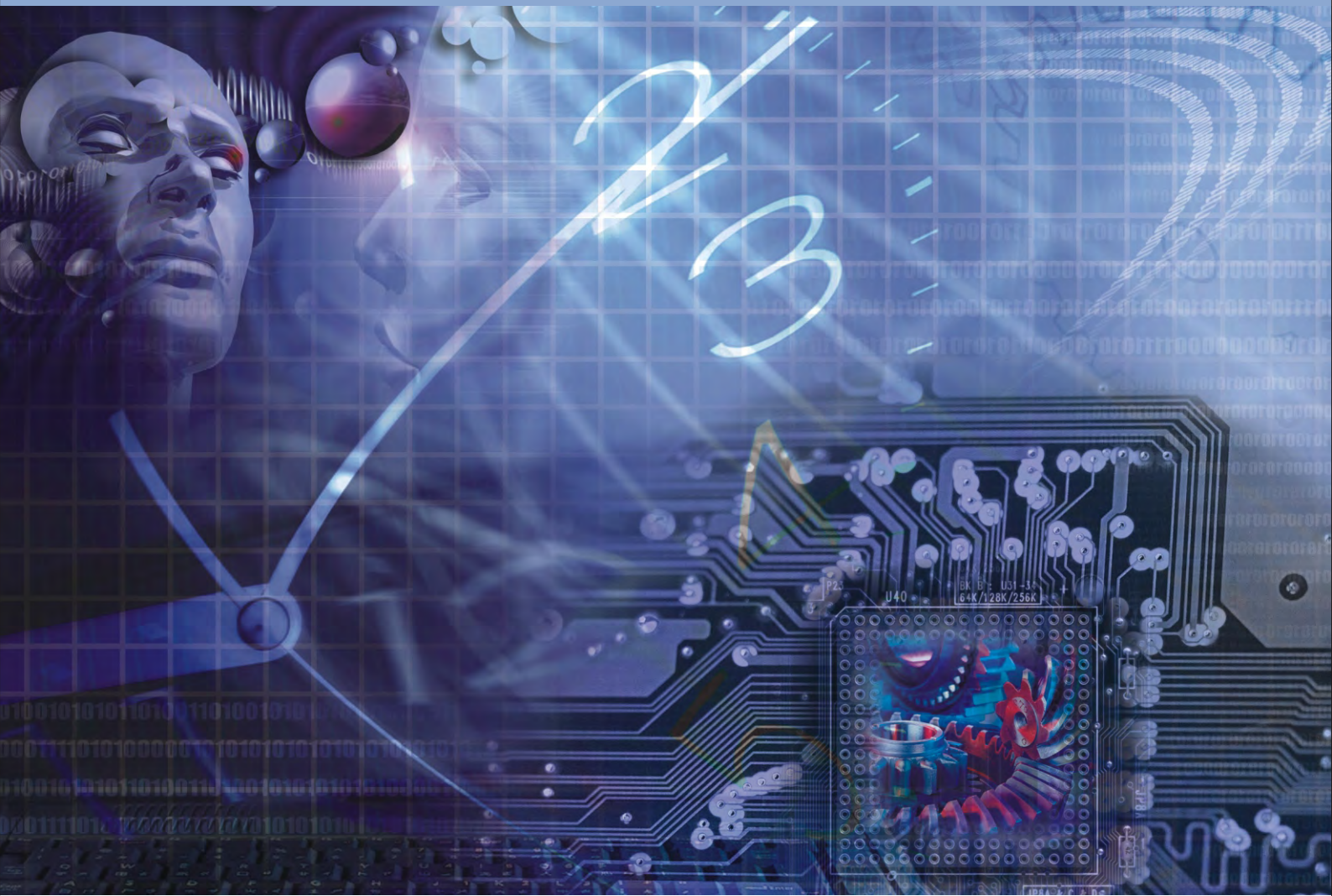


Intelligent Information Management

Editor-in-Chief: Dr. Bin Wu



Journal Editorial Board

ISSN: 2150-8194 (Print), 2150-8208 (Online)

<http://www.scirp.org/journal/iim>

Editor-in-Chief

Dr. Bin Wu

National Library of China, China

Editorial Board

Prof. Micheal Berry

Department of Electrical Engineering and Computer Science,
University of Tennessee, USA

Dr. Xilin Chen

Chinese Academy of Sciences, China

Prof. Zhoujun Li

Beihang University, China

Prof. Zhendong Niu

Beijing Institute of Technology, China

Prof. Riadh Robbana

Tunisia Polytechnic School, Tunisia

Prof. Yigang Sun

National Library of China, China

Dr. Junyong You

Norwegian University of Science and Technology, China

Dr. P. Vasant

PCO GLOBAL, Malaysia

Dr. Gabriele Milani

Technical University of Milan, Italy

Prof. Dilip Kumar Pratihar

Indian Institute of Technology, India

Prof. Ali Kaveh

Iran University of Science & Technology, Iran

Dr. M. Sohel Rahman

Bangladesh University of Engineering & Technology, Bangladesh

Dr. Leonardo Garrido

Monterrey Institute of Technology, Mexico

Dr. Daya Shanker

Indian Institute of Technology Roorkee, India

Dr. Chang-Hwan Lee

DongGuk University, Korea (South)

Dr. Pierluigi Siano

University of Salerno, Italy

Dr. Babak Forouraghi

Saint Joseph's University, USA

Dr. Alireza Givvehchi

Charite - University, Germany

Prof. Bingru Yang

University of Science and Technology Beijing, China

Dr. Jyh-Horng Chou

National Kaohsiung First University, Taiwan (China)

Dr. Maria do Carmo Nicoletti

University of New South Wales, Australia

Dr. Elizabeth Jacob

National Institute for Interdisciplinary Science and Technology, India

Prof. Dilip Kumar Pratihar

Indian Institute of Technology, India

Dr. Rashid Jayousi

Al-Quds University, Palestine

Prof. Damon Shing-Min Liu

National Chung Cheng University, Taiwan (China)

Dr. Weiru Liu

Queen's University Belfast, UK

Dr. Dalila B. M. M. Fontes

University of Porto, Portugal

Dr. George L. Caridakis

National Technical University of Athens, Greece

Dr. Yu Zhang

Rutgers University, USA

Dr. Wai-keung Wong

Hong Kong Polytechnic University, China

Dr. Marco Sacco

National Research Council), Milan, Italy

Table of Contents

Prototypicality Gradient and Similarity Measure: A Semiotic-Based Approach Dedicated to Ontology Personalization

X. Aimé, F. Furst, P. Kuntz, F. Trichet.....65

A Nonlinear Control Model of Growth, Risk and Structural Change

P. E. Petrakis, S. Kotsios.....80

Distortion of Space and Time during Saccadic Eye Movements

M. Suzuki, Y. Yamazaki.....90

Probabilistic Verification over $GF(2^m)$ Using Mod2-OBDDs

J. L. Imaña.....95

Theoretic and Numerical Study of a New Chaotic System

C. X. Zhu, Y. H. Liu, Y. Guo.....104

Sensing Semantics of RSS Feeds by Fuzzy Matchmaking

M. W. Yuan, P. Jiang, J. Zhu, X. N. Wang.....110

Text Extraction in Complex Color Document Images for Enhanced Readability

P. Nagabhushan, S. Nirmala.....120

Existence and Uniqueness of the Optimal Control in Hilbert Spaces for a Class of Linear Systems

M. Popescu.....134

Signed (b,k)-Edge Covers in Graphs

A. N. Ghameshlou, A. Khodkar, R. Saei, S. M. Sheikholeslami.....143

On the Mechanism of CDOs behind the Current Financial Crisis and Mathematical Modeling with Lévy Distributions

H. W. Du, J. L. Wu, W. Yang.....149

Intelligent Information Management (IIM)

Journal Information

SUBSCRIPTIONS

The *Intelligent Information Management* (Online at Scientific Research Publishing, www.SciRP.org) is published monthly by Scientific Research Publishing, Inc., USA.

E-mail: service@scirp.org

Subscription rates: Volume 2 2010

Print: \$50 per copy.

Electronic: free, available on www.SciRP.org.

To subscribe, please contact Journals Subscriptions Department, E-mail: service@scirp.org

Sample copies: If you are interested in subscribing, you may obtain a free sample copy by contacting Scientific Research Publishing, Inc. at the above address.

SERVICES

Advertisements

Advertisement Sales Department, E-mail: service@scirp.org

Reprints (minimum quantity 100 copies)

Reprints Co-ordinator, Scientific Research Publishing, Inc., USA.

E-mail: service@scirp.org

COPYRIGHT

Copyright© 2010 Scientific Research Publishing, Inc.

All Rights Reserved. No part of this publication may be reproduced, stored in a retrieval system, or transmitted, in any form or by any means, electronic, mechanical, photocopying, recording, scanning or otherwise, except as described below, without the permission in writing of the Publisher.

Copying of articles is not permitted except for personal and internal use, to the extent permitted by national copyright law, or under the terms of a license issued by the national Reproduction Rights Organization.

Requests for permission for other kinds of copying, such as copying for general distribution, for advertising or promotional purposes, for creating new collective works or for resale, and other enquiries should be addressed to the Publisher.

Statements and opinions expressed in the articles and communications are those of the individual contributors and not the statements and opinion of Scientific Research Publishing, Inc. We assume no responsibility or liability for any damage or injury to persons or property arising out of the use of any materials, instructions, methods or ideas contained herein. We expressly disclaim any implied warranties of merchantability or fitness for a particular purpose. If expert assistance is required, the services of a competent professional person should be sought.

PRODUCTION INFORMATION

For manuscripts that have been accepted for publication, please contact:

E-mail: iim@scirp.org

Prototypicality Gradient and Similarity Measure: A Semiotic-Based Approach Dedicated to Ontology Personalization

Xavier Aimé^{1,3}, Frédéric Furst², Pascale Kuntz³, Francky Trichet³

¹*Société TENNAXIA, Paris, France*

²*MIS - Laboratoire Modélisation, Information et Système University of Amiens, Amiens, France*

³*LINA - Laboratoire d'Informatique de Nantes Atlantique (UMR-CNRS 6241), University of Nantes,
Team "Knowledge and Decision", Nantes, France*

Email: xaime@tennaxia.com, frederic.furst@u-picardie.fr, {pascale.kuntz,francky.trichet}@univ-nantes.fr

Abstract

This paper introduces a new approach dedicated to the Ontology Personalization. Inspired by works in Cognitive Psychology, our work is based on a process which aims at capturing the user-sensitive relevance of the *categorization process*, that is the one which is really perceived by the end-user. Practically, this process consists in decorating the Specialization/Generalization links (*i.e.* the *is-a* links) of the hierarchy of concepts with 2 gradients. The goal of the first gradient, called Conceptual Prototypicality Gradient, is to capture the user-sensitive relevance of the categorization process, that is the one which is perceived by the end-user. As this gradient is defined according to the three aspects of the semiotic triangle (*i.e.* intentional, extensional and expressional dimension), we call it Semiotic based Prototypicality Gradient. The objective of the second gradient, called Lexical Prototypicality Gradient, is to capture the user-sensitive relevance of the lexicalization process, *i.e.* the definition of a set of terms used to denote a concept. These gradients enrich the initial formal semantics of an ontology by adding a pragmatics defined according to a context of use which depends on parameters like culture, educational background and/or emotional context of the end-user. This paper also introduces a new similarity measure also defined in the context of a semiotic-based approach. The first originality of this measure, called SEMIOSEM, is to consider the three semiotic dimensions of the conceptualization underlying an ontology. Thus, SEMIOSEM aims at aggregating and improving existing extensional-based and intentional-based measures. The second originality of this measure is to be context-sensitive, and in particular user-sensitive. This makes SEMIOSEM more flexible, more robust and more close to the end-user's judgment than the other similarity measures which are usually only based on one aspect of a conceptualization and never take the end-user's perceptions and purposes into account.

Keywords: Semantic Measure, Conceptual Prototypicality, Lexical Prototypicality, Gradient, Ontology Personalization, Semiotics

1. Introduction

This paper deals with Subjective knowledge, that is knowledge which is included in the semantic and episodic memory of Human Being [1]. Objective knowledge, which can be expressed through textual, graphic or sound documents, corresponds to what must be captured within a Domain Ontology, as it is specified by the consensual definition of T. Gruber [2]: "an ontology is a formal and explicit specification of a shared conceptualization". The

advent of the Semantic Web and the standardization of a Web Ontology Language (OWL) have led to the definition and the sharing of a lot of ontologies dedicated to scientific or technical fields. However, with the current emergence of Cognitive Sciences and the development of Knowledge Management applications in Social and Human Sciences, subjective knowledge becomes an unavoidable subject and a real challenge, which must be integrated and developed in Semantic Web, and more generally in Ontology Engineering. Our work aims at

providing measures dedicated to the personalization of a Domain Ontology (which by definition only includes Objective knowledge). This personalization process mainly consists in adapting the content of an ontology to its context of use. This latter usually implies an end-user and therefore mobilizes *Subjective knowledge* which is defined according to several parameters such as culture, educational background and emotional state. Our approach of ontology personalization aims at taking these parameters into account in order to reflect the relevance users of ontologies perceive on the is-a hierarchies and to what extent the terms associated to the concepts are representative.

Our work is based on the semiotic triangle, as defined by the linguists Ogden and Richard [3]. The three corners of this triangle are 1) the reference (*i.e.* the intentional dimension) which is an unit of thought defined from abstraction of properties common to a set of objects (*i.e.* the concepts of an ontology and their properties), 2) the referent (*i.e.* the extensional dimension) which corresponds to any part of the perceivable or conceivable world (*i.e.* the instances of concepts) and 3) the term (*i.e.* the expressional dimension) which is a designation of an unit of thought in a specific language (*i.e.* the linguistic expressions used to denote the concepts). The goal of our first gradient is to capture the user-sensitive relevance of the categorization process, that is the one which is perceived by the end-user. As this gradient is defined according to the three aspects of the semiotic triangle (*i.e.* intentional, extensional and expressional dimension), we call it *Semiotic-based Prototypicality Gradient*. The goal of our second gradient, called *Lexical Prototypicality Gradient*, is to capture the user-sensitive relevance of the lexicalization process, *i.e.* the definition of a set of terms used to denote a concept.

This paper also introduces a new similarity measure also defined in the context of a semiotic-based approach. The first originality of this measure, called SEMIOSEM, is to consider the three dimensions of the conceptualization underlying an ontology: the intention (*i.e.* the properties used to define the concepts), the extension (*i.e.* the instances of the concepts) and the expression (*i.e.* the terms used to denote both the concepts and the instances). Thus, SEMIOSEM aims at aggregating and improving existing extensional-based and intentional-based measures, with an original expressional one. The second originality of this measure is to be context-sensitive, and in particular user-sensitive. Indeed, SEMIOSEM exploits multiple information sources: 1) a textual corpus, validated by the end-user, which must reflect the domain underlying the ontology which is considered, 2) a set of

instances known by the end-user, 3) an ontology enriched with the perception of the end-user on how each property associated to a concept *c* is important for defining *c* and 4) the emotional state of the end-user. The importance of each source can be modulated according to the context of use and SEMIOSEM remains valid even if one of the sources is missing. This makes our measure more flexible, more robust and more closer to the end-user's judgment than the other similarity measures.

The rest of this paper is structured as follows. Section 2 presents the formal definition of the Semiotic-based Prototypicality Gradient (SPG) which corresponds to the aggregation of three components related to the *intentional*, the *expressional* and the *extensional* dimension of a conceptualization. Section 3 introduces the formal definition of the Lexical Prototypicality Gradient (LPG). Section 4 briefly introduces some well-known similarity measures and describes in detail SEMIOSEM: the basic foundations, the formal definitions, the parameters of the end-user and their interactions. Section 5 presents the tool TOOPRAG which implements our approach of ontology personalization; it also introduces some experimental results defined in two contexts: 1) a case study dedicated to the analysis of texts describing the Common Agricultural Policy (CAP) of the European Union and 2) a case study dedicated to Legal Intelligence within regulatory documents related to the domain "Health, Safety and Environment" (HSE).

2. Semiotic-Based Prototypicality Gradients (SPG)

Defining an ontology *O* of a domain *D* at a precise time *T* consists in establishing a consensual synthesis of individual knowledge belonging to a specific endogroup; an endogroup is a set of individuals which share the same distinctive signs and, therefore, characterize a community. For the same domain, several ontologies can be defined by different endogroups. We call *Vernacular Domain Ontologies* (VDO) this kind of resources¹. This property is also described by E. Rosch as *ecological* [4, 5], in the sense that although an ontology belongs to an endogroup, it also depends on the context in which it evolves. Thus, given a domain *D*, an endogroup *G* and a time *T*, a VDO depends on three factors, characterizing a precise context: 1) the culture of *G*, 2) the educational background of *G* and 3) the emotional state of *G*. In this way, a VDO can be associated to a pragmatic dimension. Indeed, a same VDO can be viewed (and used) from multiple points of view, where each point of view, although not reconsidering the formal semantics of *D*, allows us to adapt 1) the degrees of truth of the is-a links defined between concepts and 2) the degrees of expressivity of the terms used to denote the concepts. We call *Personalized Vernacular Domain Ontologies* (PVDO) this kind of resources. Our work is based on the funda-

¹*Vernacular*, which comes from the latin word *vernaculus*, means native. For instance, vernacular architecture, which is based on methods of building which use locally available resources to address local needs, tends to evolve over time to reflect the environmental, cultural and historical context in which it exists.

mental idea that all the sub-concepts of a decomposition are not *equidistant* members, and that some sub-concepts are more representative of the super-concept than others. This phenomenon is also applicable to the set of terms used to denote a concept. This assumption is validated by works in Cognitive Psychology [1,6]. Formally, our gradient is based on a Vernacular Domain Ontology (VDO), given a field D and an endogroup G . This type of ontology is defined by the following tuple: $O_{(D,G)} = \{C, P, I, \Omega_{(D,G)}, \leq^C, \sigma^P, L\}$ where:

- C, P, I represent respectively the disjointed sets of concepts, properties² and instances;
- $\Omega_{(D,G)}$ is a set of documents (e.g. text, graphic or sound documents) related to a domain D and shared by the members of the endogroup G ;
- $\leq^C: C \times C$ a partial order on C defining the hierarchy of concepts ($\leq^C(c_1, c_2)$ means that the concept c_1 subsumes the concept c_2);
- $\sigma^P: P \rightarrow C \times C$ defines the domain and the range of a property;
- $L = \{L_C, ftermc\}$ is the lexicon related to the dialect of G where a) L_C represents the set of terms associated to C , b) the function $ftermc: C \rightarrow (L_C)^n$ which returns the tuple of terms used to denote a concept.

We define $spg_{G,D}: C \times C \rightarrow [0; 1]$ the function which, for all pairs of concepts $c_f, c_p \in C$ such as it exists an is a link between the super-concept c_p and the sub-concept c_f , returns a real (null or positive value) which represents the conceptual prototypicality gradient of this link, in the context of a PVDO dedicated to a domain D and an endogroup G . For two concepts c_p and c_f , this function is formally defined as follows:

$$\begin{aligned} spg_{G,D}(c_p, c_f) = & [\alpha * \text{intentional}(c_f, c_p) \\ & + \beta * \text{expressional}_{G,D}(c_f, c_p) \\ & + \gamma * \text{extensional}_{G,D}(c_f, c_p)]^\delta \end{aligned} \quad (1)$$

with 1) $\alpha + \beta + \gamma = 1$, where $\alpha \geq 0$ is a weighting of the intentional component, $\beta \geq 0$ is a weighting of the expressional component, $\gamma \geq 0$ is a weighting of the extensional component, and 2) $\delta \geq 0$ is a weighting of the mental state of the endogroup G . Our gradient is based on the three dimensions introduced by Morris and Peirce in their theory of semiotics [7]: 1) the signified, *i.e.* the concept defined in intention, 2) the referent, *i.e.* the concept defined in extension via its instances, and 3) the signifier, *i.e.* the terms used to denote the concept. The main advantage of our approach is that 1) it integrates both the intentional, extensional and expressional dimension of a conceptualization for defining how a sub-concept is *representative/typical* of its super-concept and the influence of these dimensions can be modulated via the α, β and γ parameters and 2) it allows us to modulate this representativeness according

to an emotional dimension via the δ parameter. The values of α, β and γ are defined manually according to the context of the ontology personalization process. Indeed, when no instances (or few) are associated to the ontology then it is relevant to minimize the influence of the extensional dimension by assigning a low value to γ . In a similar way, when the ontology does not include properties then it is relevant to minimize the influence of the intentional dimension by assigning a high value to α . And when the ontology is associated to a huge and rich textual corpora then it is relevant to maximize the influence of the expressional dimension by assigning a high value to β given that $\alpha + \beta + \gamma = 1$. The value of δ is used to modulate the influence of the emotion state on the perception of the conceptualization. Multiple works on the influence of emotions on human evaluation have been done in psychology [8, 9]. The conclusion of these works can be summarized as follows: when we are in a negative mental state (e.g. fear or nervous breakdown), we tend to focus us on what appears to be the more important from an emotional point of view. In our context, this consists in reducing the universe to what is very familiar; for instance, our personal dog (or the one of a neighbor) which at the beginning is inevitably the most characteristic of the category becomes the quasi unique and quasi unique dog. Respectively, in a positive mental state (e.g. love or joy), we are more open in our judgment and we accept more easily the elements which are not yet be considered as so characteristic. According to [10], a *negative* mental state leads to the reduction of the value of representation, and conversely for a *positive* mental state. Thus, we characterize: 1) a *negative* mental state by a value $\delta \in]1, +\infty[$, 2) a *positive* mental state by a value $\delta \in]0, 1[$, and 3) a *neutral* mental state by the value 1. When the value of δ is low, the value of the gradients associated to the concepts which are initially not considered as being so representative increases considerably, because a positive state favours openmindedness, self-actualization, etc. Conversely, when the value of δ is high (*i.e.* a strongly negative mental state), the effect is to only *select* the concepts which own a high value of typicality, eliminating *de facto* the other concepts.

2.1. Intentional Component

The intentional component of our gradient aims at taking 1) the structure of a conceptualization and 2) the intentional definition of its components into account. In order to compare two concepts from an intentional point of view, we propose a measure based on the properties shared by the sub-concepts as developed in [11,12]. For each concept $c \in C$, we define a *Characteristic Vector* (CV) $vc = (v_{c1}, v_{c2}, \dots, v_{cn})$ with $n = |P|$, and $v_{ci} \in [0, 1]$, $\forall i \in [1, n]$ a weight assigned to each

² Properties include both attributes of concepts and domain relations.

property p_i of P . A concept is defined by the union of all the properties whose weight is not null. The set of concepts corresponds to a point cloud defined in a space with $|P|$ dimensions. Weights associated with the properties have to satisfy a constraint related to the *isa* relationship: a concept c is subsumed by a concept d (noted $\leq^C(d, c)$) if and only if $v_{ci} \geq v_{di}$, $\forall i \in [1, n]$ with $n = |P|$. For any $c \in C$, we define a prototype concept from all the sub-concepts of c . This prototype concept is characterized by a *Prototype Vector* (PV) $tc = (t_{c1}, t_{c2}, \dots, t_{cn})$. Prototype concepts correspond to summaries of semantic features characterizing categories of concepts. They are stored in the episodic memory, and are used in the process of categorization per comparison. In our work, we consider that a prototype concept of a concept c corresponds to the barycenter of the point cloud formed by the set of the CV of all concepts belonging to the descent of c ³. Thus, the Prototype Vector of a concept c is formally defined as follows:

$$\vec{t}_c = \frac{1}{\sum_{s \in S} \lambda(s)} \sum_{s \in S} \lambda(s) \vec{v}_s \quad (2)$$

where:

- $\lambda(s)$ is equal to $(depthtree(c) - depth(s) + 1) / depthtree(c)$ with 1) $depthtree(c)$, the depth of the sub-tree having for root c and 2) $depth(s)$, the depth of s in the sub-tree having for root c ;
- S , the set of concepts belonging to the descent of c .

The objective of the coefficient $\lambda(s)$ (for a concept s) is to relativize the properties which are hierarchically distant from the super-concept (cf. the use of the ratio of depths)⁴. Note that if we consider only the sub-concepts of c (i.e. only one level of hierarchy), then $\lambda(s) = 1$, $\forall s \in S$ and we find the formula of the PV defined by [12]:

$$\vec{t}_c = \frac{1}{|S|} \sum_{s \in S} \vec{v}_s \quad (3)$$

In our work, we advocate the following principle: the more a concept is close to the prototype concept, the more it is representative of its super-concept. We consider this value as being the Euclidean normalized distance between 1) the PV of the super-concept and 2) the CV of the subconcept which is considered; it corresponds to the normalized distance between a point and the barycenter of the point cloud. The function *intentional*: $C \times C \rightarrow [0, 1]$ is formally defined as follows:

$$intentional(c_f, c_p) = 1 - dist(\vec{t}_{c_p}, \vec{v}_{c_f}) \quad (4)$$

The more the value of this function is near to 1, the more the concept c_f is *representative/typical* of the concept c_p , from an intentional point of view.

2.2. Expressional Component

The expressional component of our gradient aims at taking the expressional view of a conceptualization into account, through the terms used to denote the concepts. This approach is based on the observation frequency of a concept related to a domain D , in an universe of the endogroup G . In this way, the more an element is frequent in the universe, the more it is considered as *representative/typical* of its category. This notion of typicality is introduced in the work of E. Rosch [4,5]. In our context, the universe of an endogroup is composed of the set of documents identified by $\Omega_{(D,G)}$. Our approach is inspired by the idea of Information Content introduced by Resnik [13]. Indeed, this is not because an idea is often expressed that it is really true and objective. Psychologically, it is acknowledged that the more an event is presented (in a frequent way), the more it is judged probable without being really true for an individual or an endogroup; this is one of the ideas supported by A. Tversky in its work on the evaluation of uncertainty [14]. The function *expressional*: $C \times C \rightarrow [0, 1]$ is formally defined as follows⁵:

$$expressional_{G,D}(c_f, c_p) = \frac{Info(c_f)}{Info(c_p)} \quad (5)$$

where :

$$Info(c_f) = \sum_{term \in world(c)} \left(\frac{count(term)}{N} * \frac{count(doc, term)}{count(doc)} \right) \quad (6)$$

with:

- *Info(c)* defines the information content of the concept c ;
- *count(term)* returns the weighting number of term occurrences in the documents of $\Omega_{(D,G)}$. Note that this function takes the structure of the documents into account. Indeed, in the context of a scientific article, an occurrence of a term t located in the keywords section is more important than another occurrence of t located in the summary or the body of text. Thus, this function is formally defined as follows:

$$count(term) = \sum_{i=1}^m M_{term,i} \quad (7)$$

³We understand by descent all the sub-concepts of c , from generation 1 to n (i.e. the leaves).

⁴Contrary to [12], we propose to extend the calculation of the prototype to all the descent of a concept, and not only to its direct sub-concepts (i.e. only one level of hierarchy). Indeed, we think that all the concepts belonging to the descent (and in particular the leaves) contribute to the definition of the prototype from a cognitive point of view.

⁵This function is only applicable if it exists 1) a direct *isa* link between the super-concept c_p and the sub-concept c_f , with an order relation $c_f \leq c_p$, or 2) an indirect link composed of a serie of *isa* links between the c_p and c_f .

where $M_{term,i} \in Z$ is the hierarchical coefficient relating to the position (in the structure of the document) of the i th occurrence of the term. The values of these coefficients are fixed in a manual and consensual way by the members of the endogroup.

- $count(doc, term)$ returns the number of documents of $\Omega_{(D,G)}$ where the term appears;
- $count(doc)$ returns the number of documents of $\Omega_{(D,G)}$;
- $world(c)$ returns all the terms concerning the concept c via the function f_{termC} and all its sub-concepts from generation 1 to generation n ;
- N is the sum of all the weighting numbers of occurrence of all the terms contained in $\Omega_{(D,G)}$.

Intuitively, the function $Info(c)$ measures “the ratio of use” of a concept in an universe, by using first the terms directly associated to the concept and then, by using the terms associated to all its sub-concepts, from generation 1 to generation n . We balance each frequency by the ratio between the number of documents where the term is present and the global number of documents. An idea which is frequently presented in few documents is less relevant than an idea which is perhaps less supported in each document but which is presented in a lot of documents of the endogroup’s universe.

2.3. Extensional Component

The extensional component of our gradient aims at taking the extensional view of a conceptualization into account, through the instances. This approach is based on the quantity of instances of a concept related to a domain D , in an universe of the endogroup G . In this way, the more a concept is frequent in the universe (because it owns a lot of instances), the more it is considered as representative/typical of its category. The function $extensional: C \times C \rightarrow [0, 1]$ is formally defined as follows:

$$extensional_{G,D}(c_f, c_p) = \frac{1}{1 - \log\left(\frac{count_i(c_f)}{count_i(c_p)}\right)} \quad (8)$$

Where the function $count_i(c): C \times I \rightarrow Z$ return the number of instances $i \in I$ of a concept $c \in C$. The form $(1/1 - \log(x))$ has been adopted in order to obtain a non-linear behavior which is more close to human judgment.

3. Lexical Prototypicality Gradient (LPG)

The goal of the Lexical Prototypicality Gradient (LPG) is to evaluate the fact that the terms used to denote a concept have not the same representativeness within the endogroup. Indeed, the question is the following: “why do we more frequently name the concept x with the term y

rather than z ?” To define these lexical variations, we propose to adapt the expressional component of our conceptual gradient previously defined. Thus, our formula is based on the Information Content of a concept, by using the ratio between the frequency of use of the term and the sum of the appearance frequencies of all the terms related to the concept in $\Omega_{(D,G)}$. We define $lpg_{G,D}: L_C \times C \rightarrow [0,1]$ the function, which for all concept $c \in C$ and the term $t \in L_C$ such as $t \in f_{termC}(c)$, returns a positive or null value representing the lexical prototypicality gradient of this term, and this for a domain D and an endogroup G . This function is formally defined as follows:

$$lpg_{G,D}(t, c) = \frac{1}{1 - \log\left(\frac{count(t)}{\sum count(f_{termC}(c))}\right)} \quad (9)$$

4. SEMIOSEM: A Semiotic-Based Similarity Measure

Currently, the notion of similarity has been highlighted in many activities related to Ontology Engineering such as ontology learning, ontology matching or ontology population. In the last few years, a lot of measures for defining concept (dis-)similarity have been proposed. These measures can be classified according to two approaches: 1) extensional-based measures such as [13,15–16] or [17] and 2) intentional based measures such as [18,19] or [20]. Most of these measures only focus on one aspect of the conceptualization underlying an ontology, mainly the intention through the structure of the subsumption hierarchy or the extension through the instances of the concepts or the occurrences of the concepts in a corpus. Moreover, they are usually sensitive to the structure of the subsumption hierarchy (because of the use of the more specific common subsumer) and, therefore, they are dependent on the modeling choices. Finally, these measures never consider the end-user’s perceptions of the domain which is considered [21]. Our goal is to provide a measure more flexible, more robust and more closer to the end-user’s judgment than the other similarity measures.

4.1. Current Intentional-Based Similarity Measures

Intentional-based measures are founded on the analysis of the structure of a semantic network. In the field of ontology engineering, the hierarchy of concepts is considered as a directed graph (where the nodes correspond to the concepts and the edges correspond to taxonomic links). Intuitively, these works are based on the following principle: concept A is more similar to concept B than concept C , if the distance from A to B (in the graph) is shorter than the one from A to C . This distance can be

calculated following different ways. [18] considers this distance, noted $dist_{edge}(c_1, c_2)$, as being the length of the shortest path between two concepts. The similarity between c_1 and c_2 is defined as follows:

$$Sim_{Rad}(c_1, c_2) = \frac{1}{dist_{edge}(c_1, c_2)} \quad (10)$$

[13] enhances this definition by introducing the maximum depth of the hierarchy, noted max .

$$Sim_{Res}(c_1, c_2) = \frac{2 * max}{dist_{edge}(c_1, c_2)} \quad (11)$$

[19] standardizes this latter measure in order to obtain results in the interval $[0, 1]$.

$$Sim_{Lea}(c_1, c_2) = -\log\left(\frac{dist_{edge}(c_1, c_2)}{2 * max}\right) \quad (12)$$

[20] suggests another measure which takes the depth of the concepts into account. The similarity between c_1 and c_2 , with $depth(c_i)$ the depth of the concept c_i in the hierarchy and c the Most Specific Common Subsumer (MSCS) of c_1 and c_2 , is defined as follows:

$$Sim_{wu}(c_1, c_2) = \frac{2 * depth(c)}{depth(c_1) + depth(c_2)} \quad (13)$$

These measures only consider the is-a links between concepts and not the semantics of concepts. Thus, they can be incorrect (concepts with high similarity can be semantically far from each other) or incomplete (for concepts semantically similar but not very close in the hierarchy, the measure can be very low). Another intention-based approach for defining a similarity measure consists in analyzing and comparing the properties of the concepts. For illustration, let us consider two objects on which we can sit down: an armchair and a chair. These two objects share common properties and there are other properties which differentiate them without ambiguity. From the comparison of these properties, it is possible to state that a chair is more close to an armchair than a stool. Indeed, in set theory, we can state that two concepts are close if the number of common properties is greater than the number of distinct properties. [14] suggests the following function:

$$Sim_{Tversky}(c_1, c_2) = \alpha.comm(c_1, c_2) - \beta.diff(c_1, c_2) - \lambda.diff(c_2, c_1) \quad (14)$$

where α, β, γ are constants.

4.2. Current Extensional-Based Similarity Measures

Extensional-based measures have first been inspired by the measure of Jaccard [22]:

$$Sim_{Jaccard}(c_1, c_2) = \frac{|I_{c_1} \cap I_{c_2}|}{|I_{c_1}| + |I_{c_2}| - (|I_{c_1} \cap I_{c_2}|)} \quad (15)$$

where $|I_c|$ is the number of instances of the concept c .

According to [17], this approach is not really appropriate to ontologies because two concepts can be similar without having common instances. [17] proposes a new measure which does not evaluate the extension overlap but the variation of the cardinality of the extensions of the considered concepts w.r.t. the cardinality of the extension of their MSCS.

$$Sim_{Amd}(c_1, c_2) = \frac{\min(|I_{c_1}|, |I_{c_2}|)}{|I_{gcs(c_1, c_2)}|} \left(1 - \frac{|I_{gcs(c_1, c_2)}|}{|I|}\right) \left(1 - \frac{\min(|I_{c_1}|, |I_{c_2}|)}{|I_{gcs(c_1, c_2)}|}\right) \quad (16)$$

where $gcs(c_1, c_2)$ is the MSCS of c_1 and c_2 , $|I_c|$ is the number of instances of the concept c , and $|I|$ is the total number of instances of the ontology.

Most of the current extensional-based measures are founded on the notion of "Information Content" (IC) introduced by Resnik [13]. The main idea consists in measuring the similarity of concepts on the ground on the amount of information that they share. [13] approximates the IC of a concept c to the probability $p(c)$ to have occurrences of c in a given corpus. Thus, IC is defined as follows:

$$\psi(c) = -\log(p(c)) \quad (17)$$

where :

$$p(c) = \frac{\sum_{n \in word(c)} count(n)}{N}$$

N represents the number of the occurrences of the terms of all the concepts in the corpus, $words(c)$ represents the set of terms used to denote the concept c and all its sub-concepts. This measure assumes that each term is associated to one and only one concept. [23] deals with this problem by modifying the calculation of the appearance frequency of a term by:

$$p(c) = \frac{\sum_{n \in word(c)} \frac{count(n)}{nb_{classe}(n)}}{N} \quad (18)$$

The similarity measure advocated by [13] is based on the most informative common subsumer of c_1, c_2 (i.e. the common subsumer which owns the most important IC; this is not necessary the MSCS of c_1, c_2). The similarity between c_1 and c_2 , where $S(c_1, c_2)$ is the set of concepts which subsume both c_1 and c_2 , is defined as follows:

$$Sim_{Res2}(c_1, c_2) = \max_{c \in S(c_1, c_2)} \psi(c) \quad (19)$$

[15] suggests another measure based on the common IC of the concepts. The similarity between c_1 and c_2 , with ppc the concept in $S(c_1, c_2)$ which minimizes $\psi(c)$, is

defined as follows:

$$Sim_{Lin}(c_1, c_2) = \frac{2 * \psi(ppc)}{\psi(c_1) + \psi(c_2)} \quad (20)$$

Based on this approach of IC, [16] advocates the following measure:

$$Sim_{jiang}(c_1, c_2) = \sum_{c \in path(c_1, c_2) - MSCS(c_1, c_2)} [\psi(c) - \psi(SC(c))] * TC(c, SC(c)) \quad (21)$$

where $TC(c_i, c_j)$ is a weighting of the edge connecting c_i to c_j such as $c_j = SC(c_i)$, and $SC(c)$ is the super-concept of c .

4.3. Semioseme: A Semiotic-Based Similarity Measure

SEMIOSEM is a conceptual similarity measure defined in the context of a PVDO; it takes as input a PVDO and the three additional resources:

- a set of **instances** supposed to be representative of the end-user's conceptualization (for instance, in the case of a business information system, these instances are the customers the end-user deals with, the products he/she sells to them, etc);
- a **corpus** given by the end-user and supposed to be representative of its conceptualization (for instance, this corpus can be the documents written by the end-user on a blog or a wiki);
- a **weighting of the properties** associated to each concept. The weights quantify the importance of the end-user associates to the properties in the definition of the concept. These weightings are fixed by the end-user as follows: for each property $p \in P$, the user ordines, on a 0 to 1 scale, all the concepts having p , in order to reflect its perception on how p is important for defining c . For instance, for the property has an author, the concept *Scientific Article* will be put first, secondly the concept *Newspaper Article*, for which the author is less important, thirdly the concept *Technical Manual*.

Thus, SEMIOSEM corresponds to an aggregation of three components:

- an *intentional* component based on the comparison of the properties of the concepts;
- an *extensional* component based on the comparison of the instances of the concepts;
- an *expressional* component based on the comparison of the terms used to denote the concepts and their instances.

Each component of SEMIOSEM is weighted depending on the way the end-user apprehends the domain he/she is working on (e.g. giving more importance to the intentional component when the end-user better apprehends the domain via an intentional approach rather than an extensional one). These differences of importance are conditioned by the domain, the cognitive universe of the end-user and the context of use (e.g. ontology-based in-

formation retrieval). For instance, in the domain of the animal species, a zoologist will tend to conceptualize them in intention (via the biological properties), whereas the majority of people use more extensional conceptualizations (based on the animals they have met during their life). Formally, the function SEMIOSEM: $C \times C \rightarrow [0, 1]$ is defined as follows:

$$SemioSem(c_1, c_2) = [\alpha * intens(c_1, c_2) + \beta * extens(c_1, c_2) + \gamma * express(c_1, c_2)]^{\frac{1}{\delta}} \quad (22)$$

The following sections present respectively the functions *intens* (cf. Section 4–C1), *extens* (cf. Section 4–C2), *express* (cf. Section 4–C3) and the parameters α , β , γ , δ (cf. Section 4–C4).

1) *Intentional component*: From an intentional point of view, our work is inspired by [12] and is based on the representation of the concepts by vectors in the space of the properties. Formally, to each concept $c \in C$ is associated a vector $!vc = (v_{c1}, v_{c2}, \dots, v_{cn})$ where $n = |P|$ and $v_{ci} \in [0; 1]; \forall i \in [1; n]$. v_{ci} is the weighting fixed by the end-user which precises how the property i is important for defining the concept c (by default, v_{ci} is equal to 1). Thus, the set of concepts corresponds to a point cloud defined in a space with $|P|$ dimensions. We calculate a prototype vector of c_p , which was originally introduced in [12] as the average of the vectors of the subconcepts of c_p . However [12] only considers the direct subconcepts of c_p , whereas we extend the calculation to all the sub-concepts of c_p , from generation 1 to generation n . Indeed, some properties which can only be associated to indirect subconcepts can however appear in the prototype of the superconcept, in particular if the intentional aspect is important. Thus, the prototype vector pcp is a vector in the space properties, where the importance of the property i is the average of the importances of the properties of all the sub-concepts of c_p having i . If for $i \in P$, $S_i(c) = \{c_j \leq^C c, c_j \in dom(i)\}$, then:

$$\vec{p}_{cp}[i] = \frac{\sum_{c_j \in Si(cp)} \vec{v}_{c_j}[i]}{|S_i(c_p)|} \quad (23)$$

From an intentional point of view, the more the respective prototype vectors of c_1 and c_2 are close in terms of euclidean distance (i.e. the more their properties are close), the more c_1 and c_2 are similar. This evaluation is performed by the function *intens*: $C \times C \rightarrow [0, 1]$, which is formally defined as follows:

$$intens(c_1, c_2) = 1 - dist(\vec{p}_{c_1}, \vec{p}_{c_2}) \quad (24)$$

2) *Extensional component*: From an extensional point of view, our work is based on the Jaccard's similarity (cf. Section 4–B). Formally, the function *extens*: $C \times C \rightarrow [0, 1]$

is defined as follows:

$$\text{extens}(c_1, c_2) = \frac{|\sigma(c_1) \cap \sigma(c_2)|}{|\sigma(c_1)| + |\sigma(c_2)| - (|\sigma(c_1) \cap \sigma(c_2)|)} \quad (25)$$

This function is defined by the ratio between the number of common instances and the total number of instances minus the number of having common instances. Thus, two concepts are similar when they have a lot of instances in common and few distinct instances.

3) *Expressional component*: From an expressional point of view, the more the terms used to denote the concepts c_1 and c_2 are present together in the same documents of the corpus, the more c_1 and c_2 are similar. This evaluation is carried out by the function *express*: $C \times C \rightarrow [0, 1]$ which is formally defined as follows:

$$\text{express}(c_1, c_2) = \sum_{t_1, t_2} \left(\frac{\min(\text{count}(t_1), \text{count}(t_2))}{N_{\text{occ}}} * \frac{\text{count}(t_1, t_2)}{N_{\text{doc}}} \right) \quad (26)$$

With :

- $t_1 \in \text{words}(c_1)$ and $t_2 \in \text{words}(c_2)$ where $\text{words}(c)$ returns all the terms denoting the concept c or one of its sub-concept (direct or not);
- $\text{count}(t_i)$ returns the number of occurrences of the term t_i in the documents of the corpus ;
- $\text{count}(t_1, t_2)$ returns the number of documents of the corpus where the term t_1 and t_2 appear simultaneously;
- N_{doc} returns the number of documents of the corpus;
- N_{occ} is the sum of the numbers of occurrences of all the terms included in the corpus.

4) *Parameters of SEMIOSEM*: α, β, γ are the (positive or null) weighting coefficients associated to the three components of SEMIOSEM. In a way of standardization, we impose that the components vary between 0 and 1 and that $\alpha + \beta + \gamma = 1$. The values of these three coefficients can be fixed arbitrarily, or evaluated by experiments. We also advocate a method to automatically calculate approximations of these values. This method is based on the following principle. As shown by Figure 1, we consider that the relationship between α, β, γ , characterizes the cognitive coordinates of the end-user in the semiotic triangle. To fix the values of α, β, γ , we propose to systematically calculate γ/α , and γ/β , and then to deduce α from the constraint $\alpha + \beta + \gamma = 1$. γ/α (resp. γ/β) is approximated by the cover rate of the concepts (resp. the instances) within the corpus. This rate is equal to the number of concepts (resp. instances) for which at least one of the terms appears in the corpus divided by the total number of concepts (resp. instances). The factor $\delta \geq 0$ aims at taking the mental state of the end-user into account. Multiple works have been done in Cognitive Psychology on the relationship between human emotions and judgments [8]. The conclusion of these works can be summarized as follows: when we are in a negative mental state on what appears to be the more important from an

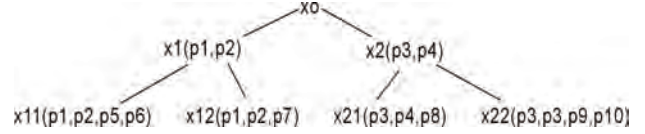


Figure 1. Intuitive example.

state (e.g. fear or nervous breakdown), we tend to focus emotional point of view. Respectively, in a positive mental state (e.g. love or joy), we are more open-minded in our judgment and we accept more easily the elements which are not yet be considered as so characteristic. According to [10], a negative mental state supports the reduction in the value of representation, and conversely for a positive mental state. In the context of our measure, this phenomenon is modelised as follows. We characterize 1) a *negative* mental state by a value $\delta \in]1, +\infty[$, 2) a *positive* mental state by a value $\delta \in]0, 1[$, and 3) a *neutral* mental state by a value of 1. Thus, a low value of δ , which characterizes a positive mental state, leads to increase the similarity values of concepts which initially would not been considered as so similar. Conversely, a strong value of δ , which characterizes a negative mental state, leads to decrease these values.

5. Experimental Results

5.1. Comparisons with Human Judgment

In order to evaluate the relevance of our gradients, we have done a first experiment with a group of 19 students of the University of Nantes. The experiment mainly consisted in 1) building by hand a vernacular domain ontology from texts, 2) personalizing this ontology for each student by using the prototypicality gradients, and 3) comparing the results which have been computed to the judgment of each student (*i.e.* its *personal categorization*). The first step of the process was focused on the construction of a vernacular domain ontology. The domain which has been adopted is delimited by the *Grenelle Environment Round Table*, an open multi-party debate in France that brings together representatives of national and local governments and organizations. The aim of this round-table is to define the key points of a public policy on ecological and sustainable development issues for the next 5 years. Each student has selected 15 texts from several web sites (e.g. industry, political party, professional associations or non-governmental organizations).

Then, for each text, each student has stated a list of the most salient terms (between 10 and 15 per text). The union of all the terms selected individually has conducted to a set of 350 terms⁶ which clearly denote the

⁶These terms are only considered as relevant in the following context: 1) a specific endogroup (composed of our 19 students), 2) a delimited domain and 3) a corpus (composed of the texts selected by the students).

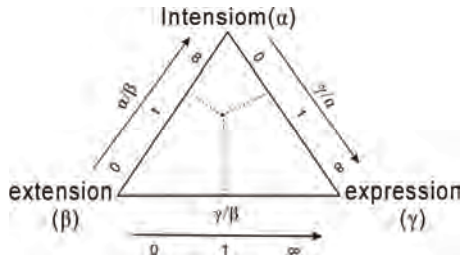


Figure 2. The weighting coefficients considered as the cognitive coordinates of the end-user in the semiotic triangle. γ/α near to 0 indicates that the end-user apprehends the domain in intention (and not in expression); the same ratio close to 1, indicates the opposite and a ratio close to 1 indicates that the intentional and expressional approaches are equilibrated. A similar interpretation is adopted for the other ratios. When the three approaches are equilibrated, $\alpha=\beta=\gamma=1/3$ and the cognitive coordinates of the end-user correspond to the barycenter of the semiotic triangle.

concepts included in the *Grenelle Environment Round Table*. From this terms, we have build a hierarchy composed of 130 concepts (depth = 3, max width = 9). Figure 2 presents an extract of this hierarchy. The second step the process was focused on the personalization of this ontology for each student via the calculation of the prototypicality gradients. For this purpose, we have only used the expressional component of our approach (i.e. we didn't define properties, nor instances). First, we asked to each student to consider a concept (e.g. the one which owns the most sub-concepts), and then to sort all its subconcepts by order of typicality. Then, each student has selected several texts (extracted from the web) which, for him, clearly concern the *Grenelle Environment Round Table* and clearly respect its opinions and convictions on this subject. From these texts, each student has calculated the expressional component between all the sub-concepts and the considered concept. The last step of the process was focused on the comparison of the results ($spg_{G,D}$ value) to human judgment (cf. Table 1). (89%) (17 students) have obtained results completely similar or very close to their opinion, and 11% (2 students) obtained different results because their personal corpus didn't really match with their real vision of the subject. This experiment reveals us that 1) the quality of the personalization process is dependent on the composition of the textual corpus: can we consider all the documents selected by the end-user such as mails, texts, web sites or blogs? and 2) the personalization process necessarily requires an adapted corpus. We currently deal with another experiment in collaboration with the University of Bretagne Sud. The main objective of this experiment is to measure the influence of the emotions on cognitive processes, such as for instance how arousal influences the perception of the less typical concepts.

⁷Lucene is a high-performance, full-featured text search engine library written entirely in Java. Lucene is an open source project available at <http://lucene.apache.org/>.

Table 1. Values of spg for a student, with the concept *Resource*.

c	$spg_{G,D}(resource,c)$	Human judgement
energy	0.93	1 st
water	0.41	2 nd
money	0.01	3 rd
material	0.03	4 th

5.2. TOOPRAG: A Tool Dedicated to the Pragmatics of Ontology

TOOPRAG (A Tool dedicated to the Pragmatics of Ontology) is a tool dedicated to the automatic calculation of our gradients. This tool, implemented in Java 1.5, is based on Lucene⁷ and Jena⁸. It takes as inputs 1) an ontology represented in OWL 1.0, where each concept is associated to a set of terms defined via the primitive *rdfs:label* (for instance, `<rdfs:label xml:lang="EN">farmer </rdfs:label>`) and 2) a corpus composed of text files. Thanks to the Lucene API, the corpus is first indexed. Then, the ontology is loaded in memory (via the Jena API) and the SPG values of all the is-a links of the concepts hierarchies are computed. The LPG values of all the terms used to denote the concepts are also computed.

These results are stored in a new OWL file which extends the current specification of OWL 1.0. Indeed, as shown by Figure 4, a LPG value is represented by a new attribute *xml:lpg* which is directly associated to the primitive *rdfs:label*. For instance, the LPG values of the terms "grower" and "peasant", used to denote the concept "agricultural labour force" (`<owl:Class rdf:ID="agricultural-labour-force">`), are respectively 0.375 and 0. In a similar way, a SPG is represented by a new attribute *xml:cpg*⁹ which is directly associated to the primitive *rdfs:subClassOf*. For instance, the SPG values of the is a links defined between the superconcept "working-population-engaged-in-agriculture" and its sub-concepts "agricultural-labour-force", "farmer", "for-stranger" and "agricultural-adviser" are respectively 0.0074, 0.9841, 0 and 0.

5.3. Distributional Analysis of the SPG

In order to analyze the statistical distribution of the SPG values on different types of hierarchies of concepts,

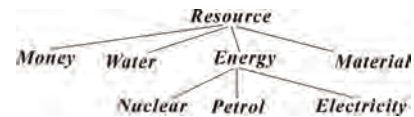


Figure 3. Extract of a hierarchy of concepts about grenelle environment round table.

⁸Jena is a Java framework for building Semantic Web applications. It provides a programmatic environment for RDF, RDFS, OWL and SPARQL and includes a rule-based inference engine. Jena is an open source project available at <http://jena.sourceforge.net/>.

⁹This attribute is called CPG for Conceptual Prototypicality Gradient; SPG and CPG are synonyms.

```

...
<owl:Class rdf:ID="agricultural labour force">
  <rdfs:label xml:lang="EN" xml:lp=0.7>farm worker</rdfs:label>
  <rdfs:label xml:lang="EN" xml:lp=0.3>agricultural labour force</rdfs:label>
  <rdfs:subClassOf rdf:resource="#working population engaged in agriculture" xml:cpg=0.0074/>
</owl:Class>

<owl:Class rdf:ID="farmer">
  <rdfs:label xml:lang="EN" xml:lp=0.375>grower</rdfs:label>
  <rdfs:label xml:lang="EN" xml:lp=0.0>peasant</rdfs:label>
  <rdfs:label xml:lang="EN" xml:lp=0.0>raiser</rdfs:label>
  <rdfs:label xml:lang="EN" xml:lp=0.625>farmer</rdfs:label>
  <rdfs:subClassOf rdf:resource="#working population engaged in agriculture" xml:cpg=0.9841/>
</owl:Class>

<owl:Class rdf:ID="forest ranger">
  <rdfs:label xml:lang="EN" xml:lp=0.0>forest ranger</rdfs:label>
  <rdfs:subClassOf rdf:resource="#working population engaged in agriculture" xml:cpg=0.0/>
</owl:Class>

<owl:Class rdf:ID="agricultural adviser">
  <rdfs:label xml:lang="EN" xml:lp=0.0>agricultural adviser</rdfs:label>
  <rdfs:subClassOf rdf:resource="#working population engaged in agriculture" xml:cpg=0.0/>
</owl:Class>
...

```

Figure 4. Extract of an OWL file produced by TOOPRAG.

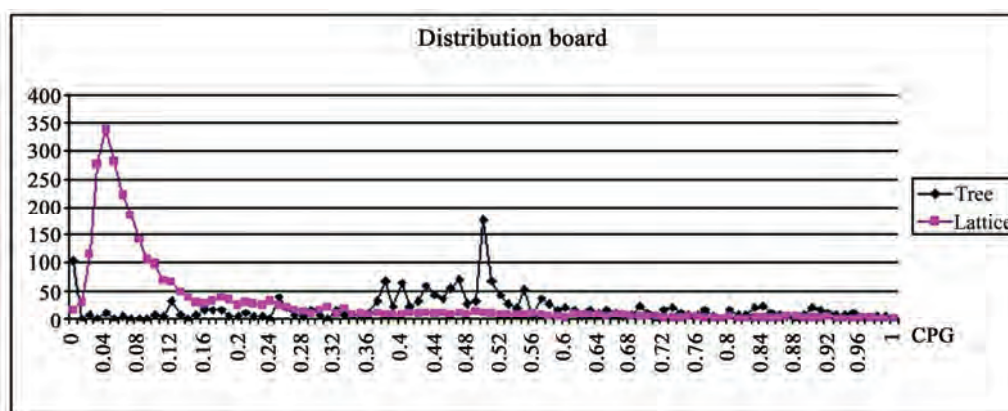


Figure 5. Influence of the number of edges (with a constant number of concepts).

we have developed a specific simulator whose parameters (given an ontology O) are: N the number of concepts of O , H the depth of O , and W the max width of O . From these parameters, the prototype automatically generates a random hierarchy of concepts. The results presented in Figure 5 computed in the following context: 1) a hierarchy O_1 based on a tree described by ($N=800$, $H=9$, $W=100$), 2) a hierarchy O_2 based on a lattice with a density of 0.5 described by ($N=800$, $H=9$, $W=100$); and (3) $\alpha=0.3$, $\beta=0.3$, $\gamma=0.3$, $\delta=1$. The results clearly attest the fact that multiple inheritance leads to a dilution of the typicality notion.

The results presented in Figure 5 have been calculated in the following context: 1) a hierarchy O_1 based on a tree described by ($N=800$, $H=9$, $W=100$); 2) a hierarchy O_2 based on a tree described by ($N=50$, $H=2$, $W=30$); and 3) $\alpha=0.3$, $\beta=0.3$, $\gamma=0.3$, $\delta=1$, of the distribution of

SPG values, proportionally to the size of the hierarchies, for a same density of graphs.

The results presented in Figure 6 have been calculated in the following context: 1) a hierarchy O based on a lattice with a density of 0.66 described by ($N=13000$, $H=7$, $W=240$), and 2) $\alpha=0.3$, $\beta=0.3$, $\gamma=0.3$, $\delta \in [0,1]$. These results clearly show the relevance of our emotional parameter: in a negative mental state, the distributional analysis focuses on strong values of SPG and in a positive mental state, the distributional analysis is more uniform.

5.4. Application of Our SPG: A Case Study in Agriculture

TOOPRAG has been used in a project dedicated to the analysis of texts describing the Common Agricultural

Policy (CAP) of the European Union. In this project, we have defined a specific ontology from the multilingual

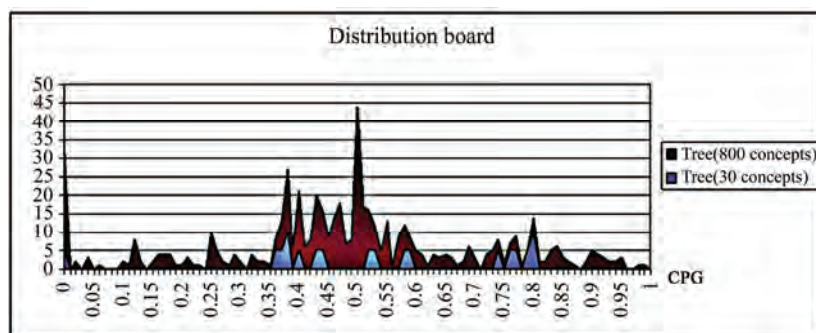


Figure 6. Influence of the number of concepts in a tree.

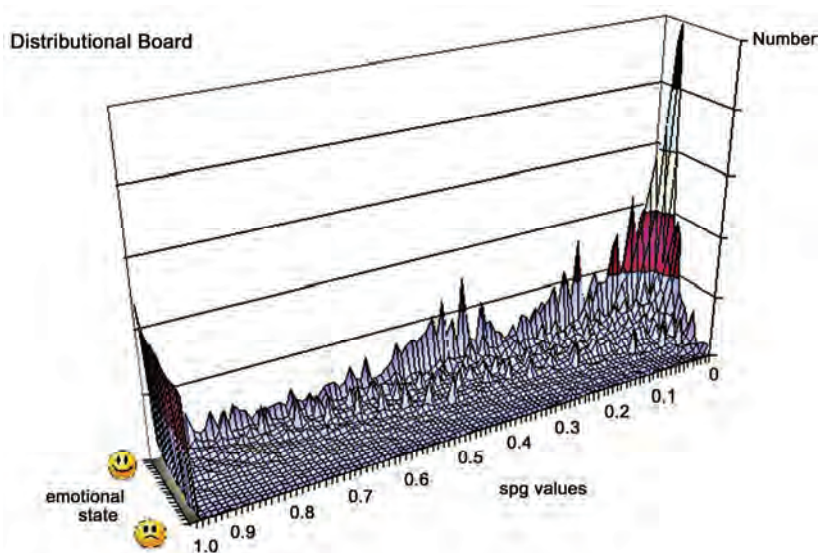


Figure 7. Emotional parameter influence.

thesaurus Eurovoc (<http://europa.eu/eurovoc/>). This thesaurus, which exists in 21 official languages of the European Union, covers multiple fields (e.g. politics, education and communications, science, environment, agriculture, forestry and fisheries, energy, etc.). It provides a means of indexing the documents in the documentation systems of the European institutions and of their users (e.g. the European Parliament, some national government departments and European organizations). From the Eurovoc field dedicated Agriculture, we have defined a first hierarchy of concepts by using the hyponymy/hyperonymy relationships (identified by the “Broader Term” links in Eurovoc) and the synonymy relationships (identified by the “Used For” links in Eurovoc). Then, this hierarchy has been modified and validated by an expert in Agriculture and Forestry. In its current version, this ontology includes a hierarchy of concepts based on a tree described by 283 concepts (depth=4 and max width=11). The lexicon of this ontol-

ogy is composed of 597 terms. In average, each concept is associated to 2,1 terms (min=1 and max=11). The corpus used for this experimentation is composed of 55 texts published in the Official Journal of the European Union (<http://eur-lex.europa.eu>) since 2005, and in particular 43 regulations, 1 directive, 8 decrees, 3 community opinions. It includes 1.360.000 words. From a statistical point of view, 61 concepts of the ontology are directly evoked in the corpus through the terms and 37 indirectly (via inheritance). Thus, the reverse ratio is 34,63%. Although the ontology considered in this project does not yet include properties and instances, the results provided by TOOPRAG (in the context of this specific corpus) are interesting because they help the expert to analyze the Common Agricultural Policy (CAP) of the European Union through regulatory texts. For instance, as shown by the Figure 8, the SPG values clearly underline that since 2005, the cultivation system which is particularly encouraged by the PAC is the organic farming.

In a similar way, the SPG values presented in Figure 4 state that the blue-collar workers of the agricultural sector (e.g. farm workers, farmers or peasants) are more supported by the PAC than the white-collar workers (e.g. agricultural advisers or head of agricultural holdings).

5.5. Application of our SPG: A Case Study in “Health, Safety and Environment”

Our approach is currently evaluated in the context of a project¹⁰ dedicated to Legal Intelligence within regulatory documents related to the domain “Health, Safety and Environment” (HSE). A first ontology of this domain has been defined¹¹. In its current version, it is composed of 3776 concepts structured in a lattice-based hierarchy (depth = 11 ; width = 1300). The calculation of our gradient has been applied on a specific corpus which includes 2052 texts related to the HSE domain from a reglementary point of view¹²: 782 administrative orders, 347 circulars, 288 parts of regulation, 178 decrees, 139 parts of legal rules, 102 european directives, 32 laws, 11 memorandum, 4 ordinances, etc. This process indicates that 1) 30.2% of the SPG values are non-null, 2) 3.34% of the SPG values are equal to 1, 3) 6.18% of the SPG values belong to [0.5, 1] and 4) 63.23% of SPG values belong to [0, 0.01]. The median value of the GPS is equal to 0.128.

5.6. Application of SEMIOSEM

SEMIOSEM is currently evaluated with the ontology of the HSE domain. In order to evaluate our measure and to compare it with related work, we have focused our study on the hierarchy presented in Figure 9: the goal is to compute the similarity between the concept *Carbon* and the sub-concepts of *Halogen*; for the expert of Tennaxia, these similarities are evaluated as follows: *Fluorine*=0.6; *Chlorine*=0.6; *Bromine*=0.3; *Iodine*=0.3; *Astatine*=0.1. We have also elaborated a specific corpus of texts composed of 1200 european regulatory documents related to the HSE domain (mainly laws, regulations, decrees and directives).

Table 2 presents the similarity values obtained with three intentional-based measures: Rada, Leacock and Wu. One can note that all the values are equal because these measures only depend on the structure of the hierarchy.

Table 3 depicts the similarity values obtained with three extensional-based measures: Lin, Jiang and Resnik. Table 4 presents the similarity values obtained with

¹⁰This ongoing research project is funded by the French company Tennaxia (<http://www.tennaxia.com>). This “IT Services and Software Engineering” company provides industry-leading software and implementation services dedicated to Legal Intelligence.

¹¹INPI June 13, 2008, Number 322.408 – SCAM-Velasquez September 16, 2008, Number 2008090075. All rights reserved.

¹²These texts have been extracted from the European Parliament and the French Parliament, mainly from LegiFrance (<http://www.legifrance.gouv.fr/>) and Eur-Lex (<http://eur-lex.europa.eu/>).

SEMIOSEM according to 6 contexts defined by the following parameters:

```
...
<owl:Class rdf:ID="organic farming">
  <rdfs:label xml:lang="EN" xml:lpq=1.0>organic farming</rdfs:label>
  <rdfs:subClassOf rdf:resource="#cultivation system" xml:cpq=0.7/>
</owl:Class>

<owl:Class rdf:ID="intensive farming">
  <rdfs:label xml:lang="EN" xml:lpq=0.0>intensive farming</rdfs:label>
  <rdfs:subClassOf rdf:resource="#cultivation system" xml:cpq=0.0/>
</owl:Class>

<owl:Class rdf:ID="single-crop farming">
  <rdfs:label xml:lang="EN" xml:lpq=0.0>single-crop farming</rdfs:label>
  <rdfs:subClassOf rdf:resource="#cultivation system" xml:cpq=0.0/>
</owl:Class>

<owl:Class rdf:ID="extensive farming">
  <rdfs:label xml:lang="EN" xml:lpq=1.0>extensive farming</rdfs:label>
  <rdfs:subClassOf rdf:resource="#cultivation system" xml:cpq=0.0333/>
</owl:Class>

<owl:Class rdf:ID="dry farming">
  <rdfs:label xml:lang="EN" xml:lpq=0.0>dry farming</rdfs:label>
  <rdfs:subClassOf rdf:resource="#cultivation system" xml:cpq=0.0/>
</owl:Class>

<owl:Class rdf:ID="crop rotation">
  <rdfs:label xml:lang="EN" xml:lpq=1.0>crop rotation</rdfs:label>
  <rdfs:subClassOf rdf:resource="#cultivation system" xml:cpq=0.2667/>
</owl:Class>
...
```

Figure 8. Extract of an OWL file produced by TOOPRAG.

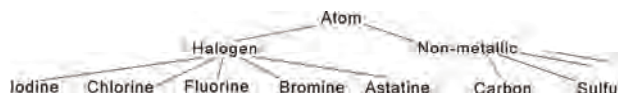


Figure 9. Extract of the hierarchy of concepts of the HSE ontology.

- A ($\alpha=0.7, \beta=0.2, \gamma=0.1, \delta=1$);
- B ($\alpha=0.2, \beta=0.7, \gamma=0.1, \delta=1$);
- C ($\alpha=0.2, \beta=0.1, \gamma=0.7, \delta=1$);
- D ($\alpha=0.33, \beta=0.33, \gamma=0.33, \delta=0.1$);
- E ($\alpha=0.7, \beta=0.2, \gamma=0.1, \delta=0.1$);
- F ($\alpha=0.7, \beta=0.2, \gamma=0.1, \delta=5.0$);

These experimental results lead to the following remarks:

- in all the contexts, SEMIOSEM provides the same order of similarities as the other measures. In a context which gives priority to the intentional component (cf. Context A), SEMIOSEM is better than the other measures. In the context B which gives priority to the extensional component (resp. the context C which gives priority to the expressional component), SEMIOSEM is close to Jiang’s measure (resp. Lin’s measure). In a context

Table 2. Similarity with Carbon (Rada, Leacock, Wu).

Halogen	Rada	Leacock	Wu
Fluorine	0.25	0.097	0.6
Chlorine	0.25	0.097	0.6
Bromine	0.25	0.097	0.6

Iodine	0.25	0.097	0.6
Astatine	0.25	0.097	0.6

Table 3. Similarity with carbon (Lin, Jiang, Resnik).

Halogen	Lin	Jiang	Resnik
Fluorine	0.31	0.14	1.43
Chlorine	0.28	0.12	1.43
Bromine	0.23	0.09	1.43
Iodine	0.22	0.09	1.43
Astatine	0	0	1.43

Table 4. Similarity with carbon (SemioSem).

Halogen	A	B	C	D	E	F
Fluorine	0.40	0.14	0.32	0.27	0.91	0.025
Chlorine	0.36	0.12	0.29	0.25	0.90	0.017
Bromine	0.29	0.10	0.23	0.20	0.88	0.007
Iodine	0.28	0.10	0.23	0.19	0.88	0.006
Astatine	0.01	2.10 ⁻⁴	2.10 ⁻⁴	3.10 ⁻⁴	0.63	1.10 ⁻⁸

which gives no priority to a specific component (cf. Context D), SEMIOSEM is between Lin's measure and Jiang's measure;

- context E and F clearly show the influence of the emotional factor: a positive mental state (cf. context E) clearly increases the similarities values and a negative mental state (cf. context F) clearly decreases similarities values;
- the concept Astatine is not evoked in the corpus, nor represented by instances. Thus, it is not considered as similar by Lin's and Jiang's measures. SEMIOSEM finds a similarity value thanks to the intentional component.

6. Discussion and Future Work

6.1. Lexical and Conceptual Prototypicality Gradients

The purpose of our work, which is focused on the notion of "Personalized Vernacular Domain Ontology", is to deal with subjectivity knowledge via 1) its specificity to an endogroup and a domain, 2) its ecological aspect and 3) the prominence of its emotional context. This objective leads us to study the pragmatic dimension of an ontology. Inspired by works in Cognitive Psychology, we have defined two measures identifying two complementary gradients called SPG and LPG which are respectively dedicated to 1) the conceptual prototypicality which evaluates the representativeness of a concept within a decomposition and 2) the lexical prototypicality which evaluates the representativeness of a term within a set of terms used to denote a concept. It is important to underline that these gradients do not modify the formal semantics of the ontology which is considered; the subsumption links remain valid. These gradients only reflect the pragmatics of an ontology for knowledge (re)-using. They can be an effective help in different activities, such as:

- *Information Retrieval*. Our conceptual and lexical prototypicality gradients can be used to classify the results of a query, and more particularly an extended query, according to a relevance criteria which consists in considering the most representative element of a given concept (resp. a given term) as being the most relevant result of a query expressed by a (set of) term(s) denoting this concept (resp. corresponding to this term). This approach permits a classification of the extended results from a qualitative point of view. Moreover, our approach also allows us to proportion the number of results according to the value of the gradients (i.e. a quantitative point of view). Thus, information retrieval becomes customizable, because it is possible to adapt the results to the pragmatics of the ontology, i.e. privileging the intentional dimension (and not the extensional and expressional one) or conversely, working with different mental states, etc. In this way, *Ontology Personalization* is used as a means for Web Personalization and Semantic Web Personalization. We currently evaluate this approach in the context of the SweetWiki Semantic Web platform [24].

- *Ontological Analysis* of text Corpora. As introduced in Section 5-D, our Semiotic-based Prototypicality Gradient can be used to evaluate the ontological contents of text corpora: which are the main concepts involved in a text corpus? By using the same ontology applied on different corpora related to the same domain, it is possible to compare, at the conceptual level, the Information Content of these corpora. We currently evaluate this approach in the context of an experimentation which aims at making a comparative analysis of the health-care preoccupations of different populations, in particular French, English and American population. For this purpose, we currently define a multilingual ontology from the MeSH¹³. In order to really deal with the preoccupations of people, we have selected the three most popular and complete medical websites where french, english and american people can find and exchange all the information they are looking about their health care needs: Doctissimo in France¹⁴, Health-care Republic in United Kingdom¹⁵ and HealthCare.com in USA¹⁶. These websites will be considered as three distinctive corpora from which our gradients will be calculated, by using the same multilingual ontology.

Our work is currently in progress towards the improvement of the gradients according to many works related to Cognitive and Social Psychology. We also study how to enrich the intentional component by taking the axiomatic part of an ontology into account. From an application point of view, we currently evaluate our approach in the context of a project dedicated to Legal Intelligence within regulatory documents related to the areas "Hygiene, Safety and Environment".

¹³MeSH is the U.S. National Library of Medicine's controlled vocabulary (<http://www.nlm.nih.gov/mesh/meshhome.html>). MeSH terminology provides a consistent way to retrieve information that may use different terminology for the same concepts.

¹⁴<http://www.doctissimo.fr>

¹⁵<http://www.healthcarepublic.com>¹⁶<http://www.healthcare.com>

6.2. SemioSem

SEMIOSEM is particularly relevant in a context where the perception (by the end-user) of the domain which is considered (and which is both conceptualized within an ontology and expressed by a corpus and instances) can have a large influence on the evaluation of the similarities between concepts (e.g. ontology-based information retrieval). We advocate that such an user-sensitive context, which de facto includes subjective knowledge (whereas ontologies only includes objective knowledge), must be integrated in a similarity measure since ontologies co-evolve with their communities of use and human interpretation of context in the use. Formally, SEMIOSEM respects the properties of similarity measure reminded in [17]: *positiveness*¹⁷, *reflexivity*¹⁸ and *symmetry*¹⁹. But, SEMIOSEM is not a semantic distance since it does not check simultaneously the strictness property²⁰ and the triangular inequality²¹. For the extensional component, our first choice was the Amato measure. But one of our goal is to be independent from the modeling structure and this measure clearly depends on the Most Specific Common Subsumer (MSCS). Moreover, in the case of our experiment, it does not really provide more relevant results. This is why we have selected the Jaccard measure for its simplicity and its independence from the MSCS but we currently study the use of the Dice measure. For the expressional component, the Latent Semantic Analysis could be adopted but since it is based on the tf-idf approach, it is not really appropriated to our approach: we want to keep the granularity of the corpus in order to give more importance to concepts which perhaps appears less frequently in each document, but in an uniform way in the whole corpus (than concepts which are frequently associated in few documents). Then, as we simply compare the terms used to denote the concepts in the corpus, our approach is clearly limited since it can not deal with expressions such as t_1 and t_2 are opposite. To deal with this problem, we plan to study more sophisticated computational linguistic methods. Finally, for the intentional component, our approach can be timeconsuming (when the end-user decides to weight the properties of the concepts²²), but, as far as we know, it is really innovative and it provides promising results. The parameters alpha, beta, gamma and delta are used to adapt the measure to the context which is related to the end-user perception of the domain according to the intentional, extensional, expressional and emotional dimension. We consider that this is really a new approach

¹⁷ $\forall x, y \in C : \text{SEMIOSEM}(x, y) \geq 0$ ¹⁸ $\forall x, y \in C : \text{SEMIOSEM}(x, y) \leq \text{SEMIOSEM}(x, y)$ ¹⁹ $\forall x, y \in C : \text{SEMIOSEM}(x, y) = \text{SEMIOSEM}(y, x)$ ²⁰ $\forall x, y \in C : \text{SEMIOSEM}(x, y) = 0 \Rightarrow x = y$ ²¹ $\forall x, y, z \in C : \text{SEMIOSEM}(x, y) + \text{SEMIOSEM}(y, z) \geq \text{SEMIOSEM}(x, z)$

and this is why we call our measure SEMIOSEM (Semi-otic-based Similarity Measure). Moreover, the aggregation we advocate does not just correspond to a sum: these parameters are used both to adapt the influence of each dimension and/or to adapt the calculus according to the resources which are available. Thus, when no corpus is available, the expressional component can not be used ($\gamma=0$). A similar approach is adopted for the intentional component (an ontology without properties leads to $\alpha = 0$) and the extensional component (no instances leads to $\beta = 0$). The value of delta (emotional state) can be defined according to a questionnaire or the analysis of data given by physical sensors such as the speed of the mouse, the webcam-based facial recognition, etc. To sum up, SEMIOSEM is more flexible (since it can deal with multiple information sources), more robust (since it performs relevant results under unusual conditions as shown by the case Astatine of the experimental results) and more user-centered (since it is based on the domain perception and the emotional state of the end-user) than all the current methods.

7. References

- [1] S. Harnad, "Categorical perception," Encyclopedia of Cognitive Science, Vol. LXVII, No. 4, 2003. [Online]. Available: <http://cogprints.org/3017/>.
- [2] T. Gruber, "Toward principles for the design of ontologies used for knowledge sharing," in Formal Ontology in Conceptual Analysis and Knowledge Representation, N. Guarino and R. Poli, Eds. Dordrecht, The Netherlands: Kluwer Academic Publishers, 1993.
- [3] C. K. Ogden and L. Richards, "The meaning of meaning: A study of the influence of language upon thought and of the science of symbolism," Harcourt, ISBN-13: 978-0156584463, 1989.
- [4] D. L. M. Gabora, D. E. Rosch, and D. D. Aerts, "Toward an ecological theory of concepts," Ecological Psychology, Vol. 20, No. 1-2, pp. 84-116, 2008. [Online]. Available: <http://cogprints.org/5957/>.
- [5] E. Rosch, "Cognitive reference points," Cognitive Psychology, No. 7, pp. 532-547, 1975.
- [6] M. McEvoy and D. Nelson, "Category norms and instance norms for 106 categories of various sizes," American Journal of Psychology, Vol. 95, pp. 462-472, 1982.
- [7] C. Morris, "Foundations of the theory of signs," Chicago University Press, 1938.
- [8] S. Bluck and K. Li, "Predicting memory completeness and accuracy: Emotion and exposure in repeated autobiographical recall," Applied Cognitive Psychology, No. 15, pp. 145-158, 2001.
- [9] J. Park and M. Nanaji, "Mood and heuristics: The influence of happy and sad states on sensitivity and bias in stereotyping," Journal of Personality and Social

²²Again, by default, all the weightings are equal to 1 and the function Intens remains valid. In the case of our experiment, the results obtained in this context for the concept Fluorine are: A - 0.59; B - 0.19; C - 0.38 D - 0.37; E - 0.95; F - 0.12.

Psychology, No. 78, pp. 1005–1023, 2000.

- [10] M. Mikulincer, P. Kedem, and D. Paz, “Anxiety and categorization-1, the structure and boundaries of mental categories,” *Personality and Individual Differences*, Vol. 11, No. 11, pp. 805–814, 1990.
- [11] C. M. Au Yeung and H. F. Leung, “Formalizing typicality of objects and context-sensitivity in ontologies,” in *AAMAS '06: Proceedings of the fifth international joint conference on Autonomous Agents and Multiagent Systems*. New York, NY, USA: ACM, ISBN 1-59593-303-4, pp. 946–948, 2006.
- [12] C. M. Au Yeung and H. F. Leung, “Ontology with likeliness and typicality of objects in concepts,” in *Proceedings of the 25th International Conference on Conceptual Modeling–ER 2006*, S. B. Heidelberg, Ed., Vol. 4215/2006, ISSN 0302-9743 (Print), 2006.
- [13] P. Resnik, “Using information content to evaluate semantic similarity in a taxonomy,” in *14th International Joint Conference on Artificial Intelligence (IJCAI 95)*, Montrai, Vol. 1, pp. 448–453, August 1995.
- [14] A. Tversky and D. Kahneman, “Judgment under uncertainty: Heuristics and biases,” *Science*, No. 185, pp. 1124–1131, 1974.
- [15] D. Lin, “An information-theoretic definition of similarity,” in *Proceedings of the 15th International Conference on Machine Learning*, pp. 296–304, 1998.
- [16] J. Jiang and D. Conrath, “Semantic similarity based on corpus statistics and lexical taxonomy,” in *International Conference on Research in Computational Linguistics*, pp. 19–33, 1997.
- [17] C. d’Amato, S. Staab, and N. Fanizzi, “On the influence of description logics ontologies on conceptual similarity,” in *EKAU 2008, International Conference on Knowledge Engineering and Knowledge Management Knowledge Patterns*, pp. 48–63, October 2008.
- [18] R. Rada, H. Mili, E. Bicknell, and M. Blettner, “Development and application of a metric on semantic nets,” *IEEE Transactions on Systems, Man and Cybernetics*, Vol. 19, No. 1, pp. 17–30, 1989.
- [19] C. Leacock and M. Chodorow, “WordNet: An electronic lexical database,” Cambridge, MA, The MIT Press, 1998, ch. Combining local context and wordnet similarity for word sense identification, pp. 265–283.
- [20] Z. Wu and M. Palmer, “Verb semantics and lexical selection,” in *Proceedings of the 32nd Annual Meeting of the Association for Computational Linguistics*, pp. 133–138, 1994.
- [21] E. Blanchard, M. Harzallah, and P. Kuntz, “A generic framework for comparing semantic similarities on a subsumption hierarchy,” in *Proceedings of the 18th European Conference on Artificial Intelligence (ECAI’2008)*. IOS Press, pp. 20–24, 2008.
- [22] P. Jaccard, “Distribution de la flore alpine dans le bassin des dranses et dans quelques rgions voisines,” *Bulletin de la Socit Vaudoise de Sciences Naturelles*, Vol. 37, pp. 241–272, 1901, (in French).
- [23] M. Sanderson and W. Croft, “Deriving concept hierarchies from text,” in *Proceedings of the 22nd International ACM SIGIR Conference*, pp. 206–213, 1999.
- [24] M. Buffa, F. Gandon, G. Ereteo, P. Sander, and C. Faron, “Sweetwiki: A semantic wiki,” *Special Issue of the Journal of Web Semantics on Semantic Web and Web 2.0*, Vol. 6, pp. 84–97, February 2008.

A Nonlinear Control Model of Growth, Risk and Structural Change

P. E. Petrakis, S. Kotsios

National and Kapodistrian University of Athens, Athens, Greece

Email: ppetrak@cc.uoa.gr, skotsios@di.uoa.gr

Abstract

Uncertainty is perceived as the means of removing the obstacles to growth through the activation of Knightian entrepreneurship. A dynamic stochastic model of continuous-time growth is proposed and empirically tested, including equilibrating and creative entrepreneurial activity. We find that uncertainty affects economic growth and the rate of return, and causes structural changes in portfolio shares for the two types of entrepreneurial events. Structural change depends mainly on the intertemporal rate of substitution, productivity ratios, and finally intersectoral difference in return and risk.

Keywords: Growth, Risk, Entrepreneurship, Structural Change

1. Introduction

This paper examines the relationship between growth and risk through structural change. Structural change is analyzed through the examination of growth, since it relates to entrepreneurship and uncertainty. Uncertainty is treated as the means of removing barriers to growth through the activation of Knightian entrepreneurship.

We assume that growth is the result of equilibrating and creative entrepreneurial events. Equilibrating entrepreneurial events (adaptive behaviour) are the most common ones and bring demand and supply to an equilibrium [1]. On the other hand, creative entrepreneurial events (innovative Schumpeterian behaviour) are those that result from the creation of new (innovative) products and services.

Considering structural change issues prompts discussion of 'why industries grow at different rates and which industries come to have an increasing weight in the total output while others decline and eventually wane' [2]. In search for an answer, researchers usually bring up differences in income elasticities of domestic demand, supply-side productivity differences, or different productivity growth rates, which is the result of selection mechanisms within the general evolutionary process [2–10]. Changes in the relationship between intersectoral conditions of risk and return are believed to have implications for structural changes and economic growth. Thus, risk and uncertainty are examined in the framework of structural change and economic growth. Intuitively one can expect to find a causal relationship between uncertainty and

structural change. However, this has not yet been verified using a growth structural change model.

A dynamic stochastic model of continuous-time growth is proposed, including two basic types of entrepreneurial events, based on the work of Turnovsky [11]. It includes three distinct 'crucial' individual concepts: growth rates, portfolio shares, and rates of return. Thus, our analysis includes the performance indexes (growth rates), 'incentives' (rates of returns) and 'results' (portfolio shares) of entrepreneurial behaviour. This paper therefore contributes to the analysis of uncertainty, entrepreneurship, and risk. It also contributes to the analysis of structural change patterns with regards to risk.

The rest of the paper is organized as follows. In Section 2 of the paper the sustainable growth conditions are discussed. Section 3 provides a short literature review on structural change. In Section 4 we introduce the model and its implications. Section 6 concludes.

2. Sustainable Growth and Uncertainty

The purpose of this study is to establish a theoretically acceptable relationship between growth and uncertainty taking into consideration the role of entrepreneurship. The argument runs as follows: uncertainty activates entrepreneurship, which stimulates social capital and influences growth. Thus, the focus is on the relationship between growth, entrepreneurship, and uncertainty. Social capital is the common ground where entrepreneurship and uncertainty operate. Do the above dynamics work towards the elimination of obstacles to sustainable growth?

Lucas [12] rearranges the neoclassical model and establishes that our attention should be on human capital and its externalities on labour (L). Romer [13,14] augments this theory, arguing that additional investment in research could result in increasing returns through knowledge spillover embodied in human capital. What is important here is alertness ([15], *i.e.* the 'knowledge' of where to find market data [16]). The process of finding entrepreneurial opportunities relates to the stock of knowledge (social capital) inherent in everyday life experience [17].

Yu [17], utilising a) Kirzner's [14] theory of entrepreneurial discovery, b) Schumpeter's [18] two types of economic responses (extraordinary and adaptive), and c) the Austrian theory of institutions as building blocks, constructs an entrepreneurial theory of institutional change and social capital accumulation. Yu [17,19], as well as other researchers in the field, do not use the concept of social capital as an alternative for institutions. However, social capital as defined by Westlund and Bolton [20] is greatly comparable to the concept of institutions, as described in Yu [17,19]. The process of institutional change is the continuous interaction between entrepreneurial exploitation and exploitation of opportunities [21]. Institutions (stores of knowledge) emerge as a consequence of the attempt to reduce structural (as opposed to neo-classical static) uncertainty. Therefore, entrepreneurship expands institutional development and social capital accumulation. Evidently in this process there are second-round effects. Social capital accumulation boosts entrepreneurship through externalities. These externalities promote the distribution of information and generate asymmetric information. At the same time, institutions reinforce entrepreneurial alertness and the process of discovering new entrepreneurial opportunities [19]. Thus, entrepreneurship increases social capital. Thus entrepreneurship affects growth positively through social capital accumulation.

According to Brouwer [22], in Knight's [23] view, true uncertainty is the only source of profits because they vanish as soon as change becomes predictable, or they become costs if uncertainty is hedged. Brouwer [22] shows that diminishing returns to investment in innovation can be avoided with the use of Knightian uncertainty. This can be achieved through R&D cooperation; that is, by creating social capital through R&D networks. We can therefore suggest that uncertainty makes perpetual innovation more likely. Thus, growth and uncertainty are positively related. Knight [23] supports that rates of return on entrepreneurial investment vary around an average, and it is the relative entrepreneurial ability that is rewarded.

Entrepreneurs also create a great deal of uncertainty through Schumpeterian innovation, which creates confusion in the market. A lack of entrepreneurship indicates an over-reliance on old structures, interpretations, and understandings [17]. Thus, entrepreneurial activation is positively related with uncertainty.

From the above analysis, we can conclude that entrepreneurship and uncertainty are related, with the latter positively affecting entrepreneurship. Therefore, growth and uncertainty are related, with the latter positively affecting growth.

Montobbio [2], in his critical and concentrated literature review, presents three main trends stemming from studies on structural change:

a) Endogenous growth models assess the determinants of aggregate growth in a multi-sectoral economy [14,24] but they incur difficulties in explaining major processes of structural change.

b) Industry life-cycle models examine growth, maturity and decline [25] but do not address 1) demand pressures and 2) the relationship between growth and sector that decline.

c) The supply and demand side factors approach seems to attract most of the recent work done in the field.

The supply side was first proposed by Schumpeter [18]. Kuznets [6] also stressed the importance of different impact of 1) technological innovations and 2) a selection mechanism based on competitive advantage. Pasinetti [9,10] demonstrates that growth rates depend on productivity rates. Montobbio [2] shows aggregate productivity growth can be achieved without technological change at the firm level.

This paper follows the supply and demand side factors approach to structural change. The model includes basic characteristics of the supply side, especially the influence of uncertainty, productivity, and social capital and networks on the portfolio shares. In addition, it takes into account the intertemporal elasticity of substitution. It evidently shows how growth and structural change can be achieved simultaneously, without the use of further assumptions.

3. The Model of Growth, Creative Equilibrating Events and the Role of Uncertainty

The preceding analysis sets the basis for the introduction of a representative agent model based on three fundamental concepts: economic growth, and the two types of entrepreneurial events, including their basic characteristics. Obviously, this points to a stochastic growth model, which will include stochastic capital accumulation, capital return specification, and consumer utility maximisation procedures. The model is based on Turnovsky's [11] stochastic growth model with an entrepreneurial event.

We consider an economy, where the household and production sectors are consolidated. The representative agent consumes output over the period $(t, t + dt)$ at a non-stochastic rate Cdt .

The agent distributes his resources between the two types of entrepreneurial events. This means that he func-

tions within an environment of perfect information, with no costs or limitations regarding the initiation of a creative or equilibrating event.

The two types of entrepreneurial events influence growth rates in different ways. In particular:

a) The role of creative vs. equilibrating events, with regards to the accumulated flow of output over the period $(t, t + dt)$, is rather different (see description of Equation (3) below).

b) The two types of entrepreneurial activities face different technologies. Each activity adds to the total production flow in the same way.

Each agent maximizes expected lifetime utility captured by a standard concave utility function:

$$E_0 \int_0^{\infty} U(C) e^{-bt} dt \text{ with } U'(C) > 0 \text{ and } U''(C) < 0 \quad (1)$$

subject to the stochastic accumulation equation,

$$dK^c + dK^e = dY - Cdt \quad (2)$$

where:

K^c = stock of physical capital devoted to creative entrepreneurial events at time t ;

K^e = stock of physical capital devoted to equilibrating events at time t ;

dY = flow of output (from both entrepreneurial events) over the period $(t, t + dt)$.

The initial stocks of capital are given by K_0^c and K_0^e .

4. The Mechanics of Growth, Risk and Structural Change

This section outlines the model. Subsection 4.1 gives the model's assumptions; 4.2 is concerned with the stochastic process; 4.3 analyzes the determinants of risk; 4.4 summarizes the findings concerned with the growth process; last, 4.5 examines the portfolio shares and the rates of return. Finally, Table 1 presents the findings of our theoretical analysis.

4.1. The Model's Assumptions

Besides the basic assumptions, which are:

a) Linearity in production equations;
b) Individuals are risk averse, which implies that $\gamma - 1 < 0$, *i.e.* relatively large elasticity of intertemporal substitution.

The model adopts the following two additional hypotheses:

c) $r_c > r_e$; *i.e.* the rate of return on creative entrepreneurial events is larger than that of equilibrating events;

d) $\sigma_e^2 > \sigma_{ce}$, and $\sigma_c^2 > \sigma_{ce}$; the risk of equilibrating and creative entrepreneurial events is greater than the covariance of risk between the two types of events. This hypothesis is based on the fundamental principle of portfolio structuring according to which the risk of each portfolio component is greater than the total portfolio risk. In other words, since the agent is risk averse, it always makes sense to reduce risk by composing portfolios that include both types of entrepreneurial events.

e) We assume that the production functions are linear in their components; that is, that $Y_e = \theta_e K^e$ and $Y_c = \theta_c K^c$, where θ_e, θ_c the TFP variables in the two production technologies respectively. There is no a priori reason to assume that $\theta_e \neq \theta_c$; thus the model does not necessarily assume heterogeneity of the sectors (firms) in terms of productivity. However, for simplicity reasons the two productivities are different in notation. The consequences of identical intersectoral productivity values will be examined later.

Output is assumed to be generated from capital through the following stochastic process:

$$dY(t) = F(K^e)dt + F(K^c)dt + H(K^e)dy^e + H(K^c)dy^c \quad (3)$$

Equation (3) has the following interpretation: the change in total output depends on deterministic events and stochastic episodes stemming from equilibrating and creative entrepreneurial activity. $F(K^c) + F(K^e)$ represent the deterministic effects of creative and equilibrating events respectively, whereas $H(K^e)dy^e + H(K^c)dy^c$ represent the stochastic events in each category of entrepreneurship.

The stochastic terms in the production function are based on two crucial assumptions: 1) $H(K^i)$, $i = e, c$ is constant and the shocks enter the production function additively, 2) or $H(K^i) = hK^i$, $i = e, c$ where the disturbances are proportional to the aggregate capital stock, thus entering the production function in a multiplicative way. The assumptions regarding the stochastic disturbance terms are crucial for obtaining tractable, closed-form solutions to the optimisation problem.

Stochastic disturbances of creative events are additive to the model. Equilibrating events, however, enter the model in a multiplicative function. Thus, according to this specification, equilibrating events are assumed to depend on the existing level of capital, whereas creative events are independent of the stock of capital.

Total capital stock held by the representative agent is

denoted in the following way:

$$K \equiv K^e + K^c \equiv W \quad (4)$$

with the corresponding portfolio shares being

$$n_e \equiv \frac{K^e}{W}, n_c = 1 - n_e \equiv \frac{K^c}{W} \quad (5)$$

From relationship (5) we see that

$$K^e = n_e K, \quad K^c = n_c K, \quad n_e + n_c = 1,$$

where n_e, n_c are the corresponding portfolio shares.

4.2. The Stochastic Processes

Capital (K) follows a continuous time stochastic process:

$$dK = \psi K dt + K dk, \quad \psi \in \mathfrak{R} \quad (6)$$

where $\psi = \frac{F(K)}{K} - \frac{C}{K}$ is the growth rate of capital and

dk denotes a stochastic component. The properties of this stochastic component are: $E(dk) = 0$ and

$$\text{Var}(dk) = \sigma_K^2 dt, \quad \sigma_K^2 = \frac{H^2}{K^2} \sigma_y^2$$

• The rate of return on capital occupied in equilibrating entrepreneurship is determined through the following stochastic process:

$$dR_e = r_e dt + du_e \quad (7)$$

The variance of the stochastic part of (7) is given by $\text{Var}(du_e) = \sigma_e^2 dt$

Stochastic real rate of return on capital occupied in creative events is described through a similar stochastic process:

$$dR_c = r_c dt + du_c \quad (8)$$

With the variance of the stochastic part being equal to $\text{Var}(du_c) = \sigma_c^2 dt$.

The deterministic parts $r_e dt$ and $r_c dt$ denote the rates of return on the two types of capital. The stochastic parts are normally distributed with $E(du) = 0$ and $\text{Var}(du) = \sigma_u^2 dt$. They represent the risks that the agent undertakes when he employs capital on equilibrating and creative entrepreneurship.

4.3. The Determinants of Risk

Theorem 4.3.1: *The following relations hold:*

$$\sigma_e^2 = \theta_e^2 \sigma_K^2, \quad \sigma_c^2 = \theta_c^2 \sigma_K^2 \quad (9)$$

(Proof available upon request)

Theorem 4.3.1 suggests that the two levels of risk associated with entrepreneurial equilibrating and creative events are directly dependent upon the corresponding productivity ratio. They are also directly related to the stochastic part of capital accumulation, described by the variance. This is consistent with economic intuition since it implies that economies with small capital accumulation variance (*i.e.* low density business cycles) are characterized by low levels of entrepreneurial risk. The positive relationship between entrepreneurial risk and productivity is also an anticipated outcome since according to theorem 4.3.1 entrepreneurial risk, and thus the rate of return, has a positive relationship with productivity.

4.4. The Growth Process

The functional form portraying the relationship between individual entrepreneurial risk and the growth rate of capital, establishes the mechanism and the necessary conditions through which the two types of risk influence the rate of growth in our model. In order to address these two issues we introduce the following optimisation problem:

$$\max_C E_0 \int_0^{+\infty} \frac{1}{\gamma} C^\gamma e^{-bt} dt \quad (10)$$

s.t.

$$dK = \psi K dt + \rho K dz \quad (11)$$

In other words, we want to maximise the expected value of the following utility function

$$U(C) = \frac{1}{\gamma} C^\gamma e^{-bt}, \quad -\infty < \gamma < 1 \quad (12)$$

where C is the agent's consumption and b is the discount rate.

The following result plays an essential role in the study of the comparative statics.

Theorem 4.4.1: *In equilibrium the following hold:*

$$\begin{aligned} \frac{\partial \psi}{\partial \sigma_e^2} &= -\frac{1}{2} \left(\frac{1}{\theta_e} \right)^2 (\gamma - 1), \\ \frac{\partial \psi}{\partial \sigma_c^2} &= -\frac{1}{2} \left(\frac{1}{\theta_c} \right)^2 (\gamma - 1) \end{aligned} \quad (13)$$

(Proof available upon request)

The effect of entrepreneurial risk on growth depends on the intertemporal elasticity of substitution. When large (*i.e.* $\gamma < 1$) the effect of risk on growth is positive. When $\gamma > 1$ then the effect is negative. This effect also depends on the reciprocal of productivity ratio, where the larger

the productivity the smaller its influence on growth. This fact is demonstrated in the following theorem:

Theorem 4.4.2: *The following inequalities are valid:*

$$\frac{\partial \psi}{\partial \sigma_e^2} > 0, \quad \frac{\partial \psi}{\partial \sigma_c^2} > 0 \quad (14)$$

Moreover, in both cases, entrepreneurial risk has a direct relationship with the aggregate capital stock and an inverse relationship with productivity. The greater the quantity of aggregate capital, the larger the influence of risk on growth. On the other hand, the larger the productivity in the economy, the weaker the influence of risk on growth. In other words the greater the average ratio in the economy, the greater is the risk influence, as the financial theory suggests. On the other hand, the larger the productivity in the economy is, the smaller the effect of risk on growth. Different sector productivity levels imply differences on their impact on growth. Thus, if $\theta_e > \theta_c$, that is the productivity ratio in the equilibrating sector is greater than the productivity ratio in the creative sector, then if risk increases (decreases) by the same amount in both sectors, the impact on growth will be smaller (larger) for the equilibrating events vs. the creative events.

4.5. Portfolio Shares and Rates of Return of Entrepreneurial Events

By using the stochastic accumulation equation we have successively shown:

$$dK = dY - Cdt = dY_e + dY_c - Cdt \quad (15)$$

Dividing by K we get

$$\begin{aligned} \frac{dK}{K} &= \frac{dY_e}{K} + \frac{dY_c}{K} - \frac{C}{K}dt = \\ &= n_e \frac{dY_e}{K^e} + n_c \frac{dY_c}{K^c} - \frac{C}{K}dt = n_e(dR_e) + n_c(dR_c) - \frac{C}{K}dt = \\ &= r_e n_e dt + n_e du_e + r_c n_c dt + n_c du_c - \frac{C}{K}dt = \\ &= \left(r_e n_e + r_c n_c - \frac{C}{K} \right) dt + n_e du_e + n_c du_c = \left(r_e n_e + r_c n_c - \frac{C}{K} \right) dt + dk \end{aligned} \quad (16)$$

with, $n_e + n_c = 1$. We can easily prove that the variance of this stochastic process satisfies the relation

$$\sigma_K^2 = n_e^2 \sigma_e^2 + n_c^2 \sigma_c^2 + 2n_e n_c \sigma_{ce} \quad (17)$$

where σ_{ce} is the covariance of dR_e and dR_c .

This is a standard portfolio construction statement. As we have already seen above, the variances of equilibrating and creative events are connected with the covariance of the σ_{ce} with the basic relation of $\sigma_e^2 > \sigma_{ce}$

and $\sigma_c^2 > \sigma_{ce}$

The original optimisation problem is transformed to

$$\max_{C, n_c, n_e} E_0 \int_0^{+\infty} \frac{1}{\gamma} C^\gamma e^{-bt} dt, \quad -\infty < \gamma < 1 \quad (18)$$

s.t.

$$\frac{dK}{K} = \left(r_e n_e + r_c n_c - \frac{C}{K} \right) dt + dk \quad (19)$$

with

$$n_e + n_c = 1 \quad (20)$$

$$\sigma_K^2 = n_e^2 \sigma_e^2 + n_c^2 \sigma_c^2 + 2n_e n_c \sigma_{ce} \quad (21)$$

Theorem 4.5.1: *The first-order conditions for the optimisation problem can be written as follows*

$$(\delta\gamma)^{1/(\gamma-1)} = \frac{b - (r_e \hat{n}_e + r_c \hat{n}_c) \gamma - \frac{1}{2} \gamma (\gamma-1) \sigma_K^2}{1-\gamma} \quad (22)$$

$$\delta\gamma = \left(\frac{\hat{C}}{K} \right)^{\gamma-1}$$

$$(\gamma-1)(\hat{n}_e \sigma_e^2 + \hat{n}_c \sigma_{ce}) + r_c = \frac{\mu}{\delta\gamma K^\gamma}$$

$$(\gamma-1)(\hat{n}_e \sigma_e^2 + \hat{n}_c \sigma_{ce}) + r_e = \frac{\mu}{\delta\gamma K^\gamma}$$

$$\hat{n}_c + \hat{n}_e = 1$$

where $\hat{n}_e, \hat{n}_c, \hat{C}$ are the maximum achieved values of n_e, n_c, C , and μ is the Lagrange multiplier, related to the equation $n_e + n_c = 1$.

(Proof available upon request)

These above equations describe the solution of our optimisation problem. For the sake of the appearances, in the next paragraphs we will omit the hat symbol. The reader should keep in mind that we refer to the optimal values. Let us now assume that the quantities n_e, n_c, r_e, r_c are functions of σ_e^2 . Differentiating the above equations, (proof on request), with respect to σ_e^2 , we get:

$$\frac{\partial n_e}{\partial \sigma_e^2} = -\frac{(\gamma-1)n_e^2}{2(r_c - r_e)}, \quad (24)$$

$$\frac{\partial n_c}{\partial \sigma_e^2} = -\frac{(\gamma-1)n_e^2}{2(r_c - r_e)} \quad (25)$$

$$\frac{\partial r_e}{\partial \sigma_e^2} = -(\gamma-1)n_e + \frac{(\gamma-1)^2 n_e^2 (\sigma_e^2 - \sigma_{ce})}{2(r_c - r_e)} \quad (26)$$

$$\frac{\partial r_c}{\partial \sigma_e^2} = \frac{(\gamma-1)^2 n_e^2 (\sigma_{ce} - \sigma_c^2)}{2(r_c - r_e)} \quad (27)$$

Accordingly for the special case where the quantities n_e, n_c, r_e, r_c are considered as functions of the variable σ_c^2 , we get:

$$\frac{\partial n_e}{\partial \sigma_c^2} = -\frac{(\gamma-1)n_c^2}{2(r_c-r_e)}, \quad (28)$$

$$\frac{\partial n_c}{\partial \sigma_c^2} = -\frac{(\gamma-1)n_c^2}{2(r_c-r_e)} \quad (29)$$

$$\frac{\partial r_e}{\partial \sigma_c^2} = \frac{(\gamma-1)^2 n_c^2 (\sigma_e^2 - \sigma_{ce})}{2(r_c-r_e)} \quad (30)$$

$$\frac{\partial r_c}{\partial \sigma_c^2} = -(\gamma-1)n_c + \frac{(\gamma-1)^2 n_c^2 (\sigma_{ce} - \sigma_c^2)}{2(r_c-r_e)} \quad (31)$$

From the comparative statics presented above we conclude that there are three variables that play a significant role: 1) the difference between the rates of return on the two entrepreneurial activities, which is reciprocal to the degree of structural change, 2) the intertemporal elasticity of substitution, which is directly related to structural change, and 3) the existing portfolio share, which has a strong impact because of the exponent. The larger the difference between the two rates of return (creative vs. equilibrating activities), the smaller the effect of structural change is expected to be. It is also noted that the larger the intertemporal rate of substitution, the larger the value of return $(\gamma-1)$ and thus, the larger the extent of structural change. Finally, the greater the portfolio share, the larger the effect of entrepreneurial risk on portfolio shares will be.

The involvement of portfolio share on structural change requires further elaboration. The influence of existing portfolio shares refers to the concept of network effects and social capital accumulation. The larger the portfolio shares, the greater the effects on entrepreneurial activity, the larger is the social capital employed in the production function, and the greater the influences on sectoral growth.

Regarding the impact of entrepreneurial risk on the rates of return, the analysis becomes more complicated. The factors mentioned above still play a significant role, as in the case of portfolio shares. In addition, there are two more terms which appear to affect this relationship. These are the difference $\sigma_e^2 - \sigma_{ce}$ for the equilibrating events and the difference $\sigma_{ce} - \sigma_c^2$ for the creative events. In both cases, the central principle of portfolio construction holds. That is, the risk involved in equilibrating and creative entrepreneurial events is greater than the covariance between the two types of risk. The greater the magnitude of these differences the larger the degree to which risk affects the rates of return.

We are now ready to obtain our next results:

Theorem 4.5.2:

If $r_c > r_e$ then

$$\frac{\partial n_e}{\partial \sigma_e^2} > 0, \quad \frac{\partial n_c}{\partial \sigma_e^2} < 0 \quad (32)$$

$$\frac{\partial n_e}{\partial \sigma_c^2} > 0, \quad \frac{\partial n_c}{\partial \sigma_c^2} < 0 \quad (33)$$

Proof: Obvious from the above formulas and the fact that $\gamma-1 < 0$.

These results are particularly interesting regarding the impact of risk on portfolio shares in the two entrepreneurial activities. Due to the difference in the rates of return between creative and equilibrating events, the higher level of risk involved in creative events directs total entrepreneurial activity more towards equilibrating activities than creative ones. The opposite also holds. So the level of risk essentially determines the nature of the entrepreneurial activity adopted (creative-equilibrating). In the case where both types of risk are in high levels, one should expect equilibrating entrepreneurship to dominate creative entrepreneurship. This fact will also have a positive effect on economic growth. This result influences the intertemporal elasticity of substitution, making agents more willing to give up consumption stemming from the equilibrating sector rather than the creative sector. As a matter of fact, the two conditions, $r_c > r_e$ and $\gamma < 1$, should coincide if the particular type of structural change takes place. If one of the above two relationships changes direction, then one condition could offset the other. Thus, for a steady-state situation, the larger the reward of creative entrepreneurship the less the required intertemporal elasticity of substitution.

Finally, we test the relationship between entrepreneurial risk and the economy's growth rate, without distinguishing between the two types of entrepreneurship. In essence, we accept that

Theorem 4.5.3: If $r_c > r_e$ and $\sigma_e^2 > \sigma_{ce}, \sigma_c^2 > \sigma_{ce}$

$$\text{then:} \quad \frac{\partial r_e}{\partial \sigma_e^2} > 0, \quad \frac{\partial r_c}{\partial \sigma_e^2} < 0 \quad (34)$$

Proof: Obvious from the above formulas and the fact that $\gamma-1 < 0$.

Theorem 4.5.4: If $r_c > r_e$ and $\sigma_e^2 > \sigma_{ce}, \sigma_c^2 > \sigma_{ce}$

$$\text{then:} \quad \frac{\partial r_e}{\partial \sigma_c^2} > 0, \quad \frac{\partial r_c}{\partial \sigma_c^2} > 0 \quad (35)$$

Proof: Obviously by the above formulas and the fact that $\gamma-1 < 0$.

Theorem 4.5.5: If $\gamma < 1$ we have at the equilibrium:

$$\frac{\partial \psi}{\partial \sigma_K^2} > 0 \quad (36)$$

Proof: Indeed, from equation (37) we have:

$$\begin{aligned} & \frac{1}{\gamma} C^\gamma - b \delta K^\gamma + \psi K \delta \gamma K^{\gamma-1} + \\ & \frac{1}{2} \sigma_K^2 K^2 \delta \gamma (\gamma - 1) K^{\gamma-2} = 0 \end{aligned} \quad (37)$$

which means that

$$\frac{\partial \psi}{\partial \sigma_K^2} = -\frac{1}{2} (\gamma - 1) > 0 \quad (38)$$

Thus it is proved that risk and growth are positively related.

Table 1 presents the findings from the preceding analysis. We concentrate on the results regarding the effect of both types of entrepreneurial events on growth rate, portfolio shares, and rates of return.

a) The effect of a change in equilibrate risk

In equilibrating events, risk has a positive relationship with the corresponding rate of return. In turn, it has a negative relationship with the rate of return of the equilibrating events. In this case, the portfolio share of equilibrating entrepreneurial events will increase against the portfolio share of creative events and eventually the growth rate will increase. Otherwise, a decrease in equilibrating entrepreneurship risk, decreases the rate of return and the portfolio share of equilibrating events. The total growth rate in this case will fall.

b) The effect of a change in creative risk

An increase (decrease) in the risk of a creative entrepreneurial event will increase (decrease) the corresponding rate of return while having the opposite effect on the rate of return of an equilibrating event. However the increase (decrease) in the risk of a creative event will shrink the corresponding portfolio share and will increase (decrease) the portfolio share of equilibrating events. Eventually, the total rate of growth will increase (decrease). The above comments hold when the creative risk influences positively the rate of return of creative events.

c) How the portfolio share of creative entrepreneurial events could be increased

Evidently, any kind of increase in risk will increase the portfolio share of equilibrating events. Then the question arises regarding the conditions that need to hold in order for the portfolio share of creative events to increase. Two mutually exclusive conditions can provide the conditions for the creative portfolio share to be increased. The first is that the rate of return of creative events is less than the rate of return of equilibrating events. In other words, the creative sector expands when the risk becomes smaller than the level of the risk of equilibrating events. The second refers to the intertemporal elasticity

of substitution. The lower its value, the larger the portfolios share of the creative events. This last result should be evaluated in the light of the fact that the model pinpoints the direction in the change of the basic variables as we depart from the equilibrium point and for very small changes. Thus, an increase in the *high* (by definition) risk of creative events reduces their portfolio share.

d) The question of possible uniformity of growth

Equations (24–25) and (26–31) imply that uniform growth can be achieved when the existing portfolios of the two entrepreneurial events are equal. Since this is a rare situation, we conclude that risk will exercise non-uniform influences on the different portfolio shares.

5. Conclusions

This paper explores the fundamental growth question regarding which forces and under what conditions the obstacles to sustainable growth can be removed. The answer focuses on the role of uncertainty. The outcome of the analysis underlines the fact that uncertainty affects growth and structural change. In turn, structural change impacts growth, if the agent follows the basic principles of portfolio construction.

6. References

- [1] P. Petrakis, "Entrepreneurship and growth: Creative and equilibrating events," *Small Business Economics*, Vol. 9 No. 5, pp. 383–402, 1997.
- [2] F. Montobbio, "An evolutionary model of industrial growth and structural change," *Structural Change and Economic Dynamics*, No.13, pp. 387–414, 2002.
- [3] W. J. Baumol, "Macroeconomics of unbalanced growth: The anatomy of urban crisis," *American Economic Review* 57, pp. 415–426, 1967.
- [4] W. J. Baumol, S. A. B. Blackman, and E. N. Wolff, "Unbalanced growth revisited: Asymptotic stagnancy and new evidence," *American Economic Review*, No. 75, pp. 806–817, 1985.
- [5] S. Kuznets, "Economic growth and nations: Total output and production structure," Cambridge University Press, Cambridge, 1971.
- [6] S. Kuznets, "Economic development, the family and income distribution. selected essays," Cambridge University Press, Cambridge, 1988.
- [7] J. S. Metcalfe, "Evolutionary economics and creative destruction," Routledge, London, 1998.
- [8] J. S. Metcalfe, "Restless capitalism: Increasing returns and growth in enterprise economics," Mimeo, CRIC, Manchester, March 1999.
- [9] L. Passinetti, "Structural change and economic growth: A theoretical essay on the dynamics of the wealth of nations,"

- Cambridge University Press, Cambridge, 1981.
- [10] L. Passinetti, "Structural change and economic dynamics," Cambridge University Press, Cambridge, 1993.
 - [11] S. J. Turnovsky, "Methods of macroeconomic dynamics," MIT Press, Cambridge, 2000.
 - [12] R. J. Lucas, "On the mechanics of economic development," *Journal of Monetary Economics*, No. 22, pp. 3–42, January 1988.
 - [13] P. M. Romer, "Increasing returns and long-run growth," *Journal of Political Economy*, Vol. 94, No. 5, pp. 1002–1037, 1986.
 - [14] P. M. Romer, "Endogenous technological change," *Journal of Political Economy*, Vol. 98, No. 5, pp. S71–102, 1990.
 - [15] I. M. Kirzner, "Competition and Entrepreneurship," University of Chicago Press, Chicago, 1973.
 - [16] R. G. Holcombe, "Entrepreneurship and economic growth," *Quarterly Journal of Austrian Economics*, Vol. 1, No. 2, pp. 45–62, Summer 1998.
 - [17] T. F. L. Yu, "Entrepreneurial alertness and discovery" *Review of Austrian Economics*, Vol. 14, No. 1, pp. 47–63, 2001a.
 - [18] J. Schumpeter, "The instability of capitalism," *The Economic Journal*, No. 38, pp. 361–386, Vol. 123, No. 2, pp. 297–307, 1928.
 - [19] T. F. L. Yu, "An entrepreneurial perspective of institutional change," *Constitutional Political Economy*, Vol. 12, No. 3, pp. 217–236, 2001b.
 - [20] H. Westlund and R. Bolton, "Local social capital and entrepreneurship," *Small Business Economics*, No. 21, pp. 77–113, 2003.
 - [21] G. Dosi and F. Malebra, "Organizational learning and institutional embeddedness. In: Organization and strategy in the evolution of the enterprise," Macmillan, London, pp. 1–24, 1996.
 - [22] M. Brouwer, "Entrepreneurship and uncertainty: innovation and competition among many," *Small Business Economics*, Vol. 15, No. 2, pp. 149–160, 2000.
 - [23] F. H. Knight, "Risk, uncertainty and profit," Houghton Mifflin, New York, 1921.
 - [24] P. Aghion and P. Howitt, "Endogenous growth theory," MIT Press, Cambridge, 1998.
 - [25] D. Audretsch, "An empirical test of the industry life cycle," *Weltwirtschaftliches Archiv*, 1987.

Appendix A. Tables

Table 1. Summary of theoretical findings the relation among growth, portfolio shares, rates of return, and risk of two types of entrepreneurial events.

α/α	Theorems	Variables	Risk	
			σ_e^2	σ_c^2
(1)	5.4.2	ψ	> 0	> 0
(2)	5.5.2	n_e	> 0	> 0
(3)	5.5.2	n_c	< 0	< 0
(4)	5.5.3, 5.5.4	r_e	> 0	> 0
(5)	5.5.3, 5.5.4	r_c	< 0	$? 0$
(6)	5.5.5	$\frac{\partial \psi}{\partial \sigma_K^2} > 0$		

Appendix B. Proof of theorem 4.3.1 (for the reviewers only)

We shall work with the first relation, $dK = dY - Cdt$. We suppose that $Y_e = \theta_e K^e$, θ_e a constant during the period $(t, t+dt)$ and thus θ_e is the productive ratio for the equilibrating events. This means that $Y_e = \theta_e n_e K$.

We know that $dR_e = \frac{dY_e}{K^e}$. Using the Ito's Lemma and equation (6) we get:

$$dY_e = \theta_e n_e \psi K dt + \theta_e n_e K dk \quad \text{and thus}$$

$$\frac{dY_e}{K^e} = \frac{dY_e}{n_e K} = \theta_e \psi dt + \theta_e dk \quad (40)$$

We know that $dR_e = \frac{dY_e}{K^e}$ and thus by means of equation 7 we have: $Var(du_e) = Var(\theta_e dk)$ or $\sigma_e^2 dt = \theta_e^2 \sigma_K^2 dt$ which means that $\sigma_e^2 = \theta_e^2 \sigma_K^2$. Working similarly, we take the second equation, too.

Appendix C. Proof of theorem 4.4.1 (for the reviewers only)

We have to deal with the following problem:

$$\max_C E_0 \int_0^{+\infty} \frac{1}{\gamma} C^\gamma e^{-bt} dt, \quad -\infty < \gamma < 1 \quad (41)$$

s.t.

$$\frac{dK}{K} = \psi dt + dk \quad (42)$$

This is a classical stochastic optimal control problem.

We define $V(K, t) = \max_C E_t \int_t^{+\infty} \frac{1}{\gamma} C^\gamma e^{-bt} dt$. The corresponding expression to be maximised with respect to C is

$$e^{-bt} \frac{1}{\gamma} C^\gamma + \frac{\partial V}{\partial t} + \psi K \frac{\partial V}{\partial K} + \frac{1}{2} \sigma_K^2 K^2 \frac{\partial^2 V}{\partial K^2} \quad (43)$$

Assuming that the unknown function $V(K, t)$ has the separable form $V(K, t) = e^{-bt} X(K)$, Equation (43) becomes:

$$\frac{1}{\gamma} C^\gamma - bX + \psi K X_K + \frac{1}{2} \sigma_K^2 K^2 X_{KK} \quad (44)$$

The value of C where the maximum is achieved, denoted by $\hat{C}$, must make (44) equal to zero. This is the well-known Bellman equation with $X(K)$ as the unknown function. In order to solve it, we postulate a solution of the form: $X(K) = \delta K^\gamma$. Substituting this solution for (44), we have:

$$\frac{1}{\gamma} \hat{C}^\gamma - b\delta K^\gamma + \psi K \delta \gamma K^{\gamma-1} + \frac{1}{2} \sigma_K^2 K^2 \delta \gamma (\gamma-1) K^{\gamma-2} = 0 \quad (45)$$

Substituting now for this equation the relation (9), we get

$$\begin{aligned} & \frac{1}{\gamma} \hat{C}^\gamma - b\delta K^\gamma + \psi K \delta \gamma K^{\gamma-1} + \\ & \frac{1}{2} \left(\frac{K}{\theta_e} \right)^2 \sigma_e^2 \delta \gamma (\gamma-1) K^{\gamma-2} = 0 \end{aligned} \quad (46)$$

Finally, by differentiating (46), with respect to σ_e^2 and rewriting ψ instead of $\hat{\psi}$ we get

$$\frac{\partial \psi}{\partial \sigma_e^2} \delta \gamma K^\gamma + \frac{1}{2} \delta \gamma \left(\frac{1}{\theta_e} \right)^2 (\gamma-1) K^\gamma = 0 \quad (47)$$

which means that

$$\frac{\partial \psi}{\partial \sigma_e^2} = -\frac{1}{2} \left(\frac{1}{\theta_e} \right)^2 (\gamma-1) \quad (48)$$

Similarly, we can prove that

$$\frac{\partial \psi}{\partial \sigma_c^2} = -\frac{1}{2} \left(\frac{1}{\theta_c} \right)^2 (\gamma-1) \quad (49)$$

and the theorem has been established.

Appendix D. Proof of theorem 4.5.1 (for the reviewers only)

We have to solve the following optimisation problem.

$$\max_{C, n_e, n_c} E_0 \int_0^{+\infty} \frac{1}{\gamma} C^\gamma e^{-bt} dt, \quad -\infty < \gamma < 1 \quad (50)$$

s.t.

$$\frac{dK}{K} = \left(r_e n_e + r_c n_c - \frac{C}{K} \right) dt + dk \quad (51)$$

with

$$n_e + n_c = 1$$

$$\sigma_K^2 = n_e^2 \sigma_e^2 + n_c^2 \sigma_c^2 + 2n_e n_c \sigma_{ce} \quad (52)$$

This is a classical stochastic optimal control problem; to solve it we shall follow Turnovsky [11]. The corresponding Lagrangian expression to be maximised is:

$$e^{-bt} \frac{1}{\gamma} + \frac{\partial V}{\partial t} + \left(r_e n_e + r_c n_c - \frac{C}{K} \right) K \frac{\partial V}{\partial K} + \frac{1}{2} \sigma_K^2 K^2 \frac{\partial^2 V}{\partial K^2} + \mu [1 - n_c - n_e] \quad (53)$$

We assume now that the unknown function $V(K, t)$, has the separable form $V(K, t) = e^{-bt} X(K)$, we substitute this value into (53) and we put then, the partial derivatives of the resulting expression, with respect to the variables C, n_e, n_c, μ , equal to zero. We shall take:

$$\begin{aligned} C^{\gamma-1} &= X_K \\ r_c K X_K + (n_c \sigma_c^2 + n_e \sigma_{ce}) K^2 X_{KK} - \mu &= 0 \\ r_e K X_K + (n_e \sigma_e^2 + n_c \sigma_{ce}) K^2 X_{KK} - \mu &= 0 \\ n_e + n_c &= 1 \end{aligned} \quad (54)$$

These equations determine the optimal values: $\frac{C}{K}, n_e, n_c, \mu$ as functions of X_K, X_{KK} . Furthermore, substituting once more, the values we got above, for the relation (53) and cancelling the term e^{-bt} we take the Bellman equation:

$$\begin{aligned} \frac{1}{\gamma} \hat{C}^\gamma - bX(K) + \left(r_e \hat{n}_e + r_c \hat{n}_c - \left(\frac{\hat{C}}{K} \right) \right) K X_K + \\ \frac{1}{2} \sigma_K^2 K^2 X_{KK} = 0 \end{aligned} \quad (55)$$

where denotes optimised value. To solve the Bellman equation we postulate a solution of the form $X(K) = \delta K^\gamma$, where the coefficient δ can be determined. Substituting this into the Bellman equation and the equations (54), we perform some manipulations:

$$\begin{aligned} (\delta \gamma)^{1/(\gamma-1)} &= \frac{b - (r_e \hat{n}_e + r_c \hat{n}_c) \gamma - \frac{1}{2} \gamma (\gamma-1) \sigma_K^2}{1 - \gamma} \\ \delta \gamma &= \left(\frac{\hat{C}}{K} \right)^{\gamma-1} \end{aligned} \quad (22)$$

$$(\gamma-1)(\hat{n}_c \sigma_c^2 + \hat{n}_e \sigma_{ce}) + r_c = \frac{\mu}{\delta \gamma K^\gamma}$$

$$(\gamma-1)(\hat{n}_e \sigma_e^2 + \hat{n}_c \sigma_{ce}) + r_e = \frac{\mu}{\delta \gamma K^\gamma}$$

$$\hat{n}_c + \hat{n}_e = 1$$

and the theorem has been established.

Distortion of Space and Time during Saccadic Eye Movements

Masataka Suzuki¹, Yoshihiko Yamazaki²

¹Department of Psychology, Kinjo Gakuin University, Nagoya, Japan

²Department of Computer Science and Engineering Graduate School of Engineering
Nagoya Institute of Technology, Nagoya, Japan
Email: msuzuki@kinjo-u.ac.jp

Abstract

The space-time distortion perceived subjectively during saccadic eye movements is an associative phenomenon of a transient shift of observer's visual frame of reference from one position to another. Here we report that the lines of subjective simultaneity defined as two spatially separated flashes perceived during saccades were nearly uniformly tilted along the physical time-course. The causality of the resulting space-time compression may be explained by the Minkowski space-time diagram in physics.

Keywords: Saccade Space Time Compression

1. Introduction

In vision, an observer's frame of reference is one of the necessary elements of perceiving an event in space (where) and time (when). During saccadic eye movements, a space once recognized in one frame of reference is distorted toward the target as the eye fixates it in the new frame of reference [1,2]. Morrone *et al.* [3] demonstrated that the resulting space compression accompanies an underestimation of the perceived time interval between temporally separated stimuli. Since these two illusory effects occurred in nearly the same time range, they suggested the possibility of a single unifying mechanism of both the space and time compressions. Moreover, they speculated that the space-time compression is originated from anticipatory repositioning of visual receptive fields [4–6], leading to an immediate relativistic consequence in perceptual space and time [7]. To approach this mechanism, however, it is necessary to probe into the transient dynamics of these illusions [8]. In this study, we provide psychophysical evidence that the space-time compression *during* saccadic eye movements could be attributable to the backward temporal shift of time-course by which the observer perceives the saccade target in a new frame of reference. These effects on visual percepts of space, time and simultaneity may be explained along the framework of a 'thought experiment' of

special relativity theory [9].

2. Methods

In this study, observers made a judgment as to the simultaneity of two briefly flashed stimuli appearing at different times during the course of a horizontal saccade (Figure 1(a)). The observers were seated in the dark, with their head fixed by a chin- and forehead-rest. Two-colored LEDs positioned on the black board in front of the observers (viewing distance: 40cm), were used as a central fixation target (FT) and a saccade target (ST) (panel 1 in Figure 1a). After fixating to the fixation target (FT) for a period of time between 1500 and 2500 ms, both the FT and saccade targets (ST) were turned off simultaneously (pane 2 in Figure 1(a)), and the observers made a 30 degrees (°) horizontal saccade to the remembered ST as soon as both targets disappeared. Either of the FT or ST was also used to provide a standard stimulus (SS) or a comparison stimulus (CS), respectively. The early rising phase of electro-oculographic (EOG) signal less than 15% relative to its maximal value was used as a common triggering source (CTS) of SS and CS, each of which provides a flash with short exposure time (1 ms) at different latencies from CTS (e.g., see panel 3 in Figure 1(a)). The short exposure (1 ms) of flashed stimuli effectively minimized motion blur during saccades, allowing for the

observers to specify the apparent positions of the two flashes.

In the first condition (Cond. 1), FT and ST provide SS and CS, respectively (e.g., see panel 3 in Figure 1(a)), while in the second condition (Cond. 2) they provide CS and SS, respectively. The SS was flashed at different constant latencies of 20, 33, 50 or 100 ms from a CTS signal. In each latency condition, the experimenter initially set the time interval of CS and SS far below or far above the observer's threshold, and the observers were asked to adjust the variable timing of the CS on a trial-by-trial basis until it appeared equal to the latency of the SS, and to report corresponding spatial positions of the two flashes (e.g., see panel 4 in Figure 1(a)). In each of these ascending or descending sessions, the timing of the CS relative to that of SS was varied by the observer continuously via a dial on a pulse generator, and apparent position of two flashes perceived as simultaneous during the course of the saccade was pointed by the observer on a trial-by-trial basis, by adjusting the sensor head of a linear potentiometer (Novotechnik, TLH1000, length 0.8 m), fixed sideways along a black board. In Cond.1 and 2, the subjective simultaneity of two flashes was estimated from the mean of five sets of ascending and descending sessions in each latency condition.

Horizontal eye position was recorded using an EOG system (AVH-10, Nihon-Koden) by placing Ag-AgCl skin electrodes at the outer canthi of both eyes. A ground electrode was placed just above the eyebrows in the center of the forehead. As described earlier, the early rising phase of the EOG signal was used to trigger two pulse generators for the LED flashes of SS and CS. The base line adjustment of EOG signal was carried out carefully on a trial-by-trial basis. Target presentation and data collection were controlled using custom software programmed in LabVIEW (National Instruments). The eye position signals were digitally low-passed filtered at 50 Hz, using a second-order Butterworth filter implemented in MatLab (The Mathworks). The onsets of eye movement were scored on the basis of 5% of the peak velocity of their position signals. Seven and six observers were used in Cond. 1 and Cond. 2, respectively.

3. Results

The spatial relationship between the perceived positions of the SS flashes and the corresponding eye positions was out-of-phase in both conditions (Figure 1(b), (c)). In Cond. 1, just after the instant of the saccade the SS appearing on FT was greatly mislocalized once nearest to the ST, but appeared near to the initial FT position as the eye fixated on the ST. In Cond. 2, the SS triggered 20 ms after the onset of CTS was invisible on ST, but in other

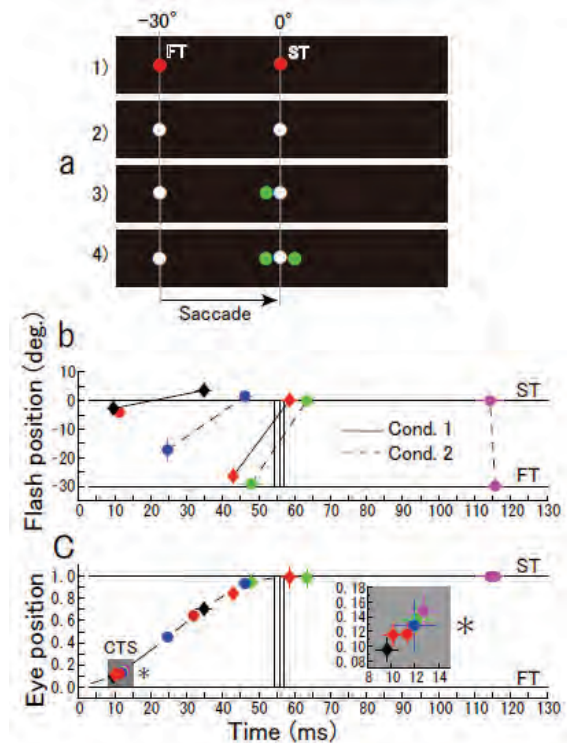


Figure 1. Subjective simultaneity of two flashes. (a) Spatial layout of two targets (FT, ST) and perceived flash stimuli specified on a black board. The same two-colored LEDs were used for both FT and ST targets. 1) FT (red) and ST (red) presented simultaneously at -30° right and 0° (screen center), respectively. 2) latency period from the simultaneous disappearance of two targets. White circles are not real, but are to refer the spatial position of the two targets. 3) a typical example of Cond. 1, showing that the spatial position of a FT flash (green) is greatly mislocalized toward the ST, and the ST flash triggered at the same time with the FT flash is invisible. 4) spatial positions of ST and FT flashes perceived as simultaneous. (b) Space-time diagram of two flashes perceived as simultaneous. The lines of subjective simultaneity in all latency conditions are represented by the solid and dashed lines in Cond.1 and Cond. 2, respectively. In Cond. 1, a pair of black or red diamonds represents the averaged estimate for the tasks when the latencies of SS flashes from the CTS signals are set at 0 or 33 ms, respectively. In Cond. 2, similarly, red, blue, green and magenta circles represent the averaged estimate for the tasks when latencies of SS flashes from the CTS signals were set at 20, 33, 50 and 100 ms, respectively. Note when SS latency from the CTS signal was set at 20 ms, the flash on the ST was invisible and thus the CTS signal (red) is depicted alone. (c) The corresponding amplitudes of the EOG signal to the flash times of a pair of SS and CS in panel b. Each value is normalized to the magnitude at 100 ms after the movement ends. Small panel on the right (asterisk) shows enlarged representation of latencies and amplitudes of the CTS for all tasks. Note that the variations of CTS measures are limited in time and amplitudes across all tasks. Three vertical bars in panels b) and c) represent mean±s.d. of eye movement times for all tasks and subjects.

tasks it was clearly identified, having a negative relationship to eye position, similar to the Cond.1.

In both conditions, when two stimuli were presented simultaneously, the ST flashes were perceived to occur earlier than the FT flashes across the saccadic period. Therefore, the observers estimated the simultaneity of the two flashes by delaying the onset time of the ST flash relative to that of FT in Cond. 1 (e.g., see panel 4 in Fig 1(a)), or by preceding the onset time of FT flash relative to the ST in Cond. 2 (Figure 1(b)). In Cond. 1, the time intervals of two flashes as an estimate of the subjective-simultaneity averaged $25 (\pm 2)$ and $16 (\pm 1)$ ms, for SS set at 0 and 33 ms from the onset of CTS, respectively (solid lines). In Cond. 2 the ST at less than 30 ms after saccade onsets was invisible, so the same measures of subjective simultaneity of the two flashes were limited to the other three cases, averaging $22 (\pm 1)$, $16 (\pm 4)$ and $2 (\pm 1)$ ms, for SS set at 33, 50 and 100 ms from the onset of CTS, respectively (dashed lines). In both conditions, therefore, the subjective simultaneity of two flashes can be referred to as the rightward tilt of the lines of simultaneity and as their directional uniformity across the saccadic period.

The effect of target eccentricity on perceiving simultaneity of two flashes was examined under static conditions without a saccade. The observers gazed -30° , -15° or 0° relative to the ST, and flash stimuli (SS/CS) were provided by two targets positioned at $-5^\circ/5^\circ$, $-15^\circ/5^\circ$ or $-30^\circ/0^\circ$, respectively. These spatial relationships between gaze directions and two flashes were roughly analogous to those in Cond. 2 (Figure 1 (b), (c)). Both stimuli were given by the experimenter, while the observers were asked to synchronize them by adjusting their interval on a trial-by-trial basis. Two flashes were apparently perceived in all tasks, and their intervals perceived as simultaneous averaged $2 (\pm 2)$, $1 (\pm 3)$ and $2 (\pm 5)$ ms at -30° , -15° and 0° conditions, respectively. In the static condition, therefore, the effect of target eccentricity on estimating the subjective simultaneity of two flashes could be minor.

4. Discussions

In this study, we found that the lines of subjective simultaneity defined as two spatially separated flashes during saccades were nearly uniformly distorted on the physical time-course (Figure 1(b)). When interpreting this in perceptual space-time, however, two simultaneous events (flashes) must be on a line parallel to the space axis. This corresponds to the backward temporal shift of the time-course of ST flashes, relative to the FT flashes (Figure 2). When this shift component, herein termed Δt ,

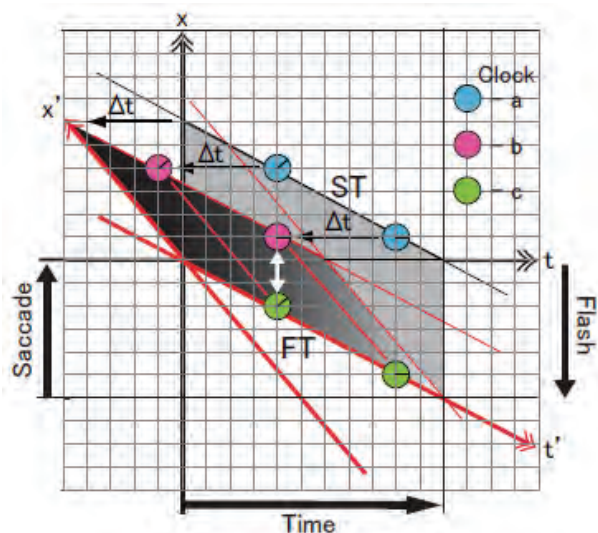


Figure 2. Schematic illustration of the space-time compression based on our results. Two coordinates, (x, t) and (x', t') , are superimposed, analogues to the Minkowski space-time diagram. The first is *Newtonian* space-time coordinates to define position (x) and time (t) for a moving object(s) in real (physical) space-time, and the second is perceptual space-time coordinates (x', t') hypothesized based on the present experiments. For details, see text.

is applicable over a saccadic period, the invisibility of the ST flashed during the initial half of the saccadic period (Figure 1(b)) could be ascribed to the backward temporal shift of the flash's percept beyond the extent of the conscious time window [10] by the amount Δt . Similarly, an earlier recovery time of the flash position to the ST than to the FT near the movement end may be explained by the same scheme. As for the latter, it is well established [11] that the target percept at the end of the saccade is referred backward in time to compensate for the time lost during saccadic suppression [12,13]. Of more importance in Figure 2 is the strong dependence of both the space and time compressions on the backward shift of the time-course of the ST flashes. This scheme is different from the convergent type of compression proposed previously [1–3].

Among these, Morrone *et al.* [3] have shown that time compression perceived subjectively during saccadic eye movements is an associative phenomenon of space compression, which was evidenced as a convergent type of mislocalization of visual stimuli toward the endpoint of the saccade. They suggested that the resulting alternation of spatial and temporal metrics of perceptual space-time would lead to relativistic-like effects on the visual percept [7]. Our results, by contrast, indicate that space and time compressions could be attributable to the shift component of mislocalization of visual stimuli. To illustrate this, as shown in Figure 2, we define *Newtonian*

space-time coordinates (x, t') , where the tilted time axis t' is to specify the corresponding position (x) of two flashes, both moving in the opposite direction to the saccadic eye movement. Since in *Newtonian* space-time the geometry of space is *Euclidian* and the time is universal for all observers, the simultaneity of two flashes is specified on the line parallel to the space axis x (e.g., see clock a and c). However, this was not the case in the perceptual space-time of our observers. According to the rightward tilt of the lines of simultaneity over a saccadic period in figure 1b, the sequence of events from the viewpoint of our observer may be illustrated graphically by shifting the timescale on the ST in the diagram backward by an amount Δt . This corresponds to a tilt in the space axis from vertical (x) to leftward (x') (red lines). Note this lead us to define another space-time coordinates (x', t') , in which a hypothetical observer moving with this frame of reference, sees all events occurring on a line parallel to the space axis x' as simultaneous (e.g., see clock b and c). The superimposed representation of two coordinates, (x, t') and (x', t') , and their interrelationship are analogues to the Minkowski space-time diagram in special relativity theory [9]. Thus, considering our results from the viewpoint of our observer, three relativistic-like effects, the lack of absolute simultaneity, space contraction (compression) and time compression, can be expected. First, the fact that the observers estimated the simultaneity of two flashes by delaying the time of the ST flash relative to that of FT by the amount Δt suggests that during saccadic periods the observers see the two flashes occurring in the space-time coordinates (x', t') . As a result, the two flashes perceived as simultaneous in (x', t') are not simultaneous in (x, t') . Second, if this was the case, as shown by the distance of two white arrow heads in the figure, the spatial distances of the two flashes appear to contract in the direction of motion (space compression). Third, for this observer, all events happening on the moving ST flashes are compressed relative to that on the FT (time compression), or the time passed on the clock c is dilated relative to time passed on the clock b , by the amount Δt . Taken together, what is novel here is to present a single unifying mechanism of space and time compression using the Minkowski diagram [9], in which the space-time compression during saccades could be ascribed to the homogeneous distortion of space along a time scale, rather than the convergent type of compression proposed by Morrone *et al.* [3].

5. Conclusions

In the present study, the pattern of space-time distortion perceived subjectively during saccadic eye movements was studied in order to gain insights into the nature of

corresponding space-time compression inherent in human visual perception. We found that the lines of subjective simultaneity defined as two spatially separated flashes perceived during saccades were nearly uniformly tilted along the physical time-course. This tempted us to speculate that vision may be subject to relativistic effects, similar to physical relativistic effects that occur at speeds approaching the speed of light. It is well established that neuron's receptive fields or their representation of space are not static entities but that they start to change perisaccadically to bring a visual stimulus defined in pre-saccadic frame of reference into a post-saccadic frame of reference [4–6]. When this dynamic coordinate transformation is rapid, approaching the physical limit of neural information transfer, the relativistic consequences may be expected.

6. References

- [1] M. Lappe, H. Awater and B. Krekelberg, "Postsaccadic visual references generate presaccadic compression of space," *Nature* 403, pp. 892–895, 2000.
- [2] J. Ross, M. C. Morrone, and D. C. Burr, "Compression of visual space before saccades," *Nature* 384, pp. 598–601, 1997.
- [3] M. C. Morrone, J. Ross, and D. C. Burr, "Saccadic eye movements cause compression of time as well as space," *Nature Neuroscience* 8, pp. 950–954, 2005.
- [4] J. R. Duhamel, C. L. Colby, and M. E. Goldberg, "The updating of the representation of visual space in parietal cortex by intended eye movements," *Science* 255, pp. 90–92, 1992.
- [5] M. Kusunoki and M. E. Goldberg, "The time course of perisaccadic receptive field shifts in the lateral intraparietal area of the monkey," *Journal of Neurophysiology* 89, pp. 1519–1527, 2003.
- [6] M. M. Umeno and M. E. Goldberg, "Spatial processing in the monkey frontal eye field. I. Predictive visual responses," *Journal of Neurophysiology* 78, pp. 1373–1383, 1997.
- [7] M. C. Morrone, J. Ross, and D. C. Burr, "Keeping vision stable: rapid updating of spatiotopic receptive fields may cause relativistic-like effects," In R. Nijhawan (Ed.), *Space and time in perception and action*, Cambridge: Cambridge University Press, 2008.
- [8] D. M. Eagleman, "Distortion of time during rapid eye movements," *Nature Neuroscience* 8, pp. 850–851, 2003.
- [9] A. Einstein, "Relativity: The Special and General Theory," New York: Henry Holt, 1920.
- [10] B. Libet, E. W. J. Wright, B. Feinstein, and D. K. Pearl, "Subjective referral of the timing for a conscious sensory experience: a functional role for the somatosensory specific projection system in man," *Brain* 102, pp. 193–224, 1979.
- [11] K. Yarrow, P. Haggard, R. Heal, P. Brown, and J. C. Roth-

- well, "Illusory perceptions of space and time preserve cross-saccadic perceptual continuity," *Nature* 414, pp. 302–305, 2001.
- [12] M. R. Diamond, J. Ross, and M. C. Morrone, "Extraretinal control of saccadic suppression," *Journal of Neuroscience* 20, pp. 3442–3448, 2000.
- [13] M. C. Morrone, J. Ross, and D. C. Burr, "Apparent position of visual targets during real and simulated saccadic eye movements," *Journal of Neuroscience* 17, pp. 7941–7953, 1997.

Probabilistic Verification over $GF(2^m)$ Using Mod2-OBDDs

José Luis Imaña

Department of Computer Architecture, Faculty of Physics, Complutense University, Madrid, Spain
Email: jluimana@dacya.ucm.es

Abstract

Formal verification is fundamental in many phases of digital systems design. The most successful verification procedures employ Ordered Binary Decision Diagrams (OBDDs) as canonical representation for both Boolean circuit specifications and logic designs, but these methods require a large amount of memory and time. Due to these limitations, several models of Decision Diagrams have been studied and other verification techniques have been proposed. In this paper, we have used probabilistic verification with Galois (or finite) field $GF(2^m)$ modifying the CUDD package for the computation of signatures in classical OBDDs, and for the construction of Mod2-OBDDs (also known as $\oplus$ -OBDDs). Mod2-OBDDs have been constructed with a two-level layer of $\oplus$ -nodes using a positive Davio expansion (pDE) for a given variable. The sizes of the Mod2-OBDDs obtained with our method are lower than the Mod2-OBDDs sizes obtained with other similar methods.

Keywords: Verification, Probabilistic, OBDD, Mod2-OBDD, Galois Field $GF(2^m)$

1. Introduction

One of the most important aspects during circuit design is the verification, *i.e.*, checking for functional equivalence. The translation of a circuit design from a high-level specification to a physical implementation depends on the correct transformation of its description at higher levels of abstraction to equivalent descriptions at more detailed levels. At logic level, verification consists of checking the equivalence of a Boolean function specification and its logic implementation. Most of the current successful equivalence checkers use Binary Decision Diagrams (BDDs) [1,2] or their derivatives as a core of the equivalence deduction engine. OBDDs [3,4] are used as canonical representations for both Boolean circuit specifications and logic designs. While OBDD-based methods have been quite successful in verifying combinational and sequential circuits, they have significant limitations. For many circuits, verification systems that represent functions as OBDDs require a large amount of memory and time, and for some circuits the resource requirements are unacceptably large. Furthermore, OBDD sizes are quite sensitive to the ordering of their Boolean

variables [5,6]. Various techniques have been proposed to reduce the memory complexity of BDDs by exploiting the structural and functional similarities of the circuits [7–9]. Some alternative to BDD-based equivalence checkers use Boolean Satisfiability (SAT) [10,11] or SAT-like methods (ATPG [12], recursive learning [13]) as a principal engine.

In spite of the considerable advances in the area, the growing complexity of the verification instances motivates exploring the alternative approaches. Verification performed using OBDDs can be considered as a Deterministic Verification, in which if the OBDDs of the functions are the same (isomorphic), then these functions are said to be equivalent [14]. Another method that can be considered in order to circumvent the above limitations is the Probabilistic Verification, based on the theory of Blum, Chandra and Wegman [15], in which numeric codes or signatures represent the functions to be verified. This method is based on algebraic transforms of Boolean functions, so that a function can be substituted by its algebraic representation. A signature representing the function is then obtained assigning numeric codes (randomly selected from a finite field) to the variables in the algebraic representation, and then evaluating the result. The comparison is then performed over signatures, not over

This work was partially supported by the Spanish Government Research Grant TIN2008-00508.

data structures representing the functions [16]. Signatures can be computed more efficiently than graph-based representations, consume less time and space, and distinguish any pair of Boolean functions with a very high probability of success. By performing several such runs with different random input variable assignments, the resultant algebraic simulation has a probability of error in verification that decreases exponentially from an initially small value [16]. The probabilistic approach presents significant advantages over deterministic methods using OBDDs. In deterministic verification, OBDDs are canonical representations of functions to be verified, and provide a basis for efficient computation. In probabilistic verification, functions are represented by signatures, not by a large data structure. Graph-based data structures are used only in intermediate evaluation steps. Therefore, more general OBDD-based models have been studied which can be used as such intermediate data structures [17–20].

In this paper, we consider Mod2-OBDDs [21] (also known as $\oplus$ -OBDDs) which are extensions of OBDDs. Mod2-OBDDs are non canonical representations of Boolean functions. For canonical representations as OBDDs, testing the equivalence of two OBDDs simply reduces to the comparison of their pointers. For non canonical representations as Mod2-OBDDs, a deterministic equivalence test requires time cubic in the number of nodes [22], and thus, it seems not to be suitable for practical purposes. In [21], a fast probabilistic equivalence test for Mod2-OBDDs that requires only a linear number of arithmetic operations is used.

In this contribution, a very efficient OBDD package (CUDD package from Colorado University) has been modified in order to construct Mod2-OBDD representations of Boolean functions given in multilevel BLIF, and probabilistic verification based on Galois field $GF(2^m)$ arithmetic for the equivalence test (signatures comparison) has been used. The addition of two elements from a binary extension field $GF(2^m)$ is simply a bitwise XOR of their corresponding binary representations, the subtraction is exactly the same as addition, and the complexity of the multiplication depends on the irreducible generating polynomial or the basis selected to represent the field elements [23–25]. These properties justify the use of $GF(2^m)$ as finite field.

Firstly, OBDDs with signatures have been constructed. The signature-inclusion is carried out by including a 32 bit signature field on each OBDD-node and computing the signature in the synthesis process. This signature-OBDD obtained has the same number of nodes as the original OBDD, but with the signature computed for the Boolean function that it represents. Secondly, Mod2-OBDDs have been constructed by the inclusion of a two-level layer of $\oplus$ -nodes. This layer is created—using the positive Davio expansion (pDE) for a selected variable—in the synthesis of the BLIF file for the function.

Signature is so computed for the Mod2-OBDD and is compared with signature obtained with the signature-OBDD for the equivalence test. Times and sizes are compared for the three used structures (OBDDs, signature-OBDDs and Mod2-OBDDs).

The paper is structured as follows: In Section 2, some basic concepts concerning Mod2-OBDDs are introduced. In Section 3, probabilistic equivalence test using Galois field $GF(2^m)$ is presented. In Section 4, modifications on CUDD package for signature computation are outlined. Section 5 deals with the introduction of $\oplus$ -nodes for the Mod2-OBDD construction. In Section 6, experimental results are presented. Finally, some conclusions are included in Section 7.

2. Mod2-OBDDs

Mod2-OBDDs have been defined in [21]. A Mod2-OBDD (also known as $\oplus$ -OBDDs) over a set $X_n = \{x_1, x_2, \dots, x_n\}$ of Boolean variables is a directed acyclic connected graph P where each node has out-degree 2 or 0. There is a distinguished non-terminal node, the *root*, which has in-degree 0. The two terminal nodes with out-degree 0, the 0-sink and the 1-sink, are labeled with the Boolean constants 0 and 1, respectively. The remaining non-terminal nodes v are either labeled with Boolean variables $x_i \in X_n$ (denoted as branching or decision nodes), or with the binary Boolean operation XOR ($\oplus$ -nodes). Let $l(v)$ denote the label of the node v . The two edges starting in a non-sink node are labeled with 0 and 1. The 0-successor and 1-successor nodes of v are denoted by v_0 and v_1 , respectively. If v_2 is a successor of v_1 in P and $l(v_1), l(v_2) \in X_n$ then $l(v_1) < l(v_2)$ according to a given ordering on the set of input variables.

The function f_P associated with a Mod2-OBDD P is determined as follows. Given an input assignment $a = (a_1, a_2, \dots, a_n) \in \{0, 1\}^n$, the Boolean values assigned to the leaves extend to Boolean values associated with all nodes of P as follows:

- If the successor nodes v_0, v_1 of a node v of P carry the Boolean values b_0, b_1 , respectively, and if $l(v) = x_i$, then v is associated with the value b_0 or b_1 according to $x_i = 0$ or $x_i = 1$.
- If $l(v) = \oplus$, then the value $\oplus(b_0, b_1) = (b_0 + b_1) \bmod 2$ is associated with v .

The value $f_P(a)$ of the Boolean function f_P represented by a Mod2-OBDD P is the value 1 or 0 associated with the root of P under the assignment a . A more compact representation can be obtained by using complemented edges [26].

3. Probabilistic Verification Using Galois Fields

Probabilistic verification consists of the generation of a

code or signature for a Boolean function. The probabilistic comparison of two functions can be carried out by evaluating their representations on a randomly chosen vector of values selected from a finite field and comparing the results (signatures) obtained from the evaluation. If the signatures are different, then the functions are inequivalent with certainty. If they are equal, then the functions are equivalent with a small probability of error.

The signature for a Boolean function $f(x_{n-1}, \dots, x_0)$ is generated by selecting random numerical values for each x_i . These values can be selected from any field, but we will use the finite field $GF(p^m)$, with p prime and $m \in \mathbb{N}$, representing a Galois field with p^m elements. The evaluation of the function (with these randomly selected values assigned to its inputs) is carried out by the replacement of $f(x_{n-1}, \dots, x_0)$ by an equivalent arithmetic function defined over the finite field. This replacement is determined by an algebraic transformation of the Boolean function in terms of polynomials over the finite field. If we assign the polynomial $p_x = x$ to a Boolean variable x , we can transform [16] the Boolean functions $\neg f$ and $f_1 \wedge f_2$ into the arithmetic expressions $1 - p_f$ and $p_{f_1} \cdot p_{f_2}$, respectively, where p_f represents the polynomial assigned to the Boolean function f . By using the law of DeMorgan and idempotence, we can also transform $f_1 \vee f_2$ and $f_1 \oplus f_2$ into $p_{f_1} + p_{f_2} - p_{f_1} \cdot p_{f_2}$ and $p_{f_1} + p_{f_2} - 2 \cdot p_{f_1} \cdot p_{f_2}$, respectively. It must be noted that the above arithmetic operations are carried out on the selected finite field.

In this paper, we consider the Galois field $GF(2^m)$ which is a characteristic 2 finite field with 2^m elements, each of them represented as an m -bit vector. $GF(2^m)$ is an extension field of the ground field $GF(2) = \{0, 1\}$. The nonzero elements of $GF(2^m)$ are generated by a primitive element α , where α is a root of a primitive irreducible polynomial $g(x) = x^m + g_{m-1}x^{m-1} + \dots + g_0$ over $GF(2)$. The nonzero elements of $GF(2^m)$ can be represented as the powers of α , i.e., $GF(2^m) = \{0, \alpha^1, \alpha^2, \dots, \alpha^{2^m-2}, \alpha^{2^m-1} = 1\}$. Since α is a root of $g(x)$, $g(\alpha) = 0$, and $\alpha^m = g_{m-1}\alpha^{m-1} + \dots + g_1\alpha + g_0$. Therefore, an element of $GF(2^m)$ can also be expressed as a polynomial of α with degree less than m , that is, $GF(2^m) = \{a_{m-1}\alpha^{m-1} + \dots + a_1\alpha + a_0 \mid a_i \in GF(2) \text{ for } 0 \leq i \leq m-1\}$. Arithmetic in a field of characteristic 2 is essentially modulo arithmetic. Therefore, the addition of two polynomials becomes the bitwise XOR of the corresponding binary representations, and subtraction is the same as addition. Multiplication of two polynomials is the most important and one of the most complex and time-consuming operations. Complexity depends on many factors, such as the selection of the irreducible polynomial or the basis selected to represent the field elements: polynomial, dual or normal bases [25]. Because of its characteristic 2, the product $2 \cdot p_{f_1} \cdot p_{f_2}$ in $GF(2^m)$ is zero, and the above polynomial for the XOR is simplified to the $GF(2^m)$ addition. Complement of a field ele-

ment is simply the complementation of its least significant bit (in polynomial basis), and multiplication can be carried out in any of the representation bases. It must be noted that the multiplication operation in $GF(2^m)$ requires $O(m)$ elementary operations [27].

When the signatures $[H_1]$ and $[H_2]$ of two functions f_1 and f_2 to be checked for equivalence are computed (using the above transformations, and evaluating them over the Galois field), then the signatures are compared: if $[H_1] \neq [H_2]$, then the functions are inequivalent with certainty; but if $[H_1] = [H_2]$, then the functions are equivalent, with small probability of error.

The probability of error [16] is bounded by $1 - e^{-n/p}$ (where n is the number of variables and p is the cardinality of the finite field) and if $p \gg n$, then $1 - e^{-n/p} \approx n/p$. Therefore the error probability associated with the probabilistic equivalence check can be reduced by either increasing the size of the field, or by making multiple runs (using on each run an independent set of random assignments and computing the signatures for the functions). If the signatures are different, we are sure that the two functions are not the same. If they are equal, we choose a new set of input assignments and reevaluate. The probability of erroneously deciding that the functions is equal decreases exponentially with the number of such runs.

Applying these concepts, the probabilistic equivalence test of two functions represented by their Mod2-OBDDs is determined by the algebraic transformation of the Mod2-OBDDs in terms of polynomials over $GF(2^m)$. A polynomial $p_v: (GF(2^m))^n \rightarrow GF(2^m)$ can be associated with each node v of the Mod2-OBDD as follows [28]:

$$p_v(x_{n-1}, \dots, x_0) = \begin{cases} 0(1), & v \text{ is } 0(1) - \text{sink} \\ x \cdot p_{v_l}(x_{n-1}, \dots, x_0) + (1-x) \cdot p_{v_r}(x_{n-1}, \dots, x_0), & l(v) = x \in X_n \\ p_{v_0}(x_{n-1}, \dots, x_0) + p_{v_1}(x_{n-1}, \dots, x_0), & l(v) = \oplus \end{cases}$$

The polynomial associated with a Mod2-OBDD P is the polynomial associated with the root of P , so the equivalence test of two functions is the signature comparison of their Mod2-OBDD roots polynomials evaluated at the randomly chosen set of input variables selected from $GF(2^m)$.

Let P_f and P_g be two Mod2-OBDDs representing the Boolean functions f and g , and assume that $a_0, \dots, a_{n-1} \in GF(2^m)$ are generated independently and uniformly at random. For the Boolean signatures p_f and p_g computed for the Mod2-OBDDs P_f and P_g it holds [28] that $p_f(a_0, \dots, a_{n-1}) = p_g(a_0, \dots, a_{n-1})$, if $f = g$, and $\text{Prob}(p_f(a_0, \dots, a_{n-1}) = p_g(a_0, \dots, a_{n-1})) < 1/2$, if $f \neq g$. Therefore, if two given signatures $p_f(a_0, \dots, a_{n-1})$ and $p_g(a_0, \dots, a_{n-1})$ are equal, then the functions f and g are equal only with a certain probability. An estimation of the probability that the signatures for two nodes representing different Boolean functions in a

Mod2-OBDD P are equal can be found in [28]. By using s different signatures per node the error probability computes to at most

$$\text{error} < \frac{\text{size}(P)^2 \cdot n^s}{2 \cdot |GF|^s} \quad (1)$$

where $\text{size}(P)$ denotes the number of nodes of the Mod2-OBDD P , n is the number of variables, and $|GF|$ the cardinality of the finite field $GF(2^m)$. For our experiments given in Section 6, we use the field $GF(2^{16})$. Therefore, if we have, for example, a Mod2-OBDD with 10^7 nodes depending on 100 variables, we should use 6 different signatures per node in order to obtain an error probability of less than $6.31 \cdot 10^{-4}$. It must be noted that for the circuit c1908 studied in section 6, for example, we should use 4 different signatures in order to obtain an error probability of less than $1.25 \cdot 10^{-5}$. However, this is an error bound. In fact, in our experiments and for only one signature per node, all the experiments performed in order to check the equivalence of two different circuits, were successfully detected by obtaining two different signatures. Furthermore, the work here presented is a first approach that can be further improved.

4. Including Signatures on CUDD Package

We have used for our implementations the CUDD package, one of the most efficient OBDD packages [29]. CUDD package has been modified in order to include signatures into OBDD nodes, and compute the signature of the Boolean function represented by an OBDD. The signature of the function will be the signature of the root node(s) of the OBDD. We name these decision diagrams signature-OBDDs (or s-OBDDs for simplicity). When the s-OBDD for a function has been constructed, then it can be verified comparing its signature with the signature computed by the construction of the Mod2-OBDD.

CUDD modifications consist of the inclusion of two new fields for each OBDD node: a 1-bit decision- x or field used to distinguish a decision node from a $\oplus$ -node, and a 32-bit signature field used for the signature of the node. The most significant 16-bits of this field store the initial signature randomly assigned to the variable associated with the node, and the least significant 16-bits store the signature of the (sub)function represented by the node. Therefore, $GF(2^{16})$ is used for signature representation.

In the synthesis process, Boolean operations are implemented using the *ite* operator [30] defined as a ternary Boolean function for three inputs F, G, H by “If F then G else H ”. This is equivalent to $\text{ite}(F, G, H) = F \cdot G + \bar{F} \cdot H$ and can be evaluated by recursive application of the Shannon decomposition $f = x \cdot f|_{x=1} + \bar{x} \cdot f|_{x=0}$, where positive and

negative cofactors $f|_{x=1}$ and $f|_{x=0}$ are the function f evaluated at $x=1$ and $x=0$, respectively. Positive and negative cofactors of a function associated with a decision node v , $l(v)=x$, are derived by simply returning the 1-successor and the 0-successor of v , respectively. In our approach, signature computation is performed in the synthesis process using the Shannon decomposition. Let $[x]$ be the signature initially assigned to x variable (the most significant 16-bits of v 's signature field), and $[v_0]$ and $[v_1]$ the signatures associated to 0- and 1-successor of v node (signatures of the functions represented by v_0 and v_1), respectively. Then the signature $[v]$ associated to the function represented by v can be computed by the following arithmetic operations over $GF(2^{16})$:

$$\begin{aligned} [v] &= \bar{[x]} \cdot [v_0] + [x] \cdot [v_1] = (1 - [x]) \cdot [v_0] + [x] \cdot [v_1] \\ &= (1 + [x]) \cdot [v_0] + [x] \cdot [v_1] \end{aligned} \quad (2)$$

where the properties and transformations of Boolean functions into arithmetic expressions given in Section 3 have been used. If the polynomial basis for the representation of the field elements is used, then the signature $1+[x]$ is simply the complementation of the least significant bit of $[x]$. When the signature $[v]$ of the function represented by v has been computed, then it is stored at the least significant 16-bits of the signature field of v . Therefore, the signature of a Boolean function represented by an OBDD can be computed in the synthesis process (*ite* operation). An OBDD with signatures is named signature-OBDD (s-OBDD).

5. Introduction of $\oplus$ -Nodes into the OBDD

Mod2-OBDD for a Boolean function is created by the introduction of an upper two-level layer of $\oplus$ -nodes while the OBDD is being constructed. The main advantage of a data structure with OBDDs and $\oplus$ -nodes is that the signature for a $\oplus$ -node can be directly computed by simply performing the bitwise exclusive-or of the signature associated to its 1- and 0-successors (if Galois field is used).

We have constructed Mod2-OBDDs for combinational circuits given in BLIF format. For the introduction of the two-level layer of $\oplus$ -nodes, we have used the positive Davio expansion. The positive Davio expansion (pDE) of a function f with reference to a variable x_i is given by the following expression:

$$\begin{aligned} f(\dots, x_i, \dots) &= f|_{x_i=0} \oplus x_i \cdot (f|_{x_i=0} \oplus f|_{x_i=1}) \\ &= f|_{x_i=0} \oplus (x_i \cdot f|_{x_i=0} \oplus x_i \cdot f|_{x_i=1}) \end{aligned} \quad (3)$$

Using this expansion, an upper layer of two $\oplus$ -nodes is constructed for the function as shown in Figure 1(a). In this representation, the upper $\oplus$ -node of the layer has its

else edge pointing to the function $f|_{k_i=0}$, while its then edge points to the lower $\oplus$ -node. The else edge of this lower $\oplus$ -node points to the function $x_i \cdot f|_{k_i=0}$ and the then edge points to the function $x_i \cdot f|_{k_i=1}$. This structure is given in Figure 1(b) for a circuit with many outputs and when a given variable x_i is selected for the pDE. This expansion could be used for the whole circuit construction in Mod2-OBDD form, but we have restricted the $\oplus$ -nodes inclusion to a two-level layer of $\oplus$ -nodes and for a selected input variable for the pDE. This general structure has been selected because the number of $\oplus$ -nodes in the Mod2-OBDD impacts on its size. For many circuits, a few $\oplus$ -nodes lead to small sizes of Mod2-OBDDs, but too many $\oplus$ -nodes lead to large Mod2-OBDDs [28]. The introduction of a limited number of $\oplus$ -nodes in a mainly OBDD structure, tries to take advantage of both approaches (Mod2-OBDDs and OBDDs).

This two-level layer structure is constructed in the synthesis of the BLIF circuit. In BLIF, there are primary inputs, outputs and internal gates. A gate can have only primary inputs as fanins, or can have primary inputs and/or internal gates as inputs. A gate is represented with several lines, each line representing a cube, and the disjunction of these cubes provides the function of that gate. In general, the cubes are not disjoint. Each entry of a line, gives the value (0, 1 or $-$) that an input to that gate takes.

For the construction of the Mod2-OBDD with two-level layer $\oplus$ -nodes, we select a primary input x_i with reference to which the positive Davio decomposition is performed. We will use the same variable for the pDE in all gates of the circuit. For each gate, we will get two functions $F|_{k_i=0}$ and $F|_{k_i=1}$ that represent the negative and positive cofactors, respectively, with reference to the selected variable x_i . For any gate of the circuit, two cases can be distinguished: a) the gate only has primary inputs as fanins, and b) the gate has primary inputs and/or internal gates as inputs.

In case a), for each line (cube) of the gate we have to perform the conjunction of all the primary inputs different from x_i (we name the result of this conjunction as product). Then the value that the input x_i has for that cube is checked. The possible cases we can take for computing the negative and positive cofactors for each cube ($C|_{k_i=0}$ and $C|_{k_i=1}$) are the following:

- $x_i=0 \Rightarrow (C|_{k_i=0} = \text{product})$ and $(C|_{k_i=1} = 0)$.
- $x_i=1 \Rightarrow (C|_{k_i=0} = 0)$ and $(C|_{k_i=1} = \text{product})$
- $x_i = - \Rightarrow (C|_{k_i=0} = \text{product})$ and $(C|_{k_i=1} = \text{product})$

These operations are carried out for each line of the gate, and finally we compute the conjunction of the $C|_{k_i=0}$'s functions (getting $F|_{k_i=0}$) and the conjunction of

the $C|_{k_i=1}$'s functions (getting $F|_{k_i=1}$) obtained for all the lines. After that, we get the two functions $F|_{k_i=0}$ and $F|_{k_i=1}$ that represent the gate.

In case b), we have that some inputs to the gate (the internal gates) have been already decomposed with respect to the x_i variable. If we name the conjunction of the $F|_{k_i=0}$'s functions of these internal gates as $N|_{k_i=0}$, and the conjunction of the $F|_{k_i=1}$'s of the internal gates as $N|_{k_i=1}$, then we will have the following three possibilities in order to get the cofactors $C|_{k_i=0}$ and $C|_{k_i=1}$ for each cube:

- $x_i=0 \Rightarrow (C|_{k_i=0} = \text{product} \cdot N|_{k_i=0})$ and $(C|_{k_i=1} = 0)$.
- $x_i=1 \Rightarrow (C|_{k_i=0} = 0)$ and $(C|_{k_i=1} = \text{product} \cdot N|_{k_i=1})$.
- $x_i = - \Rightarrow (C|_{k_i=0} = \text{product} \cdot N|_{k_i=0})$ and $(C|_{k_i=1} = \text{product} \cdot N|_{k_i=1})$

As in case a), these operations are performed for each cube of the gate. Finally, the conjunction of all the $C|_{k_i=0}$'s functions obtained for all the lines gives $F|_{k_i=0}$, and the conjunction of all the $C|_{k_i=1}$'s functions obtained for all the lines gives $F|_{k_i=1}$, so $F|_{k_i=0}$ and $F|_{k_i=1}$ represent the gate.

The graph so created for internal gates has a structure similar to Figure 1(a), except that then and else edges of the lower $\oplus$ -node are not multiplied by the x_i variable. This is because all decomposition in the circuit is carried out with respect to this variable and internal gates are used only as an intermediate step in order to get the outputs of the circuit, so it is not necessary. We must perform these multiplications if the gate is an output, obtaining the structure given in Figure 1(a). The Mod2-OBDD obtained with this method, is not a canonical representation. As we are interested in the signature computation of the circuit using this data structure, the graph can be simplified not including x_i variable in the Mod2-OBDD. The contribution of the variable is considered multiplying the signature of the lower $\oplus$ -node in Figure 1(a) by the signature initially assigned to the variable x_i . By this way the size of the final Mod2-OBDD is even more reduced. The structure finally obtained is given in Figure 1(b), with an upper two-level layer of $\oplus$ -nodes and classical OBDDs (in fact, signature-OBDDs) at the lower part of the graph. The $\oplus$ -nodes used are introduced in the CUDD package by their specification in the decision-xor field of the nodes, and signatures are stored in their signature field. All arithmetic operations involved above must be performed over $GF(2^m)$ using the signature field of the nodes (decision and $\oplus$ -nodes).

For the experimental results given in the following

section, the variable x_i selected to perform the positive Davio decomposition has been the first variable given in the initial ordering for the benchmarks. It must be noted that several heuristics for variable ordering could be used, such as those based on topological or logic information of the circuit [31,33]. Therefore, the variable x_i selected to perform the Davio decomposition could be the first variable obtained in these new orderings.

6. Experimental Results

Experiments have been carried out in a Sun Ultra1 workstation with 128 Mbytes of main-memory for the verification of LGSynth91 Benchmark multilevel BLIF circuits, and the modified CUDD has been used for the construction of OBDDs, s-OBDDs and two-level layer Mod2-OBDDs. Arithmetic operations over the finite field $GF(2^{16})$ generated by the irreducible pentanomial $f(x)=x^{16}+x^5+x^3+x^2+1$ have been performed using a

polynomial basis for representation of the field elements, and the multiplication algorithm used has been the one given in [32].

Firstly, OBDDs and s-OBDDs for some of the given benchmarks have been constructed, and their construction times have been compared. In the 2nd column of Table 1, the ratios between the times needed to construct s-OBDDs and OBDDs of the benchmarks are given. The number of nodes (sizes) of both graphs is the same. The difference is that s-OBDD computes the signature of the circuit. From the results, it can be observed that the s-OBDD construction is 2.3 times slower in average than OBDD construction (with mux.blif as worst case). Despite of this, the construction of s-OBDDs has not an excessive cost in time in comparison with OBDDs, because we can get canonical representation in OBDD form together with the signature(s) of the circuit, so we could perform deterministic and/or probabilistic verifications. It is important to note that the more time needed for the s-OBDD construction is mainly due to the $GF(2^m)$ multiplication of signatures. Therefore, the use of more efficient multiplication algorithms [23] will reduce the time needed for the s-OBDD construction.

Secondly, we have constructed two-level layer Mod2-OBDDs for the benchmarks, as given in Section 5. The times needed for their construction have been compared with the times needed for s-OBDD construction, and their signatures have been compared in order to check the correctness of the results (probabilistic verification). In the 3rd column of Table 1, the ratios between the times

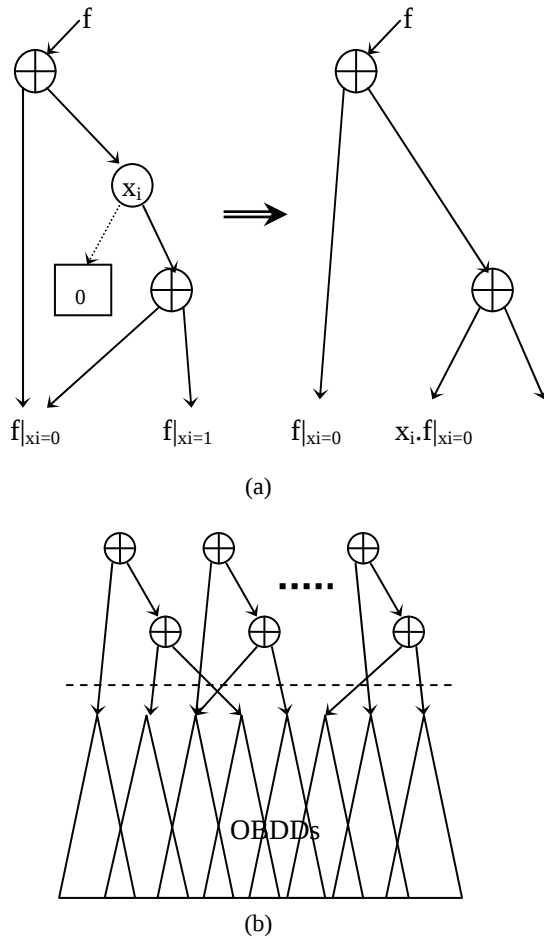


Figure 1. (a) Positive Davio expansion for variable x_i ; (b) Structure for multiple outputs.

Table 1. Comparison of time ratios for LGSynth91 multi-level benchmarks.

Circuit	s-OBDD/OBDD	Mod2-OBDD/s-OBDD
alu2	2.4	0.80
apex6	2.1	0.66
apex7	2.4	0.22
c1355	2.4	0.99
c1908	2.3	0.44
cm151a	2.5	0.05
cordic	1.5	0.77
count	1.5	0.90
des	2.9	0.96
example2	1.7	0.83
frg2	2.3	0.93
i2	2.3	0.70
k2	2.9	0.16
mux	3.1	0.94
pcler8	2.0	0.70
term1	2.5	0.73
too_large	3.0	0.61
ttt2	2.0	0.87
vda	2.9	0.22
x3	2.3	0.51
x4	1.8	0.84

Table 2. Sizes for OBDDs, two-level layer Mod2-OBDDs, and ratios between sizes of Mod2-OBDD and OBDDs.

Circuit	OBDD	Mod2-OBDD	Mod2-OBDD/OBDD
alu2	231	231	1.00
apex6	2760	1712	0.62
apex7	1660	566	0.34
c1355	45922	45953	1.00
c1908	36007	19697	0.55
cm151a	511	49	0.10
cordic	45	50	1.11
count	234	232	0.99
des	73919	74028	1.00
example2	469	553	1.18
frg2	6471	6216	0.96
i2	335	268	0.80
k2	28336	7351	0.26
mux	131071	131071	1.00
pcler8	139	146	1.05
term1	580	459	0.79
too_large	7096	5129	0.72
ttt2	223	228	1.02
vda	4345	1919	0.44
x3	2760	1712	0.62
x4	891	917	1.03
Total size	344005	298487	0.87

needed to construct two-level layer Mod2- OBDDs and s-OBDDs of the benchmarks are given. It can be observed that in all cases, Mod2-OBDD construction is faster than s-OBDD construction (0.66 times faster in average), even with hard benchmarks as c1908.blif (0.44 times faster).

The sizes (number of nodes) of OBDDs and two-level layer Mod2-OBDDs have been also compared. In the 2nd column of Table 2 the OBDD sizes for the benchmarks are given, in the 3rd column the two-level layer Mod2-OBDD sizes are also given, and in the 4th column the ratios between the Mod2-OBDD and OBDD sizes are represented. From the experimental results, it can be observed that Mod2-OBDD sizes are in average 0.87 times smaller than OBDD sizes. Significant examples are apex7, c1908, k2, vda and x3 blif files, which Mod2-OBDD sizes are 0.34, 0.55, 0.26, 0.44 and 0.62 times smaller than OBDD sizes, respectively.

We can also compare our method with other similar approaches given in the literature. The comparison of the sizes obtained by our method with the best results obtained in the work of Meinel and Sack [34] for some Benchmarks is given in Table 3. In [34], Mod2-OBDDs are constructed depending on a threshold factor, and dynamic variable reordering techniques are neither used. From this comparison, it can be observed that the Mod2-OBDD sizes obtained by Meinel and Sack (with threshold factor equal to 1.0) are 0.99 times smaller than OBDD sizes, while the two-level layer Mod2-OBDD

sizes obtained using our approach is 0.83 times smaller than OBDD sizes. In this table, ratio represents the ratio between the Mod2-OBDD and OBDD sizes.

Finally, from the experimental results, it can be observed that the probabilistic verification procedure using signature-OBDDs and two-level layer Mod2-OBDDs seems to be a promising approach for verification. The times needed for the construction of signature-OBDDs are not very large compared with the times needed for the construction of classical OBDDs, and the times needed for two-level layer Mod2-OBDD construction are always smaller than the times needed for s-OBDDs. Furthermore, the two-level layer Mod2-OBDDs sizes are in average smaller (in some cases, much smaller) than classical OBDDs sizes. The comparison of our experimental results with similar work done in the literature shows the suitability of our approach in order to obtain Mod2-OBDDs with reduced sizes. The achieved results could be further improved by trying heuristics for initial variable ordering, using dynamic variable reordering [34], [35], and using efficient multiplication algorithms over $GF(2^m)$, which is part of our ongoing work.

7. Conclusions

Probabilistic approach seems to be a promising alternative for circuit verification. For probabilistic methods

Table 3. Results of the comparison of our approach with the work done by Meinel and Sack for some benchmarks.

Circuit	OBDD	Meinel & Sack		Our approach	
		Mod2-OBDD	ratio	Mod2-OBDD	ratio
alu2	231	237	1.02	231	1.00
apex7	1660	1570	0.94	566	0.34
c1355	45922	45921	1.00	45953	1.00
c1908	36007	35869	0.99	19697	0.55
count	234	226	0.96	232	0.99
example2	469	456	0.97	553	1.18
frg2	6471	6348	0.98	6216	0.96
i2	335	334	1.00	268	0.80
k2	28336	26361	0.93	7351	0.26
mux	131071	131072	1.00	131071	1.00
term1	580	584	1.01	459	0.79
too_large	7096	7091	1.00	5129	0.72
vda	4345	4214	0.97	1919	0.44
x3	2760	2429	0.88	1712	0.62
Total size	265517	262712	0.99	221357	0.83

where Galois fields $GF(2^m)$ are used, Mod2-OBDDs are suitable data structures for signature computation due to the properties of finite fields with characteristic 2. In this work, a highly optimised package (CUDD) for OBDD construction has been modified to compute signatures in the synthesis process (signature-OBDD) and in order to construct two-level layer Mod2-OBDDs using positive Davio expansion with reference to a selected variable. Signatures obtained from signature-OBDDs and those obtained from Mod2-OBDDs are then compared for checking correctness (probabilistic verification). Experimental results have proven that the signature computation has not an excessive time cost and that Mod2-OBDDs with controlled number of $\oplus$ -nodes provide reduced sizes compared with classical OBDDs, so probabilistic verification can be a suitable alternative to classical verification methods. Comparisons with experimental results obtained by other similar approaches found in the literature have been also given, proving that our method is very suitable for the construction of reduced Mod2-OBDDs.

8. References

- [1] C. Y. Lee, "Representation of switching circuits by binary-decision programs," *Bell Systems Technology Journal*, Vol. 38, pp. 985–999, 1959.
- [2] S. B. Akers, "Binary decision diagrams," *IEEE Transactions on Computers*, Vol. C-27, pp. 509–516, 1978.
- [3] L. Fortune, J. Hopcroft, and E. M. Schmidt, "The complexity of equivalence and containment for free single variable program schemes," in: Goos, Hartmanis, Ausiello, Baum (Eds.), *Lecture Notes in Computer Science*, Springer-Verlag, New York, Vol. 62, pp. 227–240, 1978.
- [4] R. E. Bryant, "Graph based algorithms for Boolean function representation," *IEEE Transactions on Computers*, Vol. C-35, pp. 677–690, 1986.
- [5] R. Drechsler, B. Becker, and N. Göckel, "A genetic algorithm for variable ordering of OBDDs," *International Workshop on Logic Synthesis*, pp. P5c:5.55–5.64, 1995.
- [6] P. W. C. Prasad, A. Assi, A. Harb, and V. C. Prasad, "Binary decision diagrams: An improved variable ordering using graph representation of boolean functions," *International Journal of Computer Science*, Vol. 1, No. 1, pp. 1–7, 2006.
- [7] J. R. Burch and V. Singhal, "Tight integration of combinational verification methods," *IEEE/ACM International Conference on CAD*, pp. 570–576, 1998.
- [8] A. Kuehlmann and F. Krohm, "Equivalence checking using cuts and heaps," *Proceedings of Design Automation Conference*, pp. 263–268, 1997.
- [9] V. Paruthi and A. Kuehlmann, "Equivalence checking using cuts a structural SAT-solver, BDDs and simulation," *International Conference Computer Design*, 2000.
- [10] E. Goldberg, M. R. Parasad, and R. K. Brayton, "Using SAT for combinational equivalence checking," *IEEE/ACM Design, Automation and Test in Europe, Conference and Exhibition'01*, pp. 114–121, 2001.
- [11] J. Marques-Silva and T. Glass, "Combinational equivalence checking using satisfiability and recursive learning," *IEEE/ACM Design, Automation and Test in Europe*, pp. 145–149, 1999.
- [12] D. Brand, "Verification of large synthesized designs," *IEEE/ACM International Conference on Computer-Aided Design*, pp. 534–537, 1993.
- [13] W. Kunz, "HANNIBAL: An efficient tool for logic verification based on recursive learning," *IEEE/ACM International Conference on CAD*, pp. 538–543, November 1993.
- [14] R. E. Bryant, "Symbolic Boolean manipulation with ordered binary decision diagrams," *ACM Computing Surveys*, Vol. 24, No. 3, pp. 293–318, 1992.
- [15] M. Blum, A. K. Chandra, and M. N. Wegman, "Equivalence of free Boolean graphs can be decided probabilistically in polynomial time," *Information Processing Letters*, Vol. 10, No. 2, pp. 80–82, 1980.
- [16] J. Jain, J. Bitner, D. Fussell, and J. Abraham, "Probabilistic verification of Boolean functions," *Formal Methods in System Design*, Kluwer, Vol. 1, pp. 61–115, 1992.
- [17] R. E. Bryant and Y. Cheng, "Verification of arithmetic functions with binary moment diagrams," *Carnegie Mellon University Technical Report: CMU-CS-94-160*, May 1994.
- [18] R. Drechsler, B. Becker, and S. Ruppertz, "K*BMDs: A new data structure for verification," *IEEE European Design and Test Conference*, pp. 2–8, 1996.
- [19] U. Keschull, E. Schubert, and W. Rosentiel, "Multilevel logic based on functional decision diagrams," *European Design Automation Conference*, pp. 43–47, 1992.
- [20] Y. T. Lai and S. Sastry, "Edge-valued binary decision diagrams for multi-level hierarchical verification," *29th Design Automation Conference*, pp. 608–613, 1992.
- [21] J. Gergov and C. Meinel, "Mod2-OBDDs: A data structure that generalizes exor-sum-of-products and ordered binary decision diagrams," *Formal Methods in System Design*, Kluwer, Vol. 8, pp. 273–282, 1996.
- [22] S. Waack, "On the descriptive and algorithmic power of parity ordered binary decision diagrams," *Proceedings of 14th Symposium on Theoretical Aspects of Computer Science*, LNCS 1200, Springer, 1997.
- [23] J. L. Imaña, J. M. Sánchez, and F. Tirado, "Bit-parallel finite field multipliers for irreducible trinomials," *IEEE Transactions on Computers*, Vol. 55, No. 5, pp. 520–533, May 2006.
- [24] Ç. K. Koç and B. Sunar, "Low-complexity bit-parallel canonical and normal basis multipliers for a class of finite fields," *IEEE Transactions on Computers*, Vol. 47, No. 3, pp. 353–356, March 1998.

- [25] R. Lidl and H. Niederreiter, "Finite fields," Addison-Wesley, Reading, MA, 1983.
- [26] J. C. Madre and J. P. Billón, "Proving Circuit correctness using formal comparison between expected and extracted behaviour," Proceedings of 25th ACM/IEEE Design Automation Conference, pp. 308–313, 1988.
- [27] P. A. Scott, S. E. Tavares, and L. E. Peppard, "A fast VLSI multiplier for $GF(2^m)$," IEEE Journal on Selected Areas in Communications, Vol. 4, pp. 62–66, 1986.
- [28] C. Meinel and H. Sack, " $\oplus$ -OBDDs—A BDD structure for probabilistic verification," Pre-LICS Workshop on Probabilistic Methods in Verification (PROBMIV'98), Indianapolis, IN, USA, pp. 141–151, 1998.
- [29] F. Somenzi, "CUDD: CU decision diagram package," University of Colorado at Boulder. <http://vlsi.colorado.edu/~fabio/CUDD/>.
- [30] K. S. Brace, R. L. Rudell, and R. E. Bryant, "Efficient implementation of a BDD package," 27th ACM/IEEE Design Automation Conference, pp. 40–45, 1990.
- [31] K. M. Butler, D. E. Ross, R. Kapur, and M. R. Mercer, "Heuristics to compute variable orderings for efficient manipulation of ordered binary decision diagrams," Proceedings of 28th ACM/IEEE DAC, pp. 417–420, June 2001.
- [32] C. S. Yeh, I. S. Reed, and T. K. Truong, "Systolic multipliers for finite fields $GF(2^m)$," IEEE Transactions on Computers, Vol. 33, No. 4, pp. 357–360, April 1984.
- [33] M. Fujita, H. Fujisawa, and Y. Matsunaga, "Variable ordering algorithms for ordered binary decision diagrams and their evaluation," IEEE Transactions on CAD, Vol. 12, No. 1, pp. 6–12, January 1993.
- [34] C. Meinel and H. Sack, "Improving XOR-Node placement for $\oplus$ -OBDDs," 5th International Workshop on Applications of the Reed-Muller Expansion in Circuit Design, Starkville, Mississippi, USA, pp. 51–56, 2001.
- [35] C. Meinel and H. Sack, "Variable Reordering for $\oplus$ -OBDDs," International Symposium on Representations and Methodology of Future Computing Technologies (RM2003), Trier, Germany, pp. 135–144, 2003.

Theoretic and Numerical Study of a New Chaotic System

Congxu Zhu, Yuehua Liu, Ying Guo

School of Information Science and Engineering, Central South University, Changsha 410083, China
Email: zhucongxu@126.com

Abstract

This paper introduced a new three-dimensional continuous quadratic autonomous chaotic system, modified from the Lorenz system, in which each equation contains a single quadratic cross-product term, which is different from the Lorenz system and other existing systems. Basic properties of the new system are analyzed by means of Lyapunov exponent spectrum, Poincaré mapping, fractal dimension, power spectrum and chaotic behaviors. Furthermore, the forming mechanism of its compound structure obtained by merging together two simple attractors after performing one mirror operation has been investigated by detailed numerical as well as theoretical analysis. Analysis results show that this system has complex dynamics with some interesting characteristics.

Keywords: Chaotic System, Lyapunov Exponent, Poincaré Mapping, Fractal Dimension, Power Spectrum

1. Introduction

Chaos is found to be useful or has great potential application in many disciplines. Recently, it has been noticed that purposefully creating chaos can be a key issue in many technological applications such as communication, encryption, information storage, etc. Since Lorenz found the first chaotic attractor in a three-dimensional (3-D) autonomous system in 1963 [1], Lorenz system has become a paradigm for chaos research, and many new Lorenz-like chaos system have been constructed [2–6]. In finding of a new system, one can construct and determine the system parameter values such that the system can become chaotic following some basic ideas of chaoticification [7], namely:

- 1) It is dissipative.
- 2) There exist some unstable equilibria, especially some saddle points.
- 3) There are some cross-product terms, so there are dynamical influence between different variables.
- 4) The system orbits are all bounded.

This motivates the present study on the problem of generating new chaotic attractors. Under these guidelines, though not sufficient, a new chaotic system is generated by modifying from the Lorenz system. It turns out that the new system has five equilibria, therefore is not topologically equivalent to the Lorenz system. This paper has also briefly studied and analyzed its forming mechanism. The compound structure of the attractor obtained by

merging together two simple attractors after performing one mirror operation is explored here. Simulation results support brief theoretical derivations. The proposed approach in finding a new chaotic system has advantages of intuitiveness, simpleness, and convenience over some existing methods. But it has a shortcoming of uncertainty.

In the rest of the paper, we use the lowercase letters x, y, z denote the state variables of the new 3D chaotic system, the uppercase letter X denotes the vector of state variables, F denotes the name of function, J denotes the Jacobian matrix, λ denotes the eigenvalue of a matrix, L denotes the Lyapunov exponents of the system.

2. The New Chaotic System and Its Properties

The new chaotic system introduced in this paper is described as the following autonomy differential equations:

$$\begin{cases} \dot{x} = -x - ay + yz, \\ \dot{y} = by - xz, \\ \dot{z} = -cz + xy. \end{cases} \quad (1)$$

Here, a, b, c are constant parameters of the system. When $a = 1.5; b = 2.5; c = 4.9$, system (1) is chaotic and belongs Lorenz system family. In system (1) each equation contains a single quadratic cross-product term, and the linear terms in the first and second equations are dif-

ferent from the Lorenz system, and therefore is a new Lorenz-like attractor, but not equivalent chaotic attractor in the topological structure [1].

2.1. Symmetry and Dissipativity

System (1) has a natural symmetry under the coordinates transform $(x, y, z) \rightarrow (-x, -y, z)$, which persists for all values of the system parameters. So, system (1) has symmetry about the z -axis.

system (1) can be expressed as the form of vectors

$$\dot{X} = F(X) = [f_1(X), f_2(X), f_3(X)]^T. \quad (2)$$

where $X = [x(t), y(t), z(t)]^T$. The divergence of the vector field $F(X)$ on $\mathbf{R}^3$ is given by

$$\nabla F = \frac{\partial f_1(X)}{\partial x} + \frac{\partial f_2(X)}{\partial y} + \frac{\partial f_3(X)}{\partial z}. \quad (3)$$

The divergence of F measures how fast volumes change under the flow Φ_t of F . Suppose D is a region in $\mathbf{R}^3$ with a smooth boundary, and let $D(t) = \Phi_t(D)$, the image of D under the time t map of the flow. Let $V(t)$ be the volume of $D(t)$. Then Liouville's theorem asserts that

$$\frac{dV(t)}{dt} = \int_{D(t)} (\nabla F) dx dy dz \quad (4)$$

For the system (1), we compute immediately that the divergence is the negative constant $-(c+1-b) = -3.4$, so that volume decreases at a constant rate

$$\frac{dV(t)}{dt} = -(c+1-b) \int_{D(t)} dx dy dz = -(c+1-b)V \quad (5)$$

Solving this simple differential Equation (5) yields

$$V(t) = V(0)e^{-(c+1-b)t} = V(0)e^{-3.4t} \quad (6)$$

So that any volume must shrink exponentially fast to 0, and dynamical system described by (1) is a dissipative system. Hence, all orbits of system (1) are eventually confined to a specific subset that have zero volume, the asymptotic motion settles onto an attractor of the system (1) [8].

2.2. System Equilibria

In this section, assume that $b \neq 0$, $bc > 0$, and $a^2 + 4b > 0$. The equilibria of system (1) are found by solving the three equations $\dot{x} = \dot{y} = \dot{z} = 0$. It is found that system (1) has five equilibria, which are respectively described as follows:

$O(0,0,0)$; $E_1(x_0, y_+, z_+)$; $E_2(x_0, y_-, z_-)$; $E_3(-x_0, -y_-, z_-)$; and $E_4(-x_0, -y_+, z_+)$, where $x_0 = \sqrt{bc}$,

$$y_+ = \sqrt{bc}(a + \sqrt{a^2 + 4b})/(2b),$$

$$y_- = \sqrt{bc}(a - \sqrt{a^2 + 4b})/(2b), \quad z_+ = (a + \sqrt{a^2 + 4b})/2,$$

and $z_- = (a - \sqrt{a^2 + 4b})/2$. When $a = 1.5$; $b = 2.5$; $c = 4.9$, the five equilibria are $O(0,0,0)$; $E_1(3.5, 3.5, 2.5)$; $E_2(3.5, -1.4, -1)$; $E_3(-3.5, 1.4, -1)$; and $E_4(-3.5, -3.5, 2.5)$.

The Jacobian matrix J of system (1) is

$$J = \begin{bmatrix} \frac{\partial f_1(X)}{\partial x} & \frac{\partial f_1(X)}{\partial y} & \frac{\partial f_1(X)}{\partial z} \\ \frac{\partial f_2(X)}{\partial x} & \frac{\partial f_2(X)}{\partial y} & \frac{\partial f_2(X)}{\partial z} \\ \frac{\partial f_3(X)}{\partial x} & \frac{\partial f_3(X)}{\partial y} & \frac{\partial f_3(X)}{\partial z} \end{bmatrix} = \begin{bmatrix} -1 & z-a & y \\ -z & b & -x \\ y & x & -c \end{bmatrix} \quad (7)$$

For equilibrium point $O(0, 0, 0)$, the Jacobian matrix has the following result:

$$J_0 = \begin{bmatrix} -1 & -1.5 & 0 \\ 0 & 2.5 & 0 \\ 0 & 0 & -4.9 \end{bmatrix} \quad (8)$$

To gain its eigenvalues, let $[\lambda I - J_0] = 0$. Then the eigenvalues that corresponding to equilibrium $O(0, 0, 0)$ are respectively obtained as follows: $\lambda_1 = -4.9$, $\lambda_2 = 2.5$, $\lambda_3 = -1$. Here λ_2 is a positive real number, λ_1 and λ_3 are two negative real numbers. Therefore, the equilibrium $O(0, 0, 0)$ is a saddle point. So, this equilibrium point $O(0, 0, 0)$ is unstable.

In the same way, the eigenvalues that corresponding to equilibrium point E_1 are obtained as: $\lambda_1 = -6.5337$, $\lambda_2 = 1.5 + 3.2664i$, $\lambda_3 = 1.5 - 3.2664i$, where i denote the unit of imaginary number. The eigenvalues that corresponding to equilibrium point E_2 are obtained as: $\lambda_1 = -4.4632$, $\lambda_2 = 0.5316 + 2.7208i$, $\lambda_3 = 0.5316 - 2.7208i$. The eigenvalues that corresponding to equilibrium point E_3 and E_2 are the same, and the eigenvalues that corresponding to equilibrium point E_4 and E_1 are also the same.

For each of the four equilibrium points E_1 , E_2 , E_3 and E_4 , the results show that λ_1 is a negative real number, λ_2 and λ_3 become a pair of complex conjugate eigenvalues with positive real parts. Therefore, equilibrium points E_1 , E_2 , E_3 and E_4 are all saddle-focus points; so, these equilibrium points are all unstable.

2.3. Observation of Chaotic and Complex Dynamics

The initial values of the system are selected as $(-0.5, 2.0, 3.5)$. Using MATLAB program, the numerical simulation have been completed. This nonlinear system exhibits the complex and abundant chaotic dynamics behaviors, the strange attractors are shown in Figures 1. Apparently, the strange attractors in this nonlinear-system are different

from Lorenz chaos attractor.

The waveforms of $x(t)$ in time domain are shown in Figure 2. The waveforms of $x(t)$ is aperiodic.

In order to discriminate between a multiple periodic

motion that can show also a complicated behavior and a chaotic motion, the power spectrum of nonlinear system (1) is also studied, and it is continuous as shown in Figure 3.

The Poincaré section for chaotic attractors shows a certa-

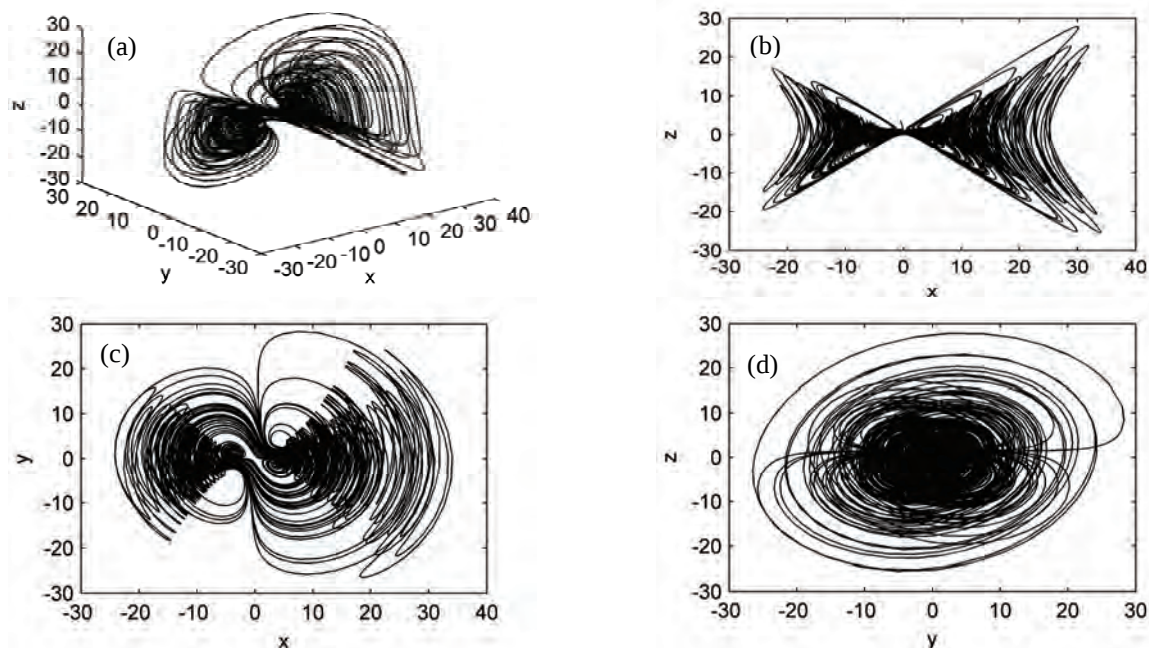


Figure 1. Phase plane strange attractors. (a) Three-dimensional view; (b) x - z phase plane; (c) x - y phase plane; (d) y - z phase plane

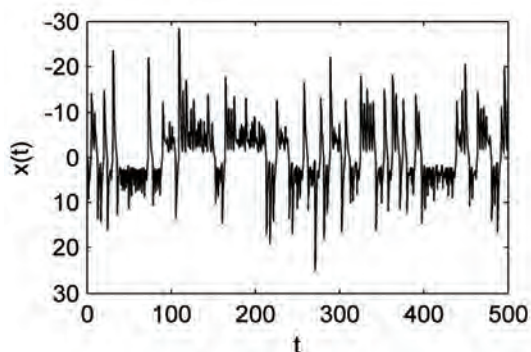


Figure 2. $x(t)$ waveform of system (1)

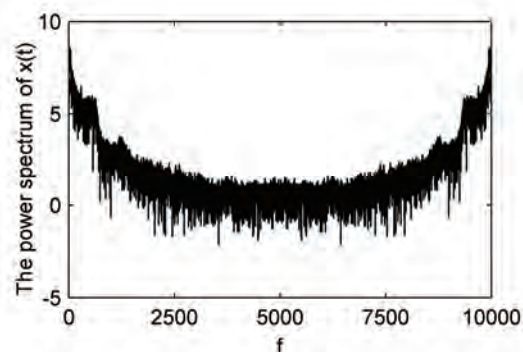


Figure 3. Power spectrum of $x(t)$

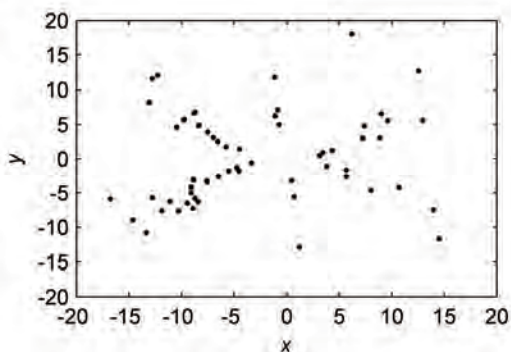


Figure 4. The Poincaré map of x - y plane

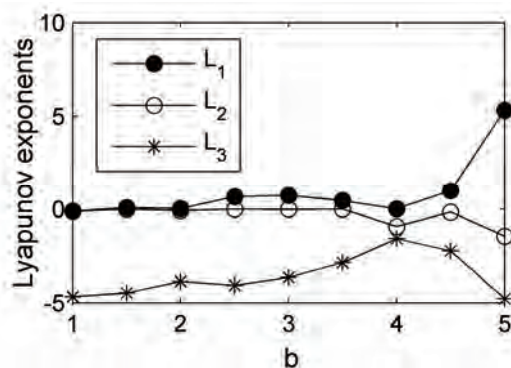


Figure 5. Spectrum of Lyapunov exponents

in type of organization but with a complex geometry. The section can be composed by an infinity of points (in contrast with the finite number of points for the periodic and quasi-periodic case) irregularly scattered, in one line or many arcs of curves. Poincaré mapping of system (1) is shown in Figure 4, its structure becomes better and better defined in time after the accumulation of points.

As is well known, the Lyapunov exponents measure the exponential rates of divergence or convergence of nearby trajectories in phase space, according to the detailed numerical as well as theoretical analysis, the largest value of positive Lyapunov exponents of this chaotic system is obtained as $L_1=0.6747$. It is related to the expanding nature of different direction in phase space.

Another one Lyapunov exponent is $L_2 = 0$. It is related to the critical nature between the expanding and the contracting nature of different direction in phase space.

While negative Lyapunov exponent is $L_3 = -4.0738$. It is related to the contracting nature of different direction in phase space.

Fix parameters $a = 1.5$, $c = 4.9$, and let b vary. The spectrum of Lyapunov exponents of system (1) versus b are computed, and the results are shown in Figure 5.

The Lyapunov dimension of chaos attractors of this nonlinear system is of fraction dimension, it is described as

$$D_L = j + \frac{1}{|L_{j+1}|} \sum_{i=1}^j L_i = 2.166 \quad (9)$$

The fractal nature of an attractor does not merely imply non-periodic orbits; it also causes nearby trajectories to diverge. As all strange attractors, orbits that are initiated from different initial conditions soon reach the attracting set, but two nearby orbits do not stay close to each other [9]. They soon diverge and follow totally different paths in the attractor. Therefore, there is really chaos in this nonlinear system.

3. Forming Mechanism of This New Chaotic Attractor Structure

In order to reveal the forming mechanism of this new chaotic attractor structure, its controlled system is proposed, the autonomous differential equations of its controlled system are expressed as

$$\begin{aligned} \dot{x} &= -x - 1.5y + yz, \dot{y} = 2.5y - \\ & xz + k, \dot{z} = -4.9z + xy. \end{aligned} \quad (10)$$

In this system, k is the parameter of control, the value of it can be changed within a certain range.

When the parameter k is changed, the chaos behavior of this system can effectively be controlled. So it is a controller. Here, the initial values of the system are selected as $(-0.5, 0, 0.5)$.

Let $k = 1.9$, the corresponding strange attractors are shown in Figure 6(a), the attractor evolves into partial but is still bounded in this time.

Let $k = 2.1$, the corresponding strange attractors are shown in Figure 6(b). Moreover the strange attractors are evolved into single left scroll attractor; it is only one half the original chaotic attractors in this time.

Let $k = 2.5$, the strange attractor evolves into the period-doubling bifurcations; period-doubling bifurcations are shown in Figure 6(c).

While k is a negative value, the chaos of this system can also be affected. The initial values of the system are still selected as $(-0.5, 0, 0.5)$.

Let $k = -1.9$, the strange attractors are shown in Figure 6(d), the attractor evolves also into partial but is still bounded in this time.

Let $k = -2.1$, the corresponding strange attractors are shown in Figure 6(e). Moreover the strange attractors are evolved into single right scroll attractor; it is also only one half the original chaotic attractors in this time.

Let $k = -2.5$, the strange attractor evolves into the period-doubling bifurcations; period-doubling bifurcations are shown in Figure 6(f).

In the controller, one can see when $|k|$ is large enough, chaos attractor disappears; when $|k|$ is small enough, a complete chaos attractor appears. So $|k|$ is an important parameter to control chaos in the nonlinear-system [10].

This means the attractor is a compound structure obtained by merging together two simple attractor after performing one mirror operation [11].

4. Conclusions

This paper has reported and analyzed a new three-dimensional continuous autonomous chaotic system, in which each equation has a single cross-product term. Basic properties of the system have been analyzed by means of Lyapunov exponents, Poincaré mapping, fractal dimension, power spectrum and chaotic behaviors. This new attractors proposed can be also realized with an electronic circuit and have great potential for communication and electronics [12]. Apparently there are more interesting problems about this new system in terms of dynamics, complexity, control and synchronization, among others, leaving rooms for further studies.

5. Acknowledgments

This work was supported by the National Natural Science Foundation of China under Grant 60902044 and the First Batch Fostering Foundation in Training Graduate Students of Central South University.

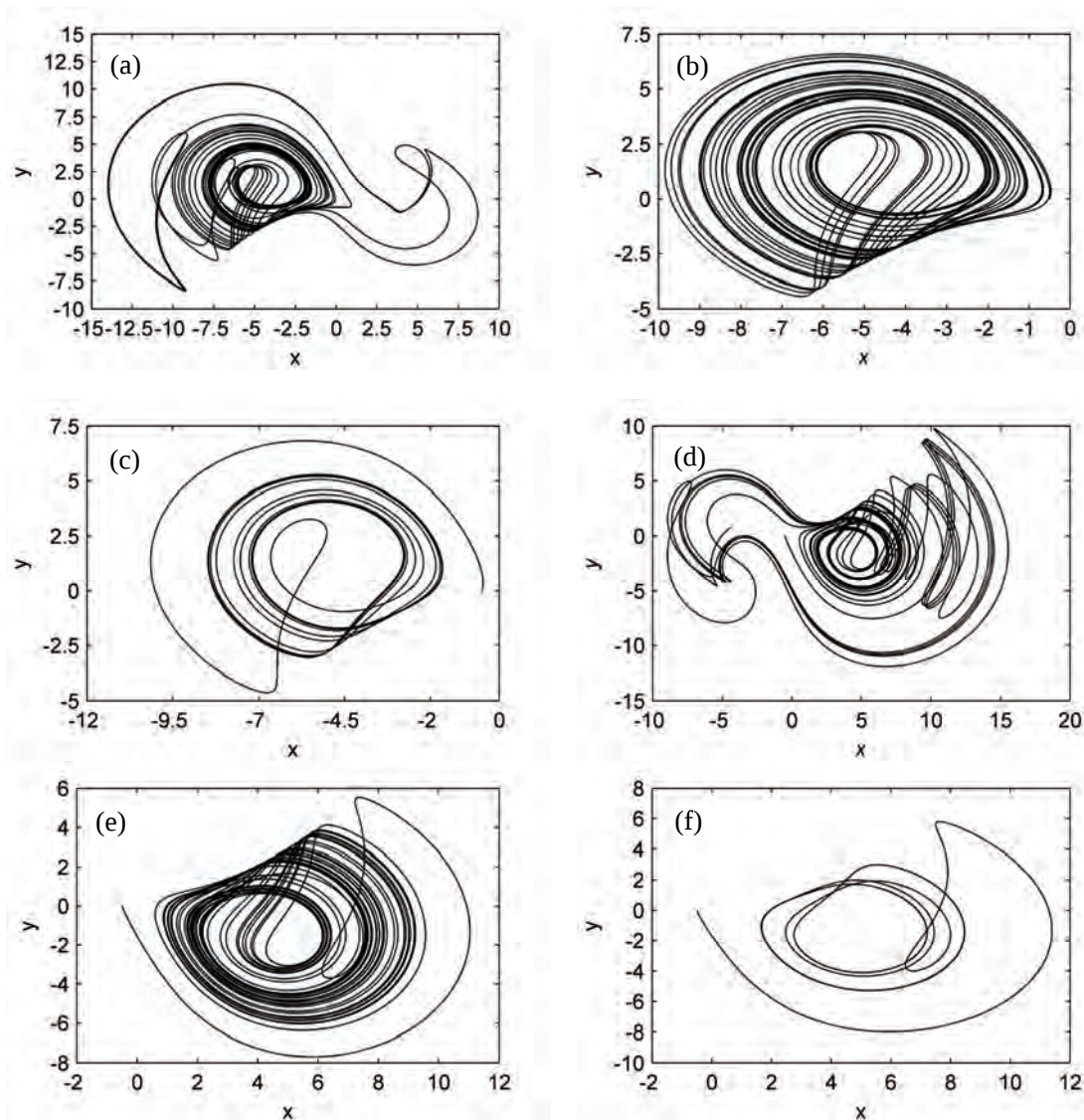


Figure 6. Dynamics behavior of the system changes with the control parameter k . a) x - y phase plane strange attractors ($k=1.9$); b) x - y phase plane strange attractors ($k=2.1$); c) x - y phase plane period-doubling bifurcations ($k=2.5$); d) x - y phase plane strange attractors ($k=-1.9$); e) x - y phase plane strange attractors ($k=-2.1$); f) x - y phase plane period-doubling bifurcations ($k=-2.5$).

6. References

- [1] C. Sparrow, "The Lorenz equations: Bifurcations chaos and strange attractors," Springer, New York, 1982.
- [2] G. R. Chen and T. Ueta, "Yet another chaotic attractor," International Journal of Bifurcation and Chaos, Vol. 9, No. 7, pp. 1465–1466, 1999.
- [3] J. H. Lü and G. R. Chen, "A new chaotic attractor coined," International journal of bifurcation and chaos, Vol. 12, No. 3, pp. 659–661, 2002.
- [4] J. H. Lü, G. R. Chen, D. Cheng, and S. Celikovsky, "Bridge the gap between the Lorenz system and the Chen system," International Journal of Bifurcation and Chaos, Vol. 12, No. 12, pp. 2917–2926, 2002.
- [5] G. Qi, G. R. Chen, and S. Du, "Analysis of a new chaotic system," Physica A, Vol. 352, No. 2–4, pp. 295–308, 2005.
- [6] C. X. Liu, L. Liu, and T. Liu, "A new butterfly-shaped attractor of Lorenz-like system," Chaos, Solitons & Fractals, Vol. 28, No. 5, pp. 1196–1203, 2006.
- [7] X. Wang, "Chaos control," Springer, New York, 2003.
- [8] J. H. Lü, G. R. Chen, and S. Zhang, "Dynamical analysis of a new chaotic attractor," International Journal of Bifurcation and Chaos, Vol. 12, No. 5, pp. 1001–1015, 2002.

- [9] A. Wolf, J. B. Swift, H. L. Swinney, and J. A. Vastano, "Determining Lyapunov exponents from a time series," *Physica D*, Vol. 16, No. 3, pp. 285–317, 1985.
- [10] C. X. Liu, L. Liu, and K. Liu, "A new chaotic attractor," *Chaos, Solitons & Fractals*, Vol. 22, No. 5, pp. 1031–1038, 2004.
- [11] J. H. Lü, G. R. Chen, and S. Zhang, "The compound structure of a new chaotic attractor," *Chaos, Solitons & Fractals*, Vol. 14, No. 5, pp. 669–672, 2002.
- [12] G. Q. Zhong and W. K. S. Tang, "Circuitry implementation and synchronization of Chen's attractor," *International Journal of Bifurcation and Chaos*, Vol. 12, No. 6, pp. 1423–1427, 2002.

Sensing Semantics of RSS Feeds by Fuzzy Matchmaking

Mingwei Yuan¹, Ping Jiang^{1,2}, Jin Zhu¹, Xiaonian Wang¹

¹*Department of Information and Control Engineering, Tongji University, Shanghai, China*

²*Department of Computing, University of Bradford, Bradford, UK*

Email: yuan_mingwei@yahoo.com.cn, p.jiang@bradford.ac.uk, {zhujintj, dawnyear}@tongji.edu.cn

Abstract

RSS feeds provide a fast and effective way to publish up-to-date information or renew outdated contents for information subscribers. So far RSS information is mostly managed by content publishers but Internet users have less initiative to choose what they really need. More attention needs to be paid on techniques for user-initiative information discovery from RSS feeds. In this paper, a quantitative semantic matchmaking method for the RSS based applications is proposed. Semantic information is extracted from an RSS feed as numerical vectors and semantic matching can then be conducted quantitatively. Ontology is applied to provide a common-agreed matching basis for the quantitative matchmaking. In order to avoid semantic ambiguity of literal statements from distributed and heterogeneous RSS publishers, fuzzy inference is used to transform an individual-dependent vector into an individual-independent vector. Semantic similarities can be revealed as the result.

Keywords: RSS Feeds, Matchmaking, Multi-Agent, Semantics

1. Introduction

Internet is a complex environment with dynamically changing contents and large-scale distributed users. An incessant research topic for the web based applications is how to acquire information more efficiently and effectively from the Internet. Nowadays there are mainly two approaches. One is user-active that a user visits websites to find interests manually. In order to improve its searching efficiency, favourite websites could be bookmarked for later usage. However, manual search or research could be a tedious process for information acquisition. An alternative approach is publisher-active that a user subscribes relevant websites and waits for updates from publishers. It is obvious that the latter mode is more convenient and instant in terms of variant interests and effortless information retrieval. RSS (RDF Site Summary or Really Simple Syndication) feeds are such sources to support automatic information acquisition from the Internet. A user can select which RSS feeds to monitor and avoid unnecessary visit.

RSS is a term used by two independent camps for web content publishing. Therefore, there are two translations about what the acronym RSS stands for: RDF (Resource Description Framework) Site Summary or Really Simple

Syndication. Currently the latest versions of the two formats are RSS1.0 (<http://web.resource.org/rss/1.0/spec>) and RSS2.0 (<http://blogs.law.harvard.edu/tech/rss>), respectively. The two formats are both XML-based (Extensible Markup Language) and provide similar functions for information updating. RSS1.0 is RDF-compliant so it has more flexible and extendable features than RSS2.0, but RSS2.0 is simpler and more widely used today. In general, the RSS is a metadata language for describing web content changes. Nowadays RSS is adopted by almost all mainstream websites for web content publishing, e.g., web site modifications, news, wiki, and blog updates. But research shows that awareness of RSS is still quite low in the Internet users, 12% of users are aware of RSS, and only 4% have knowingly used RSS [1]. Therefore, it is necessary to develop the RSS based applications to be more conveniently accessible by the ordinary Internet users.

Although a RSS feed is a standardized format with some simple semantics, such as authorship, published date and summary, filtering the received RSS documents using such simple attributes is not able to reflect a user's complex intention. For efficient and effective information acquisition, a user may want to further narrow his/her focus and ignore irrelevant information. This requires a new RSS reader with the features of:

1) semantic awareness

In a distributed web environment, the published information and knowledge can be very complex and diverse. Information acquisition requires a semantics oriented approach that knowledge is represented in a hierarchical data structure. Traditionally agent based information matchmaking often flattens the structure of knowledge into a *free text vector* with simple valued attributes e.g. with keywords, price, delivery time. Semantic relations of knowledge could be lost [2–4].

2) fuzzy sensing

Information acquisition from RSS feeds requires more capability to deal with uncertainties because there is no prior agreement on how information is represented by heterogeneous publishers [4]. Logic based approaches have been widely used to support rule-based matchmaking or consistency checking by proving subsumption and (un) satisfiability [5–6]. However, distributed web applications usually cannot retain a closed-world knowledge base for logic based inference. It would be more realistic to recognise the degree of similarity for flexible matches [7]. It requires more intelligence but less precision, *i.e.* fuzzy sensing.

Today RSS research rests mostly on effective ways to aggregate/syndicate content [8,9] and to improve its applicability [10]. For information acquisition from RSS feeds, current studies often adopt classical text mining methods. In paper [11], fuzzy concept based [12] and word sequence kernels based [13–14] text classifications were applied to measure the similarity of RSS-formatted documents. They intended to reveal similarity by direct comparison of literal texts but ignoring the inherent correlation of individual words, *i.e.* semantic similarity. Due to the autonomy of heterogeneous RSS publishers in an open environment it is impossible to force them to use strictly consistent terminologies and sentences. This introduces ambiguities into text mining of RSS feeds and makes the keyword based vector space model [15] difficult for RSS based applications. Similar text mining approach was used for learning of user preference without consideration of text semantics [16]. Statistical feature selection methods in text mining were also evaluated for classification of RSS feeds corpus [17] and the authors pointed out that topic detection [18] and automatic text classification methods [19] were important to RSS based applications.

It is undoubted that such classical textual analysis methods can be applied to RSS documents since they are formatted textual documents. It should not be ignored that RSS has a close relationship with the semantic web; especially RSS1.0 builds on the RDF. Hence the performance of RSS document classification can be improved by the semantic web technology [20], which takes into account correlations among concepts. Considering semantics, paper [21] developed a weighted schema graph for semantic search of RSS feeds. However, the weights need to be assigned manually. In fact, ontology defines the concepts and correlations in a domain. It provides a powerful tool for semantics based classification,

taking into account meaning behind words. The architecture of Personalized News Service (PNS) [22] was proposed, consisting of RSS News Feed Consolidation Service, Ontology Reasoner and Personalized News Service. A user can query interests from RSS feeds based on ontology reasoning. A logic based approach was proposed for implementation.

There have been increasing research interests in recent years addressing the semantic search in distributed applications, especially in peer-to-peer environments e.g. Bluetooth service discovery [22], grid computing [23] and the electronic marketplace [24]. Although description logic can be used for similarity ranking by counting missing/not-implied concept names and loose characteristics between documents [24], the distinguishable granularity is usually coarse and the ability to handle fuzziness and uncertainty is limited. Fuzzy logic has been extended to description logic for representing fuzzy knowledge using continuous membership, such as f-Shin [25] and rule-based f-SWRL [26]. However reasoning for fuzzy description logic still relies on a consistent fuzzy knowledge basis. Ontology can become the knowledge basis for fuzzy reasoning [27].

This paper proposes a quantitative method for information acquisition from RSS feeds with the aid of the semantic web technique. It is an intelligent agent to detect interests for a user. Ontology is used as a semantic bridge linking RSS feeds with a user's intention. Fuzzy matchmaking is carried out for ranking RSS feeds. The method proposed in this paper is for general formats of RSS feeds, and hence RSS 1.0, RSS 2.0 or any other RSS-like formats (e.g. Atom) can be applied. It acts as a real-time sensor of RSS feeds in the Internet with the capability of semantic awareness and fuzzy sensing. First, it transfers received RSS feeds into numerical vectors, *feature vectors*, underpinned by an ontology. Semantic matching of information is conducted by correlation computation. Because distributed publishers may express their opinions using different jargons or words, the obtained numerical vectors are usually individual-dependent. To solve the inherent semantic ambiguity of RSS feeds, fuzzy inference is introduced to transform an individual-dependent vector into an individual-independent vector, so that semantic matchmaking of RSS feeds is accomplished.

This paper is organized as follows, section two proposes the algorithms to extract semantic distance from a domain ontology; section three discusses the method to formulate RSS feeds into ontology instance for facilitating semantics based matchmaking; a job finding agent is developed using the proposed approach in section four. Section five summarizes the proposed method.

2. Semantic Distance between Concepts in Ontology

Information providers publish their information in RSS.

A web user subscribing favorite RSS feeds is often interested in some specific topics. Therefore, irrelevant RSS items should be filtered out. A virtual sensor can be developed for this purpose, which connects with RSS channels and monitors incoming RSS items for those close to the topics semantically.

An illustrative RSS feed from a job publishing site is shown below:

```
<rss version="2.0">
<channel>
<title> Yahoo! HotJobs:DVR</title>
<link>http://hotjobs.yahoo.com/jobs/USA/All/All-jobs</link>
<description>Top HotJobs results for jobs matching:
DVR</description>
<webMaster>webmaster-rss@hotjobs.com</webMaster>
<language>en-us</language>
...
<item>
<title>Java Developers - Beta Soft Systems - Fremont, CA USA</title>
<link>http://pa.yahoo.com/*http://us.rd.yahoo.com/hot-jobs/rss/evt=23685/*http://hotjobs.yahoo.com/jobseeker/jobsearch/job_detail.html?job_id=J987065YO</link>
<description> ... to work with our top clients in USA belonging to any industry ... .BS/MS/MBA degree/Eng. (CS, MIS,CIS,IS,IT,CS... </description>
</item>
<item>
<title>Software Engineer - MPEG, Video, Compression - Sigma Designs, Inc. - Milpitas, CA USA</title>
<link>http://pa.yahoo.com/*http://us.rd.yahoo.com/hot-jobs/rss/evt=23685/*http://hotjobs.yahoo.com/jobseeker/jobsearch/job_detail.html?job_id=J497490PV</link>
<description> ... </description>
</item>
</channel>
</rss>
```

From the example, an RSS feed is composed of a channel and a series of items. Within an item tag pair, a summary of an article or a story is presented. A virtual sensor needs to detect relevant items according to subscriber's interests. In order to simplify the presentation, only the title of an item is taken into account in this paper, which is a condensed abstract of the content and is the foremost factor influencing information selection. However, the method is applicable to other tags, such as the description tags.

Selecting relevant RSS feeds relies on semantic matchmaking rather than textual matchmaking. A textual or word-to-word comparison makes little sense because distributed and heterogeneous RSS publishers may have totally different writing styles. In fact, words and concepts used in RSS feeds have certain correlations, which can be described by ontology as domain knowledge. The

domain knowledge is generally defined as a meta-data model by domain consortia or standard bodies, for example, STEP(Standard for the Exchange of Product model data) AP203[28] and DIECoM (Distributed Integrated Environment for Configuration Management) meta-data model[29] in the domain of product manufacturing. The domain knowledge can then be formulated as a domain ontology in XML using semantic web technologies, e.g. OWL (Web Ontology Language). RSS feeds from distributed sources can be understandable and interchangeable by software agents if the domain ontology is provided.

Suppose domain ontology is defined as

$$\Omega \equiv R(e_1, e_2, \dots, e_N) \quad (1)$$

where e_i ($i = 1, \dots, N$) denotes entities (concepts, terminologies, properties, attributes) used in a domain; R is the set of relationships between the entities and can be represented as a graph. For example, a domain ontology for DVR (Digital Video Recorder) development is shown in Figure 1, which is edited using Protégé editor (<http://protege.stanford.edu/plugins/owl/index.html>) and expressed in OWL(<http://www.w3.org/TR/2004/REC-owl-ref-20040210/>). A DVR is a consumer video/ audio product that can record and play video/audio encoded by using various video/ audio compression standards. The storage media includes hard disks and recordable CD/DVD disks.

Domain ontology provides a semantic bridge for mutual understanding between publishers and subscribers. Concepts defined in ontology may be semantically relevant. Hence *semantic distance* is introduced as a measure of semantic difference or similarity between two concepts, which has been applied in semantic web matchmaking [30] and conceptual clustering of database schema [31]. In general, a semantic distance is defined as an application of $E \times E$ into R^+ , where E is a set of entities in a domain ontology Ω with the following properties:

- 1) $\forall x \in E, \forall y \in E, \quad \text{Dis}(x, y) = 0 \Leftrightarrow x = y$
- 2) $\forall x \in E, \forall y \in E, \quad \text{Dis}(x, y) = \text{Dis}(y, x)$
- 3) $\forall x \in E, \forall y \in E, \forall z \in E \quad \text{Dis}(x, y) \leq \text{Dis}(x, z) + \text{Dis}(z, y)$

From Figure 1, it can be observed that an ontology exhibits a hierarchical structure. According to the *visual distance* [31], the semantic distance between two concepts can be calculated as the shortest path length, in which the unit length is assumed to be 1 if two nodes have a direct link. A concept distance matrix to represent semantic differences between two concepts can be obtained by processing the ontology.

First a concept vector is defined, which consists of all entities in the ontology.

$$V(\Omega) = [e_1, e_2, \dots, e_N]^T \quad (2)$$

To simplify computation, the elements in a concept vector are arranged in order by taking a "breadth-first" scan of the ontology hierarchy. A higher level concept

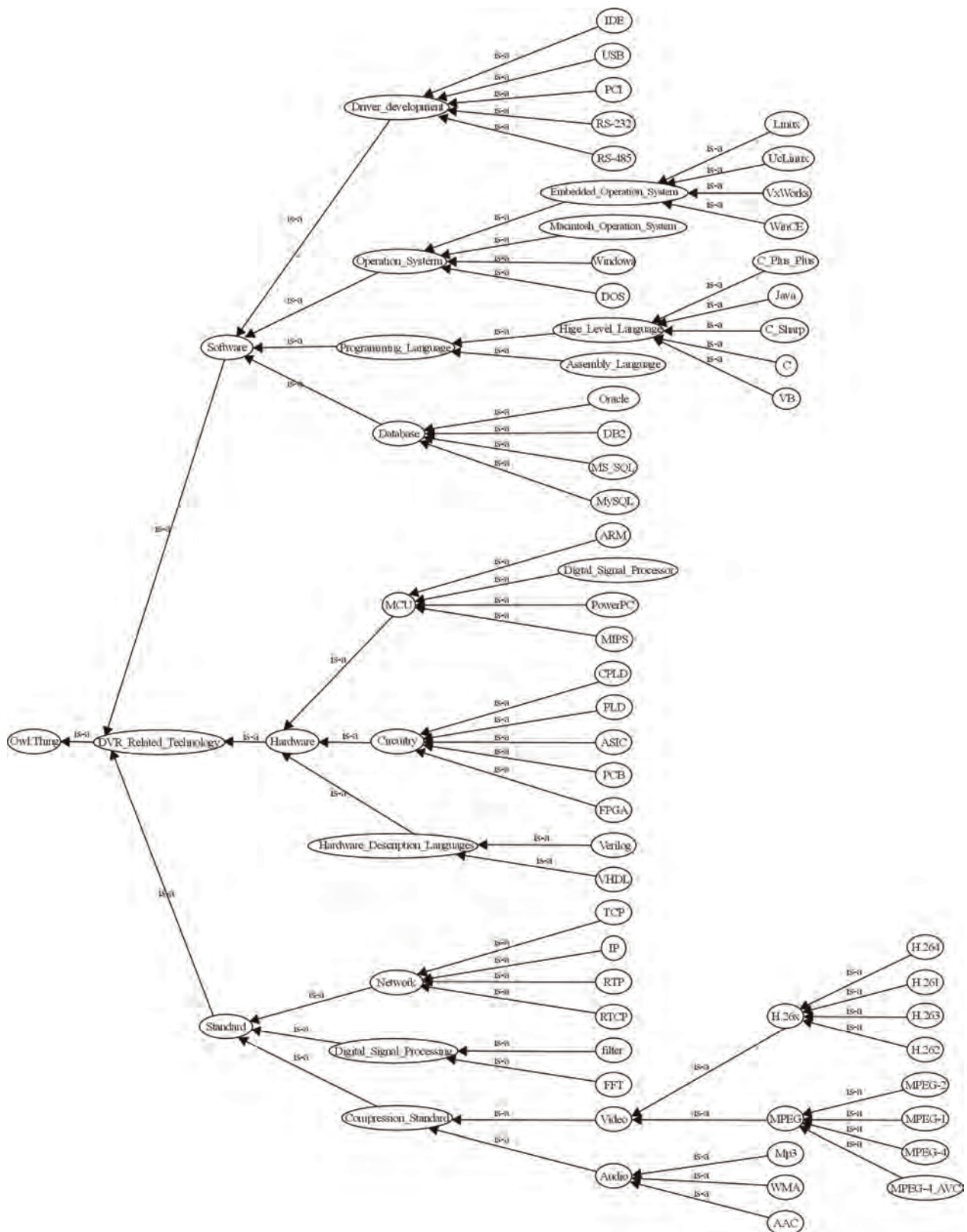


Figure 1. A DVR development ontology represented by OWLViz of protégé.

will be allocated more ahead in the vector. If a concept has multiple father concepts, the first appearance of a concept by following the “breadth-first” scan will be indexed. The following algorithm can be used to convert an OWL document into a concept vector.

// Algorithm 1: Obtain a conceptVector from an OWL document and list entities in the Breadth-First order.

BEGIN

Get root;

conceptVector (0) = root.Name ;

pos = 0; // the index of first concept in every level
count the childNum of root;

While childNum > 0

k=0; //count the children numbers of level l

For each entity i in level l

For each child j of entity i

k++;

conceptVector(pos+k) = j.Name;

Endfor

Endfor

pos = pos + childNum;

childNum = k;

Endwhile

After flattening an ontology graph into a concept vector, a semantic distance matrix can be obtained to depict semantic difference between any two concepts defined in ontology Ω .

//Algorithm 2: Calculate Semantic Distance Matrix

Step 1 Obtain an $N \times N$ initialMatrix = { $e(i,j)$ } : $e(i,j)$ denotes the distance value between the i^{th} element and the j^{th} element which have a direct link.

The initial relation matrix (initialMatrix = { $e(i,j)$ }, $i=1..N$, $j=1..N$) denotes the semantic distance between two concepts, e_i and e_j , which have a direct link in the ontology. If two concepts have a direct relationship linked by “subClassOf” or “ObjectProperty”, a distance 1 or distance 2 is assigned respectively.

$$e(i,j) = \begin{cases} 0 & \text{if } i = j \\ 1 & \text{if there exists } rdfs:subClassOf \text{ between } i \text{ and } j \\ 2 & \text{if there exists owl:ObjectProperty between } i \text{ and } j \\ X & \text{else} \end{cases}, \text{ and}$$

$$e(i,j) = e(j,i).$$

In this initial relation matrix, only the direct links are counted and X represents indirect links that will be calculated in step 2 by updating the initial matrix recursively.

Step 2: Obtain a semantic distance Matrix, DisMatrix = { $d(i,j)$ } : $d(i,j)$ denotes the distance value between the i^{th} element and the j^{th} element.

For $\forall$ concepts A, B and C,

if B = ChildOf(A) and C = $\neg$ DescendantOf(B)

Distance(B,C) = Min(Distance(B,A) + Distance(A,C));

if there is a “SubClassOf” relation between B and A

Distance(B,A) = 1 and

Distance(B,C) = Min(1 + Distance(A,C));

if there is an “ObjectProperty” relation between B and A

Distance(B,A) = 2.

where Distance(i, j) denotes the graph distance between node i and node j . Therefore, the following recursive algorithm can produce semantic distances of indirect links between concepts in an ontology graph.

BEGIN

Input initialMatrix;

For each column j in initialMatrix

If ((the element in i^{th} row) > 0)

fatherRow(k) = i; // record the index of k-th father

if

// there is multiple fathers

Endif

For each $i < j$;

For each father concept k

$d(i,j) = \min_k (d(\text{fatherRow}(k), i)$

$+ d(\text{fatherRow}(k), j))$;

$d(j,i) = d(i,j)$;

Endfor

Endfor

Endfor

END

The semantic distance matrix can be extracted from an OWL document as:

$$\text{DisMatrix} = \begin{bmatrix} 0 & d(1,2) & \dots & d(1,N) \\ d(2,1) & 0 & \dots & d(2,N) \\ & & \dots & \\ d(N,1) & d(N,2) & \dots & 0 \end{bmatrix}, \quad (3)$$

where $d(i, j)$ is the semantic distance between concept e_i and concept e_j .

3. Semantic Matchmaking of RSS Feeds

Since the Internet publishers are distributed and heterogeneous, the literal announcements in RSS have inherent ambiguity and uncertainty due to, for instance, the way to utilize synonyms and jargons. The ambiguity and uncertainty in RSS documents can be alleviated by ontology, which provides a common basis for understanding of frequently used terms and concepts. An RSS feed needs to be parsed into ontology instance first, which is composed of only the concepts defined in the ontology. Hence matchmaking of RSS feeds is transformed to comparison of ontology instances.

At first, the received RSS feeds are preprocessed with the Jena RSS package. The title and the link of each item are obtained. Then concepts defined in the ontology are found out from the title. If a word in the title does not exist in the ontology but is part of a concept in the on-

tology, the concept will be used as a replacement of this word. For example, *Software*, *MPEG*, *Video*, and *Compression_Standard* are captured from the title of the last item in the example of Section 2, where *Compression_Standard*, a concept defined in the ontology, replaces *Compression* found in the title.

Secondly, the concepts extracted from each title are arranged into a hierarchical concept graph, i.e. an ontology instance. It has its URI (link) as the root and reorganizes the concepts by retaining all ancestor descendant and sibling relationships in the ontology. For example, “*Software, MPEG, Video, Compression_Standard*” can be transformed into a concept graph as.

Assume that a subscriber of RSS feeds intends to detect information which he/she is interested in. An expression of the interest could be written into an ontology instance in OWL. For example, a job hunter expresses himself as “*Software Engineer, experiences in C language, Video Compression standard such as H.264, MPEG*”; the ontology instance can be obtained as shown in Figure 3.

Now the information from an RSS feed and a subscriber has been represented formally in accordance with the ontology definition for facilitating semantic matchmaking. As a quantitative description, an ontology instance can be transformed to a numerical vector, i.e. feature vector.

Algorithm 3: A feature vector of an ontology instance can be represented as:

$$V(i)=[s_1, s_2, \dots, s_N]^T \quad (4)$$

The element s_i in $V(i)$ has a one-to-one correspondence to the concept e_i defined in the ontology concept vector (2). The $s_i \in [0,1]$ indicates the semantic closeness between a concept, e_i , and the root of an ontology instance.

$$s_i = \begin{cases} e^{-\alpha \text{Dis}(e_i, \text{root})} & \text{if } e_i \text{ appears in the instance} \\ 0 & \text{if } e_i \text{ does not appear in the instance} \end{cases} \quad (5)$$

where $\text{Dis}(e_i, \text{root})$ is a semantic distance between entity

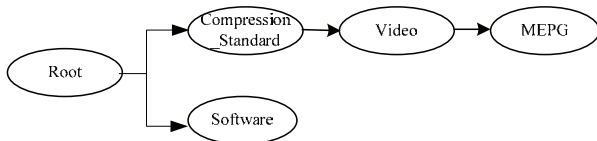


Figure 2. The concept graph of a publisher's ontology instance.

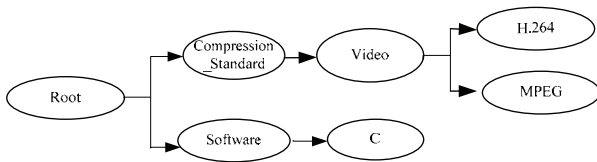


Figure 3. The concept graph of a subscriber's ontology instance.

e_i and the *root*. α is a steepness measure [32] for fuzzy modeling, which is often selected to be $-7/\text{MAX}(\text{Dis})$ because $e^{-7} \approx 0$ when $\text{Dis}(e_i, \text{root})$ reaches its maximum. The semantic distance is computed according to Algorithm 2.

Assume that a user's interest and a received RSS feed can be represented by two feature vectors, V_{user} and V_{rss} . It is expected that the similarities between a user's interest and received RSS feeds are measured by semantic matchmaking, rather than literal matching only. For example an RSS feed from a job-publishing website has information to find “*a Linux developer*”. If a user is interested in a job relevant to “*Embedded Operation System*”, the published job could be missed by using literal matchmaking. In fact “*Embedded Operating System*” and “*Linux*” have a tight semantic relation that can be observed from the ontology in Figure 1.

Due to the distributed nature of Internet applications, it is impossible to force all information publishers using strictly consistent terminologies and sentences. In fact using words in their RSS feeds relies on their own viewpoint and understanding. So the feature vectors are individual-relevant.

In fact the words or concepts used by publishers are fuzzy; any concept implies some extent of others due to the semantic correlations that can be defined by memberships in fuzzy set theory [33]. Suppose the entities of $\{e_1, e_2, \dots, e_N\}$ in ontology form a universe of discourse in a community. Any announcement i in an RSS feed, such as “*design driver program using C*”, can be considered as a linguistic variable. Then the corresponding feature vector $V(i)=[s_1, s_2, \dots, s_N]^T$ in (4) is a fuzzy representation of i from an individual's point of view, where $s_i, i=1 \dots N$, is a grade of membership corresponding to the i^{th} entity in the universe.

Now an individual-dependent $V(i)$ can be transformed into a fuzzy variable $VI(i)$ that becomes individual-independent by taking account of semantic relationships among concepts ($e_1 \dots e_N$)

$$V(i) \wedge r(\Omega) \Rightarrow VI(i) \quad (6)$$

where $r(\Omega)$ is a fuzzy relation matrix and each element of r_{ij} reflects correlation or similarity between entity e_i and entity e_j based on ontology Ω . In this case, similar entities can be taken into account even though they are not explicitly cited in an RSS document.

The fuzzy relation $r(\Omega)$ can be obtained from the ontology definition, e.g. in Figure 1. It is the inverse of the distance matrix in (3):

$$r(\Omega_c) = \begin{bmatrix} 1 & r(1,2) & \dots & r(1,N) \\ r(2,1) & 1 & \dots & r(2,N) \\ \dots & \dots & \dots & \dots \\ r(N,1) & r(N,2) & \dots & r(N,N) \end{bmatrix} \quad (7)$$

where $r(i,j)=e^{-\alpha d(i,j)}$ and α is a steepness measure; $d(i,j)$ is calculated by Algorithm 2. As an inverse of semantic distance matrix (3), equation (7) tells us the closeness between two concepts in ontology Ω .

Therefore for any linguistic item i , the fuzzy inference can be applied to consider implied semantic relations described by ontology.

Algorithm 4: Obtain an individual-independent feature vector:

$$\begin{aligned} VI(i) &= V(i) \vee \wedge r(\Omega) \\ &= [s_1 \ s_2 \ \dots \ s_N] \vee \wedge \begin{bmatrix} 1 & r(1,2) & \dots & r(1,N) \\ r(2,1) & 1 & \dots & r(2,N) \\ \dots & \dots & \dots & \dots \\ r(N,1) & r(N,2) & \dots & 1 \end{bmatrix} \quad (8) \\ &= [x_1 \ x_2 \ \dots \ x_N] \end{aligned}$$

where $\vee, \wedge$ is an inner product of fuzzy relation, such as max-min composition [33]:

$$x_i = \text{Max}(\min(s_1, r(1,i)), \min(s_2, r(2,i)), \dots, \min(s_N, r(N,i))) \quad (9)$$

Now RSS feeds and a user's interest can be represented as a set of individual-independent vectors. Selecting the interested information from RSS feeds becomes a process of similarity measuring between VI_{RSS} and VI_{user} . The following fuzzy operation can be used to filter RSS information.

$$U_i = \frac{|VI_{user} \wedge VI_{RSS}|}{|VI_{user}|} \geq \rho, \text{ where } \rho \in [0,1] \text{ is a threshold} \quad (10)$$

where $VI_{user} \wedge VI_{RSS}$ is the fuzzy min-operation whose i^{th} component is equal to the minimum of $VI_{user}(i)$ and $VI_{RSS}(i)$; $|VI_{user}|$ is the norm of VI_{user} which is defined to be the sum of its components. If every element in VI_{RSS} is equal to or greater than that in VI_{user} , then $U_i=1$ that means VI_{RSS} is regarded as a perfect match of VI_{user} . If not, VI_{RSS} with U_i higher than ρ will be considered as a close match.

4. An RSS Filter Agent for Job Hunting

An RSS filter agent for job hunting is developed by using the proposed algorithms. Suppose a job hunter wants to find a job via RSS feeds from website <http://hotjobs.yahoo.com/jobs/>. The job hunter is only interested in the jobs in a specific knowledge domain, for example the DVR developing domain in Figure 1. The job hunter will provide a favorite profile to the agent. The agent will detect relevant RSS feeds and prompt automatically.

The software architecture of the RSS filter agent is shown in Figure 4. The core components are the Quantitative Module, Matchmaking Module, Service Management Module and Ontology. The Service Management Module takes charge of coordination among the modules and confi-

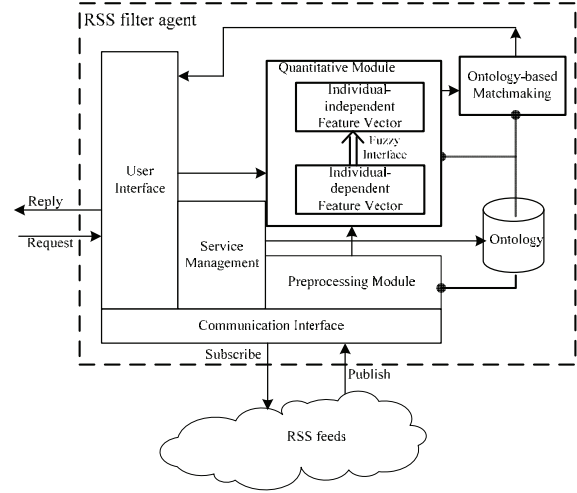


Figure 4. The software architecture of the RSS filter agent.

guration of individual modules.

Protégé(<http://protege.stanford.edu/plugins/owl/index.html>) was used to construct a domain ontology and exports the OWL document. Jena RSS package (<http://jena.sourceforge.net/>) was adopted to parse an RSS feed and Jena Ontology API was used to create and parse the OWL document.

The resulted concept vector from Algorithm 1 is as [DVR_Related_Technology, Standard, Hardware, Software, Compression_Standard, Digital_Signal_Processing, Network, Hardware_Description_Languages, Circuitry, MCU, Operation_System, Programming_Language, Driver_Development, Database, Video, Audio, FFT, filter, RTP, TCP, IP, RTCP, VH-DL, Verilog, PLD, CPLD, FPGA, ASIC, PCB, ARM, MIPS, PowerPC, Digital_Signal_Processor, Embedded_Operation_System, Macintosh_Operation_System, Windows, DOS, High_Level_Language, Assembly_Language, RS-232, USB, RS-485, IDE, PCI, DB2, MS_SQL, Oracle, MySQL, MPEG, H.26x, AAC, MP3, WMA, WinCE, ucLinux, VxWorks, Linux, Java, C_Plus_Plus, C, VB, C_Sharp, MPEG-4, MPEG-4_AVC, MPEG-2, MPEG-1, H.263, H.264, H.261, H.262].

The corresponding initial relation matrix is obtained and the distance matrix can then be computed according to Algorithm 2, which is a 70*70 matrix and is illustrated in Figure 5.

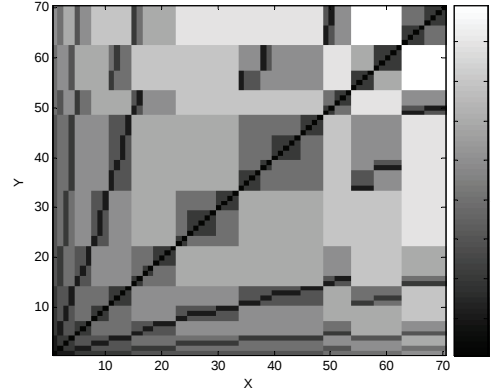


Figure 5. Distance matrix of the DVR developing ontology.

From the graph we can observe that it is a symmetry matrix. The grayscale indicates the semantic distance between an entity in x axis and an entity in y axis, which are entity indexes of the concept vector.

The following job titles were received from <http://hotjobs.yahoo.com/jobs/> RSS feeds.

J0: "Software Engineer, MPEG, Video, Compression"

J1: "Senior Firmware Engineer W/ MPEG And ARM"

J2: "Software Engineer - Programmer - Developer - C++ - Java"

J3: "C/C++/Linux/Oracle Developers"

J4: "Embedded Software Engineer -embedded OS, C, Assembly, DSP, Video"

J5: "Application Engineer, Audio/Video, Hardware, ASIC, PCB"

J6: "MPEG ASIC/Hardware Engineer"

J7: "Video Systems, H.264, MPEG, Decoder, FPGA, HDTV"

There are three users who want to find jobs and announce themselves as:

user0: "A hardware engineer, experience in using CPLD/FPGA with Verilog for design"

user1: "Software Engineer, experience in C language, Video Compression standard such as H.264, MPEG"

user2: "H.264, MPEG, Video, Assembly, FPGA, DSP"

From these statements, it is easy to observe that they are individual-dependent and do not follow a strict format.

The extracted concept graphs from J0 and J1 have been given in Figure 2 and Figure 3, respectively. The resulted feature vectors to denote the above ontology instances are listed below:

J0: $[0, 0, 0, e^{-\alpha^{*1}}, e^{-\alpha^{*1}}, 0, \dots, 0, e^{-\alpha^{*2}}, 0, \dots, e^{-\alpha^{*3}}, 0, \dots, 0]$

J1: $[0, \dots, 0, e^{-\alpha^{*1}}, 0, \dots, 0, e^{-\alpha^{*1}}, 0, \dots, 0]$

J2: $[0, 0, 0, e^{-\alpha^{*1}}, 0, \dots, 0, e^{-\alpha^{*2}}, e^{-\alpha^{*2}}, \dots, 0, 0]$

J3: $[0, \dots, 0, e^{-\alpha^{*1}}, 0, \dots, 0, e^{-\alpha^{*1}}, 0, e^{-\alpha^{*1}}, e^{-\alpha^{*1}}, 0, \dots, 0]$

J4: $[0, \dots, 0, e^{-\alpha^{*1}}, 0, \dots, 0, e^{-\alpha^{*1}}, e^{-\alpha^{*1}}, 0, \dots, 0, e^{-\alpha^{*1}}, 0, \dots, 0, e^{-\alpha^{*1}}, 0, \dots, 0]$

J5: $[0, 0, e^{-\alpha^{*1}}, 0, \dots, 0, e^{-\alpha^{*1}}, e^{-\alpha^{*1}}, 0, \dots, e^{-\alpha^{*2}}, e^{-\alpha^{*2}}, 0, \dots, 0]$

J6: $[0, 0, e^{-\alpha^{*1}}, 0, \dots, 0, e^{-\alpha^{*2}}, 0, \dots, e^{-\alpha^{*1}}, 0, \dots, 0]$

J7: $[0, \dots, 0, e^{-\alpha^{*1}}, 0, \dots, e^{-\alpha^{*1}}, \dots, e^{-\alpha^{*2}}, \dots, e^{-\alpha^{*2}}, 0, 0]$

user0: $[0, 0, e^{-\alpha^{*1}}, 0, \dots, 0, e^{-\alpha^{*1}}, e^{-\alpha^{*1}}, 0, \dots, e^{-\alpha^{*2}}, e^{-\alpha^{*2}}, 0, \dots, 0]$

user1: $[0, 0, e^{-\alpha^{*1}}, 0, \dots, e^{-\alpha^{*2}}, \dots, e^{-\alpha^{*1}}, 0, \dots, 0]$

user2: $[0, 0, 0, 0, \dots, 0, e^{-\alpha^{*1}}, 0, \dots, e^{-\alpha^{*1}}, 0, \dots, e^{-\alpha^{*2}}, 0, \dots, e^{-\alpha^{*2}}, 0, 0]$

The elements in each feature vector are corresponding to literal labels in the concept vector respectively.

α is set to 1.

According to Algorithm 4 and (10), the resulted similarities corresponding to every job J_i are:

user0: $[0.0, 0.0, 0.0, 0.0, 0.0, 0.475, 0.475, 0.175]$

user1: $[0.833, 0.045, 0.333, 0.122, 0.577, 0.122, 0.045, 0.212]$

user2: $[0.135, 0.098, 0.0, 0.0, 0.634, 0.268, 0.098, 0.465]$

The relationships between the published jobs and the announcements of users are illustrated in Figure 6.

From this figure, the agent can identify suitable jobs for users. For example, the most relevant jobs for user0 are J5 and J6.

As a general illustration, Figure 7 shows the user interface for configuration of an RSS filter agent. A user registered as Tom and subscribed RSS feeds from Yahoo hotjobs (<http://hotjobs.yahoo.com/rss/0/USA/-/-/IT/>), UK academic employment (<http://www.jobs.ac.uk/rss/disc/2516.xml>), and CareerBuilder (http://rtq.careerbuilder.com/RTQ/rss20.aspx?lr=cbcb_ct&rssid=cb_ct_rss_engine&cat=JN004&state=IL&city=chicago) for job

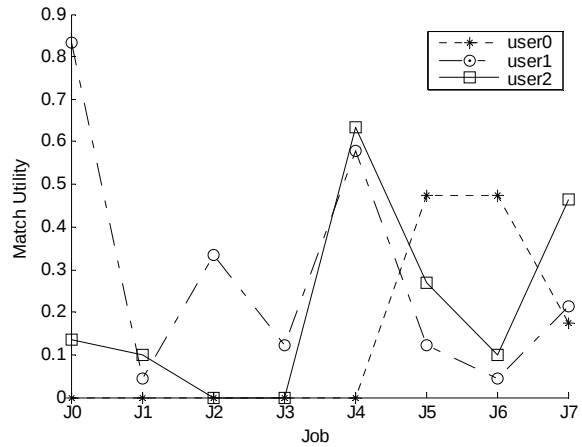


Figure 6. Jobs vs users.

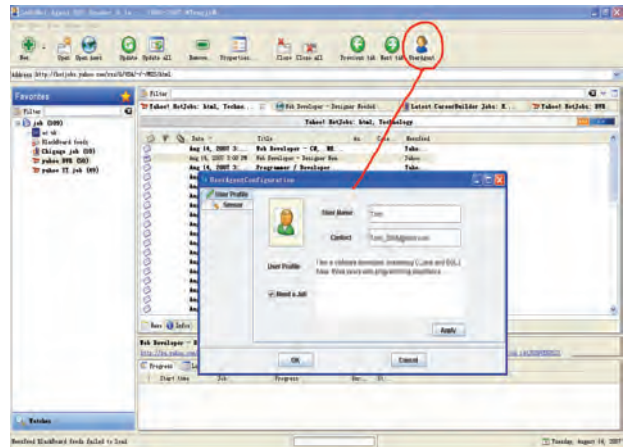
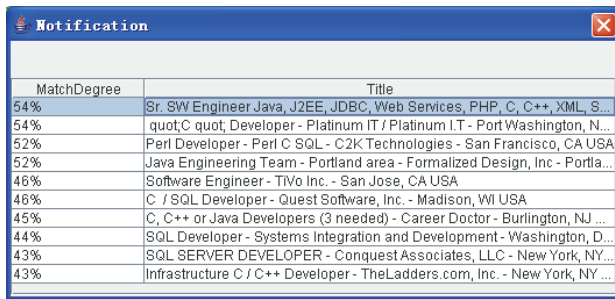


Figure 7. RSS filter agent registration.



MatchDegree	Title
54%	Sr. SW Engineer Java, J2EE, JDBC, Web Services, PHP, C, C++, XML, S...
54%	quot;C quot; Developer - Platinum IT / Platinum I.T - Port Washington, N...
52%	Perl Developer - Perl C SQL - C2K Technologies - San Francisco, CA USA
52%	Java Engineering Team - Portland area - Formalized Design, Inc - Portla...
46%	Software Engineer - Tivo Inc. - San Jose, CA USA
46%	C / SQL Developer - Quest Software, Inc. - Madison, WI USA
45%	C, C++ or Java Developers (3 needed) - Career Doctor - Burlington, NJ ...
44%	SQL Developer - Systems Integration and Development - Washington, D...
43%	SQL SERVER DEVELOPER - Conquest Associates, LLC - New York, NY...
43%	Infrastructure C / C++ Developer - TheLadders.com, Inc. - New York, NY ...

Figure 8. Virtual sensor output.

hunting. A user profile, “*I am a software developer, mastering C, Java and SQL. I have three years web programming experience*”, was announced as the interest. The parameters of the virtual sensor were set as a sensor resolution of 56% and display of top 10 candidates. The detected RSS feeds were shown in Figure 8, where candidates were identified and the match degrees were shown.

The method proposed in this paper can be applied for both content publishers and content subscribers. For example, on a business website side, the agent can be used to find potential customers and, on a user side, the agent can be used to perceive favorite information up to a certain match-level by adjusting the threshold of ρ .

5. Conclusions

In this paper, the user-oriented active choice of information from RSS feeds was discussed. A fuzzy method for matchmaking between a subscriber's interest and RSS items was proposed, which converted headlines of an RSS feed into numerical vectors and semantic closeness to the interest was measured. Ontology acted as a bridge to link heterogeneous publishers with subscribers in semantics rather than in words. Concept graphs were firstly extracted from the title of each RSS item and transformed into an individual-dependent feature vector with the aid of ontology. In order to eliminate ambiguity due to different literal expression of individuals, fuzzy inference was applied to obtain the grade of membership in terms of ontology, which became an individual independent feature vector. A job seeking agent was developed to illuminate the method and the result showed its validity.

6. References

- [1] J. Grossnickle, T. Board, B. Pickens, and M. Belmont, “RSS-crossing into the mainstream,” October 2005. http://publisher.yahoo.com/rss/RSS_whitePaper1004.pdf.
- [2] D. Kuokka and L. Harada, “Integrating information via matchmaking,” *Journal of Intelligent Information Systems*, Kluwer Academic Publishers, Vol. 6, pp. 261–279, 1996.
- [3] K. Kurbel and I. Loutchko, “A model for multi-lateral negotiations on an agent-based marketplace for personnel acquisition,” *Electronic Commerce Research and Applications*, Vol. 4, No. 3, pp. 187–203, 2005.
- [4] R. Hishiyama and T. Ishida, “Modeling e-procurement as co-adaptive matchmaking with mutual relevance feedback,” M. Barley and N. K. Kasabov (Eds.): *Intelligent Agents and Multi-Agent Systems*, the 7th Pacific Rim International Workshop on Multi-Agents (PRIMA'04), Auckland, New Zealand, pp. 67–80, 2004.
- [5] D. Trastour, C. Bartolini, and C. Preist, “Semantic web support for the business-to-business e-commerce pre-contractual lifecycle,” *Computer Networks*, Vol. 42, pp. 661–673, 2003.
- [6] J. Kopena and W. C. Regli, “DAMLJessKB: A tool for reasoning with the semantic web,” *IEEE Intelligent Systems*, Vol. 18, pp. 74–77, 2003.
- [7] S. A. Ludwig and S. M. S. Reyhani, “Introduction of semantic matchmaking to grid computing,” *Journal of Parallel and Distributed Computing*, Vol. 65, pp. 1533–1541, 2005.
- [8] D. Sandler, A. Mislove, A. Post, and P. Druschel, “Feedtree: Sharing web micronews with peer-to-peer event notification,” In *Proceedings of the 4th International Workshop on Peer-to-Peer Systems (IPTPS'05)*, Ithaca, NY, USA, pp. 141–151, 2005.
- [9] B. Hammersley, “Content syndication with RSS,” O' Reilly, ISBN: 0-596-00383-8, 2003.
- [10] E. Jung, “UniRSS: A new RSS framework supporting dynamic plug-in of RSS extension modules,” In *Proceedings of the 1st Aisan Semantic Web Conference (ASWC'06)*, Beijing, China, pp. 169–178, 2006.
- [11] K. Wegrzyn-Wolska and P. S. Szczepaniak, “Classification of RSS-formatted documents using full text similarity measures,” In *Proceedings of the 5th International Conference on Web Engineering (ICWE'05)*, Sydney, Australia, pp. 400–405, 2005.
- [12] P. S. Szczepaniak and A. Niewiadomski, “Clustering of documents on the basis of text fuzzy similarity,” Abramowicz W. (Eds.): *Knowledge-based Information Retrieval and Filtering from the Web*, pp. 219–230, Kluwer Academic Publishers, 2003.
- [13] N. Cancedda, E. Gaussier, C. Goutte, and J. Renders, “Word-sequence kernels,” *Journal of Machine Learning Research*, 3, pp. 1059–1082, 2003.
- [14] H. Lodhi, N. Cristianini, J. Shave-Taylor, and C. Watkins, “Text classification using string kernel,” *Advances in Neural Information Processing System*, Vol. 13, pp. 563–569, 2001.
- [15] G. Salton and M. J. McGill, “Introduction to modern information retrieval,” McGraw-Hill, New York, 1983.

- [16] J. J. Sampera, P. A. Castillob, L. Araujoc, J. J. Merelob, O. Cordon, and F. Tricas, "NectaRSS, an intelligent RSS feed reader," *Journal of Networks and Computer Applications*, Vol. 31, pp. 793–806, 2008.
- [17] R. Prabowo and M. Thelwall, "A comparison of feature selection methods for an evolving RSS feed corpus," *Information Processing and Management*, Vol. 42, pp. 1491–1512, 2006.
- [18] N. S. Glance, M. Hurst, and T. Tomokiyo, "BlogPulse: Automated trend discovery for weblogs," In *Proceedings of the 13th International WWW Conference: Workshop on Weblogging Ecosystem: Aggregation, Analysis and Dynamics*, New York, USA, pp.1–8, 2004.
- [19] Y. Yang and J. O. Pedersen, "A comparative study on feature selection in text categorization," In *Proceedings of the Fourteenth International Conference on Machine Learning (ICML'97)*, San Francisco, USA, pp. 412–420, 1997.
- [20] T. Berners-Lee, "Semantic web road map," 1998. <http://www.w3.org/DesignIssues/Semantic.html>.
- [21] X. Ning, H. Jin and H. Wu, "RSS: A framework enabling ranked search on the semantic web," *Information Processing and Management*, Vol. 44, pp. 893–909, 2008.
- [22] S. Avancha, A. Joshi, and T. Finin, "Enhanced service discovery in Bluetooth," *Communications*, pp. 96–99, 2002.
- [23] S. A. Ludwig and S. M. S. Reyhani, "Semantic approach to service discovery in a grid environment," *Journal of Web Semantics*, Vol. 4, pp.1–13, 2006.
- [24] S. Colucci, T. D. Noia, and E. D. Sciascio, F. M. Donini, M. Mongiello, "Concept abduction and contraction for semantic-based discovery of matches and negotiation spaces in an E-marketplace," *Electronic Commerce Research and Applications*, Vol. 4, pp. 345–361, 2005.
- [25] G. Stoilos, G. Stamou, V. Tzouvaras, J. Z. Pan, and I. Horrocks, "The fuzzy description logic f-SHIN," *International Workshop on Uncertainty Reasoning For the Semantic Web*, 2005.
- [26] J. Z. Pan, G. Stoilos, G. B. Stamou, V. Tzouvaras, and I. Horrocks, "f-SWRL: A fuzzy extension of SWRL," *Journal on Data Semantics*, Vol. 6, pp. 28–46, 2006.
- [27] P. Jiang, Q. Mair, and Z. Feng, "Agent alliance formation using ART-networks as agent belief models," *Journal of Intelligent Manufacturing*, Vol. 18, pp. 433–448, 2007.
- [28] ISO, "Application protocol: Configuration controlled design," IS 10303 – Part 203, 1994.
- [29] P. Jiang, Q. Mair, and J. Newman, "The application of UML to the design of processes supporting product configuration management," *International Journal of Computer Integrated Manufacturing*, Vol. 19, pp. 393–407, 2006.
- [30] K. P. Sycara, M. Klusch, S. Widoff, and J. Lu, "Dynamic service matchmaking among agents in open information environments," *ACM SIGMOD Record (ACM Special Interests Group on Management of Data)*, Vol. 28, pp. 47–53, 1999.
- [31] J. Akoka and I. Comyn-Wattiau, "Entity-relationship and object-oriented model automatic clustering," *Data and Knowledge Engineering*, Vol. 20, pp. 87–117, 1996.
- [32] J. Williams and N. Steele, "Difference, distance and similarity as a basis for fuzzy decision support based on prototypical decision classes," *Fuzzy Sets and Systems*, Vol. 131, pp. 35–46, 2002.
- [33] L. A. Zadeh, "Fuzzy sets," *Information and Control*, Vol. 8, pp. 338–353, 1965.

Text Extraction in Complex Color Document Images for Enhanced Readability

P. Nagabhushan, S. Nirmala

Department of Studies in Computer Science, University of Mysore, Mysore, India

Email: pnagabhushan@compsci.uni-mysore.ac.in, nir_shiv_2002@yahoo.co.in

Abstract

Often we encounter documents with text printed on complex color background. Readability of textual contents in such documents is very poor due to complexity of the background and mix up of color(s) of foreground text with colors of background. Automatic segmentation of foreground text in such document images is very much essential for smooth reading of the document contents either by human or by machine. In this paper we propose a novel approach to extract the foreground text in color document images having complex background. The proposed approach is a hybrid approach which combines connected component and texture feature analysis of potential text regions. The proposed approach utilizes Canny edge detector to detect all possible text edge pixels. Connected component analysis is performed on these edge pixels to identify candidate text regions. Because of background complexity it is also possible that a non-text region may be identified as a text region. This problem is overcome by analyzing the texture features of potential text region corresponding to each connected component. An unsupervised local thresholding is devised to perform foreground segmentation in detected text regions. Finally the text regions which are noisy are identified and reprocessed to further enhance the quality of retrieved foreground. The proposed approach can handle document images with varying background of multiple colors and texture; and foreground text in any color, font, size and orientation. Experimental results show that the proposed algorithm detects on an average 97.12% of text regions in the source document. Readability of the extracted foreground text is illustrated through Optical character recognition (OCR) in case the text is in English. The proposed approach is compared with some existing methods of foreground separation in document images. Experimental results show that our approach performs better.

Keywords: Color Document Image, Complex Background, Connected Component Analysis, Segmentation of Text, Texture Analysis, Unsupervised Thresholding, OCR

1. Introduction

Most of the information available today is either on paper or in the form of still photographs, videos and electronic medium. Rapid development of multimedia technology in real life has resulted in the enhancement of the background decoration as an attempt to make the documents more colorful and attractive. Presence of uniform or non-uniform background patterns, presence of multiple colors in the background, mix up of foreground text color with background color in documents make the documents more attractive but *deteriorates the readability*. Some of the examples are advertisements, news paper articles, decorative postal envelopes, magazine pages, decorative letter pads, grade sheets and story books of children. Further, the background patterns opted in the preparation of power point slides appear to be attractive

but cause difficulty in reading the contents during presentation on the screen. These compel to devise methods to reduce the adverse effect of background on the foreground without losing information in the foreground.

There are many applications in document engineering in which automatic detection and extraction of foreground text from complex background is useful. These applications include building of name card database by extracting name card information from fanciful name cards, automatic mail sorting by extracting the mail address information from decorative postal envelopes [1]. If the text is printed on a clean background then certainly OCR can detect the text regions and convert the text into ASCII form [2]. Several commercially available OCR products perform this; however they result in low recognition accuracy when the text is printed against shaded and/or complex background.

The problem of segmentation of text information from complex background in document images is difficult and still remains a challenging problem. Development of a generic strategy or an algorithm for isolation of foreground text in such document images is difficult because of high level of variability and complexity of the background. In the past, many efforts were reported on the foreground segmentation in document images [3–16]. Thresholding is the simplest method among all the methods reported on extraction of foreground objects from the background in images. Sezgin and Sankur [14] carried out an exhaustive survey of image thresholding methods. They categorized the thresholding methods according to the information they are exploiting, such as histogram shape based methods, clustering based methods, entropy based methods, object-attributes based methods, spatial methods and local methods. The choice of a proper algorithm is mainly based on the type of images to be analyzed. Global thresholding [7,11] techniques extract objects from images having uniform background. Such methods are simple and fast but they cannot be adapted in case the background is non uniform and complex. Local thresholding methods are window based and compute different threshold values to different regions in the image [8,14] using local image statistics. The local adaptive thresholding approaches are also window based and compute threshold for each pixel using local neighborhood information [9,13]. Trier and Jain [17] evaluated 11 popular local thresholding methods on scanned documents and reported that Niblack's method [9] performs best for OCR. The evaluation of local methods in [17] is in the context of digit recognition. Sauvola and Pietikainen [13] proposed an improved version of Niblack method especially for stained and badly illuminated document images. The approaches proposed in [9,13] are based on the hypothesis that the gray values of text are close to 0 (black) and background pixels are close to 255 (white). Leedham *et al.* [8] evaluated the performance of five popular local thresholding methods on four types of "difficult" document images where considerable background noise or variation in contrast and illumination exists. They reported that no single algorithm works well for all types of image. Another drawback of local thresholding approaches is that the processing cost is high. Still there is a scope to reduce the processing cost and improve the results of segmentation of foreground text from background by capturing and thresholding the regions containing text information. Often we encounter the documents with font of any color, size and orientation. Figure 1 shows some sample color document images where the foreground text varies in color, size and orientation. Conventional binarization methods assume that the polarities of the foreground and background intensity are known apriori; but practically it is not possible to know foreground and background color intensity in advance. This drawback of conventional thresholding methods call for specialized binarization.

Text-regions in a document image can be detected either by connected component analysis [3,18] or by texture analysis method [1,19]. The connected component based methods detect the text based on the analysis of

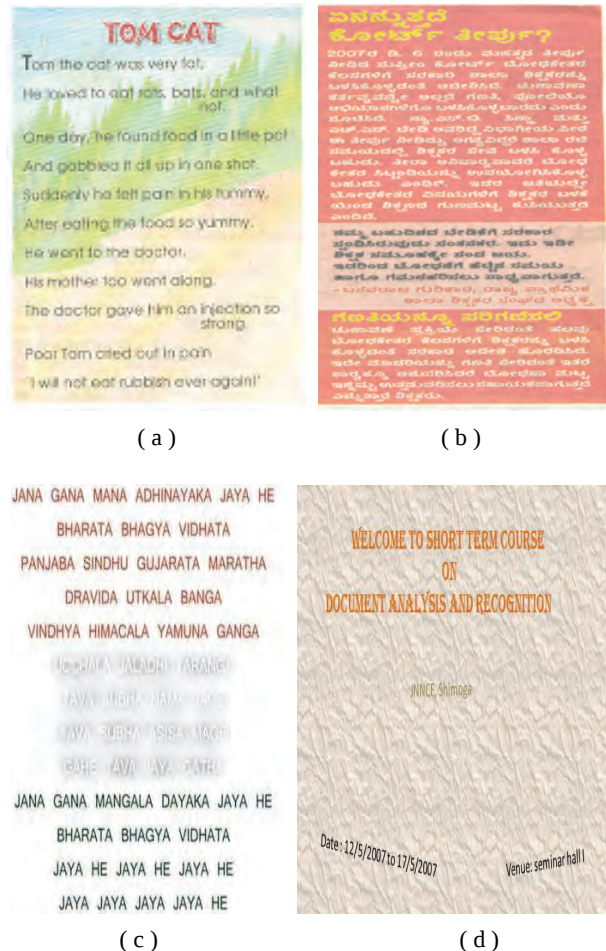


Figure 1. Color documents with printed text of different size, color and orientation.

the geometrical arrangement of the edges [16] that compose the characters. They are simple to implement and detect text at faster rate but are not very robust for text localization and also result in false text regions for images having complex background. Pietikainen and Okun [12] used edge detectors to extract the text from document images. Their method fails in extracting the tilted text lines and erroneously classifies the textured background as text. Chen *et al.* [3] proposed a method to detect the vertical and horizontal edges in an image. They used different dilation operators for these two kinds of edges. Real text regions are then identified using support vector machine. The method lacks in detecting text tilted in any orientation. Zhong *et al.* [18] used edge information to detect the text lines. Their method deals with complex color images pretty well but restricted to certain size constraints on characters. The texture based methods detect the text regions based on the fact that text and background have different textures [1]. In [19] it is assumed that text is aligned horizontally or vertically and text font size is in limited range. The method proposed in [19] uses texture features to extract text but fails in case

of small font size characters. Their method is based on the assumption that the text direction is horizontal or vertical. In [2] texture based method is proposed to detect text regions in gray scale documents having textured background. In their method text strokes are extracted from the detected text regions using some heuristics on text strings such as height, spacing, and alignment [2]. The extracted text strokes are enclosed in rectangular boxes and then binarized to separate the text from the background. Their method fails to extract the text in low contrast document images. Also they fail to extract the tilted text in document images. Most of the above methods are very restrictive in alignment and type of the text they can process. Sobotta *et al.* [15] proposed a method that uses color information to extract the text in colored books and journal covers. Their method fails to extract the isolated characters. In [4] a method is proposed to separate foreground from background in low quality ancient document images. The test documents used in their method are scanner based handwritten, printed manuscripts of popular writers. Their method fails to segment the foreground text in documents with textured background. Liu *et al.* [20] proposed a hybrid approach to detect and verify the text regions and then binarize the text regions using expectation maximization algorithm. The computation complexity of verification process of the text region is high. The performance of the algorithm proposed in [20] degrades when the documents have high complex background and fails to extract the text in low contrast document images. Kasar *et al.* [6] proposed a specialized binarization to separate the characters from the background. They addressed the degradations induced in camera based document images such as uneven lighting and blur. The approach fails to extract the text in document images having textured background. It also fails to detect the characters in low resolution document images. From the literature survey, it is evident that identifying, separating the foreground text in document images and making it smoothly readable is still a research issue in case the background of a document is highly complex and the text in foreground takes any color, font, size and tilt.

In this paper we propose a novel hybrid approach to extract the foreground text from complex background. The proposed approach is a five stage method. In the first stage the candidate text regions are identified based on edge detection followed by connected component analysis. Because of background complexity the non-text region may also be detected as text region. In the second stage the false text regions are reduced by extracting the texture feature and analyzing the feature value of candidate text regions. In the third stage we separate the text from the background in the image segments narrowed down to contain text using a specialized binarization technique which is unsupervised. In the fourth stage the text segments that would still contain noise are identified. In final stage the noise affected regions are reprocessed

to further improve the readability of the retrieved foreground text. The rest of the paper is organized as follows. Section 2 introduces our approach. In Section 3 experimental results and discussion are provided. Time complexity analysis is provided in Section 4. Conclusions drawn from this study are summarized in Section 5.

2. Proposed Approach

In this work we have addressed the problem of improving the readability of foreground text in text dominant color document images having complex background by separating the foreground from the background. The proposed work is based on the assumption that the foreground text is printed text. Two special characteristics of the printed text are used to detect the candidate text regions. They are, 1) Printed characters exhibit regularity in separation and 2) Due to high intensity gradient, a character always forms edges against its background. The sequence of the stages in proposed hybrid approach is shown in Figure 2. The proposed five stage approach is described in the subsections to follow.

2.1. Detection of Text Regions

The proposed method uses Canny edge detector to detect edges [21] because Canny edge operator has two advantages: it has low probability of missing an edge and at the same time it has some resistance to the presence of noise. We conducted experiments on both gray scale and RGB color model of source document images. It is observed from the experimental evaluations that the edge detection in gray scale document images resulted in loss of text edge pixels to certain extent. Hence edge detection in RGB color model of source document is proposed instead of transforming the color document to

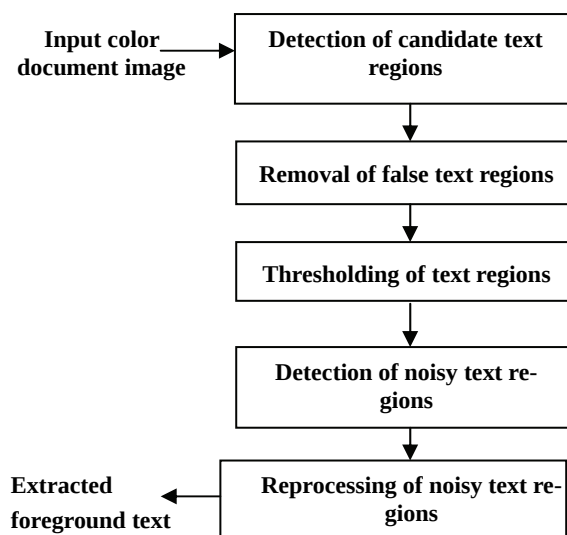


Figure 2. Stages of the proposed approach.

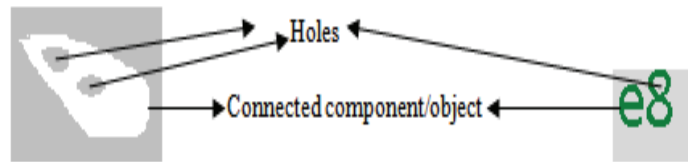


Figure 3. Holes in a connected component.

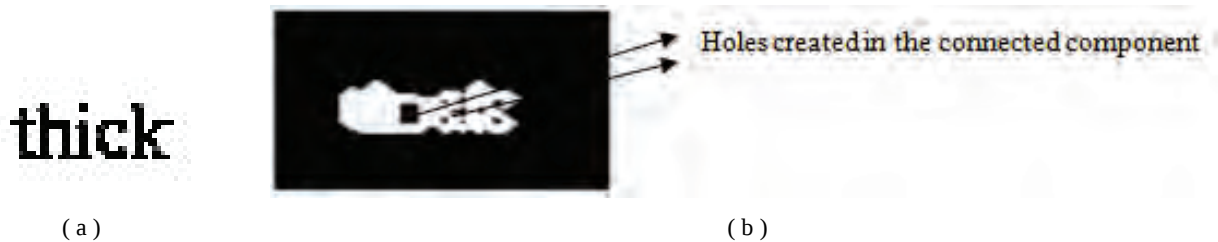


Figure 4. (a) Word that composes characters without holes, (b) holes created by connecting the characters in the word.

a gray scale document. Edge detection is carried out in each color channel separately as the foreground text may be of any color and therefore the edges could be visible in one or more of these three color channels. The results of the edge detection of all the three color channels are assimilated so that no character edge gets missed. Suppose E_R , E_G and E_B are the images after applying the Canny edge operator on red, green and blue components of the input color image, the resulting edge image “E” after assimilation is given by,

$$E = E_R \vee E_G \vee E_B \quad (1)$$

where “V” represents logical “OR” operator.

The resulting edge image “E” contains edges corresponding to character objects in the input image. When the background is highly complex and decorative, the edge image “E” might contain edges corresponding to non-text objects also. An 8-connected component labeling follows the edge detection step. The non-text components in the background such as underlines, border lines, and single lines without touching foreground characters do not contain any hole. A hole in a connected component is illustrated in Figure 3.

Generally in a document image some printed characters contain one or more holes and some other characters do not contain a hole. If a word is composed of characters without holes, using dilation operation [3] it could be possible to thicken the characters so that they get connected and additional holes are created in the space between the characters. This process is depicted in Figure 4.

As the text lines in most of the documents are conventionally aligned horizontally, we conducted experiments on dilation of the edge image “E” in horizontal direction.

It is observed from experimental evaluations that dilation of edge image “E”, only in horizontal direction is not enough to create holes in most of the connected components corresponding to character strings. Therefore we extended the dilation operation on the edge image in both horizontal and vertical directions. From the experimental results it is observed that dilation of the edge image in both horizontal and vertical directions has created holes in most of the connected components that corresponds to character strings. The size of the structuring element for dilation operation was fixed based on experimental evaluation. As no standard corpus of document images is available for this work we conducted experiments on the document images collected and synthesized by us which depict varying background of multiple colors and foreground text in any color, font, size. We dilated the edge image row-wise and column-wise with line structuring element of different sizes. Table 1(a) and Table 1(b) show the percentage loss of characters in a document image after dilating the edge image “E” with various sizes of horizontal structuring element and vertical structuring element.

Table 1(a). Percentage loss of characters for various sizes of horizontal structuring element.

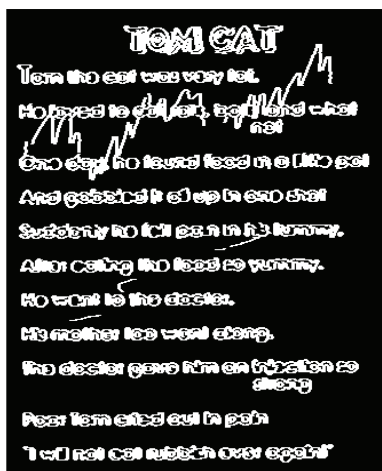
Size of vertical structuring element is 3X1, total number of characters processed=6171					
Size of the horizontal structur- ing element	1x2	1x3	1x4	1x5	1x6
Loss of characters in per- centage	1.93	1.93	2.37	2.38	4.58

Table 1(b). Percentage loss of characters for various sizes of vertical structuring element.

Size of horizontal structuring element is 1x3, total number of characters processed=6171					
Size of the vertical structuring element	2x1	3x1	4x1	5x1	6x1
Loss of characters in percentage	1.99	1.96	1.98	1.98	4.18

From Table 1(a) and Table 1(b) it is observed that with horizontal structuring element of size 1x3 and vertical structuring element of size 3x1, the percentage loss of characters in a document image is very low. This indicates that the dilation of edge image with line structuring element 1x3 in horizontal direction and line structuring element 3x1 in vertical direction creates additional holes in most of the text components which is depicted in Figure 4. Figure 5 shows the document image after assimilating the results of horizontal and vertical dilation of edge image of the input image which is shown in Figure 1(a).

The 8-connected component labeling is performed on the dilated edge image. Based on the size of the characters in the source document and spacing between the words the so labeled connected components may be composed of a single character or an entire word or part of the word or a line. The labeled component may also contain words from different lines if the words in different lines are connected by some background object. In this work the built-in function “Bwboundaries” in MATLAB image processing tool box is used to find the holes in a connected component. The connected components are analyzed to identify the object/component containing hole. We removed the connected components without hole(s). Other non-text components are eliminated by computing and analyzing the standard deviation of each connected component which is elaborated in the next subsection.

**Figure 5. Document image after dilation.**

2.2. Removal of False Text Regions

Because of background complexity certain amount of non-text region in the source document might be identified as text region in connected component analysis process. The proposed approach is based on the idea that the connected components that compose textual information will always contain holes. Holes in the connected components comprise the pixels from the background. Hence each connected component represents an image segment containing only background pixels in case there is no text information (false text region) or both foreground and background pixels in case the connected component contains text information (true text region). To remove the image segments containing only background pixels, standard deviation of gray scale values of all pixels in each image segment/connected component is calculated. The standard deviation in the image segments occupied with only background pixels (ie, image segments without text) is very low where as the standard deviation in the image segments occupied by both background and foreground pixels (ie, image segments containing text) is high [10]. Based on this characteristic property of document image it could be possible to discriminate the non-text image segments from image segments containing text. To set the value for “SD” we conducted experiments on document images having uniform/non-uniform background of multiple colors and foreground text of any font, color, size and orientation. We set the value for standard deviation from a set of 120 images (first 120 images in the corpus). The document image samples are selected randomly in multiples of 5, from the corpus of images synthesized and collected by us, to set the empirical value for standard deviation ‘SD’. The sample images selected are all distinct images from the corpus of images. From the plot shown in Figure 6, it is observed that a threshold value of 0.4 on “SD” is sufficient enough to filter out the non-text regions without loss of detected text. In addition repeating the experiment 10 times on 50 distinct samples selected randomly each time (from first 120 samples in the corpus), demonstrated that the value for standard deviation falls in the range 0.405 to 0.42. We extended the experiment on 100 more images in the corpus apart from sample images used for setting the value for “SD” and observed that SD=0.4 resulted in reduction of the false text regions without loss of text information in the document. However, although choosing a higher “SD” value reduces the false text regions it results in the loss of foreground text and choosing “SD” value lower than 0.4 leads to additional processing of more number of false text regions. Hence standard deviation of 0.4 is chosen as the threshold value.

2.3. Extraction of Foreground Text

In the proposed approach color information is not used to extract the foreground text. As already during the first

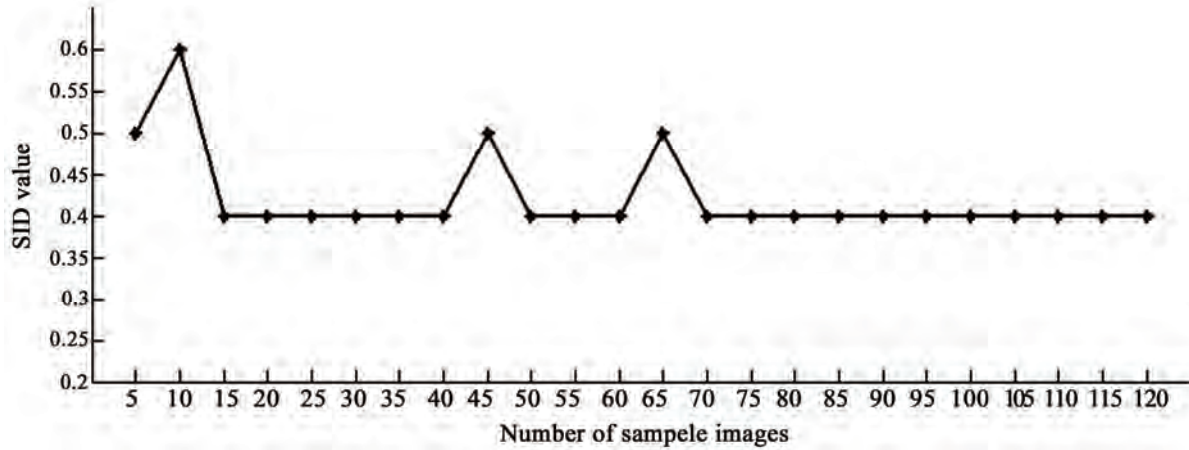


Figure 6. Plot showing the number of training sample images versus the SD value for no loss of textual information.

stage of our approach the evidences of textual edges have been drawn from intensity values of each color channel (RGB model). Also it is computationally inexpensive to threshold the gray scale of the image segment corresponding to the connected component by tightly encapsulating the segment. Figure 7 illustrates background and foreground pixels in a connected component. In each connected component average gray scale intensity value of foreground pixels and average gray scale intensity value of the background pixels are computed.

Suppose “m” and “s” are mean and standard deviation of gray scale intensities in an image segment corresponding to a connected component with hole(s), the threshold value for that segment is derived automatically from the image data as given in [9],

$$\text{threshold} = m - k * s \quad (2)$$

where (k) is a control parameter and value of (k) is decided based on the average gray scale intensity value of foreground pixels and average gray scale intensity value of background pixels. Suppose “ V_f ” is average gray scale intensity value of foreground pixels and “ V_b ” is average gray scale intensity value of background pixels. We conducted experiments on document images with varying background and foreground text of different colors. From experimental evaluations it is observed that choosing $k=0.05$ for $V_f > V_b$ and $k=0.4$ for $V_f \leq V_b$ results in a better threshold value. In this work to discriminate foreground pixels from background pixels two contrast gray

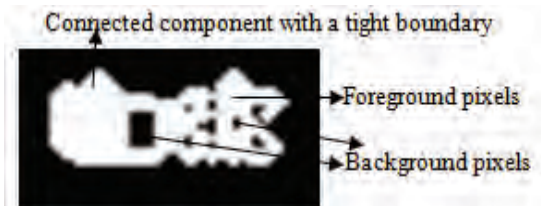


Figure 7. Illustration of foreground and background pixels in a connected component.

scale intensity values (gray value near to 0 for foreground pixels and gray value 255 for background pixels) are assigned to pixels in the output image. Irrespective of the foreground text color and background color we produced black characters on uniform white background by suitably thresholding each image segment containing text and producing the corresponding output image segment O_{bw} using the logic as given by,

$$\text{if } V_f > V_b \quad O_{bw} = \begin{cases} 10 & \text{if } I(x,y) > \text{threshold} \\ 255 & \text{if } I(x,y) \leq \text{threshold} \end{cases}$$

$$\text{if } V_f \leq V_b \quad O_{bw} = \begin{cases} 10 & \text{if } I(x,y) < \text{threshold} \\ 255 & \text{if } I(x,y) \geq \text{threshold} \end{cases}$$

Irrespective of the foreground text color and background color the extracted characters are produced in black color on uniform white background for the purpose of improving the readability of the document contents. The resulting image might contain noise in the form of false foreground. This needs reprocessing of the resulting image to further improve the readability of document contents by OCR.

2.4. Detection of Noisy Text Regions

Detection of text areas/segments that need further processing is performed using a simple method. The main idea is based on the fact that the text areas that still contain noise include more black pixels on an average in comparison to other text areas/segments. The image is divided into segments of variable sizes; each segment corresponds to one connected component. In each image segment that contains text the density of black pixels, $f(S)$ is computed. Suppose $b(S)$ is frequency of black pixels in an image segment “S” and $\text{area}(S)$ is area of image-segment “S”, the density of black pixels in “S” is given by,

$$f(S) = b(S) / \text{area}(S) \quad (3)$$

The segments that satisfy the criterion $f(S) > c * d$, are

selected for reprocessing, where “d” is the average density of black pixels of all the image segments containing text. The parameter “c” determines the sensitivity of detecting noisy text regions. High value of “c” results in less text segments to be reprocessed. Low value of “c” results in more text segments to be reprocessed which would include the text segments in which noise is already removed. Figure 8 shows the noisy areas to be reprocessed for different values of “c”. Optimal value for parameter “c” is selected based on higher character (or word) recognition rate after reprocessing noisy text regions in the output document image. We conducted experiments on the document images which we collected and synthesized. Table 2 shows the character and word recognition rates in percentage for various values of “c”. It is observed from Table 2 that character (or word) recognition rate is high for value of $c \leq 0.5$. Also it is seen from Figure 8 that number of components to be reprocessed will be less as the value of “c” increases. So 0.5 is chosen as optimal value for parameter c.

2.5. Reprocessing of Noisy Text Regions

The selected text segments containing noise in the form of false foreground pixels are reprocessed. Repeating the stage-1 on these text segments leads into text segments of smaller size. These segments are thresholded in the next stage. Since only few text segments are reprocessed instead of all the detected and verified text segments, the computation complexity of the stage-4 reduces substantially. *In fact, the entire approach can be proposed to be iterative if it is required*; but we observed that repeating stage-1 and stage-3 once on noisy regions is more than

sufficient which in turn reduces the time complexity of extracting the foreground text from complex background in document images. Figure 9 shows the results at each stage in the proposed approach.

3. Results and Discussions

3.1. Experimental Results

Since no standard corpus of document images is available for this work we created a collection of images by scanning the pages from magazines, story books of children, newspapers, decorative postal envelopes and invitation cards. In addition, one more dataset of synthesized images which are of low resolution is created by us. The details of documents in the corpus of images used for testing our proposed algorithm are depicted in Table 3. The number of document images in our datasets is 220 and they are of different resolutions (96x96 DPI, 100x100 DPI, 150x150 DPI and 200x200 DPI). The output image is obtained by depositing the black characters on the white background, irrespective of the background and foreground color in the original document. The performance of text region detection is evaluated in terms of Recall (correct detects/(correct detects + missed detects)) and Precision (correct detects/(correct detects + false alarms)). Recall is inversely proportional to missed detects whereas Precision is inversely proportional to false alarms. Missed detects indicates number of text regions incorrectly classified as non text and false alarms indicates number of non-text regions incorrectly classified as text regions. Table 4 shows the average value of precision and recall in percentage for document images in the corpus.

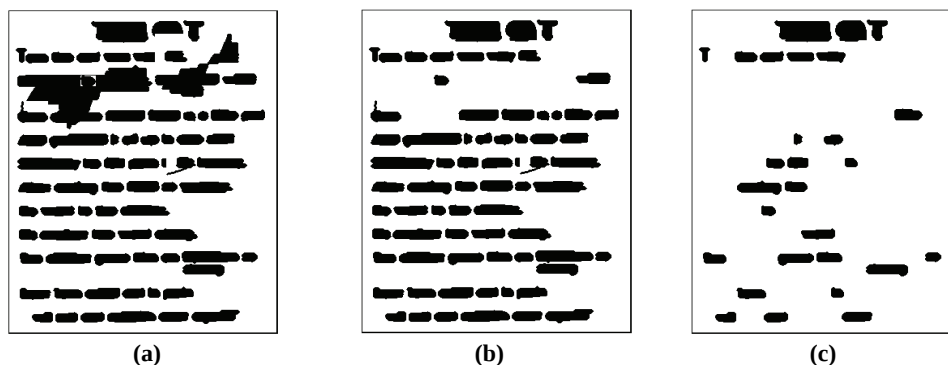


Figure 8. Noisy text segments selected based on value of c: (a) $c=0.5$ (b) $c=1.0$ (c) $c=1.5$.

Table 2. Character and word recognition rates for various values of “c”.

	$c < 0.5$	$c = 0.5$	$c = 1.0$	$c = 1.5$
Character recognition rate (%)	82.93	82.93	81.26	80.98
Word recognition rate (%)	71.33	71.33	70.10	67.38

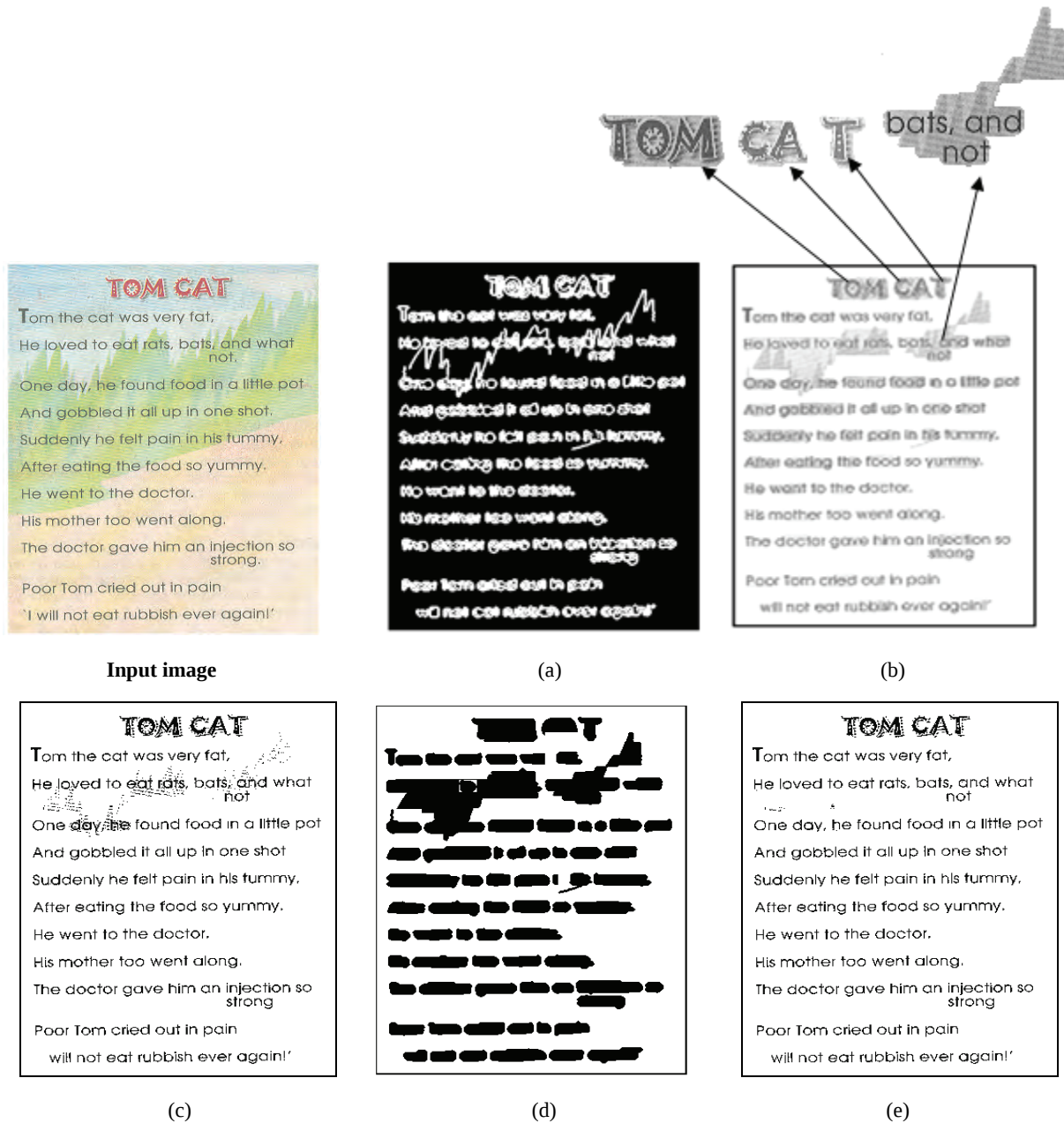


Figure 9. Result at each stage using proposed algorithm. (a) candidate text regions, (b) verified text regions, (c) extracted foreground text with noisy areas, (d) detection of noisy text segments ($c=0.5$), (e) extracted foreground text after reprocessing noisy text segments.

The proposed approach focuses on documents with English as text medium because we could quantify the performance of the improvement in the readability of document images by employing an OCR. Reading of the extracted text for documents in English as text medium is evaluated on Readiris 10.04 pro OCR. Readiris OCR converts the input document image to binary form before recognizing the characters. The Readiris 10.04 pro OCR

can tolerate a skew of 0.5 degrees on foreground text. Readability of the segmented foreground text is evaluated in terms of character and word recognition rates. OCR results for document images with printed English text are tabulated in Table 5. From Table 5 it is seen that average recognition rate at character level is higher compared to word level.

Observations from the experimental evaluation are as

Table 3. Details of document collection used for this work.

Document types	Language	Background complexity	Foreground complexity
1) Pages from Magazines	Mainly in English. Also in other languages	1) Uniform patterned background	1) Single colored and multicolored text
2) Pages from Story books of children		2) Non uniform patterned background	2) Text tilted in any orientation
3) Postal envelopes		3) Background designs from Microsoft power point	3) Text of varying sizes
4) Articles from newspapers		4) Single and multicolored background	4) Foreground with dense text and sparse text
5) Power point slides			
6) Journal cover pages			
7) Invitation cards			

Table 4. Results showing text region detection.

	Documents in English language	Documents in Kannada language	Documents in Malayalam language
Number of samples	180	30	10
Total number of characters	31784	4710	1068
Total number of words	6354	1200	317
Recall (%)	97.06	96.26	100
Precision (%)	96.78	95.1	90.23

$$\text{Character (or word) recognition rate} = \frac{\text{Number of characters (or words) correctly recognized}}{\text{Total number of characters (or words) in source document image}}$$

Table 5. OCR results for English documents.

	Average Recognition Rates (%)		
	Original document	Processed document	After further processing the noisy areas in the processed document
Character level	42.99	80.31	82.93
Word level	36.47	67.55	71.33

follows:

- For some document images the readability by the OCR without using our approach is 100% and the same is maintained even after applying our approach. [The proposed approach has not deteriorated the readability!].
- For rest of the documents due to high complexity of the background the readability through OCR is very low or even nil. After applying our approach the readability of document contents by OCR is improved to nearly 100%.

From Table 5 it is evident that the word and character recognition rates are enhanced after applying our approach. Further, it can be noted that readability is further improved after reprocessing the noisy areas in the output document images.

Many times the text lines in a document are tilted /rotated as an attempt to make the contents of the document more attractive. We extended our approach to extract the foreground text in document images with text lines tilted in any orientation. From experimental results it is evident that dilation of the edge image "E" in horizontal and vertical direction is sufficient to identify the text regions in document images having tilted text lines. Sample document images with foreground text lines tilted in any orientation and the corresponding results are

shown in Figure 10.

For documents with English as medium of text, we were able to quantify the enhanced readability through OCR and for documents in other languages we verified the extracted foreground text by visual inspection of output images, which indicates successful segmentation of foreground text from complex background. Figure 11 shows results of the proposed approach for documents in Malayalam and Kannada languages.

3.2. Discussions

Results of the proposed approach are compared with results of some existing methods of foreground separation in document images [6,9,13]. For a sample text rich document image and sparse text document image the output images obtained from the proposed method and other methods [6,9,13] are shown in Figure 12(a) and Figure 12(b) respectively.

From visual inspection of the results shown in Figure 12(a) and Figure 12(b) it is observed that, Niblack method fails to separate the foreground from complex background. Kasar method resulted in loss of foreground text information. Even though Sauvola method extracted foreground text, it introduced lot of noise compared to proposed method.



Copyright © 2010 SciRes

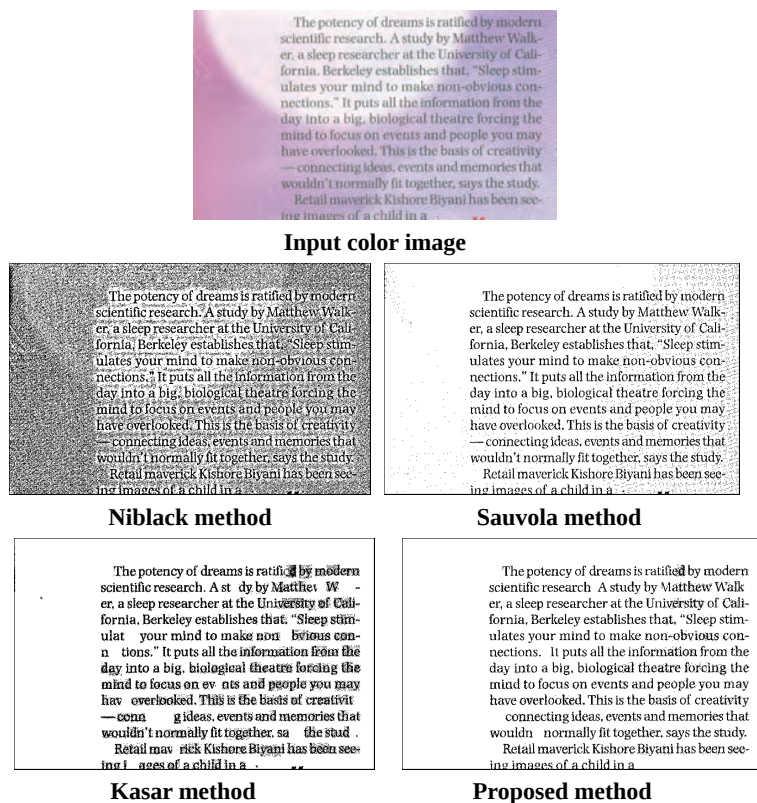


Figure 12(a). Comparison of results of foreground text separation from complex background in text rich document image.

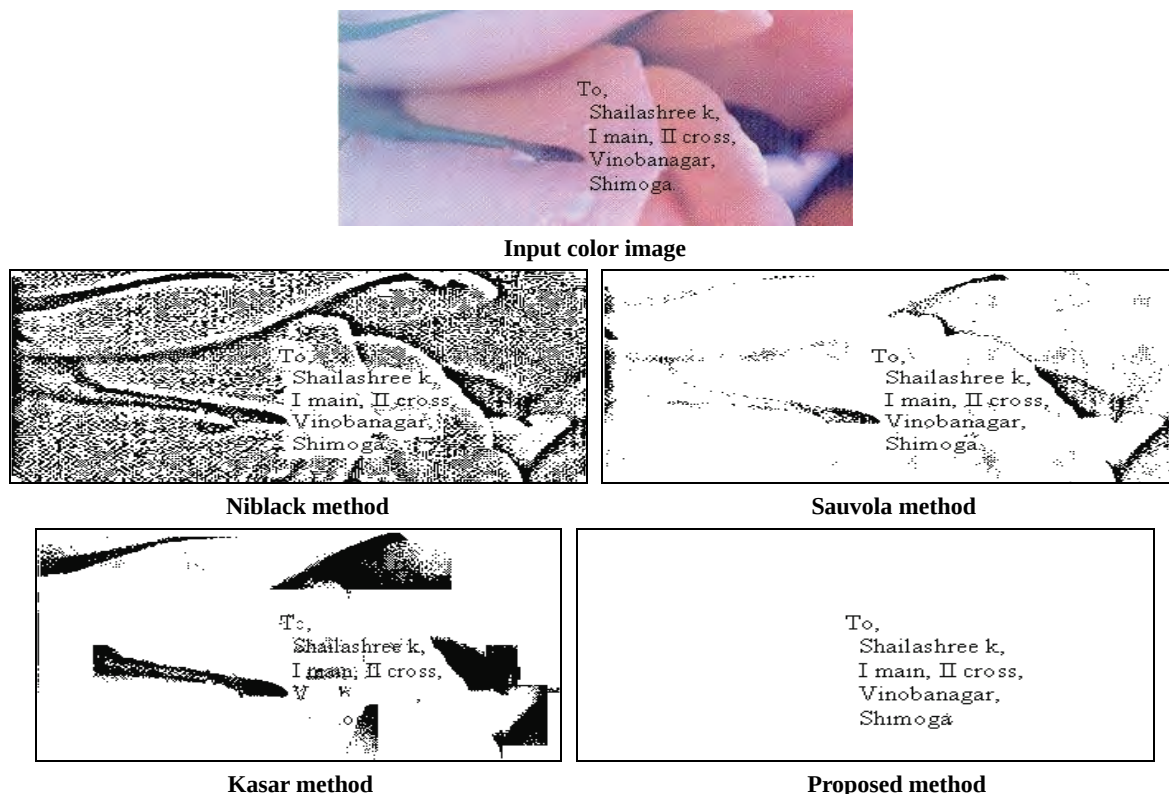


Figure 12(b). Comparison of results of foreground text separation from complex background in postal document image.

We created 20 ground truth images by selecting complex textures from Microsoft power point designs. On one set of 10 different backgrounds with varying complexities the same textual content of 540 characters is superimposed. Similarly on another set of 10 images created, each is superimposed with the same postal address shown in Figure 12(b). The outcome of the experiments on these ground truth images is shown in Table 6. As the output image produced by Niblack is too noisy the amount of characters recognized by OCR is very less. Kasar method fails to detect the foreground characters in document images having textured background which resulted in loss of text information present in the foreground of the input document image. This leads to very low character recognition accuracy by OCR. As the output images produced by Niblack, Sauvola and Kasar methods are noisier compared to the proposed method the amount of characters recognized by OCR is low which is evident from Table 6. These existing methods do not perform well when documents have textured/patterned background. This drawback is overcome by the proposed method. Yet in another experiment, a set of 10 typical document images from the corpus were tested with Niblack, Sauvola, Kasar and proposed method. The readability of extracted foreground text is evaluated on Readiris pro 10.04 OCR. The number of

characters recognized by OCR is described in Table 7.

Our approach successfully separates the foreground in document images which are of low resolution and free from degradations such as blur, uneven lighting, and wavy patterned text. From Table 7 it is evident that our approach performs well for complex background color document images compared to the methods [6,9,13] and leads to higher character recognition accuracy through OCR.

One advantage of proposed method over the existing conventional approaches is it successfully extracts the foreground text without a prior knowledge of foreground and background polarities. Another advantage over existing methods is it is less expensive as it detects the image segments containing text and extracts the text from detected text segments without using the color information. The approach is independent of medium of foreground text as it works on edge information.

4. Time Complexity Analysis

Suppose size of the document image is $M \times N$. Accordingly the size of RGB color image is $3 \times M \times N$. So the total number of pixels in input color image I is $3 \times M \times N$. The time complexity of the proposed algorithm in order notation is $O(N^2)$ if $M=N$. For the purpose of profiling

Table 6. Details of Foreground text extraction results on ground truth images by OCR.

Image type	Number of characters	Niblack method	Sauvola method	Kasar method	Proposed method
		CRR (%)	CRR (%)	CRR (%)	CRR (%)
Text rich document	540	27.89	89.63	76.33	98.53
Postal document	50	8.60	58.00	27.00	83.00

CRR–Average Character Recognition Rate when output image is OCRed.

Table 7. OCR based recognition of characters: Details for 10 test images with complex background.

Source of the document image	Number of characters	Niblack method	Sauvola method	Kasar method	Proposed method
		NCR (%)	NCR (%)	NCR (%)	NCR (%)
News paper	483	0	13.66	70.39	97.52
News paper	300	20	99.66	94.00	98.33
Magazine	144	0	0	0	92.36
Invitation card	748	0	98.93	95.45	95.32
Story book	300	0	96.66	44.66	98.66
Story book	440	0	99.09	51.59	99.55
Story book	139	0	97.12	73.38	100
Synthesized image	398	0	0	3.52	93.72
Postal doc.	47	0	0	51.06	100
Postal doc.	50	0	0	0	82
Average	305	2	50.51	48.41	95.75

NCR–Number of characters recognized by OCR.

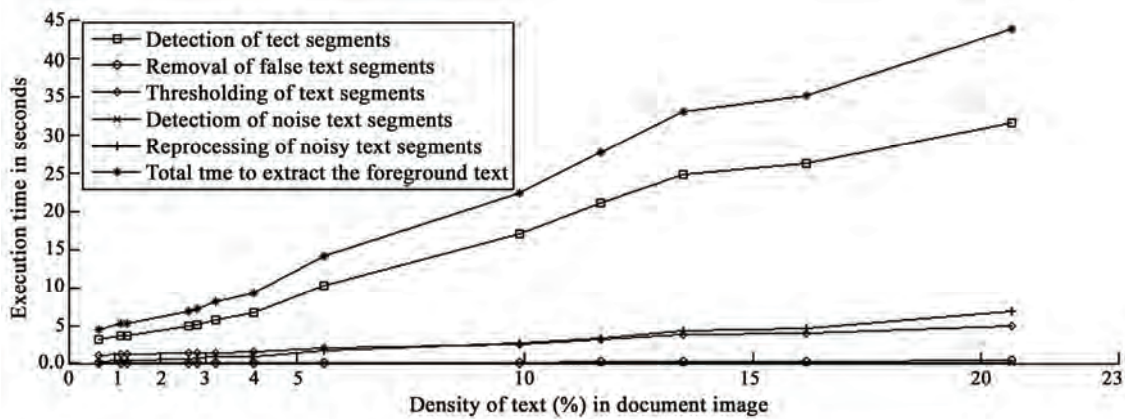


Figure 13. Plot showing execution time of each stage in proposed approach.

we have kept the size of all the test document images uniform (350x600 pixels). Postal document images contain text in sparse i.e., only printed postal address, on an average of 45 characters. Text rich document images contain text in dense, on an average of 200 characters. The algorithm was executed on a Intel(R) Core(TM) 2 Duo CPU, 2.20GHz, 1GB RAM. It is observed that total time needed to extract the foreground text is high for text rich document images compared to sparse text document images. Time needed to process the document depends on the amount of text in the foreground of the source document image. The total time of the entire process includes the time of I/O operations of document images also. Figure 13 shows the plot of execution time of each stage of the proposed approach for document images with varying density of textual information.

5. Conclusions and Future Work

In this paper a hybrid approach is presented for extraction of foreground text from complex color document images. The proposed approach combines connected component analysis and texture feature analysis to detect the segments of image containing text. An unsupervised local thresholding method is used to extract the foreground text in segments of image containing textual information. A simple and computationally less expensive method for texture analysis of image segments is proposed for reduction of false text regions. We have not used color information in extracting the text since in the first stage of our approach the evidence of all textual edges comes from intensity values in each color channel (RGB model) and this makes computations inexpensive. Threshold value to separate the foreground text is derived from the image data and does not need any manual tuning. The proposed algorithm detects on an average 97.12% of text regions in source document image. The shortcomings of the proposed approach are 1) it fails to separate the foreground text when the contrast between

the foreground and background is very poor 2) it fails to detect single letter word which neither contains a hole nor creates a hole by the dilation.

The algorithm has so far been tested on text dominant documents only which are scanned from news papers, magazines, story books of children, postal envelopes. Also we tested the proposed approach on synthesized images. The behavior of the documents containing graphic objects in foreground is considered as future extension of the present work. Design of post processing steps to recover the missed single character words without holes is another future direction of the current study.

6. References

- [1] A. K. Jain and S. K. Bhattacharjee, "Address block location on envelopes using Gabor filters," *Pattern Recognition*, Vol. 25, No 12, pp. 1459–1477, 1992.
- [2] V. Wu, R. Manmatha, and E. M. Riseman, "Textfinder: An automatic system to detect and recognize text in images," *IEEE PAMI*, Vol. 21, No. 11, pp. 1224–1229, 1999.
- [3] D. Chen, H. Bourland, and J. P. Thiran, "Text identification in complex background using SVM," *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pp. 621–626, 2001.
- [4] U. Garain, T. Paquet, and L. Heutte, "On foreground-background separation in low quality document images," *International Journal of Document Analysis and Recognition*, Vol. 8, No. 1, pp. 47–63, 2006.
- [5] H. Hase, M. Yoneda, S. Tokai, J. Kato, and C. Y. Suen, "Color segmentation for text extraction," *IJDAR*, Vol. 6, No. 4, pp. 271–284, 2003.
- [6] T. Kasar, J. Kumar, and A. G. Ramakrishnan, "Font and background color independent text binarization," *Proceedings of 2nd International Workshop on Camera Based Document Analysis and Recognition*, pp. 3–9, 2007.

- [7] E. Kavallieratou and E. Stamatatos, "Improving the quality of degraded document images," Proceedings of 2nd International Conference on Document Image Analysis for Libraries, pp. 340–349, 2006.
- [8] G. Leedham, Y. Chen, K. Takru, J. H. N. Tan, and L. Mian, "Comparison of some thresholding algorithms for text/background segmentation in difficult document images," Proceedings of 7th International Conference on Document Analysis and Recognition, pp. 859–864, 2003.
- [9] W. Niblack, "An introduction to image processing," Prentice Hall, Englewood Cliffs, 1986.
- [10] S. Nirmala, P. Nagabhushan, "Isolation of foreground-text in document images having known complex background," Proceedings of 2nd International Conference on Cognition and Recognition, pp. 99–106, 2008.
- [11] N. Otsu, "A threshold selection method from gray level histograms," IEEE Transactions on Systems, Man & Cybernetics, Vol. 9, No. 1, pp. 62–66, 1979.
- [12] M. Pietikäinen and O. Okun, "Text extraction from grey scale page images by simple edge detectors," Proceedings of the 12th Scandinavian Conference on Image Analysis (SCIA), pp. 628–635, 2001.
- [13] J. Sauvola and M. Pietikäinen, "Adaptive document image binarization," Pattern Recognition, Vol. 33, No. 2, pp. 225–236, 2000.
- [14] M. Sezgin and B. Sankur, "Survey over image thresholding techniques and quantitative performance evaluation," Journal of Electronic Imaging, Vol. 13, No. 1, pp. 146–165, 2004.
- [15] K. Sobottka, H. Kronenberg, T. Perroud, and H. Bunke, "Text extraction from colored book and journal covers," IJDAR, Vol. 2, No. 4, pp. 163–176, 1999.
- [16] C. L. Tan and Q. Yaun, "Text extraction from gray scale document image using edge information," Sixth International Conference on Document Analysis and Recognition, pp. 302–306, 2001.
- [17] O. D. Trier and A. K. Jain, "Goal directed evaluation of binarization methods," IEEE PAMI, Vol. 17, No. 12, pp. 1191–1201, 1995.
- [18] Y. Zhong, K. Karu, and A. K. Jain, "Locating text in complex color images," Pattern Recognition, Vol. 28, No. 10, pp. 1523–1536, 1995.
- [19] K. I. Kim, K. Jung, and H. J. Kim, "Texture based approach for text detection in images using support vector machines and continuously adaptive mean shift algorithm," IEEE Transactions on Pattern Analysis and Machine Intelligence, Vol. 25, No. 12, pp. 1631–1639, 2003.
- [20] Y. Liu, S. Goto, and T. Ikenaga, "A robust algorithm for text detection in color images," Proceedings of Eighth International Conference on Document Analysis and Recognition, pp. 399–403, 2005.
- [21] Z. Wang, Q. Li, S. Zhong, and S. He, "Fast adaptive threshold for Canny edge detector," Proceedings of the SPIE, pp. 501–508, 2005.

Existence and Uniqueness of the Optimal Control in Hilbert Spaces for a Class of Linear Systems

Mihai Popescu

*Institute of Mathematical Statistics and Applied Mathematics of the Romanian Academy
Bucharest, Romania
Email: ima.popescu@yahoo.com*

Abstract

We analyze the existence and uniqueness of the optimal control for a class of exactly controllable linear systems. We are interested in the minimization of time, energy and final manifold in transfer problems. The state variables space X and, respectively, the control variables space U , are considered to be Hilbert spaces. The linear operator $T(t)$ which defines the solution of the linear control system is a strong semigroup. Our analysis is based on some results from the theory of linear operators and functional analysis. The results obtained in this paper are based on the properties of linear operators and on some theorems from functional analysis.

Keywords: Existence and Uniqueness, Optimal Control, Controllable Linear Systems, Linear Operator

1. Introduction

A particular importance should be assigned to the analysis of the control of linear systems, since they represent the mathematical model for various dynamic phenomena. One of the fundamental problems is the functional optimization that defines the performance index of the dynamic product. Thus, under differential and algebraic restrictions, one determines the control corresponding to functional extremisation under consideration [1,2]. Variational calculation offers methods that are difficult to use in order to investigate the existence and uniqueness of optimal control. The method of determining the field of extremals (sweep method), that analyzes the existence of conjugated points across the optimal transfer trajectory (a sufficient optimum condition), proves to be a very efficient one in this context [3–5]. Through their resulting applications, time and energy minimization problems represent an important goal in system dynamics [6–11]. Recent results for controllable systems express the minimal energy through the controllability operator [11–13]. Also, stability conditions for systems whose energy tends to zero in infinite time are obtained in the literature [13,14]. By using linear operators in Hilbert spaces, in this study we shall analyze the existence and uniqueness of optimal control in transfer problems. The goal of this paper is to propose new methods for studying the optimal control for exactly controllable linear system. The minimization of time and energy in transfer problem is considered. The

minimization of the energy can be seen as being a particular case of the linear regulator problem in automatics. Using the adjoint system, a necessary and sufficient condition for exact controllability is established, with application to the optimization of a broad class of Mayer-type functionals.

2. Minimum Time Control

2.1. Existence

2.1.1. Problem Formulation

We consider the linear system $(\Sigma_{A,B})$:

$$\frac{d}{dt}x = Ax(t) + Bu(t), \quad x(t_0) = x_0 \in H, \quad (1)$$

where H is a Hilbert space with inner product $\langle \cdot, \cdot \rangle$ and norm $\|\cdot\|$.

$A: D(A) \subset H \rightarrow H$ is an unbounded operator on H which generates a strong semi group of operators on H , $(S(t))_{t \geq 0} = (e^{tA})_{t \geq 0}$.

$B: U \rightarrow H$ is a bounded linear operator on another Hilbert space U , for example $B \in \mathcal{L}(U, H)$.

$u: [0, \infty) \rightarrow H$ is a square integrable function representing the system control (1).

For any control function u , there exists a solution of (1) given by

$$x(t) = e^{tA}x_0 + \int_0^t e^{(t-s)A}Bu(s)ds, \quad t \geq 0 \quad (2)$$

The control problem is the following one:
Given $t_0 \in I$, $x(t_0)$, $z \in X$ and the constant $M > 0$,
let us determine $u \in U$ such that $\|u\| \leq M$ and

$$\begin{aligned} (i) \quad & x(t_1) = z \\ (ii) \quad & t_1 = \inf \{t \mid x(t) = x\} \end{aligned} \quad (3)$$

Here, X and U are Hilbert spaces.

Theorem 1. The optimal time control exists for the above formulated problem.

Proof

Let $\{t_n\}$ be a decreasing monotonous sequence such that $t_n \rightarrow t_1$. Also, let us consider the sequence $\{u_n(\tau)\}$, with $u_n = u(t_n) \in U$. We have

$$x(t_n) = S(t_n)x(t_0) + \int_{t_0}^{t_n} S(t_n - \tau)Bu_n(\tau)d\tau, \quad (4)$$

which gives

$$\begin{aligned} x(t_n) = S(t_n)x(t_0) + \int_{t_0}^{t_1} S(t_n - \tau)Bu_n(\tau)d\tau + \\ \int_{t_1}^{t_n} S(t_n - \tau)Bu_n(\tau)d\tau \end{aligned} \quad (5)$$

$S(t)$ being a continuous linear operator, it is, therefore, bounded. It follows that we have:

$$S(t_n)x(t_0) \rightarrow S(t_1)x(t_0) \quad (6)$$

$$\int_{t_1}^{t_n} S(t_n - \tau)Bu_n(\tau)d\tau \rightarrow 0 \quad (7)$$

If we consider the set of all the admissible controls

$$\Omega = \{u \in U \mid \|u\| \leq M\} \quad (8)$$

then, Ω is a weakly compact closed convex set.

As $\Omega \subset U$ is weakly compact, any sequence $(u) \in \Omega$ possesses a weakly convergent subsequence to an element $u_1 \in \Omega$. Thus, for $u_1 \in \Omega$ and any $u^* \in \Omega^*$ (the dual of Ω), we have

$$\lim_{i \rightarrow \infty} \langle u_i, u^* \rangle = \langle u_1, u^* \rangle \quad (9)$$

Since $u_1 \in \Omega$, it follows that

$$\|u_1\| \leq M \quad (10)$$

Also, we have

$$\begin{aligned} B^*: X \rightarrow U, \\ S^*(t_n - \tau)x \in U. \end{aligned} \quad (11)$$

So,

$$B^*S^*(t_n - \tau)^*x \in U. \quad (12)$$

For every $x \in X$, let us evaluate the difference

$$\begin{aligned} F = \left\langle \int_{t_1}^{t_1} S(t_n - \tau)Bu_n(\tau)d\tau, x \right\rangle - \\ \left\langle \int_{t_1}^{t_n} S(t_1 - \tau)Bu_1(\tau)d\tau, x \right\rangle \end{aligned} \quad (13)$$

By writing

$$\int_{t_1}^{t_1} S(t_k - t_0)Bd\tau = \mathcal{L}_k, \quad k = 1, n, \quad (14)$$

the Expression (13) becomes

$$\langle \mathcal{L}_n u_n, x \rangle - \langle \mathcal{L}_1 u_1, x \rangle = \langle u_n, \mathcal{L}_n^* x \rangle - \langle u_1, \mathcal{L}_1^* x \rangle \quad (15)$$

or

$$\begin{aligned} F = \langle u_n, \mathcal{L}_1^* x \rangle - \langle u_1, \mathcal{L}_1^* x \rangle - \langle u_n, \mathcal{L}_1^* x \rangle + \\ \langle u_n, \mathcal{L}_n^* x \rangle. \end{aligned} \quad (16)$$

Therefore,

$$\begin{aligned} F = \langle u_n - u_1, \mathcal{L}_1^* x \rangle + \langle u_n, (\mathcal{L}_n^* - \mathcal{L}_1^*) x \rangle = \\ \int_{t_0}^{t_1} \langle u_n(\tau) - u_1(\tau), B^*S(t_1 - \tau)^*x \rangle d\tau + \\ \int_{t_0}^{t_1} \langle u_n(\tau), B^*S(t_n - \tau)^*x - B^*S(t_1 - \tau)^*x \rangle d\tau \end{aligned} \quad (17)$$

By using the properties of the operator $S(t)$, one obtains

$$\begin{aligned} F = \int_{t_0}^{t_1} \langle u_n(\tau) - u_1(\tau), B^*S(t_1 - \tau)^*x \rangle d\tau + \\ \int_{t_0}^{t_1} \langle u_n(\tau), B^*S(t_1 - \tau)^*[S(t_n - t_1)^*x - x] \rangle d\tau. \end{aligned} \quad (18)$$

Because the sequence u_n converges weakly to u_1 , the first term in (18) tends to zero for $n \rightarrow \infty$ (see(9)). On the other hand, $S(t_1 - \tau)^*$ is a strongly continuous semigroup and, hence, from its boundedness, it follows that there exists a constant $K > 0$ such that

$$\left\| B^* S(t_1 - \tau)^* \left[S(t_n - t_1)^* x - x \right] \right\| \leq K \left\| S(t_n - t_1)^* x - x \right\| \quad (19)$$

So, it follows that the second term in (18) tends to zero as $n \rightarrow \infty$. The sequence $\{x(t_n)\} \in X$ is weakly convergent to $x(t_1) = z \in X$ if, for any $x \in X^* = X$, we have

$$\lim_{n \rightarrow \infty} \langle x(t_n), x \rangle = \langle x(t_1), x \rangle, \quad (20)$$

which gives

$$\langle z, x \rangle = \left\langle S(t_1) x(t_0) + \int_{t_0}^{t_1} S(t_1 - \tau) B u_1(\tau) d\tau, x \right\rangle \quad (21)$$

From (21), one gets

$$S(t_1) x(t_0) + \int_{t_0}^{t_1} S(t_1 - \tau) B u_1(\tau) d\tau = z \quad (22)$$

and, hence, u_1 is the optimal control.

An important result which is going to be used for proving the uniqueness belongs to A. Friedman:

Theorem 2 (bang-bang). Assuming that the set Ω is convex in a neighborhood of the origin and $u(t)$ is the control of optimal time in the problem formulated in (2.1.1), ([2, 10]) it follows that $u(t) \in \Omega$, for almost all $t \in [t_0, t_1]$.

2.2. Uniqueness of Time-Optimal Control

2.2.1. Rotund Space

Let β be the unity sphere in the Banach space U and let $\partial\beta$ be its boundary [11].

The space U is said to be rotund if the following equivalent conditions are satisfied:

- a) if $\|x_1 + x_2\| = \|x_1\| + \|x_2\|$, it follows that there exists a scalar $\lambda \neq 0$, such that $x_2 = \lambda x_1$;
- b) each convex subset $K \subset U$ has at least one element that satisfies

$$\|u\| \leq \|z\|, \quad u \in K, \quad z \in K;$$

- c) for any bounded linear functional f on U there exists at least an element $x \in \beta$ such that

$$\langle x, f \rangle = f(x) = \|f\|;$$

- d) each $x \in \partial\beta$ is a point of extreme of β .

Examples of rotund spaces:

1) Hilbert spaces.

2) Spaces l^p , L^p , $1 < p < \infty$.

3) If the Banach spaces $U_1, U_2, \dots, U_n$ are rotund, then the product space $U_1 \times U_2 \times \dots \times U_n$ is rotund, too.

4) Uniform convex spaces are rotund, but the converse implication is not true.

2.2.2. Uniqueness

We assume that u_1 and u_2 are optimal,

$u_i \in U$, $i = 1, 2$. Therefore,

$$S(t_1) x(t_0) + \int_{t_0}^{t_1} S(t_1 - \tau) B u_1(\tau) d\tau = z \quad (23)$$

$$S(t_1) x(t_0) + \int_{t_0}^{t_1} S(t_1 - \tau) B u_2(\tau) d\tau = z \quad (24)$$

By adding Equations (23) and (24), one obtains

$$S(t_1) x(t_0) + \int_{t_0}^{t_1} S(t_1 - \tau) B \frac{1}{2} (u_1(\tau) + u_2(\tau)) d\tau = z \quad (25)$$

It follows that $u_1(\tau)$, $u_2(\tau)$, $1/2(u_1(\tau) + u_2(\tau))$ are optimal. By using Theorem 2, we have

$$u_1(\tau), u_2(\tau), \frac{1}{2}(u_1(\tau) + u_2(\tau)) \in \partial\beta \Rightarrow$$

$$\|u_k\| = 1, \frac{1}{2}\|u_1 + u_2\| = 1.$$

Since the condition (i) is satisfied, U is a rotund space and $u_1 = \lambda u_2$. Thus,

$$\|u_1 + u_2\| = (1 + |\lambda|) \|u_2\| = 2 \quad (26)$$

and, therefore

$$|\lambda| = 1 \Rightarrow u_1 = u_2 \quad (27)$$

which ends the proof of the uniqueness.

3. Minimum Energy Problem

3.1. Problem Formulation

Let $\Sigma_{A,B}$ (see(1)) be a controllable system with finite dimensional state space X [6,10,12, 14].

Let $I = [0, t_1]$, $x(0) = a \in X$, $b \in X$ be given. Let U be the Hilbert space $L^2(I, U)$. The minimum norm control

problem can be formulated as follows: determine $u(t) \in U$ such that for some $t_1 \in I$,

$$1) \quad x(t_1) = b,$$

2) $\|u\|$ is minimized on $[0, t_1]$, where $\|\cdot\|$ represents the norm on $L^2(I, U)$.

Let us fix $t_1 > 0$. Consider the linear operator

$$\mathcal{L}_n : L^2(0, t_1; U) \rightarrow H_1 \quad (28)$$

defined by

$$\mathcal{L}_n u =: \int_0^{t_1} S(t_1 - s) B u(s) ds \quad (29)$$

We have

$$x(t_1) = S(t_1) a + \mathcal{L}_n u = b. \quad (30)$$

Theorem 3. Let $\mathcal{L}_t : U \rightarrow H$ be a linear mapping between a Hilbert space U and a finite dimensional space H [11]. Then, there exists a finite dimensional space $M \subset U$ such that the restriction $\mathcal{L}_t^M$ of $\mathcal{L}_t$ to M is an injective mapping.

Proof

Let $\{e_i\}_{i=1, \dots, n}$ be a basis for the range of $\mathcal{L}_t$ in H . Given any $u \in U$, $\mathcal{L}_t u$ can be written as

$$\mathcal{L}_t u = \sum_{i=1}^n \alpha_i e_i, \quad (31)$$

where each α_i can be expressed by

$$\alpha_i = \langle f_i, u \rangle \quad f_i \in U^* = U. \quad (32)$$

Then,

$$\mathcal{L}_t u = \sum_{i=1}^n \langle f_i, u \rangle e_i. \quad (33)$$

f_i are linearly independent and generate a n dimensional subspace $M \subset U$.

From the properties of Hilbert spaces (the projection theorem), it follows that $U = M \oplus M^\perp$.

Let $u \in M^\perp$. Then, $\langle f_i, u \rangle = \alpha_i = 0$. Thus,

$$M^\perp \subset \{u \in U \mid \mathcal{L}_t u = 0\} = \ker \mathcal{L}_t. \quad (34)$$

Let $u \in \ker \mathcal{L}_t$. We get

$$\mathcal{L}_t u = 0 = \sum_{i=1}^n \langle f_i, u \rangle e_i. \quad (35)$$

Since e_i are independent, $\langle f_i, u \rangle = 0, i = 1, \dots, n$ and, so, $M = \ker \mathcal{L}_t$. $\mathcal{L}_t$ maps M bijectively to the range of $\mathcal{L}_t$ and hence $\mathcal{L}_t^M$ is an injective mapping into H . Let

$$u = u_1 + u_2, \quad u_1 \in M, \quad u_2 \in M^\perp, \quad u_2 \neq 0. \quad (36)$$

Because $\mathcal{L}_t u_2 = 0$,

$$\mathcal{L}_t u = \mathcal{L}_t^M u_1 = x \in H \quad (37)$$

Since $\langle u_1, u_2 \rangle = 0$, we have

$$\|u\|^2 = \langle u_1 + u_2, u_1 + u_2 \rangle = \|u_1\|^2 + \|u_2\|^2 \quad (38)$$

Then, a unique minimum norm $(\mathcal{L}_t^M)^{-1} x$ exists and $u_1 \in \mathcal{L}_t^M x$ is the minimum norm element satisfying $\mathcal{L}_t u = x$.

Remark 1. The unique solution $u(t)$ of the equation $\mathcal{L} u = x$, with minimum norm control, is the projection of $u(t)$ onto the closed subspace $M = (f_1, \dots, f_n)$.

Hence, from (30), there exists a control $u(\cdot) \in L^2(0, t_1; U)$ transferring a to b in time t_1 if and only if $(b - S(t_1)a) \in \text{Im } \mathcal{L}_{t_1}$. The control which achieves this and minimizes the functional

$$E_{t_1} = \int_0^{t_1} \|u(s)\|^2 ds, \quad \text{called the energy functional, is}$$

$$u = \mathcal{L}_{t_1}^{-1}(b - S(t_1)a) \quad (39)$$

Define the linear operator

$$Q_t = \int_0^t S(r) B B^* S(r) dr, \quad t \geq 0 \quad (40)$$

We have the following results (see [10–12]):

Proposition 1.

1) The function $Q_t, t \geq 0$, is the unique solution of the Equation [1,13]

$$\frac{d}{dt} \langle Q_t x, x \rangle = 2 \langle Q_t^* x, x \rangle + \|B^* x\|^2 \quad (41)$$

$$x \in D(A^*), Q_0 = I$$

where

$$D(A^*) = \left\{ y \in H \mid \exists C \in \mathbb{R}^+, \left| \langle A x, y \rangle_H \right| \leq C \|x\|_H, \forall x \in D(A) \right\}. \quad (42)$$

2) If A generates a stable semigroup, then

$$\lim_{t \rightarrow \infty} Q_t = Q \quad (43)$$

exists and is the unique solution of the equation

$$2\langle QA^*x, x \rangle + \|B^*x\|^2 = 0, \quad x \in D(A^*), \quad (44)$$

The proof of Proposition 1 is given in [10,11].

The following theorem gives general results for the functionals $E_\eta(a, b)$, the minimal energy for transferring a to b in time t_1 , and $E_\infty(0, b)$, where $a, b \in H$, $t_1 > 0$ (see [12]).

Theorem 4.

1) For arbitrary $t_1 > 0$ and $a, b \in H$

$$E_\eta(a, b) = \left\| \left(Q_\eta^{1/2} \right)^{-1} (S(t_1)a - b) \right\|^2. \quad (45)$$

2) If $S(t)$ is stable and the system $(\Sigma_{A,B})$ is null controllable in time $t_0 > 0$, then

$$E_\infty(0, b) = \left\| \left(Q^{1/2} \right)^{-1} b \right\|^2, \quad b \in H \quad (46)$$

3) Moreover, there exists $C_{t_0} > 0$ such that

$$\begin{aligned} \left\| \left(Q^{1/2} \right)^{-1} b \right\|^2 &\leq E_{t_1}(0, b) \leq \\ C_{t_0} \cdot \left\| \left(Q^{1/2} \right)^{-1} b \right\|^2, \quad b \in H, \quad t_1 \geq t_0 \end{aligned} \quad (47)$$

4. Numerical Methods for Minimal Norm Control

4.1. Presentation of the Numerical Methods

Let $\Sigma_{A,B}$ be a dynamic system with $U = R^m$, $X = R_n$, $I = R^1$ [9,10,13]. Given $x_0 = 0$, t_0 , t_1 , b , determine $u(t) \in U$ such that

- 1) $x(t_1) = b$,
- 2) $\|u_p\|$ is minimized for $p \in [1, \infty)$

Now,

$$\begin{aligned} e = b &= \int_{t_0}^{t_1} S(t_1 - \tau) B u(\tau) d\tau = \\ &\int_{t_0}^{t_1} \Phi(t_1 - \tau) u(\tau) d\tau \end{aligned} \quad (48)$$

In order to make Equation (48) true, $u(t)$ must be chosen on the interval $[t_0, t_1]$. Consider the i -th component e_i of the vector e :

$$e_i = \int_{t_0}^{t_1} \varphi_i(t_1 - \tau) u(\tau) d\tau = \langle \varphi_i, u \rangle = f(\varphi_i) \quad (49)$$

where φ_i is the row of the matrix Φ and f is the unique functional corresponding to the inner product (Riesz representation theorem).

Let λ_i be an arbitrary scalar. Then,

$$\lambda_i e_i = \lambda_i f(\varphi_i) = f(\lambda_i \varphi_i), \quad (50)$$

Since R^n is a Hilbert space, the inner product

$$\sum_{i=1}^n \lambda_i e_i = \sum_{i=1}^n f(\lambda_i \varphi_i) = \langle \lambda, e \rangle \quad (51)$$

λ being a vector with arbitrary components $\lambda_1, \dots, \lambda_n$. If Equation (51) is true for at least n different linearly independent, $\lambda_i, i = 1, \dots, n$ then

Equation (48) is also true.

We have

$$\left| f \left(\sum_{i=1}^n \lambda_i \varphi_i \right) \right| = \left| \left\langle \sum_{i=1}^n \lambda_i \varphi_i, u \right\rangle \right|. \quad (52)$$

By Hölder's inequality,

$$\begin{aligned} \left| \left\langle \sum_{i=1}^n \lambda_i \varphi_i, u \right\rangle \right| &\leq \|u\|_p \left\| \sum_{i=1}^n \lambda_i \varphi_i \right\|_q, \\ \left(\frac{1}{p} + \frac{1}{q} = 1 \right). \end{aligned} \quad (53)$$

From Equations (51) and (53),

$$\begin{aligned} \|u\|_p &\geq \frac{\left| f \left(\sum_{i=1}^n \lambda_i \varphi_i \right) \right|}{\left\| \sum_{i=1}^n \lambda_i \varphi_i \right\|_q} = \frac{\sum_{i=1}^n \lambda_i e_i}{\left\| \sum_{i=1}^n \lambda_i \varphi_i \right\|_q} = \\ &\frac{\langle \lambda, e \rangle}{\left\| \sum_{i=1}^n \lambda_i \varphi_i \right\|_q}. \end{aligned} \quad (54)$$

Every control driving the system to the point b must satisfy Equation (54), while for the optimal (minimum norm) control, Equation (54) must be satisfied with equality (Theorem 2). Numerically, one must search for a n vectors λ such that the right-hand side of Equa-

tion (54) takes its maximum.

4.2. A Simple Example

We consider a single output linear dynamic system $\Sigma_{A,B}$, where

$$U = R^1, \quad X = R^1, \quad I = R^1$$

Because $t_1 \in I$ fixed, then $S(t_1 - \tau)B = \Phi(\tau)$

$$e = b = \int_{t_0}^{t_1} \Phi(\tau) u(\tau) d\tau. \quad (55)$$

Therefore,

$$|e| = b = \left| \int_{t_0}^{t_1} \Phi(\tau) u(\tau) d\tau \right| \leq \int_{t_0}^{t_1} |\Phi(\tau) u(\tau)| d\tau \quad (56)$$

and from Hölder's inequality,

$$\int_{t_0}^{t_1} |\Phi(\tau) u(\tau)| d\tau \leq \|\Phi\| \|u\|, \quad (57)$$

where we have assumed that $\Phi(t) \in L^2(I)$. The minimum norm control, if exists, will satisfy

$$\|u\| = \frac{|e|}{\|\Phi\|}. \quad (58)$$

From (56) and from (57), we have

$$\text{sign}(u(t)) = \text{sign}(\Phi(t)), \quad \forall t \in [t_0, t_1] \quad (59)$$

and, respectively,

$$|\Phi(t)|^q = h |u(t)|^p, \quad (60)$$

$$\forall t \in [t_0, t_1],$$

h arbitrary constant.

The control u satisfies the relation

$$|u(t)| = h^{-1/p} |\Phi(t)|^{q/p} = k |\Phi(t)|^{q/p} \quad (61)$$

or

$$u(t) \text{sign}(u(t)) = k |\Phi(t)|^{q/p}. \quad (62)$$

From condition (59),

$$u(t) = k \text{sign}(\Phi(t)) |\Phi(t)|^{q/p}. \quad (63)$$

By substituting Equation (63) into Equation (55), the constant k can be determined.

4.2.1. The Particular Case $p = q = 2$

In this case, the Equation (55) represents the inner product in $L^2(t_0, t_1; U)$,

$$b = \langle \Phi, u \rangle \quad (64)$$

The relation (63) becomes

$$u(t) = k \text{sign}(\Phi(t)) |\Phi(t)|. \quad (65)$$

Then,

$$b = \int_{t_0}^{t_1} k \Phi^2(\tau) d\tau = k \|\Phi\|^2. \quad (66)$$

Thus,

$$k = \frac{b}{\|\Phi\|^2} \quad (67)$$

and the optimal control is given by

$$u(t) = \frac{b \Phi(t)}{\|\Phi\|^2}. \quad (68)$$

5. Exact Controllability

5.1. Adjoint System

Let $S(t), t \in [0, t_1]$, be the fundamental solution of an homogeneous system associated to the linear control system $\Sigma_{A,B}$ (see(1)) [4,5,13].

Thus, we have

$$\frac{dS(t)}{dt} = AS(t), \quad S(0) = I, \quad t \in [0, t_1] \quad (69)$$

From

$$\frac{d}{dt} (S(t) S^{-1}(t)) = \frac{dI}{dt} = 0. \quad (70)$$

one obtains

$$\frac{d}{dt} S^{-1}(t) = S^{-1}(t) A. \quad (71)$$

which implies that

$$(S^{-1}(t))^* = -A^* (S^{-1}(t))^*. \quad (72)$$

Since

$$(S^{-1}(t))^* = (S^*(t))^{-1}. \quad (73)$$

the system becomes

$$\frac{d}{dt} [S^*(t)^{-1}] = -A^* (S^*(t))^{-1}, \quad t \in [0, t_1] \quad (74)$$

It follows that $(S^*(t))^{-1}, t \in [0, t_1]$, is the fundamental solution for the adjoint System (69).

5.2. A Class of Linear Controllable Systems

We consider the class of linear control systems in vectorial form [1,3,4,8,13,15]

$$\frac{dx}{dt} = Ax + Bu + C, \quad x(0) = x_0 \quad (75)$$

Let u_0 be any internal point of the closed bounded convex control domain U .

This domain contains the space E_m of variables $u = (u_1, \dots, u_m)$. We take

$$\bar{u} = u - u^\circ \quad (76)$$

By this transformation, system (75) becomes

$$\frac{dx}{dt} = Ax + B\bar{u} + (Bu^\circ + C). \quad (77)$$

Thus, one transfers the origin of the coordinates of space E_m in u° . At the same time, the origin of the coordinates of space E_m is an internal point inside the domain U . Let us denote by $x^\circ(t)$ the solution of system (75) which corresponds to the control $u \equiv 0$.

This control is admissible, as the origin of the coordinates belongs to the domain U and satisfies the initial condition $x^\circ(0) = x_0$. As a result,

$$\frac{dx^\circ(t)}{dt} = Ax^\circ(t) + C. \quad (78)$$

Let $u_1(t)$, $0 \leq t \leq t_1$, be any control and $x_1(t)$ be the trajectory corresponding to system (75).

We have

$$\frac{dx_1(t)}{dt} = Ax_1(t) + Bu_1(t) + C, \quad x_1(0) = x_0. \quad (79)$$

We take

$$\bar{x}(t) = x_1(t) - x^\circ(t) \quad (80)$$

Then, from (78) and (79), one obtains

$$\frac{d\bar{x}(t)}{dt} = A\bar{x}(t) + Bu_1(t).$$

Hence, for $u = u_1(t)$, the system (75) becomes

$$\Sigma_{A,B} : \begin{cases} \dot{x} = Ax + Bu \\ x(0) = 0 \end{cases} \quad x \in R^n, \quad (81)$$

Thus, the linear control system belongs to the class $\Sigma_{A,B}$ defined by (81).

5.3. Optimal Control Problems

For a given t_1 , let us determine the optimal control $\tilde{u}$

that extremises the functional of final values

$$J = F(x(t_1)) = \sum_{i=1}^n c_i x_i(t_1) \quad (82)$$

and satisfies the differential constraints represented by the system (81).

The adjoint variable $y \in R$ satisfies equation

$$\dot{y} = -\frac{\partial H}{\partial x} \quad (83)$$

where H is the Hamiltonian associated to the optimal problem

$$H = \langle y, \dot{x} \rangle \quad (84)$$

Since the final conditions $x(t_1)$ are free and the final time has been indicated from the condition of transversality, one obtains $y(t_1) = y_{t_1}$.

It follows that the adjunct system becomes

$$\Sigma_{A,B}^* : \begin{cases} \dot{y} = -A^* y \\ y(t_1) = y_{t_1} \end{cases} \quad (85)$$

with solution

$$y(t) = S^*(t_1 - t)y_{t_1}, \quad y_{t_1} \in H. \quad (86)$$

Proposition 2.

For the class of optimum problems under consideration, the following identity holds true:

$$\langle x_{t_1}, y_{t_1} \rangle_H = \int_0^{t_1} \langle u, B^* y \rangle_U dt \quad (87)$$

Proof

Assuming that $u \in C^1([0, t_1], U)$ and $y(t_1) \in D(A^*)$ it follows that $x, y \in C^1([0, t_1], U)$.

Integrating by parts, since $x(0) = 0$, $y(t_1) = y_{t_1}$, $B: U \rightarrow H$, $B^*: H \rightarrow U$, we have

$$\begin{aligned} 0 &= \int_0^{t_1} \langle \dot{x} - Ax - Bu, y \rangle_H dt = \\ &= \int_0^{t_1} \langle \dot{x}, y \rangle_H dt - \int_0^{t_1} \langle Ax, y \rangle_H dt - \\ &= \int_0^{t_1} \langle u, B^* y \rangle_U dt = \langle x, y \rangle \Big|_0^{t_1} - \\ &= \int_0^{t_1} \langle x, -A^* y \rangle_H dt - \int_0^{t_1} \langle x, A^* y \rangle_H dt - \\ &= \int_0^{t_1} \langle u, B^* y \rangle_U dt \end{aligned} \quad (88)$$

One obtains

$$\langle x_{t_1}, y_{t_1} \rangle_H - \int_0^{t_1} \langle u, B^* y \rangle_U dt = 0 \quad (89)$$

The identity has been demonstrated.

This result can be extended for arbitrary $u \in L^2(0, t_1; U)$ and $y(t_1) \in H$.

An important result referring to the exact controllability of the linear system (81) is stated in

Theorem 5. The system $\Sigma_{A,B}$ is exactly controllable if only if the following condition is satisfied

$$\int_0^{t_1} \|B^* S^*(t) y_0\|_U^2 dt \geq c \|y_0\|_U^2, \quad (90)$$

$$\forall y_0 \in H$$

Proof

" $\Rightarrow$ " We assume that $\Sigma_{A,B}$ is exactly controllable.

Let $u \in L^2(0, t_1; U)$ and $y(t_1) \in H$.

We consider that the application

$$L_{t_1} : U \rightarrow x(T) \quad (91)$$

is well defined.

Let $\Lambda : H \rightarrow L^2(0, t_1; U)$ be the inverse of L_{t_1} . From Theorem 3 it follows that there exists a finite dimensional subspace $M \subset H$, $M^\perp = \ker L_{t_1}$ such that the restriction

$$(L_{t_1})_M = (L_{t_1}) \Big|_{(\ker L_{t_1})^\perp} \quad (92)$$

is injective.

Since $L_{t_1} u = x(t_1) \Rightarrow u = L_{t_1}^{-1}(x(t_1)) = \Lambda(x(t_1))$ it follows that transfers 0 in $x(t_1)$ for the system $\Sigma_{A,B}$.

We choose $y(t_1) \in H$ and $x(t_1) = y(t_1)$, $u = \Lambda(x(t_1))$. It follows that

$$\|y_{t_1}\|_H^2 = \langle y_{t_1}, y_{t_1} \rangle = \int_0^{t_1} \langle \Lambda(x_{t_1}), B^* y \rangle_U dt \quad (93)$$

For $\Lambda(x_{t_1}) = u \in L^2$, $B^* y \in L^2 B$, using Hölders's inequality, one obtains

$$\left| \int_0^{t_1} (\Lambda(x_{t_1}))(B^* y) dt \right| \leq \int_0^{t_1} |\Lambda(x_{t_1})(B^* y)| dt \leq$$

$$\left| \int_0^{t_1} |\Lambda(x_{t_1})|^2 dt \right|^{1/2} \cdot \left| \int_0^{t_1} |B^* y|^2 dt \right|^{1/2} \quad (94)$$

From (93) and (94), for $x_{t_1} = y_{t_1}$, we have

$$\|y_{t_1}\|_H^2 \leq \|\Lambda(y_{t_1})\| \left(\int_0^{t_1} |B^* y|_U^2 dt \right)^{1/2} \leq$$

$$\|\Lambda\| \|y_{t_1}\|_H \left(\int_0^{t_1} |B^* y|_U^2 dt \right)^{1/2} \quad (95)$$

or

$$\int_0^{t_1} |B^* y|_U^2 dt \geq \frac{1}{\|\Lambda\|^2} \|y_{t_1}\|_H^2 \quad (96)$$

By changing $t \rightarrow t_1 - t \Rightarrow y_0$ becomes y_{t_1} and $S^*(t)$ is transformed into $S^*(t_1 - t)$.

Therefore, we get one equivalent relation of controllable system in which y_{t_1} substitutes y_0 .

$$\int_0^{t_1} |B^* S^*(t) y_0|_U^2 dt = \int_0^{t_1} |B^* S^*(t_1 - t) y_0|_U^2 dt \geq$$

$$\frac{1}{\|\Lambda\|^2} \|y_{t_1}\|_H^2 = c \|y_{t_1}\|_H^2 \quad (97)$$

The direct implication has been demonstrated.

" $\Leftarrow$ " We assume that condition (90) is fulfilled.

Then

$$\int_0^{t_1} |B^* y|_U^2 dt \geq c \|y_{t_1}\|_H^2 \quad (98)$$

For any $y_{t_1} \in H$ we take the set $\{u(t) = B^* y(t)\}$ ($y(t)$ is solution of $\Sigma_{A,B}^*$).

We consider the solution $x(t)$ for $\Sigma_{A,B}$ which corresponds to the above mentioned control $u(t)$. Let us define the bounded operator

$$\Gamma : y_{t_1} \in H \rightarrow x(t_1) = L_{t_1}(B^* y(\cdot)) \in H \quad (99)$$

Since $u(t) = B^* y$, the identity (89) becomes

$$\langle \Gamma y_{t_1}, y_{t_1} \rangle_H = \int_0^{t_1} |B^* y(t)|_U^2 dt \geq c \|y_{t_1}\|_H^2 \quad (100)$$

Therefore, there exists a constant $c > 0$ for which (100) is satisfied. It follows that $\langle \Gamma y_{t_1}, y_{t_1} \rangle_H$ is positively defined.

This resumes the conclusion that the system $\Sigma_{A,B}$ is controllable.

Thus, since Γ is inversable, for any $x_{t_1} \in H$, state $y_{t_1} = \Gamma^{-1}(x_{t_1})$ is such that there exists a control $u(t) = B^* y$ which transfers 0 in y_{t_1} . The theorem has been demonstrated.

6. Conclusions

The main contribution in this paper consists in proving existence and uniqueness results for the optimal control in time and energy minimization problem for controllable linear systems.

A necessary and sufficient condition of exact controllability in optimal transfer problem is obtained. Also, a numerical method for evaluating the minimal energy is presented.

In Subsection 5.2 the nonhomogeneous linear control systems are transformed into linear homogeneous ones, with a null initial condition. For the control of such a class of systems, one needs to consider the adjoint system corresponding to the associated optimal transfer problem (Subsections 5.1 and 5.3).

The above theory can be used for solving various problems in spatial dynamics (rendez-vous spatial, satellite dynamics, space pursuing), [3,4,8,9]. Also, this theory can be successfully applied to automatics, robotics and artificial intelligence problems [16–18] modelled by linear control systems.

7. References

- [1] Q. W. Olbrot and L. Pandolfi, "Null controllability of a class of functional differential systems," *International Journal of Control*, Vol. 47, pp. 193–208, 1988.
- [2] H. J. Sussmann, "Nonlinear controllability and optimal control," Marcel Dekker, New York, 1990.
- [3] M. Popescu, "Variational transitory processes and nonlinear analysis in optimal control," Technical Education Bucharest, 2007.
- [4] M. Popescu, "Sweep method in analysis optimal control for rendezvous problems," *Journal of Applied Mathematics and Computing*, Vol. 23, 1–2, pp. 249–256, 2007.
- [5] M. Popescu, "Optimal control in Hilbert space applied in orbital rendezvous problems," *Advances in Mathematical Problems in Engineering Aerospace and Science*; (ed. Sivasundaram), Cambridge Scientific Publishers 2, pp. 135–143, 2008.
- [6] S. Chen and I. Lasiecka, "Feedback exact null controllability for unbounded control problems in Hilbert space," *Journal of Optimization Theory and Application*, Vol. 74, pp. 191–219, 1992.
- [7] M. Popescu, "On minimum quadratic functional control of affine nonlinear control," *Nonlinear Analysis*, Vol. 56, pp. 1165–1173, 2004.
- [8] M. Popescu, "Linear and nonlinear analysis for optimal pursuit in space," *Advances in Mathematical Problems in Engineering Aerospace and Science*; (ed. Sivasundaram), Cambridge Scientific Publishers 2, pp. 107–126, 2008.
- [9] M. Popescu, "Optimal control in Hilbert space applied in orbital rendezvous problems," *Advances in Mathematical Problems in Engineering Aerospace and Science*; (ed. Sivasundaram), Cambridge Scientific Publishers 2, pp. 135–143, 2008.
- [10] M. Popescu, "Minimum energy for controllable nonlinear and linear systems," *Seminaire Theorie et Control Universite Savoie, France*, 2009.
- [11] M. Popescu, "Stability and stabilization dynamical systems," Technical Education Bucharest, 2009.
- [12] G. Da Prato, A. J. Pritchard, and J. Zabczyk, "On minimum energy problems," *SIAM J. Control and Optimization*, Vol. 29, pp. 209–221, 1991.
- [13] E. Priola and J. Zabczyk, "Null controllability with vanishing energy," *SIAM Journal on Control and Optimization*, Vol. 42, pp. 1013–1032, 2003.
- [14] F. Gozzi and P. Loreti, "Regularity of the minimum time function and minimum energy problems: The linear case," *SIAM Journal on Control and Optimization*, Vol. 37, pp. 1195–1221, 1999.
- [15] M. Popescu, "Fundamental solution for linear two-point boundary value problem," *Journal of Applied Mathematics and Computing*, Vol. 31, pp. 385–394, 2009.
- [16] L. Frisoli, A. Borelli, and Montagner, *et al.*, "Arm rehabilitation with a robotic exoskeleton in Virtual Reality," *Proceedings of IEEE ICORR'07, International Conference on Rehabilitation Robotics*, 2007.
- [17] P. Garrec, "Systemes mecaniques," in: Coiffet. P et Kheddar A., *Teleoperation et telerobotique*, Ch 2., Hermes, Paris, France, 2002.
- [18] P. Garrec, F. Geffard, Y. Perrot (CEA List), G. Piolain, and A. G. Freudenreich (AREVA/NC La Hague), "Evaluation tests of the telerobotic system MT200-TAO in AREVANC/La Hague hot-cells," *ENC 2007, Brussels, Belgium*, 2007.

Signed (b,k)-Edge Covers in Graphs

A. N. Ghameshlou¹, Abdollah Khodkar², R. Saei³, S. M. Sheikholeslami^{*3}

¹Department of Mathematics, University of Mazandaran Babolsar, I.R., Iran

²Department of Mathematics, University of West Georgia, Carrollton, USA

³Department of Mathematics, Azarbaijan University of Tarbiat MoallemTabriz, I.R., Iran

Email: akhodkar@westga.edu, s.m.sheikholeslami@azaruniv.edu

Abstract

Let G be a simple graph with vertex set $V(G)$ and edge set $E(G)$. Let G have at least k vertices of degree at least b , where k and b are positive integers. A function $f: E(G) \rightarrow \{-1, 1\}$ is said to be a signed (b, k) -edge cover of G if $\sum_{e \in E(v)} f(e) \geq b$ for at least k vertices v of G , where $E(v) = \{uv \in E(G) \mid u \in N(v)\}$. The value $\min \sum_{e \in E(G)} f(e)$, taking over all signed (b, k) -edge covers f of G is called the signed (b, k) -edge cover number of G and denoted by $\rho'_{b,k}(G)$. In this paper we give some bounds on the signed (b, k) -edge cover number of graphs.

Keywords: Signed Star Dominating Function, Signed Star Domination Number, Signed (b, k) -edge Cover, Signed (b, k) -edge Cover Number

1. Introduction

Structural and algorithmic aspects of covering vertices by edges have been extensively studied in graph theory. An *edge cover* of a graph G is a set C of edges of G such that each vertex of G is incident to at least one edge of C . Let b be a fixed positive integer. A *b-edge cover* of a graph G is a set C of edges of G such that each vertex of G is incident to at least b edges of C . Note that a b -edge cover of G corresponds to a spanning subgraph of G with minimum degree at least b . Edge covers of bipartite graphs were studied by König [1] and Rado [2], and of general graphs by Gallai [3] and Norman and Rabin [4], and b -edge covers were studied by Gallai [3]. For an excellent survey of results on edge covers and b -edge covers, see Schrijver [5].

We consider a variant of the standard edge cover problem. Let G be a graph with vertex set $V(G)$ and edge set $E(G)$. We use [6] for terminology and notation which are not defined here and consider only simple graphs without isolated vertices. For every nonempty subset E' of $E(G)$, the subgraph of G whose vertex set is the set of vertices of the edges in E' and whose

edge set is E' , is called the subgraph of G induced by E' and denoted by $G[E']$. Two edges e_1, e_2 of G are called *adjacent* if they are distinct and have a common vertex. The *open neighborhood* $N_G(e)$ of an edge $e \in E(G)$ is the set of all edges adjacent to e . Its *closed neighborhood* is $N_G[e] = N_G(e) \cup \{e\}$. For a function $f: E(G) \rightarrow \mathbb{R}$ and a subset S of $E(G)$ we define $f(S) = \sum_{e \in S} f(e)$. The *edge-neighborhood* $E_G(v)$ of a vertex $v \in V(G)$ is the set of all edges incident to vertex v . For each vertex $v \in V(G)$, we also define $f(v) = \sum_{e \in E_G(v)} f(e)$. Let b be a positive integer and let G have at least k vertices of degree at least b . A function $f: E(G) \rightarrow \{-1, 1\}$ is called a *signed (b, k) -edge cover* (SbkEC) of G , if $f(v) \geq b$ for at least k vertices v of G . The *signed (b, k) -edge cover number* of a graph G is $\rho'_{b,k}(G) = \min \{ \sum_{e \in E} f(e) \mid f \text{ is an SbkEC on } G \}$. The signed (b, k) -edge cover f of G with $f(E(G)) = \rho'_{b,k}(G)$ is called a $\rho'_{b,k}(G)$ -cover. For any signed (b, k) -edge cover f of G we define

*Corresponding author

$P = \{e \in E \mid f(e) = 1\}$, $M = \{e \in E \mid f(e) = -1\}$,
 $V^+ = \{v \in V \mid f(v) \geq b\}$ and $V^- = \{v \in V \mid f(v) < b\}$.

If $b=1$ and $k=n$, then the signed (b,k) -edge cover number is called the *signed star domination number*. The signed star domination number was introduced by Xu in [7] and denoted by $\gamma_{ss}'(G)$. The signed star domination number has been studied by several authors (see for example [7,10]).

If $b=1$ and $1 \leq k \leq n$, then the signed (b,k) -edge cover number is called the *signed star k -subdomination number*. The signed star k -subdomination number was introduced by Saei and Sheikholeslami in [11] and denoted by $\gamma_{ss}^k(G)$.

If b is an arbitrary positive integer and $k=n$, then the signed (b,k) -edge cover number is called the *signed b -edge cover number*. The signed b -edge cover number was introduced by Bonato *et al.* in [12] and denoted by $\rho'_b(G)$.

The purpose of this paper is to initiate the study of the signed (b,k) -edge cover number $\rho'_{b,k}(G)$. Here are some well-known results on $\gamma_{ss}'(G)$, $\gamma_{ss}^k(G)$ and $\rho'_b(G)$.

Theorem 1 [10] For every graph G of order $n \geq 4$,
 $\rho_{1,n}'(G) \leq 2n - 4$.

Theorem 2 [11] For every graph G of order $n \geq 4$ without isolated vertices, $\rho_{1,k}'(G) \leq n + k - 4$.

Theorem 3 [10] For every graph G of order n without isolated vertices, $\rho_{1,n}'(G) \geq \lceil \frac{n}{2} \rceil$.

Theorem 4 [11] For every graph G of order $n \geq 2$ without isolated vertices,

$$\rho_{1,k}'(G) \geq \lceil \frac{(\Delta(G)+1)k - n\Delta(G)}{2} \rceil.$$

Theorem 5 [12] Let b be a positive integer. For every graph G of order n and minimum degree at least b ,

$$\rho_{b,n}'(G) \geq \lceil \frac{bn}{2} \rceil.$$

We make use of the following result in this paper.

Theorem 6 [7] Every graph G with $\delta(G) \geq 3$ contains an even cycle.

2. Lower Bounds for SbkeCN of Graphs

In this section we present some lower bounds on $\rho'_{b,k}$ in

terms of the order, the size, the maximum degree and the degree sequence of G . Our first proposition is a generalization of Theorems 3, 4 and 5.

Proposition 1 Let G be a graph of order n without isolated vertices and maximum degree $\Delta = \Delta(G)$. Let b be a positive integer and let $n_0 \geq 1$ be the number of vertices with degree at least b . Then for every positive integer $1 \leq k \leq n_0$,

$$\rho'_{b,k}(G) \geq \frac{k(b+\Delta) - n_0(\Delta-b+1) - n(b-1)}{2}.$$

Proof. Let f be a $\rho'_{b,k}(G)$ -cover. We have

$$\begin{aligned} \rho'_{b,k}(G) &= \sum_{e \in E(G)} f(e) = \frac{1}{2} \sum_{v \in V(G)} \sum_{e \in E(v)} f(e) \\ &= \frac{1}{2} \sum_{v \in V^+} \sum_{e \in E(v)} f(e) + \frac{1}{2} \sum_{v \in V^-} \sum_{e \in E(v)} f(e) \\ &\geq \frac{kb}{2} - \frac{(n_0-k)\Delta + (n-n_0)(b-1)}{2} \end{aligned}$$

$$= \frac{k(b+\Delta) - n_0(\Delta-b+1) - n(b-1)}{2}.$$

Theorem 2 Let G be a graph of order n , size m , without isolated vertices and with degree sequence $(d_1, d_2, \dots, d_n)$, where $d_1 \leq d_2 \leq \dots \leq d_n$. Let b be a positive integer and let $n_0 \geq 1$ be the number of vertices with degree at least b . Then for every positive integer $1 \leq k \leq n_0$,

$$\rho'_{b,k}(G) \geq \frac{\sum_{j=1}^k (bd_j + d_j^2)}{2d_n} - m.$$

Proof. Let g be a $\rho'_{b,k}(G)$ -cover of G and let $g(v) \geq b$ for k distinct vertices v in $B(G) = \{v_{j_1}, \dots, v_{j_k}\}$. Define $f: E(G) \rightarrow \{0,1\}$ by $f(e) = (g(e)+1)/2$ for each $e \in E(G)$. We have

$$\sum_{e \in E(G)} f(N_G[e]) = \sum_{e=uv \in E(G)} \frac{g(N_G[e]) + \deg(u) + \deg(v) - 1}{2}. \quad (1)$$

Since

$$\sum_{e \in E(G)} (g(N_G[e]) + g(e)) = \sum_{v \in V} g(E(v)) \deg(v)$$

and

$$\sum_{e=uv \in E(G)} (\deg(u) + \deg(v)) = \sum_{v \in V} \deg(v)^2,$$

by (1) it follows that

$$\begin{aligned} & \sum_{e \in E(G)} f(N_G[e]) \\ &= \frac{1}{2} \sum_{v \in V} \deg(v)(g(E(v)) + \deg(v)) - \frac{1}{2} \sum_{e \in E(G)} g(e) - \frac{m}{2} \\ &\geq \frac{1}{2} \sum_{v \in V \setminus \{v_{j_1}, \dots, v_{j_k}\}} \deg(v)(g(E(v)) + \deg(v)) + \\ &\quad \frac{1}{2} \sum_{i=1}^k (bd_{j_i} + d_{j_i}^2) - \frac{1}{2} \rho'_{b,k}(G) - \frac{m}{2} \\ &\geq \frac{1}{2} \sum_{i=1}^k (bd_{j_i} + d_{j_i}^2) - \frac{1}{2} \rho'_{b,k}(G) - \frac{m}{2} \\ &\geq \frac{1}{2} \sum_{j=1}^k (bd_j + d_j^2) - \frac{1}{2} \rho'_{b,k}(G) - \frac{m}{2}. \end{aligned} \quad (2)$$

On the other hand,

$$\begin{aligned} \sum_{e \in E(G)} f(N_G[e]) &= \sum_{v \in V} f(E(v)) \deg(v) - \sum_{e \in E(G)} f(e) \\ &\leq \sum_{v \in V} f(E(v)) d_n - \sum_{e \in E(G)} f(e) \\ &= d_n \left(2 \sum_{e \in E(G)} f(e) \right) - \sum_{e \in E(G)} f(e) \\ &= (2d_n - 1) \sum_{e \in E(G)} f(e). \end{aligned} \quad (3)$$

By (2) and (3)

$$\sum_{e \in E(G)} f(e) \geq \frac{\frac{1}{2} \sum_{j=1}^k (bd_j + d_j^2) - \frac{1}{2} \rho'_{b,k}(G) - \frac{m}{2}}{2d_n - 1}. \quad (4)$$

Since $g(E(G)) = 2f(E(G)) - m$, by (4)

$$\rho'_{b,k}(G) = \sum_{e \in E(G)} g(e) \geq$$

$$\frac{1}{2d_n - 1} \left(\sum_{j=1}^k (bd_j + d_j^2) - \rho'_{b,k}(G) - m \right) - m.$$

Thus,

$$\rho'_{b,k}(G) \geq \frac{\sum_{j=1}^k (bd_j + d_j^2)}{2d_n} - m,$$

as desired.

An immediate consequence of Theorem 2 is:

Corollary 3 For every r -regular graph G of size m ,

$$\rho'_{b,k}(G) \geq \frac{k(b+r)}{2} - m. \text{ Furthermore, the bound is sharp}$$

for r -regular graphs with $b = r$ and $k = n$.

Theorem 4 Let G be a graph of order $n \geq 2$, size m , without isolated vertices, with minimum degree $\delta = \delta(G)$ and maximum degree $\Delta = \Delta(G)$. Let b be a positive integer and $n_0 \geq 1$ be the number of vertices with degree at least b . Then for each positive integer $1 \leq k \leq n_0$

$$\rho'_{b,k}(G) \geq \frac{(\Delta^2 + b^2)k - 2(\Delta - \delta)m - (b-1)^2 n - (\Delta^2 - (b-1)^2)n_0}{2\delta}. \quad (5)$$

Furthermore, the bound is sharp for n -cycles when $b = 2$ and $k = n$.

Proof. Let $B(G) = \{v \in V(G) \mid \deg(v) \geq b\}$ and let f be a $\rho'_{b,k}(G)$ -cover. Since for each $v \in V^+$, $f(v) \geq b$, it follows that $|M \cap E(v)| \leq \lfloor \frac{\deg(v) - b}{2} \rfloor$. Thus

$$\begin{aligned} & (2\delta - 1)|M| \\ &\leq \sum_{e=uv \in M} (\deg(u) + \deg(v) - 1) \\ &= -|M| + \sum_{e=uv \in M} (\deg(u) + \deg(v)) \\ &= -|M| + \sum_{v \in V(G[M])} |M \cap E(v)| \deg(v) \\ &\leq -|M| + \sum_{v \in V^+} |M \cap E(v)| \deg(v) + \sum_{v \in V^-} |M \cap E(v)| \deg(v) \\ &\leq -|M| + \sum_{v \in V^+} \lfloor \frac{\deg(v) - b}{2} \rfloor \deg(v) + \sum_{v \in V^-} \deg(v)^2 \\ &\leq -|M| + \sum_{v \in V^+} \frac{\deg(v)^2}{2} + \sum_{v \in V^-} \deg(v)^2 - \frac{b}{2} \sum_{v \in V^+} \deg(v) \end{aligned}$$

$$\begin{aligned}
&\leq -|M| + \sum_{v \in V} \frac{\deg(v)^2}{2} + \sum_{v \in V^-} \frac{\deg(v)^2}{2} - \frac{b^2}{2} |V^+| \\
&\leq -|M| + \Delta \sum_{v \in V} \frac{\deg(v)}{2} + \sum_{v \in V^- \cap B(G)} \frac{\deg(v)^2}{2} + \\
&\quad \sum_{v \in V^- \setminus B(G)} \frac{\deg(v)^2}{2} - \frac{b^2}{2} |V^+| \\
&\leq -|M| + \Delta m - \frac{b^2}{2} k + \frac{\Delta^2}{2} |V^- \cap B(G)| + \frac{(b-1)^2}{2} |V^- \setminus B(G)| \\
&\leq -|M| + \Delta m - \frac{b^2}{2} k + \frac{\Delta^2}{2} (n_0 - k) + \frac{(b-1)^2}{2} (n - n_0)
\end{aligned}$$

Hence,

$$|M| \leq \frac{\Delta m}{2\delta} + \frac{1}{4\delta} ((b-1)^2 n + (\Delta^2 - (b-1)^2) n_0 - (\Delta^2 + b^2) k).$$

Now (5) follows by the fact that $\rho'_{b,k}(G) = m - 2|M|$.

3. An Upper Bound on SbkeCN

Bonato *et al.* in [11] posed the following conjecture on $\rho'_b(G)$.

Conjecture 5 Let $b \geq 2$ be an integer. There is a positive integer n_b so that for any graph G of order $n \geq n_b$ with minimum degree b ,

$$\rho'_b(G) \leq (b+1)(n-b-1).$$

Since $\rho'_b(K_{b+1, n-b-1}) = (b+1)(n-b-1)$, the upper bound would be the best possible if the conjecture were true. They also proved that the conjecture is true for $b = 2$. In this section we provide an upper bound for $\rho'_{b,k}(G)$, where $b = 2$ and $1 \leq k \leq n$. The proof of the next theorem is essentially similar to the proof of Theorem 5 in [11].

Theorem 6 Let G be a graph of order n , size m and without isolated vertices. Let $n_0 > 0$ be the number of vertices with degree at least 2. Then for $n \geq 7$ and $1 \leq k \leq n_0$,

$$\rho'_{2,k}(G) \leq 2n + k - 9.$$

Proof. The proof is by induction on the size m of G . By a tedious and so omitted argument, it follows that $\rho'_{2,k}(G) \leq k + 5$ if $n = 7$. We may therefore assume that

$n \geq 8$. Suppose that the theorem is true for all graphs G without isolated vertices and size less than m . Let G be a graph of order $n \geq 8$, size m and without isolated vertices. We will prove that $\rho'_{2,k}(G) \leq 2n + k - 9$ for each $1 \leq k \leq n_0$. We consider four cases.

Case 1. $\delta(G) = 1$.

Let u be a vertex of degree 1 and $v \in N(u)$. First suppose $\deg(v) = 1$. Then the induced subgraph $G[u, v]$ is K_2 . It is straightforward to verify that $\rho'_{2,k} \leq 2n + k - 9$ when $n = 8$. Hence, we may assume that $n \geq 9$. Let $G' = G - uv$. Then G' is a graph of order $n - 2 \geq 7$, size $m - 1$ and without isolated vertices. By the inductive hypothesis, $\rho'_{2,k}(G') \leq 2(n - 2) + k - 9 = 2n + k - 13$. Let f be a $\rho'_{2,k}(G')$ -cover. Define $g : E(G) \rightarrow \{-1, 1\}$ by $g(uv) = -1$ and $g(e) = f(e)$ if $e \in E(G) - uv$. Obviously, g is a S2kEC and so

$$\rho'_{2,k}(G) \leq \rho'_{2,k}(G') - 1 \leq 2n + k - 14 < 2n + k - 9.$$

Now suppose $\deg(v) \geq 2$. Consider two subcases.

Subcase 1.1 $\deg(v) \geq 3$.

By the inductive hypothesis on $G - u$, $\rho'_{2,k}(G - u) \leq 2(n - 1) + k - 9 = 2n + k - 11$. Let f be a $\rho'_{2,k}(G - u)$ -cover and define $g : E(G) \rightarrow \{-1, 1\}$ by $g(uv) = 1$ and $g(e) = f(e)$ if $e \in E(G) - uv$. Obviously, g is a S2kEC and so

$$\rho'_{2,k}(G) \leq \rho'_{2,k}(G') + 1 \leq 2n + k - 10 < 2n + k - 9.$$

Subcase 1.2 $\deg(v) = 2$.

Let $w \in N(v) - \{u\}$. If $k = 1$, then define $g : E(G) \rightarrow \{-1, 1\}$ by $g(uv) = g(vw) = 1$ and $g(e) = -1$ if $e \in E(G) \setminus \{uv, vw\}$. Obviously, g is a S2kEC of G and we have

$$\rho'_{2,k}(G) \leq g(E(G)) = 4 - m \leq 2n + k - 9.$$

Let $k \geq 2$. It follows that $n_0 \geq 2$. By the inductive hypothesis on $G - \{u\}$, $\rho'_{2,k-1}(G - \{u\}) \leq 2(n - 1) + (k - 1) - 9 = 2n + k - 12$. Let f be a $\rho'_{2,k-1}(G - \{u\})$ -cover. Define $g : E(G) \rightarrow \{-1, 1\}$ by $g(uv) = g(vw) = 1$ and $g(e) = f(e)$ if $e \in E(G) \setminus \{uv, vw\}$. Obviously, g is a S2kEC and so

$$\rho'_{2,k}(G) \leq \rho'_{2,k-1}(G - \{u\}) + 3 \leq 2n + k - 9.$$

Case 2. $\delta(G) = 2$.

Let w be a vertex of degree 2 and $N(w) = \{u, v\}$. Consider two subcases.

Subcase 2.1 $uv \notin E(G)$. Let G' be the graph obtained from $G - \{w\}$ by adding an edge uv . Then G' has order $n-1$, size $m-1$ and at least $k-1$ vertices with degree at least 2. By the inductive hypothesis,

$$\rho'_{(k-1),2}(G') \leq 2(n-1) + (k-1) - 9 = 2n + k - 12.$$

Let f be a $\rho'_{2,k}(G')$ -cover. Define $g: E(G) \rightarrow \{-1, 1\}$ by $g(uw) = g(vw) = 1$ and $g(e) = f(e)$ if $e \in E(G) \setminus \{uv, vw\}$. Obviously, g is a S2kEC and so

$$\rho'_{2,k}(G) \leq g(E(G)) \leq f(E(G')) + 3 \leq 2n + k - 9.$$

Subcase 2.2 $uv \in E(G)$. First let both u and v have degree 2. Then the induced subgraph $G[\{u, v, w\}]$ is an isolated triangle. If $1 \leq k \leq 3$, then define $f: E(G) \rightarrow \{-1, 1\}$ by

$$f(uv) = f(vw) = f(uw) = 1 \text{ and } f(e) = -1 \text{ otherwise.}$$

Then

$$\rho'_{2,k}(G) \leq f(E(G)) = 6 - m \leq 2n + k - 9.$$

Now suppose that $k \geq 4$. It is not hard to show that $\rho'_{2,k}(G) \leq 2n + k - 9$ when $n = 8$ or 9 . Hence, we may assume that $n \geq 10$. Let $G' = G - \{u, v, w\}$. Then G' is a graph of order $n-3 \geq 7$, size $m-3$ and has at least $k-3$ vertices with degree at least 2. By the inductive hypothesis, $\rho'_{2,(k-3)}(G') \leq 2(n-3) + (k-3) - 9 = 2n + k - 18$. Let f be a $\rho'_{2,(k-3)}(G')$ -cover. Define $g: E(G) \rightarrow \{-1, 1\}$ by

$$g(uv) = g(vw) = g(uw) = 1 \text{ and } g(e) = f(e) \text{ if } e \in E(G').$$

Obviously, g is a S2kEC of G and

$$\rho'_{2,k}(G) = g(E(G)) \leq f(E(G')) + 3 \leq (2n + k - 18) + 3.$$

Now let $\min\{\deg(u), \deg(v)\} \geq 3$. If $k = 1$, define $g: E(G) \rightarrow \{-1, 1\}$ by $g(uw) = g(vw) = 1$ and $g(e) = -1$ otherwise. Obviously, g is a S2kEC and so

$$\rho'_{2,k}(G) \leq g(E(G)) = 4 - m < 2n + k - 9.$$

If $k \geq 2$, then $G' = G - \{w\}$ is a graph of order $n-1$, size $m-2$ and has at least $k-1$ vertices with degree at least 2. By the inductive hypothesis, we have

that $\rho'_{2,(k-1)}(G') \leq 2n + k - 12$. Let f be a $\rho'_{2,(k-1)}(G')$ -cover. We can obtain a S2kEC g of G by assigning $g(e) = 1$ for each $e \in E(G) \setminus E(G')$ and $g(e) = f(e)$ for each $e \in E(G')$. Then we have

$$g(E(G)) = f(E(G')) + 2 = \rho'_{2,(k-1)}(G') + 2 < 2n + k - 9.$$

Hence, $\rho'_{2,k}(G) < 2n + k - 9$, as desired.

Finally, assume $\min\{\deg(u), \deg(v)\} = 2$. Let without loss of generality $\deg(u) = 2$. If $1 \leq k \leq 2$, define $g: E(G) \rightarrow \{-1, 1\}$ by $g(uw) = g(vw) = g(uv) = 1$ and $g(e) = -1$ otherwise. Obviously, g is a S2kEC and so

$$\rho'_{2,k}(G) \leq g(E(G)) = 6 - m \leq 2n + k - 9.$$

If $k \geq 3$, then $G' = G - \{w\}$ is a graph of order $n-1$, size $m-2$ and has at least $k-2$ vertices with degree at least 2. By the inductive hypothesis,

$$\rho'_{2,(k-2)}(G') \leq 2(n-1) + (k-2) - 9 = 2n + k - 13.$$

Let f be a $\rho'_{2,(k-2)}(G')$ -cover. Define $g: E(G) \rightarrow \{-1, 1\}$ by

$$g(uv) = g(vw) = g(uw) = 1 \text{ and } g(e) = f(e) \\ \text{if } e \in E(G') \setminus \{uv\}.$$

Obviously, g is a S2kEC of G and

$$\rho'_{2,k}(G) \leq g(E(G)) = f(E(G')) + 4 \leq 2n + k - 9.$$

Case 3. $\delta(G) = 3$.

Let w be a vertex with degree 3. If $k = 1$, define $g: E(G) \rightarrow \{-1, 1\}$ by $g(uw) = 1$ if $u \in N(w)$ and $g(e) = -1$ otherwise. Obviously, g is a S2kEC and so

$$\rho'_{2,k}(G) \leq g(E(G)) = 6 - m < 2n + k - 9.$$

If $k \geq 2$, then $G' = G - \{w\}$ is a graph of order $n-1$, size $m-3$ and has at least $k-1$ vertices with degree at least 2. By the inductive hypothesis, we have that $\rho'_{2,(k-1)}(G') \leq 2n + k - 12$. Let f be a $\rho'_{2,(k-1)}(G')$ -cover. We can obtain a S2kEC g of G by assigning $g(e) = 1$ for each $e \in E(G) \setminus E(G')$ and $g(e) = f(e)$ for each $e \in E(G')$. Then we have

$$g(E(G)) = f(E(G')) + 3 = \rho'_{2,(k-1)}(G') + 3 \leq 2n + k - 9.$$

Hence, $\rho'_{2,k}(G) \leq 2n + k - 9$, as desired.

Case 4. $\delta(G) \geq 4$.

Then G has an even cycle by Theorem 6. Let $C = (v_1, v_2, \dots, v_s)$ be an even cycle in G . Obviously, $G' = G - E(C)$ is a graph of order n , size $m - |E(C)|$ and has at least k vertices with degree at least 2. By the inductive hypothesis, $\rho'_{2,k}(G') \leq 2n + k - 9$. Let f be a $\rho'_{2,k}(G')$ -cover. Let $v_{s+1} = v_1$ and define $g: E(G) \rightarrow \{-1, 1\}$ by

$$g(v_i v_{i+1}) = (-1)^i \text{ if } i = 1, \dots, s \text{ and } g(e) = f(e) \text{ for } e \in E(G) \setminus E(C).$$

Obviously, g is a S2kEC and hence $\rho'_{2,k}(G) = \rho'_{2,k}(G') \leq 2n + k - 9$. This completes the proof.

4. Conclusions

In this paper we initiated the study of the signed (b, k) -edge cover numbers for graphs, generalizing the signed star domination numbers, the signed star k -domination numbers and the signed b -edge cover numbers in graphs. The first lower bound obtained in this paper for the signed (b, k) -edge cover number concludes the existing lower bounds for the other three parameters. Our upper bound for the signed (b, k) -edge cover number also implies the existing upper bound for the signed b -edge cover number. Finally, Theorem 6 inspires us to generalize Conjecture 5.

Conjecture 7 Let $b \geq 3$ be an integer. There is a positive integer n_b so that for any graph G of order $n \geq n_b$ with $n_0 \geq 1$ vertices of degree at least b , and for any integer $1 \leq k \leq n_0$, $\rho'_{b,k}(G) \leq bn + k - (b+1)^2$.

5. References

- [1] D. König, "Über trennende knotenpunkte in graphen (nebst anwendungen auf determinanten und matrisen)," Acta Litterarum ac Scientiarum Regiae Universitatis Hungaricae Francisco-Josephinae, Sectio Scientiarum Mathematicarum [Szeged], Vol. 6, pp. 155–179, 1932–1934.
- [2] R. Rado, "Studien zur kombinatorik," Mathematische Zeitschrift, German, Vol. 36, pp. 424–470, 1933.
- [3] T. Gallai, "Über extreme Punkt-und Kantenmengen (German)," Ann. Univ. Sci. Budapest. Eötvös Sect. Math., Vol. 2, pp. 133–138, 1959.
- [4] R. Z. Norman and M. O. Rabin, "An algorithm for a minimum cover of a graph," Proceedings of the American Mathematical Society, Vol. 10, pp. 315–319, 1959.
- [5] A. Schrijver, "Combinatorial optimization: Polyhedra and Efficiency," Springer, Berlin, 2004.
- [6] D. B. West, "Introduction to graph theory," Prentice-Hall, Inc., 2000.
- [7] B. Xu, "On signed edge domination numbers of graphs," Discrete Mathematics, Vol. 239, pp. 179–189, 2001.
- [8] C. Wang, "The signed star domination numbers of the Cartesian product," Discrete Applied Mathematics, Vol. 155, pp. 1497–1505, 2007.
- [9] B. Xu, "Note on edge domination numbers of graphs," Discrete Mathematics, Vol. 294, pp. 311–316, 2005.
- [10] B. Xu, "Two classes of edge domination in graphs," Discrete Applied Mathematics, Vol. 154, pp. 1541–1546, 2006.
- [11] R. Saei and S. M. Sheikholeslami, "Signed star k -subdomination numbers in graph," Discrete Applied Mathematics, Vol. 156, pp. 3066–3070, 2008.
- [12] A. Bonato, K. Cameron, and C. Wang, "Signed b -edge covers of graphs (manuscript)".

On the Mechanism of CDOs behind the Current Financial Crisis and Mathematical Modeling with Lévy Distributions

Hongwen Du¹, Jianglun Wu², Wei Yang³

¹School of Finance and Economics, Hangzhou Dianzi University, Hangzhou, China

²Department of Mathematics, Swansea University, Swansea, UK

³Department of Mathematics, Swansea University, Swansea, UK

Email: zjdhw@hdu.edu.cn, {j.l.wu,mawy}@swansea.ac.uk

Abstract

This paper aims to reveal the mechanism of Collateralized Debt Obligations (CDOs) and how CDOs extend the current global financial crisis. We first introduce the concept of CDOs and give a brief account of the development of CDOs. We then explicate the mechanism of CDOs within a concrete example with mortgage deals and we outline the evolution of the current financial crisis. Based on our overview of pricing CDOs in various existing random models, we propose an idea of modeling the random phenomenon with the feature of heavy tail dependence for possible implements towards a new random modeling for CDOs.

Keywords: Collateralized Debt Obligations (CDOs), Cashflow CDO, Synthetic CDO, Mechanism, Financial Crisis, Pricing Models, Lévy Stable Distributions

Collateralized debt obligations (CDOs) were created in 1987 by bankers at Drexel Burnham Lambert Inc. Within 10 years, the CDOs had become a major force in the credit derivatives market, in which the value of a derivative is “derived” from the value of other assets. But unlike some fairly straightforward derivatives such as options, calls, and *Credit Default Swaps* (CDSs), CDOs are not “real”, which means they are constructs, and sometime even built upon other constructs. CDOs are designed to satisfy different type of investors, low risk with low return and high risk with high return.

In early 2007, following the burst of the bubble of housing market in the United States, losses in the CDOs market started spreading. By early 2008, the CDO crisis had morphed into what we now encountered the world-wide financial crisis. CDOs are at the heart of the crisis and even extend the crisis.

1. Introduction to CDOs

Collateralized Obligations (COs) are *promissory notes* backed by collaterals or securities. In the market for COs, the securities can be taken from a very wide spectrum of alternative financial instruments, such as bonds (*Col-*

lateralized Bond Obligations, or CBO), loans (*Collateralized Loan Obligations*, or CLO), funds (*Collateralized Fund Obligations*, or CFO), mortgages (*Collateralized Mortgage Obligations*, or CMO) and others. And frequently, they source their collaterals from a combination of two or more of these asset classes. Collectively, these instruments are popular referred to as CDOs, which are bond-like instruments whose cashflow structures allocate interest income and principal repayments from a collateral pool of different debt instruments to a prioritized collection of CDO securities to their investors. The most popular life of a CDO is five years. However, 7-year, 10-year, and to a less extent 3-year CDOs now trade fairly actively.

A CDO can be initiated by one or more of the followings: banks, non-bank financial institutions, and asset management companies, which are referred to as the sponsors. The sponsors of a CDO create a company so-called the *Special Purpose Vehicle* (SPV). The SPV works as an independent entity and is usually bankruptcy remote. The sponsors can earn serving fees, administration fees and hedging fees from the SPV, but otherwise has no claim on the cash flow of the assets in the SPV.

According to how the SPV gains credit risks, CDOs are classified into two kinds: *cashflow CDOs* and syn-

thetic CDOs. If the SPV of a CDO owns the underlying debt obligations (portfolio), that is, the SPV obtains the credit risk exposure by purchasing debt obligations (eg. bonds, residential and commercial loans), the CDO is referred to as a cashflow CDO, which is the basic form in the CDOs market in their formative years. In contrast, if the SPV of a CDO does not own the debt obligations, instead obtaining the credit risk exposure by selling CDSs on the debt obligations of reference entities, the CDO is referred to as a synthetic CDO; the synthetic structure allows bank originators in the CDOs market to ensure that client relationships are not jeopardized, and avoids the tax-related disadvantages existing in cashflow CDOs. The following graph (Figure 1) illustrates a construction of a cashflow CDO.

After acquiring credit risks, SPV sells these credit risks in *tranches* to investors who, in return for an agreed payment (usually a periodic fee), will bear the losses in the portfolio derived from the default of the instruments in the portfolio. Therefore, the tranches holders have the ultimate credit risk exposure to the underlying reference portfolio.

Tranching, a common characteristic of all securitisations, is the structuring of the product into a number of different classes of notes ranked by the seniority of investor's claims on the instruments assets and cashflows. The tranches have different seniorities: *senior tranche*,

the least risky tranche in CDOs with lowest fixed interest rate, followed by *mezzanine tranche*, *junior mezzanine tranche*, and finally the first loss piece or *equity tranche*. A CDO makes payments on a sequential basis, depending on the seniority of tranches within the *capital structure* of the CDO. The more senior the tranches investors are in, the less risky the investment and hence the less they will be paid in interest. The way it works is frequently referred to as a “waterfall” or cascade of cashflows. We will give a specific illustration in Section 3.

In perfect capital markets, CDOs would serve no purpose; the costs of constructing and marketing a CDO would inhibit its creation. In practice, however, CDOs address some important market imperfections. First, banks and certain other financial institutions have regulatory capital requirements that make it valuable for them to securitize and sell some portion of their assets, reducing the amount of (expensive) regulatory capital that they must hold. Second, individual bonds or loans may be illiquid, leading to a reduction in their market values. Securitization may improve liquidity, and thereby raise the total valuation to the issuer of the CDO structure.

In light of these market imperfections, at least two classes of CDOs are popular: the *balance-sheet CDO* and the *arbitrage CDO*. The balance-sheet CDO, typically in the form of a CLO, is designed to remove loans from the

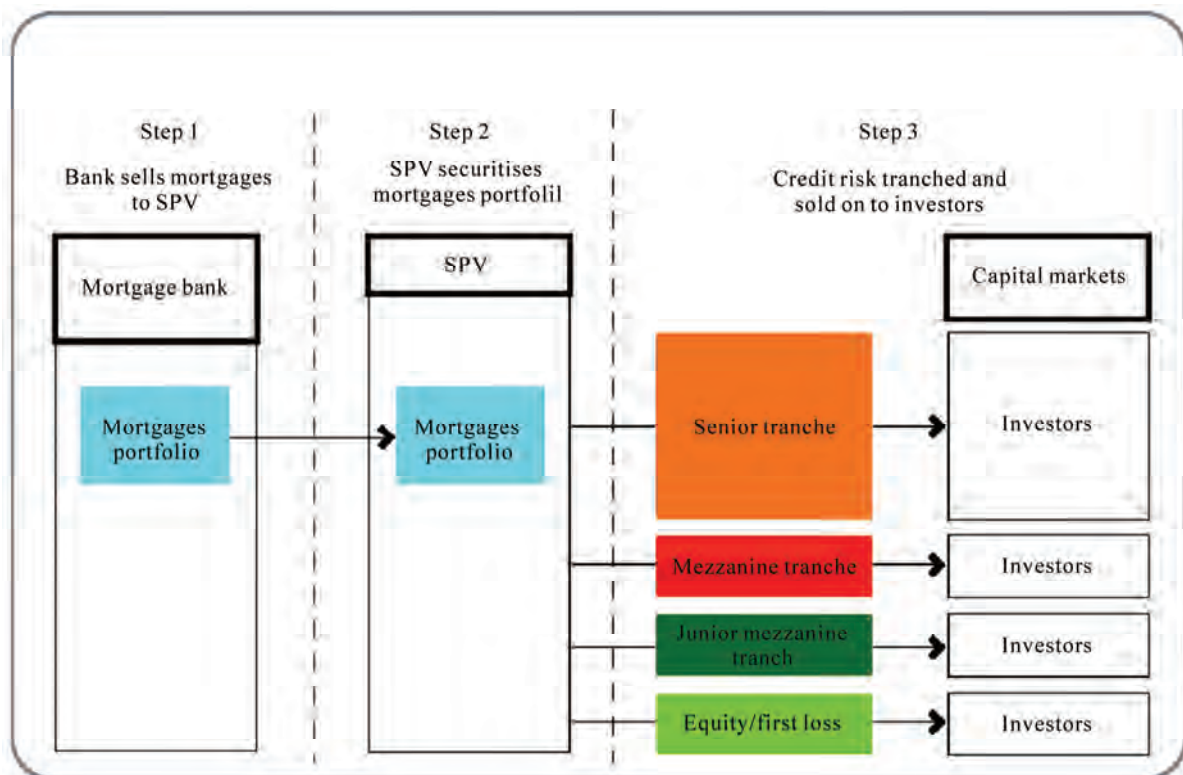


Figure 1. Cashflow collateralised mortgage obligation.

balance sheets of banks, achieving capital relief, and perhaps also increasing the valuation of the assets through an increase in liquidity. An arbitrage CDO, often underwritten by an investment bank, is designed to capture some fraction of the likely difference between the total cost of acquiring collateral assets in the secondary market and the value received from management fees and the sale of the associated CDOs structure.

2. The Development of CDOs

Although a market for CMOs—the forerunner of modern CDOs—was taking shape in the US market by the early 1980s, the market for CDOs is generally believed to date back to the late 1980s and the rapid revolution of CDOs is very much a story of the 1990s.

By the late 1990s, the structure of the international market for CDOs of all kinds was becoming characterized by a number of conspicuous and interrelated trends. Firstly, issuance volume was rising exponentially, as was understanding and acceptance of the CDO technique. Secondly, the cross-border investment flowing into CDOs were rising steeply. Thirdly, more and more asset classes were being used as security for COs. Finally, in 1999 and 2000 the concept of the COs was popularised across continental Europe with strikingly high speed and, meanwhile, changes to legislation and regulation were emerging as important sources of support for new issuance in the CDOs market in Europe.

In 2000 CDOs were made legal and at the same time were prevented from being regulated, by the Commodity Futures Modernization Act, which specifies that products offered by banking institutions could not be regulated as futures contracts. It lies at the root of America's failure to regulate the debt derivatives that are now threatening the global economy. By 2000 and 2001 globally, the most important determinant of increasing volumes in the CDO market was the explosive growth in the market for credit derivatives in general and for CDS in particular, which paved the way for an equally explosive expansion of the market for synthetic CDOs. Thereafter, the volume of traditional cashflow CDOs has been eclipsed by synthetic products.

The process of increasing diversification in the CDO market has began in 2002. An example of such diversification is the so-called CDOs of CDOs (CDOs squared): a portfolio of CDOs is assembled, tranching, and sold to investors. Other exotic CDO products include CDOs of funds (CFO), and CDOs of equity default swaps (CDOs of EDS), forward starting CDOs, options on CDO tranches, leverage super senior CDOs, and bespoke CDOs.

CDOs were originally static portfolios where the underlying names rarely changed and the static CDOs po-

ssess the advantage that they call for minimal resources in terms of management expertise and time and reduce costs involved in trading. However, when they declined in value, investors were unable to do anything to reverse that decline as credit quality began to deteriorate. Therefore, actively managed CDOs were rapidly gaining in popularity. The growth of managed products, however, was also helped by the growing maturity of the CDO market and by the increasing number of managers with proven experience in managing credit in general and credit derivatives in particular.

3. The Mechanism of CDOs

In this section, we try to describe the mechanism of CDOs in a vivid, therefore not so rigorous, way. We simplify the collaterals as mortgages. Then the process to create a CDO can be seen in the following manner: investment banks buy mortgages and then pool them into Mortgage Backed Securities (MBSs) with different ratings. Financial institutions seeking new markets purchase these MBSs, pool them with other similarly rated MBSs and sometimes derivatives, and then issue new securities. This process of buying mortgages, creating MBSs, and packaging these MBSs into CDOs is designed to apportion credit risk to those parties who are willing to take it on.

First of all, we have a CDO manager who decides to create a CDO. He/She has a bottle (SPV). In order to fill the bottle, he/she can buy collaterals (anything he/she wants: the loans, credit card debt and student loans). For example, let us assume the collaterals are \$1b mortgages paying interest rate of 10%. Then \$1m credit-linked notes (CLNs) with par value \$1k are issued based on the underlying collaterals portfolio.

Secondly, we may regard its capital structure as a 4-layer pyramid of wine glasses over a tray. Each layer (tranche) has different seniority. Into these glasses are the CLNs rated according to their riskiness. On the top layer is the senior tranche with 400k AAA-rated CLNs, the least risky tranche with lowest fixed interest rate 6%, followed by mezzanine tranche with 200k AA-rated CLNs paying fixed interest rate 7%, junior mezzanine tranche with 200k BBB-rated CLNs paying fixed interest rate 10%, and equity tranche with 200k CLNs with highest risk. Investors will get paid, at each payment day, at corresponding interest rates of tranches they are involved in.

At the payment day, because these mortgages are paying interest, the cork of the bottle pops off with much pressure. The money then flows out on the top and into the pyramid of the glasses. If all of the mortgages are paying interest, i.e., there is no default, the interest would sum up to \$100m. Because the senior tranche is the least

risky, it gets paid first (\$24m). After the senior tranche gets filled up first, the mezzanine tranche (\$14m) and then the junior mezzanine tranche (\$20m) get filled up in turn. Equity tranche on the bottom is still filled up with payment of \$42m, resulting in up to 21% return rate. (cf.

Figure 2 below)

However, if defaults happen among these mortgages in the bottle, the cashflow of interest would decrease, for instance, let us say only \$50m interests are paid (cf. Figure 3 below).

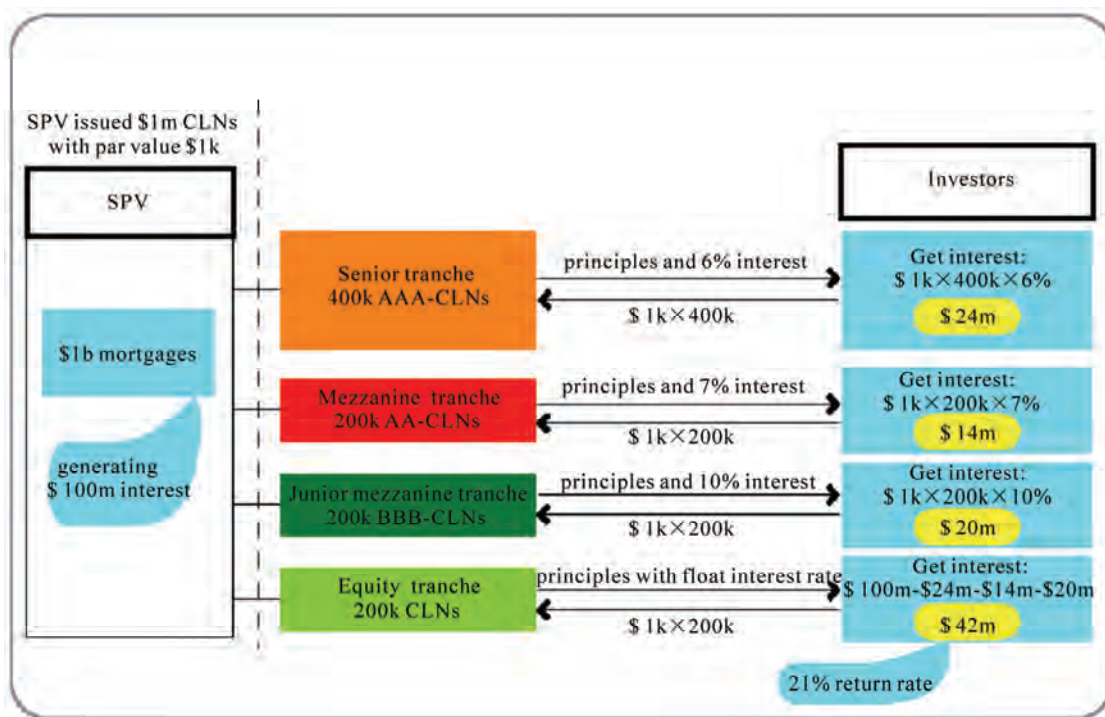


Figure 2. Cashflow of a CDO under no defaults.

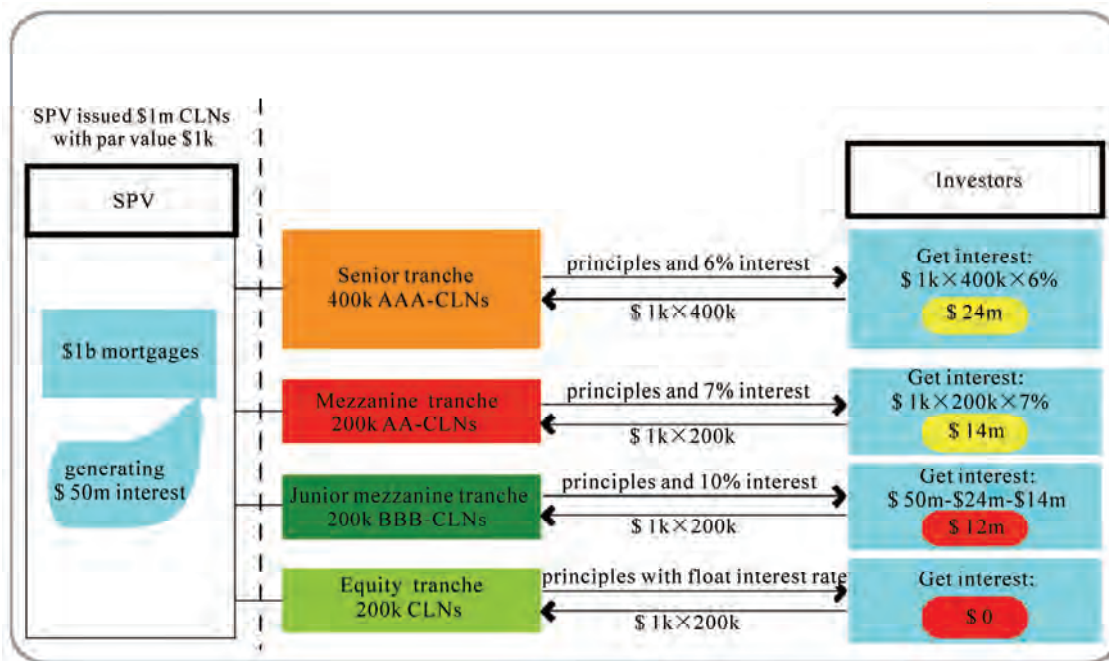


Figure 3. Cashflow of a CDO under 50% defaults.

In this situation, the senior tranche still gets paid in full first (\$24m); mezzanine tranche get paid \$14m; yet the junior mezzanine tranche only gets paid \$12m generating 6% return rate less than 10% as expected; nothing can be paid to the equity tranche holders.

If we complicate the situation further by thinking of another manager who also decides to create a CDO. Instead of filling the bottle with mortgages, he decides to fill with these MBSs (the first CDO). The second manager then takes the glasses from the bottom layer (equity tranche) in the previous CDO. In the boom, it will be no problem whenever everyone is paying their mortgages in the first bottle and these glasses are generating payment in payment day. Both pyramids of glasses are full.

However, with 50% defaults, the bottom glasses in the first pyramid are not filled up. At payment day of the second CDO, the cork pops off and generates zero. In this case, even nothing is filled at the top glasses of the second CDO, they are still rated AAA, as if they were safe as the original assets. In the situation that housing market persists weak and people default or stop paying mortgages, there will be less and less money come out from the first bottle into the pyramid of glasses, and less and less money these AAA highly rated securities in the second CDO would make. Thousands of millions of dollars have been invested into this kind of secondary CDOs.

Similarly, a third CDO can be created, which repack MBSs in the first and second CDO. Then, a fourth, a fifth, , and so on. We can easily imagine the much more complicated situation if we refer to the volume of the CDOs market. Many investment banks are involved in the enormous web by the CDOs contracts.

This is the right problem. Investors have packed a lot of funds with these securities which are now not paying anything and liable never to pay anything again. Many financial situations start to teeter because of being tied in numbers of contracts. They cannot unravel these deals. One failure in the web starts to drag down the rest of the system and suck people down in the end. Nobody knows how big it is, how far it is and who are actually involved. We only find out whenever a company began to collapse, suddenly, the second finds itself was dragged with it and then the third, the forth and so on. That's the situation we are exactly in today.

4. The Current Financial Crisis

The current unprecedented financial crisis started from the *US subprime mortgage financial crisis*, and spreaded and accelerated by the securitisation of subprime mortgages into credit derivatives, especially into kinds of CDOs. As the crisis develops, the real economy has been

been obviously seen severely affected since late 2008.

Alongside the stock bubble of mid-1990s, the US housing bubble grew up and began to burst in early 2007 as the building boom led to so much over-supply that house prices could no longer be supported, which evolved into the so-called *US subprime mortgage financial crisis*.

The over expansion of credit in US housing market led to losses by financial institutions. Initially the companies affected were those directly involved in home construction and mortgage lending such as the *Northern Rock* and *Countrywide Financial*. Take *Northern Rock*, a major British bank, as an example. It raises most of the money, which it provides for mortgages via the wholesale credit market, primarily by selling the debt on in the form of bonds. Following the widespread losses made by investors in the subprime mortgage market, these banks and investors have become wary of buying mortgage debt, including *Northern Rock's*. The highly leveraged nature of its business led the bank to request security from the *Bank of England*. News of this lead to investors panic and a bank run in mid-September 2007. *Northern Rock's* problems proved to be an early indication of the severe troubles that would soon befall other banks and financial institutions.

The crisis then began to affect general availability of credit to non-housing related businesses and to large financial institutions not directly connected with mortgage lending. It is the “*securitisation*” process, which spreads the current crisis. Many subprime mortgages were securitised and sold to investors using *asset-backed securities (ABSs)*. It has been estimated that 54% of subprime mortgages were securitised in 2001 and this rose to 75% in 2006 [1]. At the heart of the portfolios of many financial institutions were investors whose assets had been derived from bundled home mortgages.

In early 2007, when defaults were rising in the mortgage market, New York's *Wall Street* began to feel the first tremors in the CDOs world. Hedge fund managers, commercial and investment banks, and pension funds, all of which had been big buyers of CDOs, found themselves landed in trouble, as many CDOs included derivatives that were built upon mortgages—including risky, subprime mortgages. More importantly, the mathematical models that were supposed to protect investors against risk weren't working. The complicating matter was that there was no market on which to sell the CDOs. CDOs are not traded on exchanges and even not really structured to be traded at all. If one had a CDO in his/her portfolio, then there was not much he/she could do to unload it. The CDO managers were in a similar bind. As fear began to spread, the market for CDOs' underlying assets also began to disappear. Suddenly it was impossible to dump the swaps, subprime-mortgage derivatives, and other securities held by the CDOs.

In March 2008, slightly more than a year after the first indicator of troubles in the CDO market, *Bear Stearns*, which was one of *Wall Street's* biggest and most prestigious firms and which had been engaged in the securitisation of mortgages, fell prey and was acquired by *JP Morgan Chase* through the deliberate assistance from the US government. By the middle of 2008, it became clear that no one was safe; everyone—even those who had never invested in anything—would wind up paying the price. On September 15, 2008, the 158 year-old *Lehman Brothers* filed for *Chapter 11 bankruptcy protection*. The collapse of *Lehman Brothers* is the largest investment bank failure since *Drexel Burnham Lambert* in 1990 and triggered events that seemed unthinkable a year before: the high volatility of worldwide financial institutions, massive state-funded bailouts of some of the world's leading financial institutions and the disappearance of investment banks.

The 94 year-old *Merrill Lynch* accepted a purchase offer by *Bank of America* for approximately US\$ 50 billion, a big drop from a year-earlier market valuation of about US\$ 100 billion. A credit rating downgrade of the large insurer *American International Group (AIG)* led to a rescue agreement on September 16, 2008 with the *Federal Reserve Bank* for a \$ 85 billion dollar secured loan facility, in exchange for a warrants for 79.9% of the equity of *AIG*. Even, in January 2009, *HBOS*, a banking and insurance group in UK, was taken over by *Lloyds TSB Banking Group*.

The crisis is now much far beyond the *virtual economy*, the *real economy* has been severely affected. The global economy is in the midst of a deep downturn. The dramatic intensification of the financial crisis has generated historic declines in consumer and business confidence, steep falls in household wealth, and severe disruptions in credit intermediation.

In the last quarter of 2008, industrial production has fallen precipitously across both advanced and emerging economies, declining by some 15%–20% and merchandise exports have fallen by some 30%–40%, at an annual rate. Official figures show that the Britain industrial production dived at record speed, underlying how hard the global downturn has hit producers and exporters: manufacturing outputs dropped by 2.9% in January, 2009, taking annual rate of decline to 12.8%, which is the biggest decline since January, 1981; broader industrial productions, including mining and utilities, are now also falling at an annual rate of 11.4%, again the worst since 1981.

Labor markets are weakening rapidly, particularly in those advanced economies. In February, 2009, the US unemployment rate rose to 8.1%, the highest in more than 25 years and more layoffs are on the way; the Britain unemployment rate rose to 6.3%, up 1.1% on 2008, and the Euro unemployment rate rose to 8.2%, the highest level in over 2 years.

Despite production cut-backs by *OPEC*, oil prices have declined by nearly 70% since their July 2008 peak. Similarly, metals prices are now around 50% below their March 2008 peaks. Food prices have eased 35% from their peak, reflecting not only deteriorating global cyclical conditions, but also favorable harvests.

Clearly, the global economy faces a contraction in overall *Gross Domestic Product (GDP)* for the first time since the Second World War, as claimed by Dominique Strauss-Kahn, the head of the *International Monetary Fund (IMF)*.

5. Mathematical Challenges in Modeling the Mechanism of CDOs

The investment banks presented CDOs as investments in which, actually, the key factors were not the underlying assets, rather the use of mathematical calculations to create and distribute the cash flows. In other words, the basis of a CDO was not a mortgage, a bond or even a derivative, but the metrics and algorithms of quants and traders. In particular, the CDO market skyrocketed in 2001 with the invention of a formula called the *Gaussian Copula*, which made it easier to price CDOs quickly. But what seemed to be the great strength of CDOs—complex formulas that protected against risk while generating high returns—turned out to be flawed.

Normally financial institutions do not trade instruments unless they have satisfactory models for valuing them. What is surprising about the financial crisis is that financial institutions were prepared to trade senior tranches of an ABS (*i.e.*, an *asset-backed security*) or an ABS CDO (an instrument in the synthetic CDOs market) without a model [1]. The lack of a model makes risk management almost impossible and causes problems when the instrument ceases to be rated. Because models were not developed, the key role of correlation in valuing ABSs and (particularly) ABS CDOs was not well understood. Many investors and analysts assumed that CDOs were diversified, and hence made less risky, due to the large number of individual bonds that might underlie a given deal. In fact, the investments within the CDOs turned out to be more highly correlated than expected.

Pricing a CDO is mainly to find the appropriate spread for each tranche and its difficulty lies in how to estimate the default correlation in formulating models that fit market data. With the empirical evidence of the existence of mean reversion phenomena in efficient credit risk markets, mean-reverting type stochastic differential equations are considered (*cf.* *e.g.* [2]). In addition, the CDOs market has seen the phenomenon of heavy tail dependence in a portfolio, which draws the attention to use modeling with heavy tail phenomenon as a feature. Besides, the efficiency in calibrating pricing models to

market prices should be paid much attention. A well-calibrated and easily implemented model is the right goal.

The market standard model is the so-called one factor Gaussian copula model. Its origins can be found in [3,4]. The assumptions of the one factor Gaussian copula model about the characteristics of the underlying portfolio simplify the analytical derivation of CDOs premiums but are not very realistic. Thereafter more and more extensions have been proposed to pricing CDOs: homogeneous infinite portfolio is extended to homogeneous finite portfolio, and then to heterogeneous finite portfolio which represents the most real case; multi-factor models are considered other than one factor model; Gaussian copula is replaced by alternative probability distribution functions; the assumptions of constant default probability, constant default correlation and deterministic loss given default are relaxed and stochastic ones are proposed which incorporate dynamics into pricing models.

In one line of thinking, to relax the assumption of Gaussian distribution in the one factor Gaussian copula model, student-t copula [5–12], double-t copula [13,14], Clayton copula [12,15–20], Archimedian copula [21,22], Marshall Olkin copula [23–26] are studied. And default correlations are made stochastic and correlated with the systematic factor in [27,28] to relax the assumption that default correlations are constant through time and independent of the firms default probabilities. Hull and White propose the implied copula method in [29].

In the other line of thinking, many stochastic processes are applied in CDOs pricing models to describe the default dependence. Markov chains are used to represent the distance to default of single obligor (eg. [30,31]). Then correlation among obligors is introduced with nonrecombining trees [30] or via a common time change of affine type [31].

Some researchers include jumps in CDOs pricing model. Duffie and Garleanu [32], for example, propose an approach based on affine processes with both a diffusion and a jump components. To improve tractability, Chapovsky, Rennie and Tavares [33] suggest a model in which default intensities are modeled as the sum of a compensated common random intensity driver with tractable dynamics (e.g. the Cox-Ingersoll-Ross model (or CIR model in short) with jumps) and a deterministic name-dependent function.

Motivated by the possibility that price processes could be pure jump, several authors have focused their attention on pure jump models in the Lévy class. Firstly, we have the Normal Inverse Gaussian (NIG) model of Barndorff-Nielsen [34], and its generalisation to the generalised hyperbolic class by Eberlein, Keller, and Prause [35]. Kalemánova, Schmid and Werner [36] and Guégan and Houdain [37] work with NIG factor model.

Secondly, we have the symmetric Variance Gamma (VG) model studied by Madan and Seneta [38] and its asymmetric extension studied by Madan and Milne [39], Madan, Carr, and Chang [40]. Baxter [41] introduces the B-VG model where has both a continuous Brownian motion and a discontinuous variance—Gamma jump terms. Finally, we have the model developed by Carr, Geman, Madan, and Yor (acronym: CGMY) [42], which further generalises the VG model. Most of these models are special cases of the generic one-factor Lévy model supposed in [43]. Lévy models bring more flexibility into the dependence structure and allow tail dependence.

Besides default dependence, the recovery rate is also an important variable in pricing CDOs. Empirical recovery rate distributions in [44] have high variance and the certainty with which one can predict recovery is quite low. One undisputed fact about recovery rates is that average recovery rates tend to be inversely related to default rates: in a bad year, not only are there many defaults, but recoveries are also low. The loss process models involve the development of a model for the evolution of the losses on a portfolio. Graziano and Rogers [45] provide semi-analytic formulas via Markov chain and Laplace transform techniques which are both fast and easy to implement. In [46], Schoenbucher derives the loss distribution of the portfolio from the transition rates of an auxiliary time-inhomogeneous Markov chain and stochastic evolution of the loss distribution is obtained by equipping the transition rates with stochastic dynamics. Other loss process can be found in [47] where discuss a dynamic discrete-time multi-step Markov loss model and in [48] where loss follows a jump process.

6. Modeling Heavy Tail Phenomena by Lévy Distributions

From the mathematical view point, we see the highly complexity and chaotic dynamics in the system of CDOs, and, especially, the phenomenon of heavy tail dependence. In literatures, researchers have investigated quite a lot of models in pricing CDOs, but seldom incorporate Lévy stable distributions to represent the heavy tail dependence in the modeling. In this final section, we shall explicate and suggest an idea about applying Lévy stable distributions in pricing CDOs.

Historically, the application of probability distributions in mathematical modeling for the real world problems started with the use of Gaussian distributions to express errors in measurement. Concurrently with this, mathematical statistics emerged. The mean of a Gaussian distribution traditionally represents the most probable value for the actual size and the variance of it is related to the errors of the measurement. The whole distribution is in fact a prediction which is easy to check, since it was

developed for probability distributions which can be well characterized by their first two moments.

Other probability distributions have appeared in mathematical modeling where the mean and variance can not well represent the process. For example, it is well-known that all moments, of the lognormal distribution are finite but . This fact shows that a lot of weight is in the tail of the distribution where rare but extreme events can occur. This phenomenon is the so-called heavy tail dependence phenomenon and is exactly the one observed in the market for CDOs.

Moreover, probability distributions with infinite moments are also encountered in the study of critical phenomena. For instance, at the critical point one finds clusters of all sizes while the mean of the distribution of clusters sizes diverges. Thus, analysis from the earlier intuition about moments had to be shifted to newer notions involving calculations of exponents, like e.g. Lyapunov, spectral, fractal etc., and topics such as strange kinetics and strange attractors have to be investigated.

It was Paul Lévy who first grappled in-depth with probability distributions with infinite moments. Such distributions are now called Lévy distributions. Today, Lévy distributions have been expanded into diverse areas including turbulent diffusion, polymer transport and Hamiltonian chaos, just to mention a few. Although Lévy's ideas and algebra of random variables with infinite moments appeared in the 1920s and the 1930s (cf. [49,50]), it is only from the 1990s that the greatness of Lévy's theory became much more appreciated as a foundation for probabilistic aspects of chaotic dynamics with high entropy in statistical analysis in mathematical modelling (cf. [51,52], see also [53,54]). Indeed, in statistical analysis, systems with highly complexity and (nonlinear) chaotic dynamics became a vast area for the application of Lévy processes and the phenomenon of dynamical chaos became a real laboratory for developing generalizations of Lévy processes to create new tools to study nonlinear dynamics and kinetics. Following up this point, Lévy type processes and their influence on long time statistical asymptotic will be unavoidably encountered.

As a flavor on this aspect, let us finally give a brief account for modelling the risk with Lévy processes within the framework of the intensity based models.

Relative to the copula approach, intensity based models has the advantage that the parameters have economic interpretations. Furthermore, the models, by nature, deliver stochastic credit spreads and are therefore well-suited for the pricing of CDOs tranches. In the intensity based model, default is defined as the first jump of a pure jump process, and it is assumed that the jump process has an intensity process. More formally, it is assumed that a non-negative process λ exists such that the process

$$M(t) := 1_{\{\tau \leq t\}} - \int_0^t 1_{\{\tau > s\}} \lambda(s) ds$$

is a martingale. And the default correlation is generated through dependence of firms' intensities on the common factor.

Following Mortensen [55], we assumes that default of obligor is modelled as the first jump of a Cox process with a default intensity composed of a common and an idiosyncratic component in the following way

$$\lambda_i(t) = a_i X_c(t) + X_i(t)$$

where $a_i > 0$ is a constant and X_c and X_i are independent Lévy processes. Namely, the two independent processes X_c and X_i are of the following form

$$\begin{aligned} X_c(t) &= \theta_c + \sigma_c W_c(t) \\ &+ \sigma_c \int_0^t \int_{0 < |x| < 1} x [N_c(ds, dx) - ds \mu_c(dx)] \\ &+ \sigma_c \int_0^t \int_{|x| \geq 1} x N_c(ds, dx) \end{aligned}$$

and

$$\begin{aligned} X_i(t) &= \theta_i t + \sigma_i W_i(t) \\ &+ \sigma_i \int_0^t \int_{0 < |x| < 1} x [N_i(ds, dx) - ds \mu_i(dx)] \\ &+ \sigma_i \int_0^t \int_{|x| \geq 1} x N_i(ds, dx) \end{aligned}$$

where mean values θ_c , θ_i and volatilities are constants, σ_c , $\sigma_i > 0$ are constants, $W_c(t)$, $W_i(t)$ are independent Brownian motions on $(\Omega, \mathcal{F}, P; \{F_t\}_{t \geq 0})$; $N_c(t, A)$, $N_i(t, A)$ are defined to be the numbers of jumps of process X_c , X_i with size smaller than A during time period t . For fix A , $N_c(t, A)$, $N_i(t, A)$ are Poisson processes with intensity $\mu_c(A)$, $\mu_i(A)$ respectively.

Based on these assumptions, we may calculate marginal default probability, joint default probability and the characteristic function of the integrated common risk factor to get the expression of expected tranche losses. Finally, we may get the tranche spreads of CDOs. We will realize this aim in our forthcoming work towards the concrete mathematical modelling.

7. Conclusions

In this paper, we start with detailed explanation of the mechanism of CDOs and discuss the mathematical challenge in modelling the complexity systems arising from CDOs. We link the feature of CDOs with heavy tail phenomenon and then propose to use Lévy process, in particular Lévy stable process, to model risk factors in pricing CDO tranche spreads.

Our paper shows Lévy stable distribution may capture the feature of high default dependence among CDOs' underlying portfolio.

8. Acknowledgments

We would like to thank Claire Geleta of Deutsche Bank Trust Company Americas at Los Angeles for useful conversation regarding to CDOs. We also thank the referee for constructive comments on our previous manuscript.

9. References

- [1] J. Hull, "The credit crunch of 2007: What went wrong? Why? What lessons can be learned?" Working Paper, University of Toronto, 2008.
- [2] J. L. Wu and W. Yang, "Pricing CDO tranches in an intensity-based model with the mean-reversion approach," Working Paper, Swansea University, 2009.
- [3] O. Vasicek, "Probability of loss on a loan portfolio," Working Paper, KMV (Published in Risk, December 2002 with the title Loan Portfolio Value), 1987.
- [4] D. X. Li, "On default correlation: A Copula approach," Journal of Fixed Income, Vol. 9, No. 4, pp. 43–54, 2000.
- [5] L. Andersen, J. Sidenius, and S. Basu, "All your hedges in one basket," Risk, pp. 67–72, 2003.
- [6] S. Demarta and A. J. McNeil, "The t copula and related copulas," International Statistical Review, Vol. 73, No. 1, pp. 111–129, 2005.
- [7] P. Embrechts, F. Lindskog, and A. McNeil, "Modelling dependence with Copulas and applications to risk management," In Handbook of Heavy Tailed Distributions in Finance, edited by S. Rachev, Elsevier, 2003.
- [8] R. Frey and A. J. McNeil, "Dependent defaults in models of portfolio credit risk," Journal of Risk, Vol. 6, No. 1, pp. 59–92, 2003.
- [9] A. Greenberg, R. Mashal, M. Naldi, and L. Schloegl, "Tuning correlation and tail risk to the market prices of liquid tranches," Lehman Brothers, Quantitative Research Quarterly, 2004.
- [10] R. Mashal and A. Zeevi, "Inferring the dependence structure of financial assets: Empirical evidence and implications," Working paper, University of Columbia, 2003.
- [11] R. Mashal, M. Naldi, and A. Zeevi, "On the dependence of equity and asset returns," Risk, Vol. 16, No. 10, pp. 83–87, 2003.
- [12] L. Schloegl and D. O'Kane, "A note on the large homogeneous portfolio approximation with the student-t Copula," Finance and Stochastics, Vol. 9, No. 4, pp. 577–584, 2005.
- [13] J. Hull and A. White, "Valuation of a CDO and an n-th to default CDS without Monte Carlo simulation," Journal of Derivatives, Vol. 12, No. 2, pp. 8–23, 2004.
- [14] A. Cousin and J. P. Laurent, "Comparison results for credit risk portfolios," Working Paper, ISFA Actuarial School, University of Lyon and BNP-Paribas, 2007.
- [15] P. Schönbucher and D. Schubert, "Copula dependent default risk in intensity models," Working Paper, Bonn University, 2001.
- [16] J. Gregory, and J. P. Laurent, "I will survive," Risk, Vol. 16, No. 6, pp. 103–107, 2003.
- [17] E. Rogge and P. Schönbucher, "Modelling dynamic portfolio credit risk," Working Paper, Imperial College, 2003.
- [18] D. B. Madan, M. Konikov, and M. Marinescu, "Credit and basket default swaps," The Journal of Credit Risk, Vol. 2, No. 2, 2006.
- [19] J. P. Laurent and J. Gregory, "Basket default swaps, CDOs and factor copulas," Journal of Risk, Vol. 7, No. 4, pp. 103–122, 2005.
- [20] A. Friend and E. Rogge, "Correlation at first sight," Economic Notes, Vol. 34, No. 2, pp. 155–183, 2005.
- [21] P. Schönbucher, "Taken to the limit: Simple and not-so-simple loan loss distributions," Working Paper, Bonn University, 2002.
- [22] T. Berrada, D. Dupuis, E. Jacquier, N. Papageorgiou, and B. Rémillard, "Credit migration and basket derivatives pricing with copulas," Journal of Computational Finance, Vol. 10, pp. 43–68, 2006.
- [23] D. Wong, "Copula from the limit of a multivariate binary model," Working Paper, Bank of America Corporation, 2000.
- [24] Y. Elouerkhaoui, "Credit risk: Correlation with a difference," Working Paper, UBS Warburg, 2003.
- [25] Y. Elouerkhaoui, "Credit derivatives: Basket asymptotics," Working Paper, UBS Warburg, 2003.
- [26] K. Giesecke, "A simple exponential model for dependent defaults," Journal of Fixed Income, Vol. 13, No. 3, pp. 74–83, 2003.
- [27] L. Andersen and J. Sidenius, "Extensions to the Gaussian copula: Random recovery and random factor loadings," Journal of Credit Risk, Vol. 1, No. 1, pp. 29–70, 2005.
- [28] L. Schloegl, "Modelling correlation skew via mixing Copula and uncertain loss at default," Presentation at the Credit Workshop, Isaac Newton Institute, 2005.
- [29] J. Hull and A. White, "Valuing credit derivatives using an implied copula approach," Journal of Derivatives, Vol. 14, No. 2, pp. 8–28, 2006.
- [30] C. Albanese, O. Chen, A. Dalessandro, and A. Vidler, "Dynamic credit correlation modeling," Working Paper, Imperial College, 2005.
- [31] T. R. Hurd and A. Kuznetsov, "Fast CDO computations in affine Markov chains models," Working Paper, McMaster University, 2006.
- [32] D. Duffie and N. Gârleanu, "Risk and valuation of collateralized debt obligations," Financial Analysts Journal, Vol. 57, No. 1, pp. 41–59, 2001.
- [33] A. Chapovsky, A. Rennie and P. A. C. Tavares, "Stochastic intensity modeling for structured credit exotics," Working paper, Merrill Lynch International, 2006.

- [34] O. E. Barndorff-Nielsen, "Processes of normal inverse Gaussian type," *Finance and Stochastics*, Vol. 2, No. 1, pp. 41–68, 1998.
- [35] E. Eberlein, U. Keller and K. Prause, "New insights into smile, mispricing and value at risk," *Journal of Business*, Vol. 71, No. 3, pp. 371–406, 1998.
- [36] A. Kalemanna, B. Schmid and R. Werner, "The Normal Inverse Gaussian distribution for synthetic CDO pricing," *Journal of Derivatives*, Vol. 14, No. 3, pp. 80–93, 2007.
- [37] Guégan and Houdain, "Collateralized debt obligations pricing and factor models: A new methodology using Normal Inverse Gaussian distributions," *Note de recherche IDHE-MORA, ENS Cachan*, No. 007-2005, 2005.
- [38] D. B. Madan and E. Seneta, "The Variance Gamma (V.G.) model for share market returns," *Journal of Business*, Vol. 63, No. 4, pp. 511–24, 1990.
- [39] D. B. Madan and F. Milne, "Option pricing with V.G. martingale components," *Mathematical Finance*, Vol. 1, No. 4, pp. 39–55, 1991.
- [40] D. B. Madan, P. Carr and E. C. Chang, "The variance gamma process and option pricing," *European Finance Review*, Vol. 2, No. 1, pp. 79–105, 1998.
- [41] M. Baxter, "Lévy process dynamic modeling of single-name credits and CDO tranches," *Working Paper, Nomura International*, 2006.
- [42] P. Carr, H. Geman, D. Madan, and M. Yor, "The fine structure of asset returns: An empirical investigation," *Journal of Business*, Vol. 75, No. 2, pp. 305–332, 2002.
- [43] H. Albrecher, S. A. Landoucette, and W. Schoutens, "A generic one-factor Lévy model for pricing synthetic CDOs," In: *Advances in Mathematical Finance*, R. J. Elliott *et al.* (eds.), Birkhäuser, Boston, 2007.
- [44] Gupton and Stein, "Losscalc: Moody's model for predicting loss given default(LGD)," *Working Paper, Moody's Investor services*, 2002.
- [45] G. Graziano and L. C. G. Rogers, "A dynamic approach to the modeling of correlation credit derivatives using markov chains," *Working Paper, University of Cambridge*, 2005.
- [46] P. Schönbucher, "Portfolio losses and the term structure of loss transition rates: A new methodology for the pricing portfolio credit derivatives," *Working Paper, ETH Zurich*, 2006.
- [47] M. Walker, "Simultaneous calibration to a range of portfolio credit derivatives with a dynamic discrete-time multi-step loss model," *Working Paper, University of Toronto*, 2007.
- [48] F. Longstaff and A. Rajan, "An empirical analysis of the pricing of collateralized debt obligations," *Journal of Finance*, Vol. 63, No. 2, pp. 529–563, 2008.
- [49] P. Lévy, "Calcul des probabilités," *Gauthier-Villars*, 1925.
- [50] P. Lévy, "Théorie de l'addition des variables aléatoires," *Gauthier-Villars*, 1937.
- [51] G. Samorodnitsky and M. S. Taqqu, "Stable non-Gaussian random processes: Stochastic models with infinite variance," *Chapman and Hall/CRC*, Boca Raton, 1994.
- [52] M. F. Shlesinger, G. M. Zaslavsky, and U. Frisch (Eds.), "Lévy flights and related topics in physics," *Lecture Notes in Physics*, Vol. 450, Springer-Verlag, Berlin, 1995.
- [53] B. Mandelbrot, "The Pareto-Lévy law and the distribution of income," *International Economic Review*, Vol. 1, pp. 79–106, 1960.
- [54] V. M. Zolotarev, "One-dimensional stable distributions," *American Mathematical Society*, R. I. Providence, 1986.
- [55] A. Mortensen, "Semi-analytical valuation of basket credit derivatives in intensity-based models," *Journal of Derivatives*, Vol. 13, No. 4, pp. 8–26, 2006.

