

Ministère de l'Enseignement Supérieur et de la Recherche Scientifique



Université Akli Mohand Oulhadj – Bouira, Algérie

**Laboratoire d'Informatique, Mathématique,
Physique pour l'Agriculture et Forêts (LIMPAF)**



COSI'17

**Actes
du 14^{ème} Colloque sur
l'Optimisation et les Systèmes
d'Information**



Bouira, 14 – 16 Mai 2017

Actes du Quatorzième Colloque sur l'Optimisation et les Systèmes d'Information COSI'2017

du 14 au 16 Mai 2017, Bouira, Algérie

Organisation

Université Akli Mohand Oulhadj, Bouira (Algérie)

Président d'honneur

Professeur Moussa ZEREG, Recteur de l'Université de Bouira - Akli Mohand Oulhadj (Algérie)
Professeur Ahcène ARBAOUI, Doyen de la faculté des sciences et sciences appliqués, Bouira (Algérie)

Présidente du colloque

Dr. Kahina LOUADJ, Université de Bouira - Akli Mohand Oulhadj (Algérie)

Vice-Présidents du colloque

Professeur Djamel BENNOUAR, Directeur du laboratoire LIMPAF, Université de Bouira
Dr. Abdelhamid DJAIDJA, Vice-doyen de la recherche scientifique, Université de Bouira

Membres

YOUSFI Khelifa, DAHMANE Merzak, BOUDANE Khadidja, BOUDANE Massiva, DJOUDER Sofiane, ALLEM Lala Maghnia, TAKHEDMIT Baya, BOUGHANI L'hadi, BESSAD Baya, BEKRI Houria, IMINE Ouiza, SMAILI Nesserine, BOUDREF Mohamed, ABBAS Akli, HAMDACHE Sid Ali, OUKIL Nacima, WAHID Guettaf, GHEMMOUR Karim

Comité de Pilotage

Mohamed AIDENE, Université Mouloud Mammeri de Tizi-Ouzou, Algérie
Mohand-Saïd HACID, Université Claude Bernard, Lyon I, France
Nadjet KAMEL, Université Ferhat Abbas Sétif 1, Algérie.
Lhouari NOURINE, Université Clermont Auvergne, Clermont-Ferrand, France
Samia OURARI, CDTA, Alger, Algérie
Jean Marc PETIT, INSA de Lyon, France
Mohamed Said RADJEF, Université de Béjaia, Algérie
Bachir SADI, Université Mouloud Mammeri de Tizi-Ouzou, Algérie
Lakhdar SAIS, Université d'Artois, France
Farouk TOUMANI, Université Clermont Auvergne, Clermont-Ferrand, France
Rachid NOURINE, Université d'Oran 1, Algérie

Comité de programme

Président

Mourad BAIYOU, Université Clermont Auvergne - CNRS, France

Membres

MOHAMED AHMED NACER Université des Sciences et Technologie Houari Boumediene (USTHB), Alger (Algérie)
MEZIANE AIDER Université des Sciences et Technologie Houari Boumediene (USTHB), Alger (Algérie)
OTMANE AIT MOHAMED Concordia University, Montreal, Quebec (Canada)
HACENE AIT-HADDADENE Université des Sciences et Technologie Houari Boumediene (USTHB), Alger (Algérie)
ZAIA ALIMAZIGHI Université des Sciences et Technologie Houari Boumediene (USTHB), Alger (Algérie)
MAKHLOUF ALIOUAT Université Ferhat Abbas Setif 1, Sétif (Algérie)
ZIBOUDA ALIOUAT Université Ferhat Abbas Setif 1, Sétif (Algérie)
ADEL ALTI Université Ferhat Abbas Setif 1, Sétif (Algérie)
HASSAN AÏT-KACI Self
NADJIB BADACHE CERIST, Alger (Algérie)
KAMEL BARKAOUI Conservatoire national des arts et métiers (CNAM), Paris (France)
VINCENT BARRA Université Clermont Auvergne, Clermont-Ferrand (France)
LADJEL BELLATRECHE ISAE-ENSMA - École Nationale Supérieure de Mécanique et d'Aérotechnique, Poitiers (France)
KHALID BENABDESLEM Université Claude Bernard Lyon I, Lyon (France)
MOUSSA BENAÏSSA Université d'Oran, Oran (Algérie)
NACÉRA BENAMRANE Laboratoire SIMPA, Département d'Informatique, USTOMB (Algérie)
SALIMA BENBERNOU Université Paris Descartes, Paris (France)
FATIHA BENDALI Université Clermont Auvergne, Clermont-Ferrand (France)
BELAÏD BENHAMOU Université d'Aix-Marseille, Marseille (France)
ABDELHAFID BERRACHEDI Université des Sciences et Technologie Houari Boumediene (USTHB), Alger (Algérie)
ISMA BOUCHEMAKH Université des Sciences et Technologie Houari Boumediene (USTHB), Alger (Algérie)
MOURAD BOUDHAR Université des Sciences et Technologie Houari Boumediene (USTHB), Alger (Algérie)
MOHAND BOUGHANEM Université Paul Sabatier Toulouse III, Toulouse (France)
MOKRANE BOUZEGHOUB Université de Versailles Saint-Quentin-en-Yvelines - UVSQ, Versailles (France)
BRICE CHARDIN ISAE-ENSMA - École Nationale Supérieure de Mécanique et d'Aérotechnique, Poitiers (France)
BRUNO CREMILLEUX Université de Caen Basse-Normandie, Caen (France)
JEAN-PIERRE CROUZEIX Université Clermont Auvergne, Clermont-Ferrand (France)
LAURENT D'ORAZIO Université Rennes 1, Rennes (France)
BENTERKI DJAMEL Université Ferhat Abbas Setif 1, Sétif (Algérie)
HABIBA DRIAS University of Science and Technology USTHB Algiers
DIDI FEDOUA Université Abou Bekr Belkaid, Tlemcen (Algérie)

FREDERIC FLOUVAT University of New Caledonia (France)
 PIERRE FOUILHOX Université Pierre et Marie Curie Paris 6, Paris (France)
 PATRICK GALLINARI Université Pierre et Marie Curie Paris 6, Paris (France)
 ZAHIA GUESSOUM Université Pierre et Marie Curie Paris 6, Paris (France)
 MOHAND-SAÏD HACID Université Claude Bernard Lyon I, Lyon (France)
 SONIA HADDAD-VANIER Université Paris 1 (France)
 ALLEL HADJALI ISAE-ENSMA - École Nationale Supérieure de Mécanique et d'Aérotechnique, Poitiers (France)
 YOUSSEF HAMADI Microsoft Research, Cambridge (UK)
 SAÏD HANAFI Université de Valenciennes et du Hainaut Cambrésis (France)
 JIN-KAO HAO University of Angers, Angers (France)
 SOUHILA KACI Université de Montpellier II (France)
 NADJET KAMEL University Ferhat Abbas Setif 1, Sétif (Algérie)
 OKBA KAZAR Université de Biskra (Algérie)
 HAMAMACHE KHEDDOUCI INSA de Lyon, Université Lyon 1 (France)
 NACIMA LABADIE University of Technology of Troyes (France)
 ARNAUD LALLOUET Huawei Technologies Ltd, Paris (France)
 DOMINIQUE LAURENT Université Cergy-Pontoise, Paris (France)
 YAHIA LEBBAH Université d'Oran 1, Oran (Algérie)
 VINCENT LIMOUZY Université Clermont Auvergne, Clermont-Ferrand (France)
 SAMIR LOUDNI Université de Caen Basse-Normandie, Caen (France)
 SOFIAN MAABOUT University of Bordeaux 1, Bordeaux (France)
 PHILIPPE MAHEY Université Clermont Auvergne, Clermont-Ferrand (France)
 RIDHA MAHJOUB Institut des sciences de l'information et de leurs interactions (INS2i) - CNRS, Paris (France)
 ARNAUD MARY Université Claude Bernard Lyon I, Lyon (France)
 NOUREDINE MELAB Université Lille 1 / INRIA Lille / CNRS CRISTAL (France)
 ENGELBERT MEPHU NGUIFO Université Clermont Auvergne, Clermont-Ferrand (France)
 HAYET FARIDA MEROUANI University of Badji-Mokhtar Annaba (Algérie)
 ROKIA MISSAOUI UQAM - Université du Québec à Montréal (Canada)
 AIDENE MOHAMED Université Mouloud Mammeri de Tizi-Ouzou (Algérie)
 OUANES MOHAND Université Mouloud Mammeri de Tizi-Ouzou (Algérie)
 BIBI MOHAND OUAMER Université Abderrahmane Mira de Bêjaïa (Algérie)
 SAFIA NAIT BAHLOUL Université d'Oran 1, Oran (Algérie)
 RACHID NOURINE Université d'Oran 1, Oran (Algérie)
 BRAHIM OUKACHA Université Mouloud Mammeri de Tizi-Ouzou (Algérie)
 SAMIA OURARI CDTA, Alger (Algérie)
 HACÈNE OUZIA Université Pierre et Marie Curie Paris 6, Paris (France)
 JEAN-MARC PETIT Université de Lyon, INSA Lyon (France)
 AHMED-OUAMER RACHID Université Mouloud Mammeri de Tizi-Ouzou (Algérie)
 MOHAMMED-SAÏD RADJEF Université Abderrahmane Mira de Bêjaïa (Algérie)
 MICHAEL RAO CNRS - LIP - ENS Lyon (France)
 ALLAOUA REFOUFI Université Ferhat Abbas Sétif 1, Sétif (Algérie)
 BACHIR SADI Université Mouloud Mammeri de Tizi-Ouzou (Algérie)
 YAKOUB SALHI Université d'Artois, Lens (France)
 YACINE SAM Institut Universitaire de Technologie de Blois (France)
 FRÉDÉRIC SAUBION Université d'Angers (France)
 MICHEL SCHNEIDER Université Clermont Auvergne, Clermont-Ferrand (France)
 AMEUR SOUKHAL Université de Tours (France)
 PIERRE SPITERI IRT - ENSEEIHT, Toulouse (France)
 YEHIA TAHER Université de Versailles Saint-Quentin-en-Yvelines - UVSQ, Versailles (France)
 FAROUK TOUMANI Université Blaise Pascal, Clermont-Ferrand (France)
 TAKEAKI UNO National Institute of Informatics (NII), Tokyo (Japan)
 DJEMEL ZIOU Sherbrooke university, Sherbrooke, Quebec (Canada)

Relecteurs additionnels

Abdelfetah, Hentout
Ahmed, Kouider
Andrade, Ricardo
Aribi, Nouredine
Aridhi, Sabeur
Belaid, Mohammed Said
Bentounsi, Mehdi
Bouet, Marinette
De Goër De Herve, Jocelyn
Dhifli, Wajdi
Djamal, Belkasmi
Fournier-Viger, Philippe
Gastaldello, Mattia
Geniet, Annie
Grissa, Dhouha
Jelassi, Mohamed Nader
Kermia, Omar
Kouamou, Georges Edouard
Libo, Ren
Lima, Leandro I. S. De
Meziane, Hassina
Mohamed, Sayah
Mokhtari, Karima
Ndoundam, Rene
Nestor, Koueya
Palmieri, Anthony
Perret Du Cray, Henri
Pons, Luc
Rey, Christophe
Rosenfeld, Matthieu
Saidani, Fayçal Rédha
Talon, Alexandre
Tamen, Zahia
Zghal, Sami

Table des matières

Préface	page 1
----------------------	--------

Conférences invitées

— <i>Graph packing problems : theory and applications</i> Professeur Hamamache KHEDDOUCI, Université Claude Bernard - Lyon 1 (France)	page 2
— <i>Fuzzy Dual Numbers and Optimization under Uncertainty</i> Professeur Felix MORA-CAMINO, Ecole Nationale de l'Aviation Civile - ENAC, Toulouse (France)	page 3
— <i>Validation des Transformations de Modèles : Concepts et Défis</i> Professeur Allaoua CHAOUI, Université Constantine 2 - Abdelhamid Mehri (Algérie)	page 4
— <i>Modeling and Solving a Class of Real-life Bin Packing Problems in Operational Supply Chain</i> Professeur Abderrahmane AGGOUN, KLS OPTIM, Paris (France)	page 5

Articles

— <i>Characterizations of extremal graphs for somme bounds on k-independence number</i> Ahmed Bouchou, Mostafa Blidia	page 6
— <i>The b-domatic number of join and corona graphs</i> Benatallah Mohammed Ikhlef Eschouf Nouredine, Mihoubi Miloud	page 16
— <i>Vérification Formelle des Contraintes Temporelles et Spatio-temporelles des Documents SMIL</i> Fatma Zohra Mekahlia, Abdelghani Ghomari, Sami Yazid	page 24
— <i>A Mathematical Model and a Fast Lower Bound for the Open-End Bin Packing Problem with Conflicts</i> Mohamed Maiza, Mohamed Tebbal, Billal Rabia, Lakhdar Sais	page 36
— <i>Moore bipartite mixed Graphs</i> Ouiza Imine, Méziane Aïder	page 46
— <i>Méthode hybride pour le problème de contrôle optimal en temps minimal</i> Mohand Ouamer Bibi, Imane Alioua	page 56
— <i>A Comprehensive Taxonomy for Omics Data Integration and Analysis Methods</i> Amina Lemsara, Salima Ouadfel, Mohamed Batouche	page 68
— <i>GBAS : Graph Based approach for Arabic Stemming</i> Imene Boukhalifa, Sihem Mostefai	page 78

- *A hybrid direction method for solving linear programs with bounded variables* Khalil Djeloud, Mohand Bentobache, Mohand Ouamer Bibi page 86
- *Modélisation de la Consommation d’Energie dans les réseaux Ad hoc via l’outil des réseaux de Petri stochastiques* Ouiza Lekadir, Adedl-Aissaoui Karima, Bouzit Nacera et Hadidi Nadia page 97
- *Nouvelle approche pour l’optimisation continue : Optimisation par Filtres Morphologiques* Khelifa Chahinez Nour El Houda, Belmadani Abderrahim page 109
- *Fusion de minuties pour une reconnaissance efficiente des empreintes digitales* Mohamed Hedi Ghaddab, Khaled Jouini, Ouajdi Korbaa page 119
- (
- *DSMA : Diffusion Sûre des Messages d’Alerte dans les Vanets* Hadjer Goumidi, Zibouda Aliouat, Makhoulf Aliouat page 131
- *Pareto-Optimization-Based Approach for Feature Selection in Sentiment Analysis* Fayçal Saidani, Allel Hadjali, Djamal Belkasmî page 143
- *Key Managment Based AODV In Vanets Networks* Chaima Bensaid , Boukli Hacène Sofiane ,Faraoun Kame Mohamed page 155
- *Evolutionary Simulated Annealing Algorithm with SVM for Feature Selection and Classification* Brahim Sahmadi, Dalila Boughaci page 170
- *Magnification des images de documents dégradées issues des terminaux mobiles* Zakia Kezzoula, Soumia Faouci, Djamel Gaceb page 182
- *Pretrained Deep Convolutional Neural Networks for Offline Isolated Handwritten Arabic Character Recognition* Chaouki Boufenar, Adlen Kerboua, Mohamed Batouche page 194
- *Smart-Diffusion des gradients et prétraitement des images de documents en couleurs* Ryma Benabdelaziz , Djamel Gaceb , Roza Benabdelaziz page 206
- *Detection Method without crowd behavior modeling by fuzzy logic* Hocine Chebi, Dalila Acheli, and Mohamed Kesraoui page 218
- *Suivi multi-cibles* Houssam Eddine Rouabhia, Brahim Farou, Hamid Seridi, Herman Akdag page 237
- *b-chromatic number of K_m cartesian product with some particular graphs* Fatima Affif Chaouche, Abdelhafid Berrachedi page 248

- *Gradient Local Binary Patterns for Keyword Spotting in Handwritten Documents* Mohamed Lamine Bouibed, Hassiba Nemmour, and Youcef Chibanipage 259

- *A comparative analysis of context-management approaches for the Internet of Things* Farida Retima, Saber Benharzallah, Laid Kahloul, Okba Kazar page 270

Préface

Le Colloque sur l'Optimisation et les Systèmes d'Information (COSI) rentre cette année dans sa quatorzième édition, après celles de Sétif (2016), d'Oran (2015), de Bejaïa (2014), d'Alger (2013), de Tlemcen (2012), de Guelma (2011), de Ouargla (2010), d'Annaba (2009), de Tizi-Ouzou (2008), d'Oran (2007), d'Alger (2006), de Bejaïa (2005) et de Tizi-Ouzou (2004). Cette rencontre annuelle pluridisciplinaire est un modèle du genre, elle permet à des chercheurs Algériens et étrangers travaillant sur des thématiques transversales comme les systèmes d'information, l'Intelligence Artificielle, l'optimisation combinatoire et la théorie des graphes de se rencontrer et de s'imprégner des dernières avancées technologiques. Les divers thèmes abordés admettent souvent des connexions à la fois théoriques - les problèmes posés sont souvent de nature combinatoire et difficiles - et pratiques - les applications cibles font souvent appel à diverses techniques issues des différents thèmes abordés durant le colloque.

COSI est un événement chaleureux et convivial favorisant les échanges scientifiques et humains ! Cette ambiance de travail fait que tout participant en garde un souvenir mémorable. Cet événement scientifique majeur dans le paysage de recherche Algérien, Nord-africain et International continue de susciter un intérêt toujours croissant auprès de la communauté scientifique. Le maintien du niveau de soumission de papiers à un niveau assez élevé (120 papiers soumis), est une preuve indéniable de sa pérennité et de son succès. Les soumissions de COSI 2017 proviennent de différentes villes d'Algérie, de Tunisie, de France, du Canada et du Japon. Cette année le taux d'acceptation est de 20% en papiers longs et 18% en papiers posters. Les articles soumis couvrent les champs disciplinaires traditionnels de COSI : systèmes d'information (24), recherche opérationnelle (41), intelligence artificielle (42) et applications innovantes (31).

Le programme du Colloque comporte également quatre conférenciers invités de renommée internationale. La première plénière, du Professeur Hammamache Kheddouci du LIRIS - Université Claude Bernard (Lyon) porte sur le "Graph Packing Problems : Theory and Applications". Le Professeur Abderrahmane Aggoun de la société KLS OPTIM (Paris) présentera une conférence intitulée "Modeling and Solving a class of Real-life Bin Packing problems in Operational Supply Chain". La troisième, du Professeur Allaoua Chaoui de l'Université de Constantine 2 portera sur la validation des transformations de modèles : Concepts et Défis. La quatrième plénière du Professeur Felix Mora-Camino de l'ENAC (Toulouse) est intitulée "Fuzzy Dual Numbers and Optimization under Uncertainty". Merci à eux de nous avoir fait l'honneur de leur présence à COSI 2017.

Nous tenons à remercier très chaleureusement l'équipe organisatrice et plus particulièrement la Présidente du colloque Kahina Louadj pour son dévouement et son travail remarquable pour la préparation de cette quatorzième édition malgré une situation personnelle difficile. Nous tenons à lui exprimer ici, notre profonde gratitude et à lui souhaiter un prompt rétablissement. Nous sommes de tout cœur avec toi Kahina. Nos remerciements vont également aux vice-présidents du comité scientifique, le Professeur Djamel Bennouar Directeur du Laboratoire LIMPAF et le Docteur Abdelhamid Djaidja Vice-doyen de la recherche scientifique et à tous les membres du comité d'organisation (Yousfi Khelifa, Dahmane Merzak, Boudane Khadidja, Boudane Massiva, Djouder Sofiane, Allem Lala Maghnia, Takhedmit Baya. Boughani L'hadi. Bessad Baya, Bekri Houria, Imine Ouiza. Smaili Nesserine, Boudref Mohamed, Abbas Akli, Hamdache Sid Ali, Oukil Nacima, Hemad Hinda, Wahid Guettaf, Ghemmour Karim). Nous ne pouvons oublier le Professeur Moussa Zereg, Recteur de l'Université de Bouira, le Docteur Ahcène Arbaoui Doyen de la faculté des sciences et sciences appliquées pour leur fort soutien et pour avoir rendu possible l'organisation de cette manifestation scientifique. Nous tenons à remercier tous les auteurs qui ont soumis des articles montrant ainsi la vitalité de cette rencontre scientifique. Nous tenons également à exprimer notre profonde reconnaissance aux membres du comité de programme et aux rapporteurs supplémentaires qui ont fait un excellent travail, et qui ont accepté la surcharge de travail induite par le nombre important de soumissions. Enfin, n'oublions pas que l'organisation de ce colloque ne serait pas possible sans les aides de nos partenaires institutionnels et industriels. Qu'ils reçoivent ici notre profonde reconnaissance.

Tous nos vœux de réussite pour cette rencontre scientifique et longue vie à COSI !

Fait à Clermont-Ferrand, le 2 mai 2017

Mourad Baiou, Lhouari Nourine et Lakhdar Sais

Graph Packing Problems: Theory and Applications

Hamamache Kheddouci

Université Claude Bernard - Lyon 1, France

Abstract. In this talk, I review, classify and discuss several known conjectures, recent advances and results obtained on graph packing problems. In particular, I give bounds and algorithms for these problems and some of their applications in computer science.

Fuzzy Dual Numbers and Optimization under Uncertainty

Felix Mora-Camino

Ecole Nationale de l'Aviation Civile - ENAC, Toulouse, France

Abstract. In general, optimization problems assume implicitly that their parameters (cost coefficients, limit values) are perfectly known while very often for real problems this is not the case. A first approach is to perform around the nominal optimal solution a post optimization sensibility analysis. When some probabilistic information about the values of the uncertain parameters is available, stochastic optimization techniques may provide the most expected optimal solution. When these parameters are only known to remain within some intervals, robust optimization techniques will provide a robust solution. The fuzzy formalism has been also considered in this case as an intermediate approach to represent the parameter uncertainties and provide fuzzy solutions. These different approaches result in general into a very large amount of computation. In this talk, a new formalism based on fuzzy dual numbers is proposed to diminish the computational burden when dealing with uncertainty in mathematical programming problems. In this framework, a solution approach is proposed for mathematical programming problems presenting parameter uncertainty and solution diversion. The adopted formalism considers fuzzy dual numbers which are a simplified version of fuzzy numbers which adopts some elements of dual numbers calculus. The proposed special class of numbers, dual fuzzy numbers integrates the nilpotent operator Epsilon of dual numbers and considers symmetrical fuzzy numbers. Here after introducing the elements of fuzzy dual calculus useful in this study, two classes of fuzzy dual mathematical programming problems are considered: those where uncertainty relays only in the parameters of the problem and those where also the implementation of the solution is subject to uncertainty. Finally fuzzy dual dynamic programming is developed to solve sequential optimization problems under interval uncertainty.

Validation des Transformations de Modèles : Concepts et Défis

Allaoua Chaoui

Universit Constantine 2 - Abdelhamid Mehri

Abstract. L'objectif de cette confrence invite est de donner les concepts fondamentaux de la validation des transformations, un survey des diffrents travaux ainsi que quelques problmes ouverts dans ce domaine de recherche

Modeling and Solving a class of Real-life Bin Packing problems in Operational Supply Chain

Abderrahmane Aggoun

KLS OPTIM, Paris, France

Abstract. KLS OPTIM is an SME specialist in Warehouse Management Systems and Optimization in Logistics, and offers a complete range of proven packages, solutions and services in a niche of sectors in Supply Chain Management. The main problems addressed are packing, design of optimal plan of packing products in cartons and cartons in pallets, optimization in distribution by minimizing the number of pallets, optimization of vehicle/ container loading plans, optimization of assignment of containers in wagons, integration of packing and vehicle routing. Most of the problems are known in the literature as bin packing problems. In this presentation, we are interested in one hand to review a set of problems encountered in logistics, and in a second hand to highlight the contribution of Constraint Programming, Linear Programming, Meta-heuristics and numerical solvers using Non-Linear Inequalities in solving them.

Characterizations of extremal graphs for some bounds on k -independence number

Ahmed Bouchou¹ and Mostafa Blidia²

¹University of Médéa, Algeria.

E-mail: bouchou.ahmed@yahoo.fr.

²LAMDA-RO, Department of Mathematics, University of Blida, Algeria.

E-mail: m_blidia@yahoo.fr.

Abstract

For an integer $k \geq 1$ and a graph $G = (V, E)$, a subset S of V is k -independent if $\Delta(G[S]) < k$. We denote by $\beta_k(G)$, and call k -independence number, the maximum cardinality of a k -independent set of G . In this paper, we give a characterization of graphs that achieve equality in some upper bounds on β_k . Moreover, we give a constructive characterization of trees that achieve a lower bound on β_k . Finally, we provide a characterization of extremal regular graphs of a Nordhaus–Gaddum bound for the k -independence number.

Keywords: k -independence number.

AMS Subject Classification: 05C69.

1 Introduction

We consider simple graphs $G = (V(G), E(G))$ of order $|V(G)| = n(G)$ and size $|E(G)| = m(G)$. The *neighborhood* of a vertex $v \in V$ is $N_G(v) = \{u \in V : uv \in E\}$. The *degree* of a vertex v of G is $d_G(v) = |N_G(v)|$. The *maximum degree* of G is $\Delta(G) = \max\{d_G(v) : v \in V\}$ and the *minimum degree* of G is $\delta(G) = \min\{d_G(v) : v \in V\}$. If $S \subset V$, we denote by $G[S]$ the subgraph induced by S in G and write $N_S(v) = N_G(v) \cap S$. The *path* (*cycle*, *clique*, *star*, respectively) of order n is denoted by P_n (C_n , K_n , $K_{1,n-1}$, respectively). The complement of G is the graph $\overline{G} = (V(G), E(\overline{G}))$ where $E(\overline{G}) = \{uv : uv \notin E(G)\}$. We say that G is *regular* if every vertex has the same degree. If every vertex of G has degree r , we say that G is r -regular. A graph is a *bipartite graph* if its vertex sets can be partitioned in two independent sets. A *tree* is a connected graph with no cycle.

For an integer $k \geq 1$ and a graph $G = (V, E)$, a subset S of V is k -independent if $\Delta(G[S]) < k$. We denote by $\beta_k(G)$, and call k -independence number, the maximum cardinality of a k -independent set of G . A maximum k -independent set of G is also called a $\beta_k(G)$ -set.

Clearly for a graph G of order n and maximum degree Δ , $\beta_1(G)$ is the independence number $\beta(G)$ and $\beta_\Delta(G) < \beta_{\Delta+1}(G) = n$. Moreover, the sequence $(\beta_j(G))_{1 \leq j \leq n}$ is non-decreasing.

The k -independence number was introduced by Fink and Jacobson [5] in 1984. Since its introduction, more than a hundred papers have been published on various aspects of k -independence in graphs (for examples, see [1, 2, 4, 5, 6, 7]).

In this paper, we examine classes of extremal graphs for some bounds on β_k . We give a characterization of graphs that achieve equality in some upper bounds on β_k . Moreover, we give a constructive characterization of trees that achieve a lower bound on β_k . Finally, we provide a characterization of extremal regular graphs of a Nordhaus-Gaddum bound for the k -independence number.

2 Characterization of some upper bounds for β_k

We begin by recalling some important results that will be useful in our investigations. In [7], Hansberg and Pepper give upper bounds for $\beta_k(G)$ in terms of n , $m(G)$, $\delta(G)$, k and $\Delta(G)$. However, extremal graphs for these bounds are not given. Here, we characterize these extremal graphs.

Theorem 1 (Hansberg, Pepper [7]) *Let G be a graph on n vertices, $m(G) \neq 0$ edges, maximum degree $\Delta(G)$, minimum degree $\delta(G)$ and k a positive integer. Then*

$$(i) \quad \beta_k(G) \leq \frac{m(G)}{\delta(G) - k + 1} \text{ for } k \leq \delta(G).$$

$$(ii) \quad \beta_k(G) \leq \frac{n\Delta(G)}{\Delta(G) + \delta(G) - k + 1} \text{ for } k \leq \delta(G).$$

$$(iii) \quad \beta_k(G) \leq \frac{2\Delta(G)n - 2m(G)}{2\Delta(G) - k + 1} \text{ for } k \leq \Delta(G).$$

A graph is a k -semiregular bipartite graph if its vertex set can be partitioned into two subsets such that every vertex in one of the partite sets

has degree k . For subsets X and Y of $V(G)$ with $X \cap Y = \emptyset$, we denote by $m(X, Y)$ the number of edges between X and Y .

Now we give characterizations of graphs attaining bounds in Theorem 1.

Theorem 2 *Let G be a graph on n vertices, $m(G) \neq 0$ edges and k a positive integer with $k \leq \delta(G)$. Then*

$$\beta_k(G) \leq \frac{m(G)}{\delta(G) - k + 1},$$

with equality if and only if G is a δ -semiregular bipartite graph and $k = 1$.

Proof. Let S be a $\beta_k(G)$ -set. Then $m(S, V - S)$ satisfies

$$|S|(\delta(G) - k + 1) \leq m(S, V - S).$$

We deduce that $|S|(\delta(G) - k + 1) \leq m(G)$. Thus $\beta_k(G) = |S| \leq \frac{m(G)}{\delta(G) - k + 1}$.

If $\beta_k(G) = \frac{m(G)}{\delta(G) - k + 1}$, then $m(G) = m(S, V - S) = |S|(\delta(G) - k + 1)$, and so all edges of G are between S and $V - S$. Thus S and $V - S$ are independent and each vertex x of S has degree $d_G(x) = \delta(G) - k + 1$. Since $d_G(x) \geq \delta(G)$ and $k \geq 1$, $k = 1$ and $d_G(x) = \delta(G)$, that is, G is a δ -semiregular bipartite graph.

Conversely, suppose that G is a δ -semiregular bipartite graph with partite sets A and B such that each vertex in A has degree $\delta(G)$ and $k = 1$. Then A is independent and $m(G) = m(A, B) = |A|\delta(G)$. Thus $\beta_1(G) \geq |A| = \frac{m(G)}{\delta(G)}$. Equality is obtained from the fact that $\beta_1(G) \leq \frac{m(G)}{\delta(G)}$. ■

The following result in regular graphs is an immediate consequence of Theorem 2.

Corollary 1 *Let G be a r -regular graph on n vertices, $m(G) \neq 0$ edges and k a positive integer with $k \leq r(G)$. Then, $\beta_k(G) = \frac{m(G)}{r(G) - k + 1}$ if and only if $k = 1$ and G is a bipartite graph.*

Theorem 3 *If G is a graph on n vertices of maximum degree $\Delta(G)$, minimum degree $\delta(G)$ and k a positive integer with $k \leq \delta(G)$, then*

$$\beta_k(G) \leq \frac{n\Delta(G)}{\Delta(G) + \delta(G) - k + 1},$$

with equality if and only if $V = X \cup Y$, where each vertex of X has degree $\delta(G)$, each vertex of Y has degree $\Delta(G)$, $G[X]$ is $(k - 1)$ -regular and Y is independent.

Proof. Let X be a $\beta_k(G)$ -set. Then $m(X, V - X)$ satisfies

$$|X|(\delta(G) - k + 1) \leq m(X, V - X) \leq \Delta(G)|V - X|.$$

We deduce that $|X|(\delta(G) - k + 1) \leq \Delta(G)(n - |X|)$. Thus $\beta_k(G) = |X| \leq \frac{n\Delta(G)}{\Delta(G) + \delta(G) - k + 1}$.

If $\beta_k(G) = \frac{n\Delta(G)}{\Delta(G) + \delta(G) - k + 1}$, then $|X|(\delta(G) - k + 1) = m(X, V - X) = \Delta(G)|V - X|$, and so each vertex of X has exactly $\delta(G) - k + 1$ neighbors in $V - X$, and each vertex of $V - X$ has $\Delta(G)$ exactly neighbors in X , and so $V - X$ is independent. On the other hand, $|N(x) \cap X| = k - 1$, for all $x \in X$, otherwise $d_G(x) < \delta(G)$, which means that $G[X]$ is $(k - 1)$ -regular, and so each vertex of X has degree $\delta(G)$ in G .

Conversely, it is easy to see that if G is one of the graphs described above, then each vertex of X has exactly $\delta(G) - k + 1$ neighbors in Y and each vertex of Y has exactly $\Delta(G)$ neighbors in X . Thus $|X|(\delta(G) - k + 1) = m(X, Y) = \Delta(G)|Y| = \Delta(G)|V - X|$, and so $|X| = \frac{n\Delta(G)}{\Delta(G) + \delta(G) - k + 1}$. Since X is k -independent, $\beta_k(G) \geq |X| = \frac{n\Delta(G)}{\Delta(G) + \delta(G) - k + 1}$, and equality holds from the fact that $\beta_k(G) \leq \frac{n\Delta(G)}{\Delta(G) + \delta(G) - k + 1}$. ■

As an immediate consequence of Theorem 3, we have the following result for $\beta(G)$.

Corollary 2 *Let G be a graph of maximum degree $\Delta(G)$ and minimum degree $\delta(G) \geq 1$. Then, $\beta(G) = \frac{n\Delta(G)}{\Delta(G) + \delta(G)}$ if and only if G is bipartite with partite sets X (induced by vertices of degree $\delta(G)$) and Y (induced by vertices of degree $\Delta(G)$). In particular, if G is regular, then $\beta(G) = \frac{n}{2}$ if and only if G is bipartite.*

Theorem 4 *Let G be a graph of order n with $m(G)$ edges and maximum degree $\Delta(G)$, and let k a positive integer with $k \leq \Delta(G)$. Then*

$$\beta_k(G) \leq \frac{2\Delta(G)n - 2m(G)}{2\Delta(G) - k + 1},$$

with equality if and only if G is Δ -regular, $V = X \cup Y$, where $G[X]$ is $(k - 1)$ -regular and Y is independent.

Proof. Let X be a $\beta_k(G)$ -set. Then $|N(x) \cap X| \leq k - 1$ for all $x \in X$, and so $m(G[X]) \leq \left(\frac{k-1}{2}\right)|X|$. In particular

$$m(G[V - X]) + m(V - X, X) \leq \Delta(G)|V - X|.$$

Therefore

$$\begin{aligned} m(G) &= m(G[X]) + m(G[V - X]) + m(V - X, X) \\ &\leq \left(\frac{k-1}{2}\right) |X| + \Delta(G) |V - X| \\ &\leq \left(\frac{k-1}{2} - \Delta(G)\right) |X| + \Delta(G) n, \end{aligned}$$

that is $\beta_k(G) \leq \frac{2\Delta(G)n-2m(G)}{2\Delta(G)-k+1}$.

If $\beta_k(G) = \frac{2\Delta(G)n-2m(G)}{2\Delta(G)-k+1}$, then $m(G[X]) = \left(\frac{k-1}{2}\right) |X|$ and $m(G[V - X]) + m(V - X, X) = \Delta(G) |V - X|$, that is $G[X]$ is $(k-1)$ -regular, every vertex of X has exactly $\Delta(G) - k + 1$ neighbors in $V - X$, $V - X$ is independent and every vertex in $V - X$ has exactly $\Delta(G)$ neighbors in X . It is clear that G is regular.

Conversely, it is easy to see that if G is one of the graphs described above, then $\Delta(G)n = 2m(G)$, each vertex of X has exactly $\Delta(G) - k + 1$ neighbors in Y and each vertex of Y has exactly $\Delta(G)$ neighbors in X . Thus $(\Delta(G) - k + 1) |X| = m(Y, X) = \Delta(G) |Y|$, and so $|X| = \frac{\Delta(G)n}{2\Delta(G)-k+1} = \frac{2\Delta(G)n-2m(G)}{2\Delta(G)-k+1}$. Since X is k -independent, $\beta_k(G) \geq |X| = \frac{2\Delta(G)n-2m(G)}{2\Delta(G)-k+1}$, and equality holds from the fact that $\beta_k(G) \leq \frac{2\Delta(G)n-2m(G)}{2\Delta(G)-k+1}$. ■

Corollary 3 *Let G be a graph of order n with $m(G)$ edges and maximum degree $\Delta(G)$. Then $\beta(G) = n - \frac{m(G)}{\Delta(G)}$ if and only if G is a Δ -regular bipartite graph.*

Next, we improve the upper bound in Theorem 1-(ii) when $\Delta(G) + \delta(G) - k + 1 < n$. Moreover, we characterize all graphs attaining this bound.

Theorem 5 *Let G be a graph with minimum degree $\delta(G)$ and k a positive integer with $k \leq \Delta(G)$. Then*

$$\beta_k(G) \leq n - \delta(G) + k - 1$$

where equality holds if and only if $V = X \cup Y$, where $G[X]$ is $(k-1)$ -regular, each vertex x in X is joined to each vertex in Y , $d_G(x) = \delta(G)$ and $k \leq \delta(G)$.

Proof. $k \geq \delta(G) + 1$, then $n - \delta(G) + k - 1 \geq n$, and so $\beta_k(G) < n - \delta(G) + k - 1$, since $\beta_k(G) \leq n - 1$ for $k \leq \Delta(G)$. Assume that $k \leq \delta(G)$. Let

X be a maximum k -independent set in G , and let x be a vertex in X . Since x has at most $k-1$ neighbors in X and $d_G(x) \geq \delta$, x is adjacent to at least $\delta(G) - k + 1$ vertices in $Y = V - X$, which means that $|V - X| \geq \delta(G) - k + 1$, and so $\beta_k(G) \leq n - \delta(G) + k - 1$ as required.

Suppose that $|X| = n - \delta(G) + k - 1$. Then $k \leq \delta(G)$ and $|Y| = \delta(G) - k + 1$, so for every vertex x in X , x has exactly $k-1$ neighbors in X and x is adjacent to every vertex in Y , which means that $d_G(v) = \delta(G)$ and $G[X]$ is $(k-1)$ -regular.

Conversely, if G is the graph described above, then by construction, X is k -independent and $|X| = n - \delta(G) + k - 1$, so $\beta_k(G) \geq |X| = n - \delta(G) + k - 1$ and equality holds from the fact that $\beta_k(G) \leq n - \delta(G) + k - 1$. ■

Corollary 4 *Let G be a r -regular graph and k a positive integer with $k \leq r(G)$. Then $\beta_k(G) = n - r(G) + k - 1$ if and only if $V = X \cup Y$, where $G[X]$ is $(k-1)$ -regular and each vertex in X is joined to each vertex in Y .*

3 Characterization of a lower bound for β_k in trees

For nontrivial trees T , Maddox generalized the well known bound $\beta(T) \geq n/2$ and proved $\beta_k(T) \geq kn/(k+1)$ [8]. In [2], Blidia, Chellali, Favaron and Meddah gave another proof of this result which allowed one to characterize the extremal trees with a recursive construction. Let $\mathcal{F}_0(k)$ the family of trees which can be constructed recursively from $T_1 = K_{1,k}$ by adding a copy of $K_{1,k}$ attached by an edge from any vertex of $K_{1,k}$ to any vertex of T_i .

Theorem 6 (Blidia et al. [2]) *Let T be a tree of order n and maximum degree $\Delta(T)$. Then for every integer k with $2 \leq k \leq \Delta(T)$, $\beta_k(T) \geq \frac{kn}{k+1}$, with equality if and only if $T \in \mathcal{F}_0(k)$.*

If G is bipartite graph, then it is well-known and easy to see that $\beta(G) \geq \lceil \frac{n}{2} \rceil$. In [9], Volkmann presents a constructive characterization of bipartite graphs G for which $\beta(G) = \lceil \frac{n}{2} \rceil$.

Since $\beta_k(T)$ is a positive integer, the lower bound $\lceil kn/(k+1) \rceil$ for every tree can easily be obtained from Theorem 6.

Corollary 5 *Let T be a tree of order n and maximum degree $\Delta(T)$. Then for every integer k with $2 \leq k \leq \Delta(T)$, $\beta_k(T) \geq \lceil kn/(k+1) \rceil$.*

Let $n = t(k+1) + r$ for integers $t \geq 0$ and $0 \leq r \leq k$. Then $\lceil kn/(k+1) \rceil = (kn + r)/(k+1)$.

Let $\mathcal{H}(k)$ be a family of trees of order r_0 with $1 \leq r_0 \leq k$ and $\mathcal{H}_u^s(k)$ be a family of trees of order $k+1+s$ with $0 \leq s \leq k$ with a vertex u such that $\forall x \in N(u); d_G(x) \leq k, \forall x \notin N(u); d_G(x) \leq k-1$ and at least u or one of its neighbors has degree k or more.

For the purpose of characterizing the trees attaining the bound of Corollary 5, we introduce a family $\mathcal{F}(k)$ of trees T that can be obtained from a sequence $T_1, T_2, \dots, T_m$ of trees such that $T_1 \in \mathcal{H}(k) \cup \mathcal{H}_u^s(k)$, and if $m \geq 2$, then T_{i+1} can be obtained recursively from T_i by the operation $\mathcal{O}$ for $1 \leq i \leq m-1$.

- **Operation $\mathcal{O}$** : Add a copy of $H \in \mathcal{H}_u^s(k)$ attached by an edge from the center vertex u to any vertex of T_i with the condition: If $n(T_i) \equiv r_i \pmod{k+1}$, then $1 \leq r_i + s \leq k$.

Lemma 7 *Let T be a tree of order n and maximum degree $\Delta(T)$. Then for every integer k with $2 \leq k \leq \Delta(T)$, if $T \in \mathcal{F}(k)$, then $\beta_k(T) = \lceil kn / (k+1) \rceil$.*

Proof. Consider any $T \in \mathcal{F}(k)$. So T can be obtained from a sequence $T_1, T_2, \dots, T_m$ ($m \geq 1$) of trees of T , where $T_1 \in \mathcal{H}(k) \cup \mathcal{H}_u^s(k)$, $T = T_m$, and, if $1 \leq i \leq m-1$, the tree T_{i+1} is obtained from T_i by the operation $\mathcal{O}$. We proceed by induction on m .

If $m = 1$, then $T = T_1 \in \mathcal{H}(k) \cup \mathcal{H}_u^s(k)$. If $T \in \mathcal{H}(k)$, then T is a tree of order $n \leq k$, and so

$$\beta_k(T) = n = n(k+1)/(k+1) = (nk+n)/(k+1) = \lceil kn/(k+1) \rceil.$$

If $T \in \mathcal{H}_u^s(k)$, then T is a tree of order $n = k+1+s$ with $0 \leq s \leq k$, and so

$$\begin{aligned} \beta_k(T) &= n-1 = (n-1)(k+1)/(k+1) \\ &= (nk+n-k-1)/(k+1) \\ &= (nk+s)/(k+1) \\ &= \lceil kn/(k+1) \rceil. \end{aligned}$$

Assume now that $m \geq 2$ holds for T and that the result holds for all trees in T that can be constructed by a sequence of length at most $m-1$. Let $T' = T_{m-1}$. Assume that T is obtained from T' by using the operation $\mathcal{O}$, with $n(T') \equiv r' \pmod{k+1}$ and $1 \leq r' + s \leq k$. We prove that $\beta_k(T) = \beta_k(T') + k + s$. It is clear that there exists a $\beta_k(T)$ -set S containing $V(H) - \{u\}$. Thus $S - (V(H) - \{u\})$ is a k -independent set of T' . Hence $\beta_k(T') \geq$

$|S - (V(H) - \{u\})| = \beta_k(T) - k - s$. Since $V(T') \cap V(H) = \emptyset$, every $\beta_k(T')$ -set S' can be extended to a k -independent set of T by adding $V(H) - \{u\}$ and so $\beta_k(T) \geq |S' \cup (V(H) - \{u\})| = \beta_k(T') + k + s$, implying the equality. Since $\beta_k(T') = \lceil kn(T') / (k+1) \rceil$,

$$\begin{aligned} \beta_k(T) &= \lceil kn(T') / (k+1) \rceil + k + s \\ &= (kn(T') + r') / (k+1) + k + s \\ &= (k(n(T) - k - s - 1) + r') / (k+1) + k + s \\ &= (kn(T) + s + r') / (k+1). \end{aligned}$$

Since $0 \leq s + r' \leq k$, $\beta_k(T) = \lceil kn(T) / (k+1) \rceil$. ■

We now are ready to give a constructive characterization of trees T which satisfy $\beta_k(T) = \lceil kn / (k+1) \rceil$.

Theorem 8 *Let T be a tree of order n and maximum degree $\Delta(T)$. Then for every integer k with $2 \leq k \leq \Delta(T)$, $\beta_k(T) = \lceil kn / (k+1) \rceil$ if and only if $T \in \mathcal{F}(k)$.*

Proof. First suppose that $T \in \mathcal{F}(k)$. Then Lemma 7 implies that $\beta_k(T) = \lceil kn / (k+1) \rceil$.

Conversely, assume that T satisfies $\beta_k(T) = \lceil kn / (k+1) \rceil$. Let $Y(T)$ be the set of vertices of degree at least k of a tree T . We proceed by induction on $|Y(T)|$. If $Y(T) = \emptyset$, then $\beta_k(T) = n$, and so $n = \lceil (kn) / (k+1) \rceil = (kn + r) / (k+1)$ where $n \equiv r \pmod{k+1}$ with $0 \leq r \leq k$, which implies that $n = r$. Hence, T is any tree of order $n \leq k$. So, $T \in \mathcal{H}(k) \subset \mathcal{F}(k)$. If $|Y(T)| = 1$ then $\beta_k(T) = n-1$, and so $n-1 = \lceil kn / (k+1) \rceil = (kn+r) / (k+1)$ where $n \equiv r \pmod{k+1}$ with $0 \leq r \leq k$, which implies that $n = k+1+r$. Hence $T \in \mathcal{H}_u^s(k) \subset \mathcal{F}(k)$.

Suppose the property of the theorem is true for all trees with $|Y(T)| < p$ with $p \geq 2$ and let T be a tree of order n such that $|Y(T)| = p$. Root T at a leaf t . Let u be a vertex of degree at least k at maximum distance from t and under this condition, of maximum degree. Since $|Y(T)| \geq 2$, u is at distance at least 2 from t . Let v and w be the father and grandfather of u in the rooted tree.

If $d_T(u) > k$, let T' and T'' be the components of $T - uv$ respectively containing v and u . Then $d_{T''}(u) \geq k$ and from the first condition in the choice of u , all the vertices of $V(T'') - \{u\}$ have degree less than k .

If $d_T(u) = k$, let T' and T'' be the components of $T - vw$ respectively containing w and v . Then from the two conditions in the choice of u , $d_{T''}(u) = k$,

the vertices of $N_{T''}(v) - \{u\}$ have degree at most k and all the vertices of $V(T'') - N_{T''}[v]$ have degree less than k .

Every $\beta_k(T')$ -set S' can be extended to a k -independent set of T by adding $V(T'') - \{x\}$, where $x = u$ in the first case and $x = v$ in the second case, and so $\beta_k(T) \geq |S' \cup (V(T'') - \{x\})| = \beta_k(T') + |V(T'')| - 1$, and apply the inductive hypothesis to T' since $|Y(T')| < |Y(T)|$. Moreover $n(T') \equiv r' \pmod{k+1}$ where $0 \leq r' \leq k$, and $|V(T'')| = k + 1 + s$ where $s \geq 0$. Thus

$$\begin{aligned} \beta_k(T) &\geq \beta_k(T') + |V(T'')| - 1 \\ &\geq \lceil kn(T') / (k+1) \rceil + |V(T'')| - 1 \\ &= k(n(T') - |V(T'')| + r') / (k+1) + |V(T'')| - 1 \\ &= (kn(T') + |V(T'')| + r' - k - 1) / (k+1) \\ &= (kn(T') + r' + s) / (k+1) \\ &\geq \lceil kn / (k+1) \rceil, \end{aligned}$$

Since $\beta_k(T) = \lceil kn / (k+1) \rceil$, we have equality throughout this inequality chain, that is $\beta_k(T') = \lceil k(n - |V(T'')|) / (k+1) \rceil$ and $0 \leq r' + s \leq k$. So by the inductive hypothesis, $T' \in \mathcal{F}(k)$, and $T'' \in \mathcal{H}_x^s(k)$ attached by $x = u$ in the first case or $x = v$ in the second case. Since T is obtained from T' by using Operation $\mathcal{O}$, $T \in \mathcal{F}(k)$. Therefore the property is true for T . ■

4 Nordhaus-Gaddum inequality

As an extension of the known Nordhaus-Gaddum type inequality $\beta(G) + \beta(\overline{G}) \leq n + 1$, established by Chartrand and Schuster [3], Blidia, Bouchou and Volkmann showed in [1] that if $k \leq \min(\Delta(G) + 1, \Delta(\overline{G}) + 1)$, then $\beta_k(G) + \beta_k(\overline{G}) \leq n + 2k - 1$. We now give a simple proof of this result for regular graphs, and characterize the regular graphs for which equality holds.

Theorem 9 *Let G be a r -regular graph of order n and k a positive integer with $k \leq \min(r(G) + 1, r(\overline{G}) + 1)$. Then $\beta_k(G) + \beta_k(\overline{G}) \leq n + 2k - 1$ with equality, if and only if G is $(k-1)$ -regular of order $2k-1$.*

Proof. Let S be a $\beta_k(G)$ -set of G and S' a $\beta_k(\overline{G})$ -set of $\overline{G}$. Since $r(G) + r(\overline{G}) = n - 1$, Theorem 5 yields

$$\beta_k(G) + \beta_k(\overline{G}) \leq n - r(G) + k - 1 + n - r(\overline{G}) + k - 1 = n + 2k - 1.$$

If $\beta_k(G) + \beta_k(\overline{G}) = n + 2k - 1$, then equality holds throughout the calculation, and so $\beta_k(G) = n - r(G) + k - 1$ and $\beta_k(\overline{G}) = n - r(\overline{G}) + k - 1$. By Corollary

4, we have $G[S]$ is $(k-1)$ -regular and each vertex in S is joined to each vertex in $V-S$ in G , and $G[S']$ is $(k-1)$ -regular and each vertex in S' is joined to each vertex in $V-S'$ in $\overline{G}$. Hence, G and $\overline{G}$ are connected graphs. On the other hand, if $V-S \neq \emptyset$ or $V-S' \neq \emptyset$, then G or $\overline{G}$ is not connected. Therefore, the only possibility is $V-S = \emptyset$ and $V-S' = \emptyset$. So, $V = S = S'$ and $r(G) = r(\overline{G}) = k-1$, which means that G is $(k-1)$ -regular of order $n = 2k-1$. The converse is evident. ■

References

- [1] M. Blidia, A. Bouchou and L. Volkmann, *Bounds on the k -independence and k -chromatic numbers of graphs*, Ars Comb. 113 (2014) 33-46.
- [2] M. Blidia, M. Chellali, O. Favaron and N. Meddah, *On k -independence in graphs with emphasis on trees*, Discrete Mathematics 307 (2007) 2209-2216.
- [3] G. Chartrand and S. Schuster, *On the independence numbers of complementary graphs*, Trans. New York Acad. Sci., Series II, 36 : 247-251, 1974.
- [4] M. Chellali, O. Favaron, A. Hansberg and L. Volkmann, *k -Domination and k -Independence in Graphs: A Survey*, Graphs and Combinatorics 28 (2012) 1-55.
- [5] J. F. Fink and M. S. Jacobson, *n -domination in graphs*, in : Graph Theory with Applications to Algorithms and Computer. John Wiley and Sons, New York (1985) 283-300.
- [6] J. F. Fink and M. S. Jacobson, *n -domination, n -dependence and forbidden subgraphs*, in: Graph Theory with Applications to Algorithms and Computer. John Wiley and Sons, New York (1985) 301-311.
- [7] A. Hansberg and A. Pepper, *On k -domination and j -dependence in graphs*, Discrete Applied Mathematics 161(10-11) (2013) 1472-1480.
- [8] R. B. Maddox, *On k -dependent subsets and k -degenerate graphs*, Congressus Numerantium 66 (1988) 11-14.
- [9] L. Volkmann, *A characterization of bipartite graphs with independence number half their order*, Australasian Journal of Combinatorics, 41 (2008) 219-222.

The b -domatic number of join and corona graphs

Mohammed Benatallah¹, Nouredine Ikhlef Eschouf²,
and Miloud Mihoubi³

¹UZA, Faculty of Sciences, RECITS Laboratory, Djelfa, Algeria.

e-mail: m_benatallah@yahoo.fr

²UYF, Faculty of Sciences, Medea, Algeria.

e-mail: nour_eshouf@yahoo.fr

³USTHB, Faculty of Mathematics, RECITS Laboratory, Algiers, Algeria.

e-mail: mmihoubi@usthb.dz

Abstract

A partition of $V(G)$, all of whose classes are dominating sets in G , is called a domatic partition of G . The maximum number of classes of a domatic partition of G is called the domatic number of G and is denoted by $d(G)$. A domatic partition $\mathcal{P}$ of G is called *b-domatic* if no larger domatic partition of G can be obtained from $\mathcal{P}$ by transferring some vertices of some classes of $\mathcal{P}$ to form a new class. The minimum cardinality of a b -domatic partition of G is called the *b-domatic number* and is denoted by $bd(G)$.

In this paper we explore the bounds for the b -domatic numbers of the join of two graphs and corona graphs. The bounds are the best possible in the sense that there exist examples for which equalities are attained.

Keywords: domatic number, b -domatic number, join graphs, corona graphs.

1 Introduction

Let $G = (V, E)$ be a finite, simple and undirected graph with vertex-set V and edge-set E . We call $|V|$ the order of G and denote it by n . For any nonempty subset $A \subset V$. For any vertex v of G , the *neighborhood* of v is the set $N_G(v) = \{u \in V(G) \mid (u, v) \in E\}$ and the *closed neighborhood* of v is the set $N_G[v] = N_G(v) \cup \{v\}$. Let $\Delta(G)$ (respectively, $\delta(G)$) be the maximum (respectively, minimum) degree in G . Let $G[A]$ denote the subgraph of G induced by A . Through this paper, the notations P_n, C_n , and K_n always denote a path, a cycle, and a complete graph of order n , respectively, while $K_{p,q}$ ($p \geq q$) denotes a complete bipartite graph with partite sets of sizes p, q and S_n (the star with n leaves) and P (the Petersen graph). For further terminology on graphs we refer to the book by Berge [1].

A set $D \subseteq V$ is called a dominating set if every vertex in $V \setminus D$ is adjacent to some vertex in D . The minimum cardinality of a dominating set is called the domination number and is denoted by $\gamma(G)$. By analogy to the concept of chromatic partition, Cockayne and Hedetniemi [2] introduced the concept of domatic partition of a graph. A partition $\mathcal{P}$ in which each of its classes is a dominating sets is called a domatic partition of G . The domatic number $d(G)$ is defined as the largest number of sets in a domatic partition. The authors of [2] showed that

$$d(G) \leq \min\left\{\frac{n}{\gamma(G)}, \delta(G) + 1\right\} \quad (1)$$

In [3], O. Favaron introduced the b-domatic number by considering a new type of domatic partition. As defined in [3], a domatic partition $\mathcal{P}$ of G is *b-domatic* if no larger domatic partition π can be obtained by transferring some vertices of some classes of $\mathcal{P}$ to form a new class. Formally, a partition $\mathcal{P} = \{U_1, U_1, \dots, U_k\}$ is a b-domatic partition of G if there do not exist k non-empty subsets $\pi_i \subseteq U_i$, $i \in \{1, \dots, k\}$ with strict inclusion for at least one non-empty subset $(U_i \setminus \pi_i)$ for which $\pi = \{\pi_1, \pi_2, \dots, \pi_k, V \setminus \bigcup_{i=1}^k \pi_i\}$ is domatic partition of G . The minimum cardinality of a b-domatic partition of G is called the b-domatic number and is denoted by $bd(G)$.

It is observed in [3] that if $\delta(G) = 0$, then $\{V(G)\}$ is the unique domatic partition and so $bd(G) = d(G) = 1$. For this, all graphs considered in this paper are without isolated vertices. Many other properties of domatic and b-domatic partitions were given in [3]. In particular, it was observed that for

any graph G with minimum degree $\delta(G) \geq 1$,

$$2 \leq bd(G) \leq d(G) \quad (2)$$

The same author has computed the b-domestic number for some particular classes of graphs, for example $bd(K_n) = n$, $bd(C_3) = 3$ and $bd(C_n) = 2$ for $n \geq 4$, and $bd(K_{p,q}) = 2$, $bd(S_n) = 2$ and $bd(P) = 2$. Other aspects of b-domatically continuous graphs were discussed in [3].

The b-domestic spectrum of a graph G is the set of the cardinalities of all the b-domestic partitions of G .

The joint of two graph G_1 and G_2 is the graph $G_1 \vee G_2$ defined as the disjoint union of graph G_1 and G_2 with additional edges linking each vertex of G_1 with each vertex of G_2 .

The corona $G \circ H$ of two graphs G and H is the graph obtained by taking one copy of G of order n and n copies of H , and then joining the i^{th} vertex of G to every vertex in the i^{th} copy of H . For every $v \in V(G)$, denote by H^v the copy of H whose vertices are attached one by one to the vertex v . Subsequently, denote by $v \vee H^v$ the subgraph of the corona $G \circ H$ corresponding to the join $v \vee H^v$, $v \in V(G)$.

Our aim in this paper is to determine the exact values of the b-domestic number of joint graph.

2 Mains results

We make use of the following results in this paper.

Theorem 1 [7] *Let $\mathcal{P}$ be a domestic partition of a graph $G = (V, E)$. If exists a vertex $v \in V$ such that each vertex of $N_G[v]$ is either isolated in its class or has a private neighbor with respect to its class, Then $\mathcal{P}$ is b-domestic.*

Proposition 2 [7] *If v is universal vertex in G , then*

$$bd(G) = bd(G \setminus v) + 1$$

Theorem 3 [3] *Let $G_1, \dots, G_k$ be the components of a disconnected graph G without isolated vertices. Then $bd(G) = \min\{bd(G_i) : 1 \leq i \leq k\}$.*

2.1 b-Domatic number of the join $G_1 + G_2$

Theorem 4 *For any two graph G_1, G_2 we have*

$$bd(G_1) + bd(G_2) \leq bd(G_1 + G_2) \leq \min \{ \delta(G_1) + |V(G_2)|, \delta(G_2) + |V(G_1)| \} + 1$$

Proof. We put $V(G_1) = \{u_1, u_2, \dots, u_n\}$, $V(G_2) = \{v_1, v_2, \dots, v_m\}$ and assume without loss of generality that the minimum degrees $\delta(G_1), \delta(G_2)$ are realized by vertices u_k, v_l respectively. By the definition of the join $G_1 + G_2$, $\delta(G_1 + G_2) = \min \{ \delta(G_1) + |V(G_2)|, \delta(G_2) + |V(G_1)| \}$. By 1 and 2, we have $bd(G_1 + G_2) \leq \min \{ \delta(G_1) + |V(G_2)|, \delta(G_2) + |V(G_1)| \} + 1$. Consequently, the upper bound follows.

Let $\mathcal{P}_1 = \{C_1, \dots, C_{bd(G_1)}\}$ be a b-domatic partition of G_1 and $\mathcal{P}_2 = \{B_1, \dots, B_{bd(G_2)}\}$ is the b-domatic partition of G_2 . Each set C_i for $i = 1, \dots, bd(G_1)$ dominates the vertex set $V(G_2)$ of G_2 and simultaneously each set B_j , for $j = 1, \dots, bd(G_2)$ dominates the vertex set $V(G_1)$ of G_1 . Moreover, the sets C_i, B_j , for $i = 1, \dots, bd(G_1)$ and $j = 1, \dots, bd(G_2)$ are pairwise disjoint and $(\bigcup_{i=1}^{bd(G_1)} C_i) \cup (\bigcup_{j=1}^{bd(G_2)} B_j) = V(G_1 \cup G_2)$. Let $\mathcal{P} = \{U_1, \dots, U_{bd(G_1) + bd(G_2)}\}$

is a b-domatic partition of join graph $G_1 + G_2$. such that $U_i = C_i, i = 1, \dots, bd(G_1)$ and $U_{bd(G_1) + j} = B_j, j = 1, \dots, bd(G_2)$. On the other hand, if exists two vertices u and v , such that $\{u\} \in \mathcal{P}_1$ and $\{v\} \in \mathcal{P}_2$ is not isolated and has not a private neighbor with respect to its class, so the subset $\{u, v\} \subseteq V(G_1 + G_2)$ is the dominating set in the join $G_1 + G_2$. Then certainly there exists the domatic partition $\mathcal{P}' = \{U_1, \dots, U_{bd(G_1) + bd(G_2)}, \{u, v\}\}$ of the graph $G_1 \vee G_2$. Hence $\mathcal{P}$ is no b-domatic partition. So $\mathcal{P}$ is a b-domatic partition of join graph $G_1 \vee G_2$ if and only if either every vertex is isolated and has a private neighbor with respect to its class in $\mathcal{P}_1$ or $\mathcal{P}_2$. Suppose, to the contrary, that $\mathcal{P}$ is not b-domatic. Then by definition, there exist $bd(G_1) + bd(G_2)$ non-empty subsets $\pi_i \subseteq U_i, i \in \{1, \dots, bd(G_1) + bd(G_2)\}$ with strict inclusion for at least one

non-empty subset $((U_i \setminus \pi_i) \neq \emptyset)$ and $\mathcal{A} = V \setminus \bigcup_{i=1}^{bd(G_1) + bd(G_2)} \pi_i, (\mathcal{A} \neq \emptyset)$, and

for which $\pi = \{\pi_1, \pi_2, \dots, \pi_{bd(G_1) + bd(G_2)}\}$ is domatic partition of $(G_1 + G_2) - \mathcal{A}$

and $\pi = \{\pi_1, \pi_2, \dots, \pi_{bd(G_1) + bd(G_2)}, V \setminus \bigcup_{i=1}^{bd(G_1) + bd(G_2)} \pi_i\}$ is domatic partition of

G and for which $\mathcal{A}$ is not dominated by a some class of π but we have every vertices satisfies the property given in the statement of the theorem i.e every vertices in $\mathcal{P}_1$ is a isolated vertex or private neighbor in its classes, then $\mathcal{A}$ is

not contains a vertex of $\mathcal{P}_1$. So $\left(V \setminus \bigcup_{i=1}^{bd(G_1)} \pi_i\right) = \phi$, meaning the classe $\mathcal{A}$ contains only the vertices of $\mathcal{P}_2$, a contradiction thus $\mathcal{P}_2$ is b-domatic. ■

Corollary 5 *Let $\mathcal{P}_1$ and $\mathcal{P}_2$ is a two b-domatic partition. If every vertices is isolated or has private neighbor in its classes in $\mathcal{P}_1$ or $\mathcal{P}_2$, then*

$$bd(G_1 + G_2) = bd(G_1) + bd(G_2)$$

Proposition 6 *Let $P_n, n \geq 2$ is a path and G is a graph, Then*

$$bd(P_n + G) = 2 + bd(G)$$

Proof. Let $P_n (n \geq 2)$ is a path and $bd(P_n) = 2$ ie the path P_n admits a b-domatic partition contains two classes such that every vertex of two classes is isolated vertex. By Corollary 5, $bd(P_n + G) = bd(P_n) + bd(G) = 2 + bd(G)$. ■

Proposition 7 *Let T is a tree and G is a graph, then*

$$bd(T + G) = 2 + bd(G)$$

Corollary 8 *Let $\mathcal{P}_1$ and $\mathcal{P}_2$ is two b-domatic partition. If exists a vertex is not isolated or has not private neighbor in its classes in $\mathcal{P}_1$ and $\mathcal{P}_2$, then*

$$bd(G_1 + G_2) > bd(G_1) + bd(G_2)$$

Proposition 9 *Let C_n, C_m two cycle such that $n, m \geq 4$. Then*

$$bd(C_n + C_m) = \begin{cases} 5 & \text{if } C_n, C_m \text{ are odd} \\ 4 & \text{if no} \end{cases}$$

Proof. If C_n, C_m two cycle such that $n, m \geq 4$ one of two cycle is even, then $bd(C_n) = 2$ and $bd(C_m) = 2$ ie one of the cycle admits a b-domatic partition contains two classes such that every vertex of two classes is isolated vertex. By theorem 5 $bd(C_n + C_m) = bd(C_n) + bd(C_m) = 4$. Otherwise if two cycle C_n and C_m such that $n, m \geq 4$ are odd, then $bd(C_n) = 2$ and $bd(C_m) = 2$ i.e two cycle admits a b-domatic partition contains two classes such that every class of cycle contains a vertex is not isolated and not private neighbor in its classe. By corollary 8, $bd(C_n + C_m) = 5$. ■

Proposition 10 *Let P is a Peterson graph and G is a graph, then*

$$bd(P + P) = 7$$

and

$$2 + bd(G) \leq bd(P + G) \leq 5 + bd(G)$$

Corollary 11 *For $m, n \geq 2$, $bd(S_n + S_m) = 4$.*

The b-domatic spectrum of $(S_n + S_m)$ is $\{4, 2 + \min\{n, m\}\}$.

Proof. Let $m, n \geq 2$ and put $V(S_n) = \{u_0\} \cup A, V(S_m) = \{v_0\} \cup B$, let us suppose that $A = \{u_1, u_2, \dots, u_n\}$ $B = \{v_1, v_2, \dots, v_m\}$ and A, B is independent sets, with u_0, v_0 of a degree n and m . We have $v_0 \in S_n + S_m$ i.e. the vertex v_0 is adjacent to all other vertices in $S_n + S_m \setminus \{u_0\}$ and v_0 is adjacent to all other vertices in $S_n + S_m \setminus \{v_0\}$ together with $\{u_0\}, \{v_0\}$ forms a two b-domatic partition of $S_n + S_m$, we gives: $bd(S_n + S_m) = bd(S_n + S_m \setminus \{u_0, v_0\}) + 2$. We have $S_n + S_m \setminus \{u_0, v_0\} = A + B = K_{nm}$ is a complete bipartite graph, then $bd(S_n + S_m \setminus \{u_0, v_0\}) = 2$. Hence By theorem 5, $bd(S_n + S_m) = 4$, then the b-domatic spectrum of $S_n + S_m \setminus \{u_0, v_0\}$ is $\{2, \min\{n, m\}\}$. Hence the b-domatic spectrum of $S_n + S_m$ is $\{4, 2 + \min\{n, m\}\}$. ■

Corollary 12 *For $m, n, p, q \geq 2$, $bd(K_{pq} + K_{mn}) = 4$.*

The b-domatic spectrum of $K_{pq} + K_{mn}$ is $\{4, \min\{p, q\} + \min\{n, m\}\}$.

Proof. Let $K_{pq} = A_p + B_q$ and A_p, B_q independent. So By theorem 1 $\mathcal{P} = \{A_p, B_q\}$ is b-domatic partition. Therefore $bd(K_{pq}) = 2$, respectively $K_{mn} = C_n + D_m$ and C_n, D_m independent. So By theorem 1, $\mathcal{P}' = \{C_n, D_m\}$ is b-domatic partition since C_n, D_m are minimal dominating sets. Therefore $bd(K_{mn}) = 2$. The graph $K_{pq} + K_{mn} = A_p + B_q + C_n + D_m$, by consequence, $\{A_p, B_q, C_n, D_m\}$ is b-domatic partition and sets A_p, B_q, C_n, D_m independent. By theorem 2, $bd(K_{pq} + K_{mn}) = bd(K_{pq}) + bd(K_{mn}) = 4$. The b-domatic spectrum of K_{pq} is $\{2, \min\{p, q\}\}$ and the b-domatic spectrum of K_{mn} is $\{2, \min\{n, m\}\}$. Hence the b-domatic spectrum of $K_{pq} + K_{mn}$ is $\{4, \min\{p, q\} + \min\{n, m\}\}$. ■

Proposition 13

1. For $m \geq 2$, $bd(K_m + G) = m + bd(G)$.
2. For $n \geq 1$, $bd(S_n + G) = 2 + bd(G)$.

2.2 b-Domatic number of the join $G \circ H$

Theorem 14 *Let G and H be two connected graphs and G of order $n \geq 2$, then*

$$bd(G \circ H) = bd(H) + 1$$

Proof. Let $V(G) = \{v_1, \dots, v_n\}$ and let $H_1, H_2, \dots, H_n$ be n copies of H in $G \circ H$ such that v_i is adjacent to all the vertices in H_i . Let $\{C_1, C_2, \dots, C_{bd(H)}\}$ be a b -domatic partition of H . For every $v \in V(G)$, let $v_i \vee H^{v_i}, 1 \leq i \leq n$ be the components of a disconnected graph of the corona $G \circ H$. By Proposition 2, $bd(v_i + H^{v_i}) = bd(H^{v_i}) + 1 = bd(H) + 1$ and if $\bigcap_{i=1}^n (v_i + H^{v_i}) \neq \emptyset$, it

follows that $\bigcup_{i=1}^n (v_i + H^{v_i})$ is the spanning subgraph of the corona $G \circ H$, then

$bd(\bigcup_{i=1}^n (v_i + H^{v_i})) \leq bd(G \circ H)$. Otherwise by Theorem 3 $bd(\bigcup_{i=1}^n (v_i + H^{v_i})) = \min \{bd(v_i + H^{v_i}) : 1 \leq i \leq n\} = bd(H) + 1$. Hence $bd(G \circ H) \geq bd(H) + 1$.

Let $\{C_1^i, C_2^i, \dots, C_{bd(H)}^i\}$ be the corresponding domatic partition of H_i and let $U_j = \bigcup_{i=1}^n C_j^i$. Then $\mathcal{P} = \{U_1, \dots, U_{bd(H)}, U_{bd(H)+1}\}$, such that $U_{bd(H)+1} =$

$V(G)$. Suppose, to the contrary, that $\mathcal{P}$ is not b -domatic. Then by definition, there exist $bd(H) + 1$ non-empty subsets $\pi_j \subseteq U_j, j = 1, \dots, bd(H) + 1$ with strict inclusion for at least one non-empty subset $((U_j \setminus \pi_j) \neq \emptyset)$ and

$\mathcal{A} = V(G \circ H) \setminus \bigcup_{j=1}^{bd(H)+1} \pi_j, (\mathcal{A} \neq \emptyset)$, and for which $\pi = \{\pi_1, \pi_2, \dots, \pi_{bd(H)+1}\}$

is domatic partition of $(G \circ H) - \mathcal{A}$ and $\pi = \{\pi_1, \pi_2, \dots, \pi_{bd(H)+1}, V(G \circ H) \setminus \bigcup_{j=1}^{bd(H)+1} \pi_j\}$ is domatic partition of G and for which if $\mathcal{A}$ is contains at least

a vertex of $U_{bd(H)+1}$, then π is not domatic partition, i.e. $(U_{bd(H)+1} \setminus \pi_{bd(H)+1}) =$

$\emptyset$. So $\mathcal{A} = V(G \circ H) \setminus \bigcup_{j=1}^{bd(H)} \pi_j \neq \emptyset$ and we have $\bigcap_{j=1}^{bd(H)} (U_j \setminus \pi_j) = \emptyset$, it follows

that $\mathcal{A}$ is the disjoint sets. Contradicting with H is b -domatic. ■

References

- [1] C. Berge. *Graphs*. North Holland, 1985.

- [2] E.J. Cockayne and S.T. Hedetniemi, Towards a theory of domination in graphs, *Networks* 7 (1977) 247–261.
- [3] O. Favaron. The b -domatic number of a graph, *Discussiones Mathematicae, Graph Theory* 33(2013)747–757.
- [4] F. Harary, S. Hedetniemi, The achromatic number of a graph, *J. Combin. Theory* 8 (1970) 154–161.
- [5] R.W. Irving, D.F. Manlove, The b -chromatic number of graphs, *Discrete Appl. Math.* 91 (1999) 127–141.
- [6] D.F. Manlove, Minimaximal and maximinimal optimisation problems: A partial order-based approach, Ph.D. Thesis, Tech. Rep. 27, Comp. Sci. Dept., Univ. Glasgow, Scotland, 1998
- [7] M. Benatallah, N. Ikhlef Eschouf, M. Mihoubi, B -domatic partition of graphs, *Journées scientifiques du laboratoire RECITS (JSL’2016) à l’USTHB. Alger, les 20 et 24 Avril 2016, bab-ezzouar, Algérie*
- [8] Monika Kijewska, Domatic number of graph products. *Journal of Mathematics and Applications*, No 30, pp 71-81 (2008)

Vérification Formelle des Contraintes Temporelles et Spatio-temporelles des Documents SMIL

Fatma Zohra Mekahlia ¹, Abdelghani Ghomari ¹, Sami Yazid ²

¹ Laboratoire RIIR, Département d'Informatique, Faculté des Sciences Exactes et Appliquées, Université d'Oran1 Ahmed Ben Bella, Oran, Algérie.

mekahlia.fzohra@yahoo.fr mekahlia.fatima@edu.univ-oran1.dz
ghomari65@yahoo.fr

² Département d'Informatique, Faculté des Sciences Exactes et Appliquées, Université d'Oran1 Ahmed Ben Bella, Oran, Algérie.

yazid_samy@yahoo.fr

Résumé. La vérification des incohérences temporelles et spatiales des documents SMIL (Synchronized Multimedia Integration Language) est considérée dans ce papier. Dans un travail précédent, nous avons proposé un nouveau prototype basé sur une méthode formelle hybride pour le contrôle des documents SMIL. Cette méthode est inspirée par un travail connexe qui utilise la logique de Hoare pour la détection des conflits temporels dans les documents SMIL, mais avec de nombreuses améliorations et extensions. Une nouvelle règle d'inférence a été ajoutée pour vérifier la durée de répétition des éléments ainsi que les bornes supérieures et inférieures. De plus, nous avons vérifié la cohérence spatio-temporelle en proposant un algorithme de vérification des incohérences spatio-temporelles. Par la suite, nous avons utilisé la programmation par contraintes disjonctives de Marriott pour résoudre le problème des conflits spatio-temporels. Cet article étend ces précédents travaux à la vérification temporelle d'un nouvel élément SMIL.

Mots-clés. Document SMIL, méthode formelle hybride, conflit spatial, contraintes disjonctives, conflit temporel, logique de Hoare,

1 Introduction

1.1 Contexte, motivation et problématique

La conception et la création des Documents Multimédia Interactifs (DMI) sont devenues plus que nécessaire dans les réseaux de communication. En effet, les avancées technologiques en puissance de l'informatique, des appareils et des réseaux informatiques ont conduit à la création d'un grand nombre d'applications multimédias qui les utilisent, telles que : la formation assistée par ordinateur, l'apprentissage à distance, la vidéoconférence interactive, la consultation à distance, le tourisme virtuel, la vidéo à la demande, la communication sur IP, etc. SMIL (Synchronized Multimedia Integration Language) [1], est un langage qui permet la création des présentations multimédias distribuée sur le Web, en utilisant des médias comme: image, audio,

texte, vidéo, animations, et de spécifier donc les dimensions temporelle, spatiale et hypermédia de ce document. Un éditeur de SMIL est capable de vérifier seulement les erreurs syntaxiques mais pas les erreurs sémantiques. Cependant, le problème de la vérification sémantique se pose toujours. Il consiste en l'analyse et le contrôle de cohérence temporelle et spatiale des documents SMIL. Par conséquent, un outil de formalisation, d'analyse et de contrôle de cohérence temporelle et spatiale des documents SMIL est vraiment nécessaire pour assurer une bonne qualité de présentation des applications multimédias.

L'incohérence des documents SMIL est le résultat d'une mauvaise spécification des relations temporelles ou spatiales entre les médias et qui implique une mauvaise synchronisation de ces médias. La synchronisation est une caractéristique importante des documents qui matérialise la sémantique d'une présentation multimédia conçue par l'auteur. Malheureusement, l'inconvénient majeur du langage SMIL est que la vérification de ses incohérences sémantiques est difficile à réaliser, contrairement à la vérification des incohérences syntaxiques qui est prise en charge dans un éditeur SMIL. Dans ce papier, nous allons étudier les deux types de conflits dans une présentation SMIL, à savoir :

1) **Conflit spatial:** la spécification de l'organisation spatiale définit le placement des objets médias dans l'espace de présentation. Si l'on exclut l'audio, les données des médias sont immergées dans l'espace (affichées sur l'écran) et nécessitent la définition de l'organisation spatiale du document. Le conflit spatial est un problème majeur qui affecte la qualité des présentations multimédias. Nous parlons de conflit spatial lorsque deux ou plusieurs médias visuels (vidéo, image, texte) du même document détiennent en même temps la même position de présentation dans le document.

2) **Conflit temporel:** la synchronisation temporelle des éléments SMIL est une caractéristique importante qui englobe la sémantique d'une présentation multimédia conçue par un auteur. Dans le cas des objets temporels, le facteur temps apparaît comme une dimension essentielle de l'information. La synchronisation peut être définie entre les composants d'un même média (intra-média) ou entre différents médias (inter-média). On trouve dans la littérature deux cas de conflits temporels:

i. **Conflit intra-média :** représente le cas de conflit temporel au sein d'un même média. Cela est dû à un conflit entre les valeurs des attributs temporels associés au même média. Ce cas de conflit peut être détecté par le contrôle des attributs temporels associés à chaque média. Par exemple, une image qui s'affiche à l'instant 5s et disparaît après 4s, tandis que sa durée maximale est égale à 3s. Par exemple, si on considère le code SMIL du média image suivant:

`<img id="image1" begin="5" end="9" max="3" />`. On peut conclure que cette image ne peut pas démarrer à l'instant 5s et s'arrêter après 4s avec la valeur maximale de sa présentation qui est égale à 3s.

ii. **Conflit inter-média:** représente le cas de conflit entre les différents médias de la présentation. Ce cas de conflit est plus difficile à détecter, car il regroupe plusieurs médias. Cependant, la détection d'un conflit inter-média stipule une analyse de toutes les relations de synchronisation entre les différents

médias de la présentation. Par exemple, si on considère le code SMIL de l'élément *par* suivant:

```
<par id="par1" dur="15">
  <img id="image2" begin="4" end="15" />
  <audio id="audio1" begin="5" end="30" />
</par>
```

Il est impossible que sa durée totale soit égale à 15s et son fils soit joué pendant 25s.

1.2 Contribution

Suite à une étude des approches permettant la vérification des incohérences temporelles et spatiales des documents SMIL, nous nous sommes inspirés d'un travail existant de la littérature et qui est basé sur la logique de Hoare [3] pour vérifier la cohérence temporelle d'un document SMIL. Pour cela, nous avons proposé une nouvelle approche en ajoutant de nouvelles règles permettant la vérification temporelle de plusieurs éléments SMIL, et un nouvel algorithme qui détecte les incohérences spatiales des documents SMIL. En plus, nous avons utilisé les contraintes disjonctives de Marriott [4] pour corriger les conflits spatiaux. Ce choix nous a permis d'analyser les éléments SMIL sans utiliser de modèles auxiliaires comme elles l'ont fait la plupart des méthodes de l'état de l'art, de contrôler et de corriger aussi les erreurs sémantiques du document SML.

2 Travaux Connexes

Little et Ghafoor [5] proposent un modèle modifié de réseaux de Petri appelé Composition Object Petri Net (OCPN), permettant de spécifier les relations de composition temporelle binaires ou n-aires entre les médias. Toutefois, le modèle OCPN ne traite pas les contraintes de synchronisation les plus complexes en temps réel, et donc il n'est pas assez puissant pour capturer toute la sémantique d'ordonnancement et de synchronisation des médias composant un document SMIL.

Dans [6], Yang propose le modèle RTSM (Real Time Synchronization Model) qui est basé sur le modèle OCPN. L'auteur définit alors deux types de localisation et de nouvelles règles de tir sont proposées. Le modèle RTSM peut capturer la sémantique temporelle d'un document SMIL et détecte les conflits temporels des médias. Toutefois, la vérification n'est pas incrémentale et nécessite l'analyse de l'ensemble du RTSM. De plus, la traduction du modèle SMIL vers RTSM génère une perte de la structure temporelle du document SMIL.

Dans [3], les auteurs proposent un outil basé sur une sémantique formelle définissant les aspects temporels des éléments SMIL par le biais d'un ensemble de règles d'inférence proposé dans [7]. Les règles définissent comment l'exécution d'une partie du code change l'état du système par le biais de triplet de Hoare. Si un conflit temporel est détecté, le système renvoie un message à l'utilisateur en mentionnant l'élément qui a causé le conflit. Cette approche définit une sémantique inspirée par la logique de Hoare. Ce choix permet d'analyser les éléments SMIL sans aucune

transformation et de proposer ainsi une correction des erreurs trouvée. De plus, la sémantique de l'approche couvre un vaste ensemble d'éléments SMIL 3.0 et elle est facilement extensible en utilisant d'autres éléments ou attributs, et en ajoutant de nouvelles règles sans remettre en cause l'ensemble du système. Malheureusement, cette approche se concentre seulement sur la vérification des incohérences temporelles mais elle ne traite pas la vérification des incohérences spatiales. Bien que le conflit spatial soit un problème majeur qui affecte la qualité des présentations multimédias.

Dans [8], les auteurs proposent un outil de création des documents SMIL, nommé SMIL Builder qui permet à l'auteur de créer d'une façon incrémental son document SMIL. Le contrôle de cohérence temporelle est basé sur le modèle H-SMIL-Net [9] définit précédemment. L'ensemble des comportements qui peuvent être produits est encore limité, car il ne tient pas compte des éléments importants de SMIL, comme *min*, *max* et *repeatDur*.

Dans [10], les auteurs proposent un modèle basé sur les réseaux de Petri et la logique temporelle appelé SAM (Software Architecture Model). Ce dernier ne permet pas de vérifier les conflits sémantiques, mais juste la satisfiabilité des propriétés de qualité de service (QoS). Le réseau de Petri est traduit en un ensemble d'axiomes qui sont utilisés pour démontrer les formules logiques qui expriment les exigences de la QoS et non pas la vérification de la sémantique d'un document SMIL.

Dans [11], les auteurs présentent un modèle qui analyse le comportement temporel des objets SMIL. Ils offrent une modélisation modulaire qui reflète exactement les documents SMIL. Les structures temporelles sont accordées aux relations de synchronisation entre les objets ayant des références temporelles distinctes.

Dans [12], les auteurs présentent une approche pour la validation des documents multimédia réalisés avec le langage NCL. Le NCL est un langage de création important dans le développement numérique des applications TV au Brésil et dans le monde. Les documents NCL peuvent être compris comme des systèmes de transition finis. Les techniques basées sur les modèles standards sont applicables pour leur validation.

En fait, la vérification des incohérences spatiales des documents SMIL semble être un travail de recherche très peu publié dans la littérature. La preuve on trouve peu de papiers de recherche qui résolvent ce problème en utilisant l'approche basée sur les rectangles libres [13]. Dans [14], les auteurs proposent une extension de l'éditeur SMIL Builder déjà conçu [8] qui prend en charge la vérification spatiale d'un document SMIL. L'idée ici est de diviser les 81 cas de chevauchement en cinq classes, où chaque classe est définie par une formule mathématique. Cependant, pour vérifier s'il y a une incohérence spatiale ou non, l'outil exécute toutes les cinq formules, un par un, et une fois l'erreur est détectée, une correction de cette erreur est exécutée. Dans cette approche, les auteurs utilisent le formalisme des rectangles libres que nous avons déjà définis. Cependant, il existe aussi l'approche de [15] pour résoudre le problème spatio-temporel des présentations multimédias où les auteurs se concentrent sur la vérification des incohérences spatiales au cours de la présentation d'un document multimédia.

3 Vérification des Contraintes Temporelles et Spatio-temporelles des Présentations SMIL

Notre contribution consiste à proposer une méthode formelle pour contrôler et corriger les conflits temporels et spatiaux lors d'une présentation SMIL [16]. L'outil réalisé appelé « *SMIL-Verification* » fournit un ensemble de règles d'inférence et de contraintes disjonctives permettant de décrire comment une partie du code peut changer l'état d'une présentation d'un document SMIL. Notre but consiste en la détection des conflits sémantiques qui sont la cause d'un comportement inattendu d'un document SMIL. La vérification temporelle est basée sur une sémantique formelle qui est décrite par un ensemble de règles d'inférence basé sur la logique de Hoare. Nous proposons aussi un algorithme de vérification de cohérence spatio-temporelle appelé (AVIS) d'un document SMIL. Les contraintes disjonctives de Marriott sont utilisées pour la correction des conflits spatiaux.

Tous les éléments et attributs de SMIL étudiés dans notre méthode sont résumés dans le Tableau 1. En outre, un média peut être joué à différents instants.

Table 1. Liste des objets, éléments, attributs et abréviations de SMIL utilisé dans ce travail.

Abréviations	Eléments SMIL : head body
	head : attribute body: cmd cmd: media par seq excl attribute media : static cont ref static : text imag brush cont: video audio animation attribute: top left height width begin dur end min max repeatDur repeatCount
Valeurs d'attribut	valeur positive (t) ev ev+t ev : ID.begin ID.end ID : media

3.1 Présentation de “SMIL-Verification”

3.1.1 Vérification Spatio-temporelle

Nous avons opté pour la programmation par contraintes disjonctives de Marriott pour résoudre le problème des conflits spatio-temporels. Au début, nous vérifions les éléments SMIL qui sont joués en même temps, par la suite nous contrôlons s'il existe un chevauchement spatial entre ces éléments. Si c'est le cas, alors nous devons proposer à l'auteur une autre configuration de ces régions qui ne provoque pas un conflit spatio-temporel dans la présentation. L'idée est de proposer un système de quatre contraintes où chaque contrainte définit la position entre deux éléments à savoir : gauche, droite, haut et bas. Les variables utilisées par les contraintes représentent les attributs qui définissent les différentes régions de la présentation SMIL. La Figure 1 illustre la dimension spatiale d'une présentation SMIL. Nous

présentons deux régions R1 et R2, chaque région est définie par ses coordonnées (top, left, width et height).

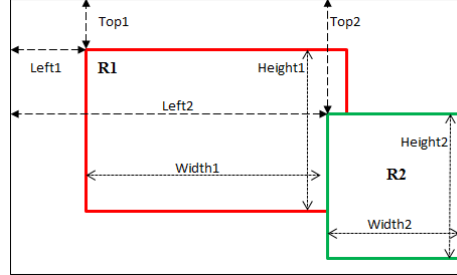


Fig. 1. Vue de l'organisation spatiale d'une présentation SMIL.

Dans ce cas, nous supposons que les deux régions (R1 et R2) sont affichées en même temps. Par l'application de l'algorithme d'AVIS (voir Algorithme 1), nous concluons la présence d'un conflit spatio-temporel dans la présentation SMIL. En effet, les contraintes disjonctives sont appelées pour résoudre ce problème et de trouver une bonne configuration des deux régions afin d'éliminer le conflit.

Le système défini est comme suit:
$$\left\{ \begin{array}{l} top_1 \geq top_2 + height_2 \vee \\ top_2 \geq top_1 + height_1 \vee \\ left_1 \geq left_2 + width_2 \vee \\ left_2 \geq left_1 + width_1 \end{array} \right\}$$

Pour résoudre le conflit, il suffit de satisfaire une seule contrainte parmi les quatre :

- La première contrainte $top_1 \geq top_2 + height_2$ propose que la région 1 (R1) soit au-dessous de la région 2 (R2).
- La deuxième contrainte $top_2 \geq top_1 + height_1$ propose que la région 1 (R1) soit au-dessus de la région 2 (R2).
- La troisième contrainte $left_1 \geq left_2 + width_2$ propose que la région 1 (R1) soit à droite de la région 2 (R2).
- La quatrième contrainte $left_2 \geq left_1 + width_1$ propose que la région 2 (R2) soit à droite de la région 1 (R1).

3.1.2 Vérification Temporelle

Notre approche est basée sur une sémantique formelle pour définir les aspects temporels d'éléments SMIL et utilise un ensemble de règles d'inférence. Les règles sont basées sur la sémantique de Hoare et utilisent un moteur d'inférence. Notre approche permet au moteur d'inférence de vérifier et valider la durée active de tous les éléments SMIL de la présentation, en appuyant sur les attributs « *min* », « *max* » et « *repeatDur* ». Nous avons étendu l'ensemble des attributs de SMIL déjà vérifiés par de nouveaux attributs tels que: *min*, *max* et *repeatDur*, pour offrir à l'auteur plus de

contrôle sur la durée active des objets média. La durée active d'un élément définit la période entière où le plan de montage chronologique de l'élément est actif. Cela prend en compte la durée simple de l'élément, qui est définie par les deux attributs *begin* et *end*.

Algorithme : Algorithme de Vérification des Incohérences Spatio-temporel(AVIS)

Entrées: start, dur, left, top, width, height, n.
 //start: le temps d'apparition des médias sur l'écran.
 //dur: la durée d'apparition de l'élément média sur l'écran.
 //left et top: les coordonnées verticales et horizontales de la région, respectivement.
 //width: la largeur de la région.
 //height: la hauteur de la région.
 //n: le nombre des médias dans le document.

Sortie : message de vérification

```

for i=1→n do // i et j représente les différents médias dans la présentation
for j= i+1→n do

  if (((starti ≥ startj) & (starti ≤ (startj+durj))) || ((startj ≥ starti) & (startj ≤ (starti+duri)))) then
    print(" il existe une interaction temporelle entre les deux élément i et j ")
  if ( (lefti ≤ (leftj+widthj)) & ( (topi-heighti) ≤ (topj-heightj) & ((topj-heightj) ≤ topi) ) then
    print(" il existe une incohérence spatiale entre les deux éléments i et j ")
  else
    print(" il n'y a pas une incohérence spatiale entre les deux éléments i et j ")
  else
    print(" il n'y a pas une interaction temporelle entre les deux éléments i et j ")
  end if
end if

end for
end for
  
```

Tout comportement de répétition est défini par les attributs *repeatDur* et *repeatCount*. L'attribut *repeatDur* spécifie la durée totale de la répétition d'un élément SMIL donné. Dans l'exemple suivant, l'audio va se répéter pendant 7 secondes. Il sera joué entièrement 2 fois et partiellement pendant 2 secondes. Cela est équivalent à un attribut « repeatCount=2.8 » (voir Figure 2) :

<audio src= « music.mp3 »begin= « 2 »end= « 4.5 »min= « 1 »max= « 12 »repeatDur= « 7 »>.

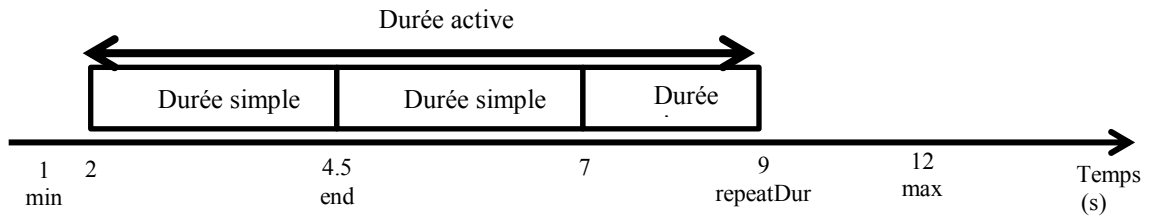


Fig. 2. Composition temporelle de l'élément « audio ».

De ce fait, nous avons proposé et implémenté une règle formelle inspirée par la logique de Hoare. Cette règle permet la vérification et le contrôle de cohérence temporelle de la durée active des éléments SMIL ainsi que leurs bornes supérieure et

inférieure. Si à la fois les attributs « *min* » et « *max* » sont définis, alors la valeur de « *max* » doit être supérieure ou égale à celle de « *min* ». Si cette condition n'est pas remplie, les deux attributs seront ignorés et aucune notification ne sera communiquée à l'auteur. Si la définition des médias est basée sur son apparition plusieurs fois dans la présentation. En d'autres termes, si l'auteur définit un média avec plusieurs « durées simple » qui sont séparées entre eux. Dans ce cas-là, l'évaluation de cet élément nécessite la définition d'une règle plus complexe (voir Figure 3).

$$\begin{aligned}
 &\text{media+begin+end+dur+min+max+repeatDur} \\
 &\{ \text{pré_condition} \} \\
 &\quad \langle \text{media } m \text{ begin} = "v1" \text{ end} = "v2" \text{ dur} \\
 &\quad \quad = "v3" \text{ min} = "v4" \text{ max} = "v5" \text{ repeatDur} \\
 &\quad \quad = "v6" \rangle \\
 &\{ \text{post_condition} \} \\
 &\text{pré_condition} = \{ t_{cr}(m) = \text{strat} - v1 \} \\
 &\text{post_condition} = \left\{ \begin{array}{l} \text{begin}(m) = \text{strat} \\ \text{end}(m) = \text{start} + v3 \\ \text{end}(m) = \text{start} - v1 + v2 \\ v4 < v5 \\ v1 \geq v4 \\ v2 \leq v5 \\ (v5 - v4) \geq v6 \end{array} \right\} \\
 &\text{conditions d'applicabilités} \\
 &= \left\{ \begin{array}{l} (\text{defined}(v2) \vee \text{defined}(v3)) \Rightarrow v3 = v2 - v1 \\ \text{repeat}(v3, n_{fois}) \Rightarrow v6 = v3 * n \\ (\text{defined}(v3)) \vee (\text{defined}(v1) \wedge \text{defined}(v2)) \Rightarrow (v3 \notin \text{notDur}) \vee ((v2 - v1) \notin \text{notDur}) \end{array} \right\}
 \end{aligned}$$

Fig. 3. Règle de preuve de la durée active (repeatDur) des éléments de la présentation SMIL.

Le moteur d'inférence utilise les règles d'inférence pour raisonner et construire l'arbre de la preuve. Si le document contient un conflit, la construction de la preuve échoue parce que l'une de ses prémisses nécessaires ne peut pas être prouvée ou les conditions d'application ne sont pas satisfaites et le système retourne l'élément qui contient l'erreur avec la traçabilité d'erreur et propose des suggestions sur la méthode de correction.

3.2 Extension de “SMIL-Verification” pour la Vérification Temporelle

Nous avons augmenté le nombre d'attributs temporels à vérifier. Pour cela, une nouvelle règle d'inférence a été ajoutée qui contrôle l'attribut « *repeatCount* » ; elle à la spécification de combien de fois l'élément SMIL soit répété (voir Figure 4).

$$\begin{aligned}
 &\text{media+begin+end+dur+min+max+repeatCount} \\
 &\quad \{ \text{pré_condition} \} \\
 &\quad \langle \text{media } m \text{ begin} = "v1" \text{ end} = "v2" \text{ dur} = "v3" \text{ min} = "v4" \text{ max} = "v5" \text{ repeatCount} = "v6" \rangle \\
 &\quad \{ \text{post_condition} \} \\
 &\text{pré_condition} = \{ t_{cr}(m) = \text{strat} - v1 \} \\
 &\text{post_condition} = \left\{ \begin{array}{l} \text{begin}(m) = \text{strat} \\ \text{end}(m) = \text{start} + v3 \\ \text{end}(m) = \text{start} - v1 + v2 \\ v4 < v5 \\ v1 \geq v4 \\ v2 \leq v5 \\ (v5 - v4) \geq v3 * v6 \end{array} \right\} \\
 &\text{conditions d'applicabilités} \\
 &= \left\{ \begin{array}{l} (\text{defined}(v2) \vee \text{defined}(v3)) \Rightarrow v3 = v2 - v1 \\ (\text{defined}(v3)) \vee (\text{defined}(v1) \wedge \text{defined}(v2)) \Rightarrow (v3 \notin \text{notDur}) \vee ((v2 - v1) \notin \text{notDur}) \end{array} \right\}
 \end{aligned}$$

Fig. 4. Règle de preuve de la durée active (repeatCount) des éléments de la présentation SMIL.

La condition d'applicabilité nous empêche d'appliquer la règle « *media + begin + end + dur + min + max + repeatCount* » dans les cas suivants : 1) La première condition d'applicabilité indique que si le début et la fin de l'élément SMIL sont indiqués alors la durée de l'élément doit être égale à leurs différences. 2) L'instant fin de l'élément SMIL qui est en train de se répéter ne peut être calculé si l'attribut *end* (ou *dur*) est égal à "indéfinie". De ce fait, une fois que nous sommes capables de décider si cet élément SMIL se termine, nous devons calculer son instant d'arrêt (*end*). 3) La dernière condition, exigé que soit la durée doit être définie obligatoirement ou bien la différence de « *begin* » et « *end* » de la durée simple doit être définie, afin de calculer la durée active de l'élément en question.

4 Implémentation et Résultats Expérimentaux

Durant la mise en œuvre, nous avons utilisé la version 1.8 de java¹ ainsi que NetBeans IDE 8.2², à cause de raison de compatibilité avec la librairie JaCoP 4.4.0 utilisée. Pour réussir la programmation par contraintes disjonctives, nous avons utilisé la bibliothèque JaCoP³ de Java. Il existe de nombreux solveurs d'instances CSP, à la fois commercial (IBM ILOG CPLEX CP Optimizer, ...) et académique (AIspace, Choco, JaCoP, ...). Pour résoudre notre problème nous avons choisi d'utiliser la bibliothèque JaCoP, qui est une API en Java. Cette bibliothèque est assez petite (moins de 3 Megaoctets) et fournit les codes sources sous licence GNU, ce qui permet de vérifier comment certaines fonctionnalités sont implémentées.

Nous avons utilisé la bibliothèque JaCoP car elle fournit un paradigme de programmation de contraintes implémenté en Java. Ainsi, elle fournit des primitives pour définir des variables de domaine fini (FD), des contraintes et des méthodes de recherche. Le programmeur doit être familiarisé avec la programmation par contraintes (logique) pour pouvoir utiliser cette bibliothèque. La bibliothèque JaCoP fournit les contraintes primitives les plus couramment utilisées, telles que l'égalité, l'inégalité ainsi que les contraintes logiques, réifiées et conditionnelles. La bibliothèque JaCoP peut être utilisée en la fournissant en tant que fichier JAR ou en spécifiant l'accès à un répertoire contenant toutes les classes JaCoP.

Dans l'application Java qui utilise JaCoP, il est nécessaire de spécifier des instructions d'importation pour toutes les classes utilisées de la bibliothèque JaCoP.

La Figure 5 présente la vue générale de l'outil développé.

¹ <http://www.oracle.com/technetwork/java/javase/downloads/jdk8-downloads-2133151.html>

² <https://netbeans-ide.fr.uptodown.com/windows>

³ <http://jacopapi.osolpro.com/>

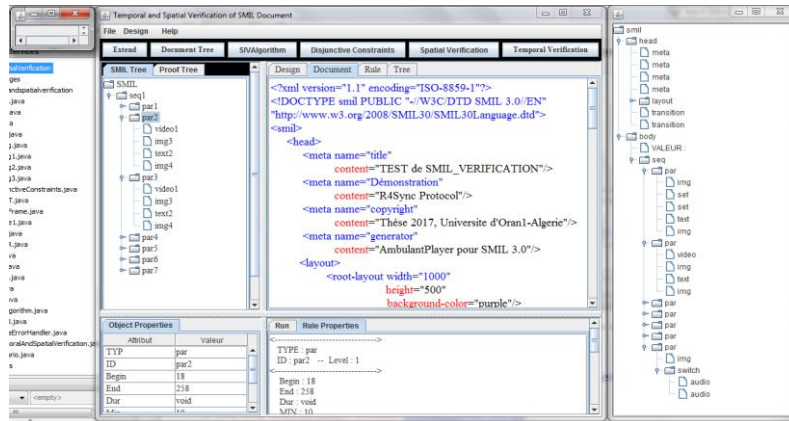


Fig. 5. Copie écran du prototype.

Pour illustrer le fonctionnement de notre prototype *SMIL-Verification*, deux cas de vérification sont présentés : spatiale et temporelle. Pour cela, nous avons créé une présentation multimédia en SMIL qui montre la précision d'un protocole de synchronisation du temps dans les réseaux de capteurs sans fil en temps réel (RCSF) [17].



Fig. 6. Copie écran qui montre une incohérence spatiale dans la présentation SMIL.

La figure 6 illustre la présentation finale de la démonstration en utilisant le lecteur Ambulant de SMIL 3.0. Cette figure décrit le comportement inattendu de la présentation. Nous remarquons que la région 6 qui présente le plan de la démonstration a été mal définie par l'auteur car elle produit un chevauchement entre les deux régions 6 et 5 (la région 5 **contient** la region6). De même, la région 3 se chevauche avec la région 4 (le text chevauche l'extrémité supérieure de la photo). Pendant la lecture de la présentation l'auteur découvre qu'il y a deux conflits spatiaux. De ce fait, nous proposons de tester le document SMIL avant sa lecture par Ambulant Player.

La Figure 7 présente le résultat de l'application de SIVA ainsi que la programmation par contraintes. Cependant, l'algorithme a détecté l'apparition de connexion spatio-temporelle entre les deux régions (3 et 4), ainsi que les régions (5 et

6). De ce fait, les contraintes disjonctives sont appelées afin de proposer une solution qui ne provoque pas de conflit spatio-temporel. La nouvelle configuration des régions (5 et 6) prend en compte la dimension de la fenêtre principale de la présentation.

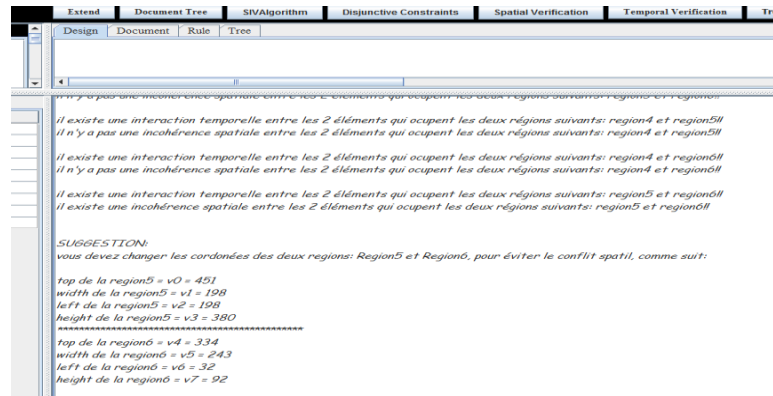


Fig. 7. Copie écran qui montre la vérification des incohérences spatiales.

5 Conclusion

Le travail présenté dans ce papier est une contribution dans le domaine du multimédia et plus spécifiquement dans la formalisation, l'analyse et le contrôle de cohérence temporelle et spatiale des présentations SMIL. Un document SMIL est un document multimédia permettant l'intégration de plusieurs médias dans une page Web. Cependant, la création d'une présentation multimédia SMIL est une tâche très difficile qui génère beaucoup d'erreurs, particulièrement quand la structure temporelle ou spatiale du document est complexe.

La contribution essentielle de ce travail est la proposition d'une approche formelle qui contrôle un grand nombre d'éléments SMIL : *begin*, *end*, *dur*, *par*, *seq*, *excl*, *min*, *max*, *repeatDur*, *top*, *left*, *width* et *height*. D'autre part, la méthode de [3] ne contrôle pas les attributs, *min*, *max*, *repeatDur*, *repeatCount*, *top*, *left*, *width* et *height*; et l'approche de [15] ne contrôle pas les attributs : *dur*, *par*, *seq*, *excl*, *min*, *max*, *repeatDur* et *repeatCount*. Nous avons conçu et réalisé ainsi un outil formel nommé *SMIL-Verification* utilisant la méthode proposée pour assister d'une façon incrémentale l'auteur dans la phase d'édition et de vérification de son document SMIL.

Il est néanmoins important de noter que notre méthode proposée est loin d'être achevée et qu'elle doit évoluer dans un futur proche. Nous suggérons donc d'améliorer notre outil pour vérifier à la fois les trois dimensions d'un document SMIL à savoir : la dimension temporelle, spatiale et la dimension hypermédia.

References

1. Consortium, W. W.-W. W., et al: Synchronized Multimedia Integration Language–SMIL 3.0 Specification, W3C Recommendation. DOI: <http://www.w3.org/TR/SMIL3>
2. Hoare, C. A. R. : An axiomatic basis for computer programming. *Communications of the ACM* (1969) 576–580
3. Gaggi, O., Bossi, A.: Analysis and verification of SMIL documents. *Multimedia Systems*, Vol.17, Num.6 (2011) 576–580
4. Marriott, K., Moulder, P., Stuckey, P. J., Borning, A.: Solving Disjunctive Constraints for interactive graphical applications. In *International conference on principles and practice of constraint programming* (2001) 361–376
5. Little, T. D., Ghafoor, A.: Synchronization and storage models for multimedia objects. *Selected Areas in Communications*, Vol.8, Num.3 (1990) 413–427
6. Yang, C.-C.: Detection of the time conflicts for SMIL based multimedia presentations. In *Workshop on computer networks, internet, and multimedia* (2000) 57–63
7. Bossi, A., & Gaggi, O.: Enriching SMIL with assertions for temporal validation. In *Proceedings of the 15th international conference on multimedia* (2007) 107–116
8. Bouyakoub, S., Belkhir, A.: SMIL Builder: An incremental authoring tool for SMIL documents. *ACM Transactions on Multimedia Computing, Communications, and Applications (TOMM)* (2011)
9. Bouyakoub, S., Belkhir, A.: H-SMIL-Net: A hierarchical Petri Net model for SMIL documents. In *Computer modeling and simulation, UKSIM* (2008) 106–111
10. Yu, H., He, X., Gao, S., Deng, Y.: Modeling and analyzing SMIL documents in SAM. In *Multimedia software engineering* (2002) 132–139
11. Dahmani, D., Mazouz, S., Boukala, M. :Efficient and modular modeling of hierarchical multimedia objects by using Time ERPn+. In *Multimedia computing and systems (ICMCS)* (2014) 1153–1158
12. dos Santos, J., Braga, C., & Muchaluat-Saade, D. C.: A rewriting logic semantics for NCL. *Science of Computer Programming* (2015) 64–92
13. Bertino E., Member S., Ferrari E; and Stolf M., MPGS: An Interactive Tool for the Specification and Generation of Multimedia Presentations, *IEEE Transactions On Knowledge And Data Engineering*, Vol.12, No.1, January/February 2000
14. Bouyakoub, S., & Belkhir, A.: A spatio-temporal authoring tool for multimedia SMIL documents. *International Journal*, Vol.2, Num.3, (2012) 83–88
15. dos Santos, J. A., Braga, C., Muchaluat-Saade, D. C., Roisin, C., & Layaida, N.: Spatio-temporal validation of multimedia documents. In *Proceedings of the ACM symposium on document engineering* (2015) 133–142
16. Mekahlia F., Ghomari A., Yazid S., Djenouri D.: Temporal and Spatial CoherenceVerification in SMIL Documents with Hoare Logic and Disjunctive Constraints: A Hybrid Formal Method. *Transactions of the SDPS: Journal of Integrated Design and Process Science*. IOS Press. DOI: 10.3233/jid-2016-0020, 21 October 2016. (Issue Preprint)
17. Djenouri D., Merabtine N., Mekahlia F., Doudou M. Fast Distributed Multi-hop Relative Time Synchronisation Protocol and Estimators for Wireless Sensor Networks. *Ad Hoc Networks Elsevier Science*, 20 June 2013

A Mathematical Model and a Fast Lower Bound for the Open-End Bin Packing Problem with Conflicts

Mohamed Maiza^{1,2}, Mohamed Tebbal¹, Billal Rabia¹, and Lakhdar Sais²

¹ Laboratoire de Mathématiques Appliquées, Ecole Militaire Polytechnique, BP17
Bordj-El-Bahri, Alger, Algérie

² Centre de Recherche en Informatique de Lens CNRS-UMR8188, Université
d'Artois, F-62300 Lens, France

Abstract. The Open-End Bin Packing Problem (OEBPP) aims to pack a set of items with varying sizes in a minimum number of identical bins such that any bin can be filled beyond its capacity as long as the removal of the last packed item returns the bin not fully filled. In this paper, we propose a new extension, called OEBPPC, obtained through additional incompatibility constraints. Such constraints are usually characterized by an undirected conflict graph whose vertices correspond to the items while edges correspond to couples of items that cannot be packed together in the same bin. We first provide an original formulation as an *Integer Linear Program*. Then we propose novel and efficient lower bound in order to guide the search in the solution space. An experimental evaluation shows the efficiency of our proposed approach in both running time and solution quality.

1 Introduction

The one dimensional open-end bin packing problem, called OEBPP, is a variant of the well known bin packing problem in which items with varying sizes are packed into identical bins such that in each bin, the total items size content before packing the last item is strictly less than the bin capacity. In other words, for each bin, only so-called *the overflow item* which correspond to the last packed item can exceed the capacity of the bin. The aim of the OEBPP is to minimize the number of bins used to pack all items. This problem arises in a wide variety of applications including manufacturing, transportation, affectation tasks. Illustrative examples of the OEBPP application in both the fare payment system in the Hong Kong Taipei subway stations and in job scheduling of the manufacturing environment have been given by Leung et al. [11] and Pornthipa [17] respectively.

In this paper we address the offline version of the OEBPP while avoiding joint assignments of some items that are incompatible. Which means that each incompatible items pair cannot be packed in the same bin. We denote this problem variant as *Open End Bin-Packing Problem with Conflicts (OEBPPC)*. Such additional constraints play an important role in several applications of bin packing

problems. Indeed, let us mention a job scheduling application in the manufacturing environment that operates one shift per day. A worker loads jobs to a machine that can operate automatically, and for example, due to the drying constraints or also the operating mode constraints some jobs are incompatible and cannot be executed in the same day. The objective is to schedule jobs to minimize the total time (days) to complete all the jobs. Intuitively for a given day, the job with the longest processing time which is compatible with already loaded jobs is scheduled until its completion. The goal is to fully occupy the automated machine when the worker leaves his job, and the worker can unload the finished job next day. Similarly, it also applies in the context of computer process scheduling.

The introduced OEBPPC extension generalizes two well-known NP-hard problems, namely OEBPP and classical Bin Packing Problem with Conflicts (BPPC). Indeed, if there is no conflicts over items then the conflicts graph is a discrete graph, so, the OEBPPC can be reduced to the OEBPP. On other hand, if no overflow is allowed and the total size of packed items cannot exceed the capacity of the bin, the OEBPPC can be reduced to the BPPC. Therefore, it is interesting to exploit the results obtained on the two previous cited problems to solve OEBPPC. As mentioned below, by considering incompatibility constraints, our goal is to enlarge the application scope to other real world problems.

To the best of our knowledge there is no published paper investigating this new variant. However, various closely related problems have been discussed. The classical one dimensional bin packing problem BPP, from which the current OEBPPC stems, has been extensively studied since 1970s. For a complete overview, we refer the interested reader to the survey chapter of Coffman et al. [2] and to the recent survey by Delorme et al. [3]. For what concerns the BPPC (*with conflicts constraint*), it is worth to note that this problem has been investigated through numerous papers, we refer for example to the works of Jansen et al. [7], Kalfakakou et al. [9], Gendreau et al. [5], Muritiba et al. [16], Khanafer et al. [10], Maiza et al. [13] and Sadykov et al. [18].

On the other hand, only few works have addressed OEBPP (*without conflicts*). Indeed, Leung et al. [11] introduced the one dimensional open-end bin packing problem for modeling the ticketing system of Hong Kong subway station. They showed that any online algorithm must have an asymptotic worst-case ratio of at least 2. Also, the authors proved that the problem is strongly NP-hard. Then, Yang and Leung [19] studied the Ordered Open Bin Packing Problem which extends the requirement corresponding to the order of passenger's itinerary in subway station problem. Thereafter, a branch and price algorithm for an exact optimization of the Ordered Open-End Bin Packing Problem has been provided by Ceselli and Righini [1].

To solve efficiently the OEBPPC, we present new algorithms to compute lower bound for the OEBPPC. We make an original use of heuristic-based techniques, not only to provide valid upper bounds, but also to strengthen the behavior of the lower bounding techniques. Generally, the heuristic algorithms are

often used to speed up the computation and improve the robustness of exact approaches.

The reminder of this paper is organized as follows. To clearly state the problem, we start by introducing in Section 2 some notations and a compact mathematical formulation. In Section 3 we present an efficient lower bound based on both clique computation and some known mathematical formulation. The proposed lower bound is tested and compared on a large set of instances in Section 4. Finally conclusions are drawn in Section 5.

2 Mathematical problem statement

Formally, the OEBPPC can be described as follows. We are given a set $V = \{1, 2, \dots, n\}$ of items and an infinite number of identical bins of capacity W . The items $i \in V$ must be packed into bins with a given non negative space requirement w_i . We are also given a conflict graph $G = (V, E)$, where E is a set of edges such that $(i, i') \in E$ expresses that i and i' are in conflict or *incompatible*. The problem is to assign items to bins, using a minimum number of bins, while ensuring that no two items that are in conflict are assigned to the same bin and the capacity of each bin can be exceeded only by the last item packed into it. Hence, the last item in each bin called the *overflow item* requires at least only one unit of capacity.

A natural and compact integer programming formulation makes use of binary variables x_{ij} taking value 1 if an item i is assigned to the bin in which the overflow item is item j and 0 otherwise, and binary variables y_j which indicates whether item j is the overflow item in its bin:

$$\text{Minimize } \sum_{j \in V} y_j \quad (1)$$

s.t.

$$y_i + \sum_{i \neq j} x_{ij} = 1 \quad \forall i \in V \quad (2)$$

$$\sum_{i \neq j} w_i x_{ij} \leq (W - 1)y_j \quad \forall j \in V \quad (3)$$

$$x_{ij} + x_{i'j} \leq 1 \quad \forall (i, i') \in E \quad (4)$$

$$x_{ij} + y_j \leq 1 \quad \forall (i, j) \in E \quad (5)$$

$$x_{ij} \in \{0, 1\} \quad (6)$$

$$y_i \in \{0, 1\} \quad (7)$$

The number of bins used is indicated in the objective function (1) by the number of binary variables y_j set to one. Constraints (2) impose that each item i , whether an overflow item or not, must be assigned to one bin. Constraints (3)

ensure that the overall size of the items assigned to a bin, excluding the overflow item, must fit into the bin and must leave at least one capacity unit available for accommodating the overflow item. Both constraints (4) and (5) make sure that each incompatible pair of items (i, i') cannot be assigned to the same bin j , excluding or including the overflow item respectively. Finally, constraints (6) and (7) restrict the domain of the decision variables to $\{0, 1\}$.

Note that, if an overflow item y_j in such a bin j is set to 1 then the binary variable x_{jj} which indicates whether this item j is assigned to the bin j which does not make sense. This is why, we only have x_{ij} variables with $i \neq j$ in constraints (2) and (3).

3 Lower bound

In this section, we provide a new fast lower bound computation algorithm exploiting the combinatorial structure of the OEBPPC. In the following we use L to define the value produced by the lower bounding procedure.

3.1 Continuous lower bound for the problem without conflicts

Continuous lower bound technique has been originally provided by Martello and Toth [14][15] for bounding solutions of the classical BPP. It consists of considering items in fractional way by the domain relaxation of the decision variables. Since, this technique has been enlarged to encompass several variants of packing and cutting problems [5][10][16][6][12]. For the OEBPP (*without conflict*), the continuous lower bound on a set V of items -denoted in what follow L_{OEBPP}^V - is a procedure based on the computing of the minimal number of largest items from V which can be considered as overflow items. So, only these items are not considered in fractional way. Intuitively, the largest items are assigned to bins as overflow items and require the last unit capacity of the bin. For that, the set V of items are considered in decreasing order of size. The lower bound value L_{OEBPP}^I is the minimal positive number λ of largest items in which the fractional assignment of remaining items (items indexed from $\lambda + 1$ to n) into bins with capacity $W - 1$ require at least λ bins. Meaning that, given a set V of items ordered in decreasing order of size, continuous lower bound L_{OEBPP}^V is given by

$$L_{OEBPP}^V = \min(\lambda \in \mathbb{N}^+ / \sum_{i=\lambda+1}^n w_i \leq \lambda(W - 1)) \quad (8)$$

3.2 New lower bounds for OEBPPC

In this section, we present our contribution to compute lower bounds for the one dimensional open-end bin packing problem with conflicts. The first one, is based on a simple relaxation of the compact integer programming formulation (see Section 2). The second lower bound based on clique computation presents interesting computational properties.

Model-based relaxation lower bound A lower bound for the OEBPPC can be obtained directly from the linear relaxation of (1)-(7), when the domain of the decision variables (6) and (7) are relaxed, we call this lower bound a *simple continuous lower bound* L_0 . Indeed, relaxing the decision variables means considering all items in fractional or continuous way. Therefore, the overflow items will never exceed the capacity of the bins and L_0 gives consequently the same value of relaxed model of the classical bin packing problem with conflicts BPPC. Let L_1 be the lower bound obtained by relaxing only the domain of the decision variables (6). The corresponding model is thereafter a mixed integer linear program (MILP) with some integer decision variables in which the overflow items can not be assumed in a fractional way. It is clear that L_1 improve values produced by L_0 . Unfortunately, its computation using MILP solvers is time consuming.

Clique-based lower bound Formally, the OEBPPC is similar to the classical BPPC in which the capacity of bins is $W - 1$ where we add as much overflow items as constructed bins. Hence, any valid lower bound for BPPC is a valid lower bound for a subclass of OEBPPC in which 1) the overflow items are not considered and 2) the capacity of bins is limited to $W - 1$.

In the sequel, we describe our lower bound for the OEBPPC, denoted L_{OEBPPC} , based on the adaptation of the BPPC clique computation lower bound.

We note that for the BPPC, many lower bounds have been developed in the literature ([7], [9],[5],[16],[10],[13],[18]). We focus on the work of Gendreau et al. [5] where the authors proposed a *constrained packing lower bound* (L_{cp} in the following). The used algorithm starts by computing a large clique K of the conflict graph G through a standard greedy algorithm and initializing one bin for each item in K . Then it fractionally assigns the items in $V \setminus K$ to the opened bins, by solving a transportation problem that takes into consideration both conflicts and sizes. All items or fractional items that do not fit in the opened bins are sorted in new bins by computing the classical continuous BPP lower bound without considering conflicts. Let q denotes the value produced by the continuous lower bound. We have

$$L_{cp} = |K| + q \quad (9)$$

L_{OEBPPC} is an adaptation of L_{cp} updated as follows: we start by computing a large clique K of the conflict graph G through the standard greedy algorithm and initializing one bin for each item in K . Then, by solving the standard assignment problem, we extract the set O from $V \setminus K$ of overflow items such that their number does not exceed the size of the obtained clique ($|O| \leq |K|$). Obviously, the obtained set O contains a selection of these largest items of $V \setminus K$ which are compatible with clique items. The items of O will be assigned directly to the opened bins. Since, each overflow item occupy one unit capacity of opened bin. Therefore, the residual capacity of each opened bin become $W_k = W - w_k - 1$ where w_k is the size of the assigned clique item. After that we fractionally assign the remaining items in $V \setminus K \setminus O$ to the residual capacity of the opened bins,

by solving a transportation problem that takes into consideration both conflicts and sizes. All items or fractional items that do not fit in the initialized bins are grouped in the set $\tilde{V}$ and stored in new bins by computing the continuous OEBPP lower bound $L_{OEBPP}^{\tilde{V}}$ (formula 8) without considering conflicts. Then we have

$$L_{OEBPPC} = |K| + L_{OEBPP}^{\tilde{V}} \quad (10)$$

Below a step by step description of the proposed lower bounding procedure:

Step 1 Determine a large clique K on $G = (V, E)$ using Johnson's [8] greedy heuristic, and assign each item of K to a different bin. Let $W_k = W - w_k$ be the residual capacity of the bin with clique item k . Set $V = V \setminus K$;

Step 2 Define the following assignment problem which consists of maximizing the overall assignment cost. For each item $i \in V$, the assignment cost c_{ik} of item i to bin k is equal to w_i if item i is non-conflicting with item already assigned to k , and to 0 otherwise. The assignment problem is then

$$\text{Maximize } \sum_{i \in V} \sum_{k \in K} c_{ik} x_{ik} \quad (11)$$

subject to

$$\sum_{k \in K} x_{ik} \leq 1 \quad \forall i \in V \quad (12)$$

$$\sum_{i \in V} x_{ik} \leq 1 \quad \forall k \in K \quad (13)$$

$$x_{ij} \in \{0, 1\} \quad (14)$$

where x_{ik} is the boolean decision variable indicating whether item i is assigned to bin k .

The solution of (11)-(14) is a set of overflow items $O \subset V$ with $|O| \leq |K|$;

Step 3 Set $W_k = W_k - 1$ and $V = V \setminus O$;

Step 4 Define the following transportation problem that optimally assign the entire or partial set of items from V to opened bins $\{b_1, b_2, \dots, b_{|K|}\}$, where we add a virtual bin b_0 with unlimited capacity to receive items or fractional items that cannot fit into any of the $|K|$ opened bins. The costs c_{ik} are defined as follows. For $k \neq 0$, c_{ik} is equal to 1 if i is non-conflicting with both clique item and overflow item already assigned to k ; ∞ otherwise. Whilst, for the virtual bin $k = 0$, $c_{i0} = n$ (n is a positive value different from 1) for all $i \in V$. The transportation problem is then

$$\text{Minimize } \sum_{i \in V} \sum_{k \in K} c_{ik} x_{ik} \quad (15)$$

subject to

$$\sum_{k \in K} x_{ik} = w_i \quad \forall i \in V \quad (16)$$

$$\sum_{i \in V} x_{ik} \leq w_k \quad \forall k \in K \quad (17)$$

$$x_{ik} \in \mathbb{Z}^+ \quad (18)$$

where x_{ik} is the portion of the item i assigned to bin k . Let $\tilde{V}$ be the set of items and fractional items assigned to the virtual bin b_0 ;

Step 5 Compute the continuous OEBPP lower bound $L_{OEBPP}^{\tilde{V}}$ without considering conflicts using formula (8), and deduce the lower bound value L by using formula (10);

4 Computational Results

We report in this section the computational experiments performed with proposed lower bound on a broad set of test problem. Our algorithms were coded in C++ and run on an *Intel Core i7 CPU 2.7GHz* processor with *4GB* of RAM under a Windows7 operating system. The ILP models were optimally solved using ILOG CPLEX 12.5.

We consider three data-sets containing the classical BPP instances called in what follows DS1, DS2 and DS3 respectively, on which we generate conflict graphs with distinct densities. The two first data-sets are inspired from the literature library. The last data-set contains our generated instances.

DS1 and **DS2** have been taken from the OR-Library, originally generated by Faulkenauer [4], and includes 8 classes of 20 different instances. These instances are widely recognized as being rather difficult. The 4 first classes known as *Uniform instances*, include bins of capacity $W = 150$ and items with integer size uniformly distributed in the interval $[20-100]$. The number of items n for each class is 120, 250, 500 and 1000 respectively. The 4 last classes called *Triplet instances* include bins of capacity $W = 100$. The number of items n in each class is respectively 60, 120, 249 and 501. Triplet instances consider items with integer size uniformly distributed in the interval $[25-50]$ with one decimal digit; therefore each bin cannot be filled with more than three items. Since the OEBPPC is an NP-hard problem, solving optimally its ILP formulation with a CPLEX solver is excepted for small problems. We have therefore used the 10 first instances for the two first classes of triplet instances as DS1 and DS2 respectively, that is to say a total of 20 instances.

DS3 contains a class with 10 randomly generated instances with a problem size of $n = 50$ where items sizes are generated in the interval of $[20-160]$ according to a discrete uniform distribution, the bins have a fixed capacity of $W = 200$.

A conflict graph G was generated with 10 different densities for each of the 30 described instances. For each vertex which represents an item i in the conflicts graph, an uniform value v_i on $[0-1]$ is assigned. Then, an edge (i, i') is created if $(v_i + v_{i'})/2 \leq d$ where d is the expected density of G . We used $d = \{0\%, 10\%, \dots, 90\%\}$, yielding to an overall number of 300 OEBPPC instances.

Table 1 summarizes results obtained by the application of proposed lower bound, where for each cited data-set, we present the performances in term of deviation from the optimal solution ($\text{gap}(\%)$) calculated as $\text{gap}(\%) = 100 \cdot (s^* - LB)/s^*$ (where s^* is the optimal solution); and CPU time. For each density

$d(\%)$	<i>DS1</i>		<i>DS2</i>		<i>DS3</i>	
	<i>Gap%</i>	<i>Sec.</i>	<i>Gap%</i>	<i>Sec.</i>	<i>Gap%</i>	<i>Sec.</i>
0	6.66	0.0015	6.66	0.0000	5.33	0.0000
10	3.33	0.0000	5.66	0.0000	4.83	0.0000
20	3.33	0.0016	5.66	0.0030	4.66	0.0037
30	0.74	0.0031	0.00	0.0016	1.83	0.0024
40	0.00	0.0000	0.00	0.0000	0.00	0.0000
50	0.00	0.0031	0.00	0.0000	0.00	0.0000
60	0.00	0.0000	0.00	0.0016	0.00	0.0000
70	0.00	0.0000	0.00	0.0015	0.00	0.0001
80	0.00	0.0000	0.00	0.0063	0.00	0.0012
90	0.20	0.0016	0.00	0.0088	0.00	0.0028
<i>Avg</i>	1.42	0.0010	1.79	0.0022	1.66	0.0011

Table 1. Computational tests of proposed lower bound

$d(\%)$ of conflicts graph, we report an average of the results over ten OEBPPC instances. Results in boldfaced characters indicate that on those instances the lower bound give rise to an optimal solution. Results reported in Table1 clearly indicate the high quality of the proposed lower bound. In most cases, for each tested classes, an optimal solution is found on all the 10 considered instances, especially, when the density is larger than 30%. Otherwise, the gap is almost small and does not exceed 6%. These results present the same behavior and do not depend on the problem size. Computing times are in all cases negligible and slowly increases with the density, this is due to the time required to solve the assignment model and the transportation model when the clique becomes large. In overall, the proposed lower bound seems highly successful and typically yields optimal solutions within very low computing times.

5 Conclusion

The Open End Bin-Packing Problem with Conflicts is an *NP-hard* combinatorial problem often encountered in several practical domains. In the present paper, we discussed this important variant of the bin packing problem with arbitrary conflicts between items to be packed. We first presented a compact mathematical model using an integer programming formulation which may open an interesting research avenue. Then, we have described a new lower bound based on clique computation and on both models of assignment and transportation problems. Computational experiments based on a set of 300 instances are carried out and confirm the good performance of the proposed lower bound both in term of solution quality and run time. The proposed lower bound can be subsequently exploited in the future works to improve some exact methods like branch and bound.

References

1. Ceselli, A., Righini, G.: An optimization algorithm for the ordered open-end bin-packing problem. *Operations Research* 56(2), 425–436 (2008)
2. Coffmann, J.E.G., Garey, M.R., Johnson, D.S.: Approximation algorithms for NP-hard problems: Approximation algorithms for bin-packing-a survey, chap. 2, pp. 46–96. PWS Publishing, Boston (1997)
3. Delorme, M., Iori, M., Martello, S.: Bin packing and cutting stock problems: Mathematical models and exact algorithms. *European Journal of Operational Research* 255, 1–20 (2016)
4. Falkenauer, E.: A hybrid grouping genetic algorithm. *Journal of heuristics* 2, 5–30 (1996)
5. Gendreau, M., Laporte, G., Semet, F.: Heuristics and lower bounds for the bin packing problem with conflicts. *Computers & Operations Research* 31, 347–358 (2004)
6. Haouari, M., Serairi, M.: Relaxations and exact solution of the variable sized bin packing problem. *Computational Optimization and Application* 48, 345–368 (2011)
7. Jansen, K.: An approximation scheme for bin packing with conflicts. *Journal of Combinatorial Optimization* 3, 363–377 (1999)
8. Johnson, D.S.: Approximation algorithms for combinatorial problems. *Journal of Computer and System Sciences* 9, 256–278 (1974)
9. Kalfakakou, R., Katsavounis, S., Tsouros, K.: Minimum number of warehouses for storing simultaneously compatible products. *International Journal of Production Economics* 81–82, 559–564 (2003)
10. Khanafer, A., Clautiaux, F., Talbi, E.: New lower bounds for bin packing problems with conflicts. *European Journal of Operational Research* 206, 281–288 (2010)
11. Leung, J.Y.T., Dror, M., Young, G.H.: A note on an open-end bin packing problem. *Journal of Scheduling* 4(4), 201–207 (2001)
12. Maiza, M., Labeled, A., Radjef, M.S.: Efficient algorithms for the offline variable sized bin-packing problem. *Journal of Global Optimization* 57(3), 1025–1038 (2013)
13. Maiza, M., Radjef, M.S.: Heuristics for solving the bin-packing problem with conflicts. *Applied Mathematical Sciences* 5(35), 1739 – 1752 (2011)

14. Martello, S., Toth, P.: Knapsack Problems: Algorithms and Computer Implementations. Chichester: John Wiley & Sons (1990)
15. Martello, S., Toth, P.: Lower bounds and reduction procedures for the bin packing problem. *Discrete Applied Mathematics* 28, 59–70 (1990)
16. Muritiba, A.E.F., Iori, M., Malaguti, E., Toth, P.: Algorithms for the bin packing problem with conflicts. *INFORMS Journal on Computing* 22, 401– 415 (2010)
17. Ongkunaruk, P.: Asymptotic Worst-Case Analyses for the Open Bin Packing Problem. Ph.D. thesis, Virginia Polytechnic Institute and State University (2005)
18. Sadykov, R., Vanderbeck, F.: Bin packing with conflicts: A generic branch-and-price algorithm. *INFORMS Journal on Computing* 25(2), 244–255 (2013)
19. Yang, J., Leung, J.Y.T.: The ordered open-end bin-packing problem. *Operations Research* 51(5), 759–770 (2003)

Moore bipartite mixed Graphs

Ouiza Imine¹, Méziane Aïder²

¹ Université Bouira, Algeria

² LaROMaD, Fac. Maths, USTHB, PB 32, 16111 Bab Ezzouar, Algeria.

¹ imineouiza@yahoo.fr, ² m-aider@usthb.dz

Abstract. The degree/diameter problem consists to find and to study the largest graphs of given maximum degree and given diameter. An upper bound for the order of a graph of given maximum degree and given diameter is known as a Moore bound and a graph attaining this bound is said a Moore graph. The Moore bound is generally determined by counting the number of vertices at any possible distance from a given vertex. This problem has been studied for different cases of graphs, including mixed (or partially directed) graphs (graphs that have both arcs and edges).

In this paper, we consider the case of bipartite mixed graphs and derive a like-Moore bound. Then we prove that mixed bipartite Moore graphs (mixed bipartite graphs with order attaining this like-Moore bound) do not exist for diameter 2 but exist for diameter 3 only for graphs such that each vertex is incident with a unique edge.

Keywords: Moore bound, Moore graphs, Mixed bipartite graphs, Mixed bipartite Moore graphs.

1 introduction

The order of an undirected graph (a mixed graph admitting only edges) with maximum degree d and diameter D cannot exceed the undirected Moore bound given by $M_{d,D} = 1 + d + d(d-1) + \dots + d(d-1)^{D-1}$. It is well known that there are no undirected Moore graphs (undirected graphs with order equals the Moore bound) of degree $d \geq 3$ and diameter $D \geq 3$. For $D = 1$ and $d \geq 1$ complete graphs K_{d+1} are the only Moore graphs. For $D \geq 3$ and $d = 2$ the cycles C_{2D+1} are the only Moore graphs. For $D = 2$, apart from C_5 ($d = 2$), Moore graphs exist only when $d = 3$ (the Petersen graph), $d = 7$ (the Hoffman-Singleton graph) and possibly when $d = 57$.

The “bipartite Moore bound”, that is, the maximum number $B_{\Delta,D}$ of vertices in a bipartite graph of maximum degree Δ and diameter at most D , was given by Biggs [3]

$$B_{2,D} = 2D \text{ and } B_{\Delta,D} = \frac{(\Delta-1)^D - 1}{\Delta-2} \text{ if } \Delta > 2.$$

Bipartite graphs satisfying the equality are called bipartite Moore graphs. For degrees 1 or 2, bipartite Moore graphs are K_2 and the $2D$ -cycles, respectively. When $\Delta > 3$ the possibility of the existence of bipartite Moore graphs was settled by Feit and Higman [6] in 1964 and, independently, by Singleton [9] in 1966. They proved that such graphs exist only if the diameter is 2, 3, 4 or 6.

The order of a digraph (a graph admitting only arcs) with maximum out-degree d and diameter D cannot exceed the directed Moore bound $\vec{M}_{d,D} = 1 + d + d^2 + \dots + d^D$. It was proved [5] that the Moore digraphs do not exist for $d > 1$ and $D > 1$. The unique Moore digraphs are directed cycles of length $D + 1$ and complete graphs on $d + 1$ vertices.

Aïder [1] [2] studied bipartite digraphs with maximum in- and out-degree d (> 1) and diameter D . He showed that the order of such a digraph is at most

$$2 \frac{d^{D+1} - 1}{d^2 - 1} \text{ if } D \text{ is odd;}$$

$$2 \frac{d^{D+1} - d}{d^2 - 1} \text{ if } D \text{ is even.}$$

He then proved that bipartite digraphs of the given order ("bipartite Moore digraphs") exist only for $2 \leq D \leq 4$. Additionally, a variety of properties of such digraphs are established.

In this paper, we consider graphs which are finite and mixed (they may contain both arcs as well as edges). The mixed graphs are also called partially directed graphs. Let v be a vertex of a graph G . Denote by $r(v)$ the number of edges incident from v (i.e. the undirected degree of v) and $z(v)$ the number of arcs incident with v (i.e. the directed degree of v). Denote by $id(v)$ (respectively $od(v)$) the sum of the number of arcs incident to (respectively, from) v and the number of edges incident with v . A mixed graph G is said to be regular of degree d if $od(v) = id(v) = d$ for every vertex v of G . A regular graph G of degree d is said to be totally regular with undirected degree r and directed degree $z = d - r$ if for every pair of vertices $\{u, v\}$ of G we have $r(u) = r(v) = r$. A mixed graph G is said to be a proper mixed graph if G contains at least one arc and at least one edge. We denote by $E(G)$ (respectively $A(G)$) the set of edges (respectively arcs) of G . If $E(G) = \emptyset$ then G is a directed graph, if $A(G) = \emptyset$ then G is an undirected graph.

A mixed graph is said to be a proper mixed graph if $E(G) \neq \emptyset$ and $A(G) \neq \emptyset$. A finite sequence $u_0, e_1, u_1, e_2, \dots, u_{l-1}, e_l, u_l$ of length $l \geq 0$ (where $u = u_0, u_1, u_2, \dots, u_{l-1}, u_l = v$ are vertices of G and $e_1, e_2, \dots, e_l$ are different arcs or edges of G , and if e_j is an arc, then $e_j = (u_{j-1}, u_j)$) is called path if those vertices are all different, and called cycle if those vertices are all different except $u = v$. The distance from u to v is the length of a shortest path from u to v . Note that the distance from u to v in a proper mixed graph can be different from the distance from v to u when a shortest path between u and v involves arcs. The diameter is the maximum value of the distances over all pairs of vertices and is denoted by D . By the definition of a mixed graph, the graph under consideration are special cases.

Nguyen et al. [8] observed that one can also define a like-Moore bound for mixed graphs in the following way. The order of a mixed graph of diameter D , maximum degree d and maximum out-degree z is bounded by:

$$1 + d + zd + r(d-1) + \dots + zd^{D-1} + r(d-1)^{D-1}$$

and the graphs of order reaching this bound are called Moore mixed graphs. The same authors showed that mixed Moore graphs only exist for diameter $D = 2$.

The rest of the paper is structured as follows. In Section 2 we define the Moore bound for a mixed bipartite graph of degree $d = r + z$ and diameter D . In Section 3 we show that Moore bipartite mixed graphs do not exist for diameter 2, but exist for diameter 3 and we give a construction of these graphs in section 4 for $r = 1$. In Sections 5 and 6 we construct almost Moore bipartite mixed graphs for diameter 2 (when $r \geq 2$) and $D = 4$ respectively, since Moore graphs do not exist for these values. Finally, Section 7 presents our conclusions and some open questions.

2 Moore bipartite mixed graphs

Proposition 1. *Let $G = (V, E \cup A)$ be a bipartite mixed graph of degree $d = r + z$ with E the set of edges, A the set of arcs and V the set of vertices.*

1. *If $d > 3$ then*

$$|V| \leq \begin{cases} 2 \left[\frac{d^3 - d + zd^D - zd^2}{d^2 - 1} + r \frac{(d-1)^D - (d-1)^2}{d^2 - 2d} \right] & \text{if } D \text{ is even.} \\ 2 \left[\frac{d^2 - 1 + zd^D - zd}{d^2 - 1} + r \frac{(d-1)^D - (d-1)}{d^2 - 2d} \right] & \text{if } D \text{ is odd.} \end{cases}$$

2. *If $d = 2$ then*

$$|V| \leq \begin{cases} 2 \left[\frac{2^{D+1} + 3D - 2}{3} \right] & \text{if } D \text{ is even.} \\ 2 \left[\frac{2^{D+1} + 3D - 1}{3} \right] & \text{if } D \text{ is odd.} \end{cases}$$

Proof. Let V_1 and V_2 be the bipartition of G , with $|V_1| \geq |V_2|$. We assume that the maximum out-degree d is achieved for the directed degree z and the undirected degree r .

1. If $d > 3$, then D can either be even or odd.

a) If $D = 2p$, let v be any vertex in V_2 . Note that there are at most d vertices at distance 1 from v and for $i = 1, 2, \dots, p-1$, there are at most

$d + zd^{2i} + r(d-1)^{2i}$ vertices at distance $2i+1$ from v . Then the number of vertices in V_1 is

$$\begin{aligned}
 |V_1| &= d + \sum_{i=1}^{p-1} N_{2i+1}^+(v) \leq d + \sum_{i=1}^{p-1} (zd^{2i} + r(d-1)^{2i}) \\
 &= d + z \sum_{i=1}^{p-1} d^{2i} + r \sum_{i=1}^{p-1} (d-1)^{2i} \\
 &= \frac{d^3 - d}{d^2 - 1} + \frac{zd^{2p} - zd^2}{d^2 - 1} + r \frac{(d-1)^{2p} - (d-1)^2}{d^2 - 2d} \\
 &= \frac{d^3 - d + zd^D - zd^2}{d^2 - 1} + r \frac{(d-1)^D - (d-1)^2}{d^2 - 2d}.
 \end{aligned}$$

Since $|V| = |V_1| + |V_2| \leq 2|V_1|$, it follows:

$$|V| \leq 2 \left(\frac{d^3 - d + zd^D - zd^2}{d^2 - 1} + r \frac{(d-1)^D - (d-1)^2}{d^2 - 2d} \right)$$

- b) If $D = 2p+1$, let v be any vertex in the set V_1 . For $i = 1, 2, \dots, p$, there are at most $zd^{2i-1} + r(d-1)^{2i-1}$ vertices at distance $2i$ from v . Then the number of vertices in V_1 is:

$$\begin{aligned}
 |V_1| &= 1 + \sum_{i=1}^{p-1} |N_{2i}^+(v)| \leq 1 + \sum_{i=1}^p (zd^{2i-1} + r(d-1)^{2i-1}) \\
 &= 1 + z \sum_{i=0}^{p-1} d^{2i+1} + r \sum_{i=0}^{p-1} (d-1)^{2i+1} \\
 &= 1 + zd \frac{(d^2)^p - 1}{d^2 - 1} + r(d-1) \frac{(d-1)^{2p} - 1}{(d-1)^2 - 1} \\
 &= 1 + z \frac{(d^{2p+1}) - d}{d^2 - 1} + r \frac{(d-1)^{2p+1} - (d-1)}{d^2 - 2d} \\
 &= \frac{d^2 - 1 + zd^{2p+1} - zd}{d^2 - 1} + r \frac{(d-1)^{2p+1} - (d-1)}{d^2 - 2d} \\
 &= \frac{d^2 - 1 + zd^D - zd}{d^2 - 1} + r \frac{(d-1)^D - (d-1)}{d^2 - 2d}.
 \end{aligned}$$

Since $|V| = |V_1| + |V_2| \leq 2|V_1|$, it follows:

$$|V| \leq 2 \left(\frac{d^2 - d + zd^D - zd^2}{d^2 - 1} + r \frac{(d-1)^D - (d-1)}{d^2 - 2d} \right)$$

2. If $d = 2$, then D can also either be even or odd.

- a) If $D = 2p$, let v a vertex in V_2 . For $i = 1, 2, \dots, p$, there are at most $2 + 2^{2i} + 1$ vertices at distance $2i$ from v . Then the number of vertices

in V_1 is:

$$\begin{aligned} |V_1| &= d + \sum_{i=1}^{p-1} N_{2i+1}^+(v) \leq 2 + \sum_{i=1}^{p-1} 2^{2i} + 1 = 2 + 2^2 \frac{2^{2(p-1)} - 1}{2^2 - 1} + p - 1 \\ &= 2 + 2^2 \frac{2^{2p} - 4}{3} + p + 1 = \frac{2^{2D} - 4}{3} + \frac{D}{2} + 1 \\ &= \frac{(2^{D+1}) - 8 + 3D + 6}{6} = \frac{2^{D+1} + 3D - 2}{6} \end{aligned}$$

Thus, we have $|V| = |V_1| + |V_2| \leq 2|V_1| \leq \frac{2^{D+1} + 3D - 2}{3}$

- b) $D = 2p + 1$; let v be a vertex in V_1 . For $i = 1, 2, \dots, D - 1$, there are $1 + 2^{2i} + 1$ vertices at distance $2i$ from v . Then the number of vertices in V_1 is:

$$\begin{aligned} |V_1| &= 1 + \sum_{i=0}^{p-1} N_{2i}^+(v) \leq 1 + \sum_{i=1}^{p-1} (2^{2i+1} + 1) = 1 + 2 \sum_{i=0}^{p-1} (2^{2i}) + p \\ &= 1 + 2 \frac{2^{2p} - 1}{2^2 - 1} + p = 1 + \frac{2^{2p+1} - 2}{3} + \frac{D - 1}{2} \\ &= \frac{3 + 2^D - 2}{3} + \frac{D - 1}{2} = \frac{2^{D+1} + 3D - 1}{6} \end{aligned}$$

At last $|V| = |V_1| + |V_2| \leq 2|V_1| \leq \frac{2^{D+1} + 3D - 1}{3}$

Definition 1 We denote by $B(z, r, D)$ the bound given by Proposition 1., and call a Moore bipartite mixed graph every bipartite mixed graph with maximum degree $d = z + r$ and diameter D , of order $B(z, r, D)$.

Remark 1 Observe that if $r = 0$ then $B(z, r, D) = B(z, D)$ (the Moore bound of a bipartite digraph) and if $z = 0$ then $B(z, r, D) = B(r, D)$ (the Moore bound of a non oriented bipartite graph), where $d = r + z$.

3 Nonexistence of Moore bipartite mixed graphs with diameter 2

Proposition 2. Moore bipartite mixed graphs of diameter 2 do not exist.

Proof. Let us assume that a Moore bipartite mixed graph G exists for diameter $D = 2$. It must have $2d$ vertices (the Moore bound for mixed bipartite graph of diameter 2). Let $\{U, V\}$ be the bipartition of G . Since G is neither a digraph nor a non oriented graph, it has necessarily an arc (u, v) , where $u \in U$ and $v \in V$. Since the diameter of G is 2, v must be relied to every vertex in U either by an edge or by an arc. It must be so for u and the degree of v will exceed d , a contradiction.

Remark 2 Note that the bound for Moore bipartite mixed graphs of diameter $D = 3$ and degree $d = z + r$ is $B(z, r, 3) = 2(d^2 - r + 1)$. This bound is maximal if $r = 1$.

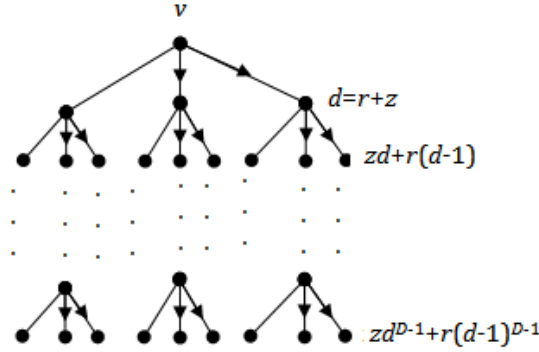


Fig. 1. A spanning subgraph of the bipartite mixed graph

4 Moore bipartite mixed graphs of diameter $D = 3$

We give a construction of Moore bipartite mixed graphs of diameter $D = 3$ based on those of bipartite Moore digraphs developed by Aïder [1, 2].

Let $G(z)$ be a graph of degree $d = z + 1$ and A the set of arcs and denote the vertices of G by $0, 1, \dots, n - 1$, where $n = 2d^2$.

We have an arc (u, v) in the following cases:

- If u is an odd vertex and for each i in $\{0, 1, \dots, d - 1\}$, $v = u + 2id + 1$.
- if u is an even vertex and for each i in $\{0, 1, \dots, d - 1\}$, $v = u + 2i + 1$.

We replace then each pair of arcs with twice orientations by an edge.

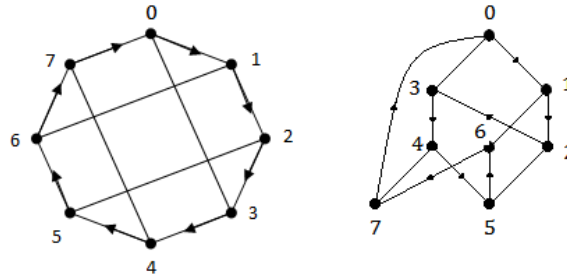


Fig. 2. Moore bipartite mixed graph of degree 2 and diameter 3

Remark 3 It is not difficult to see that $G(z)$ is a regular bipartite mixed graph of degree $d = z + r$, with $r = 1$ and diameter 3, and achieves the Moore bound

for bipartite mixed graph noted. The graph $G(z)$ is then a Moore bipartite mixed graph for every degree $d = z + 1$ and $z \geq 1$.

Remark 4 The graph $G(z)$ is vertex transitive because the image of edge (u, v) is the image of two arcs (u, v) and (v, u) and if u and u' are not of the same parity let f an application such that, if v is even $v = 2j$ then $f(v) = 2jz$ and if v is odd $v = 2j + 1$ then $f(v) = 2jz + 1$.

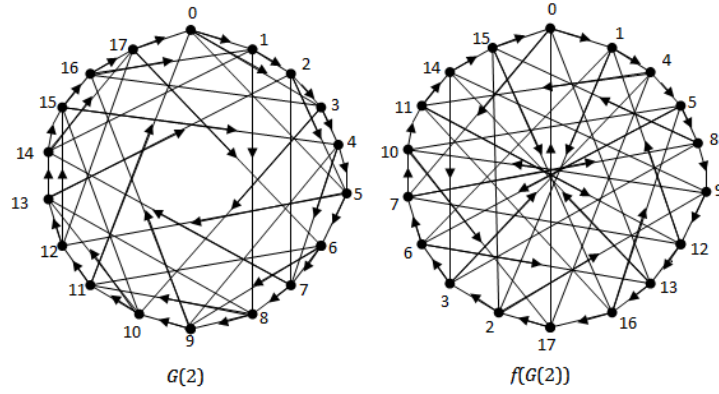


Fig. 3. The image of the graph $G(2)$

In fact the Moore bipartite mixed graphs do not exist for $r \geq 2$, we define the almost bipartite graphs of order $2(d^2 - 1) - 2$ because the number of vertices $n(z, r, 3) \leq 2(d^2 - 1)$, this bound is not affected.

5 Almost Moore bipartite mixed graphs of diameter 2 for $r \geq 2$

Since Moore bipartite mixed graphs of diameter two do not exist for $r \geq 2$, then in this case, the maximum order of a regular bipartite mixed graph of diameter two is $2(d^2 - 1) - 2$.

Here, we give a construction of graphs achieving this new bound. We call them Almost Moore bipartite mixed graphs. This construction is based on the following steps:

1. Number the vertices of G with $0, 1, \dots, n - 1$.
2. There is an arc $(u, v) \in A$ in the following cases:
 - if u is an odd vertex and for each i in $\{0, 1, \dots, d - 1\}$, $v = u + 2id - 1 \pmod{(2d^2 - 2)}$.

- if u is an even vertex and for each i in $\{0, 1, \dots, d-1\}$, $v = u + 2i + 1 \mod (2d^2 - 2)$.
- 3. We replace the arcs with twice orientation by edges (Fig. 4).

Remark 5 The graph G is vertex transitive. To see this, it suffices to consider the function f defined from G to G by:

$$\begin{aligned} f(u) &= 2jd \mod (2d^2 - 2) & \text{if } u = 2j \\ f(u) &= 2jd + 1 \mod (2d^2 - 2) & \text{if } u = 2j + 1 \end{aligned}$$

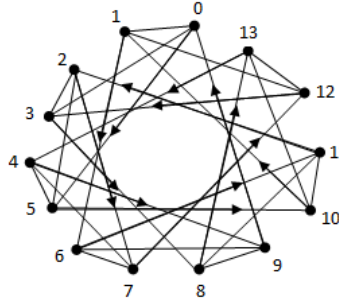


Fig. 4. Almost Moore bipartite mixed graph of degree 3 and diameter 3

6 Almost Moore bipartite mixed graphs for diameter 4

We do not know if Moore bipartite mixed graphs of diameter 4 do exist. If the answer is “no”, then the maximum order of a regular bipartite mixed graph of diameter 4 would be $2(d^3 - d^2 + (z - r + 1)d + r) - 2$.

Here, we give a construction of graphs achieving this new bound. We call them Almost Moore bipartite mixed graphs. This construction is based on the following steps:

1. Number the vertices of G with $0, 1, \dots, n-1$.
2. An arc $(u, v) \in A$ if and only if:
 - if u is an odd vertex and for each i in $\{0, 1, \dots, d-1\}$, $v = u + 3(d+1)i + 2 \mod (2(d^3 - d^2 + (z - r + 1)d + r) - 2)$.
 - if u is an even vertex and for each i in $\{0, 1, \dots, d-1\}$, $v = u + 3(d+1)i + 1 \mod (2(d^3 - d^2 + (z - r + 1)d + r) - 2)$.
3. we replace the arcs with twice orientation by edges.

The graph G is vertex transitive, it suffices to consider the function f defined by G in G by:

$$\begin{aligned} f(u) &= 2jz \bmod (2(d^3 - d^2 + (z - r + 1)d + r) - 2)) & \text{if } u = 2j \\ f(u) &= 2jz + 1 \bmod (2(d^3 - d^2 + (z - r + 1)d + r) - 2)) & \text{if } u = 2j + 1 \end{aligned}$$

The graph in Fig. 5 is an almost bipartite mixed graph of maximum degree $d = 2$ and diameter 4.

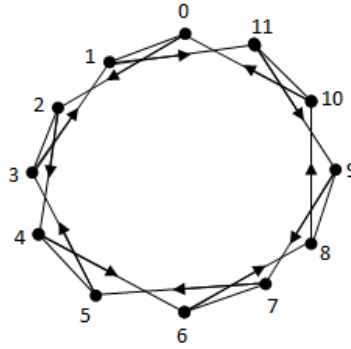


Fig. 5. Almost bipartite mixed Moore graph of degree 2 and diameter 4

7 Conclusion

In this paper, we have proved the non-existence of bipartite mixed Moore graphs of diameter $D = 2$. For diameter $D = 3$, we have proved the existence of Moore bipartite mixed graphs and constructed bipartite Moore graphs of degree $d = z + 1$ (undirected degree $r = 1$ and directed degree $z \geq 1$). The perspective in this area is to study the existence or the non-existence of bipartite Moore mixed graphs for $D \geq 4$ and to construct them if exist. It is also interesting to try to construct almost Moore bipartite mixed graphs (bipartite mixed graphs achieving an upper bound on the order) in the case of the non-existence.

References

1. M. Aïder, Réseaux d'interconnexion bipartis orientés colorations généralisées dans les graphes, Doctorat Thesis, University of Grenoble, France, 1987.
2. M. Aïder, Réseaux d'interconnexion bipartis orientés, Rev. Maghrébine Math. 1(1) (1992) 79-92.
3. N. I. Biggs, Girth, valency and excess, Linear Algebra Appl. 31 (1980) 55-59.

4. J. Bosak, Partially directed Moore graphs, *Math. Slovaca* 29(2)(1979) 181-196).
5. W. G. Bridges and S. Toueg, On the impossibility of directed Moore graphs, *J. Combinatorial Theory B* 29 (1980) 339-341.
6. W. Feit and G. Higman, The non-existence of certain generalized polygons, *J. Algebra* 1 (1964) 114-131.
7. M. Miller, J. Širáň, Moore graphs and beyond: A survey of the degree/diameter problem, *Electron. J. Combin.* 20(2), DS14v2(2013) 1-91.
8. M. H. Nguyen, M. Miller and J. Gimbert, On mixed Moore graphs, *Discrete Mathematics* 307 (2007) 964-970.
9. R. C. Singleton, On minimal graphs of maximum even girth, *Journal of Combinatorial Theory* 1 (3) (1966) 306-332.

Méthode hybride pour le problème de contrôle optimal en temps minimal

Mohand Ouamer BIBI¹ et Imane ALIOUA¹

¹ Unité de Recherche LaMOS (Modélisation et Optimisation des Systèmes),
Département de Recherche Opérationnelle, Université de Béjaia 06000, Algérie.
{mobibi.dz}@gmail.com
{alioua_imate}@yahoo.com

Résumé

Certains problèmes de contrôle optimal se modélisent avec un temps terminal fixé à l'avance, tandis que dans d'autres le temps terminal n'est pas fixé et il sera donc considéré comme une variable inconnue à trouver. On parle alors d'un problème de contrôle optimal avec un temps terminal libre.

Dans ce papier, on a présenté une nouvelle méthode hybride pour résoudre le problème de contrôle optimal en temps minimal. Cette méthode combine une méthode directe basée sur la méthode de support de programmation quadratique, et une procédure finale indirecte, s'apparentant à la méthode de tir ou shooting method.

Mots-clés : *contrôle optimal, temps minimal, principe du maximum, méthode directe de support, méthode de tir.*

1 Introduction

De nos jours, les systèmes dynamiques font complètement partie de notre quotidien, ayant pour but d'améliorer notre qualité de vie et de faciliter certaines tâches : contrôle des flux routiers [2], des réseaux informatiques [1], des systèmes mécaniques [4,5,17,30] et biologiques [9], etc.

Les applications en contrôle optimal concernent tout système dynamique sur lequel on peut agir au moyen d'une commande, avec une notion de rendement optimal. Le but du contrôle optimal est alors de faire évoluer l'état d'un système dynamique de manière à le ramener à un état désiré, tout en minimisant un certain critère donné (énergétique [25,27], économique [28,14], temporel [24,13,15,23]).

Certains problèmes de contrôle optimal se modélisent avec un temps terminal fixé à l'avance, tandis que dans d'autres le temps terminal n'est pas fixé et il sera donc considéré comme une variable inconnue à trouver. On parle alors d'un problème de contrôle optimal en temps terminal libre [18,19]. Dans ce cas, il s'agit de déterminer à la fois le contrôle et le temps terminal de manière à minimiser un critère choisi à l'avance, comme en particulier le problème de transfert en temps minimal [23].

Dans ce travail, on présente une nouvelle méthode hybride pour résoudre le problème de contrôle optimal en temps minimal. Cette méthode combine une méthode directe basée sur la méthode de support de programmation quadratique, et une procédure finale indirecte, s'apparentant à la méthode de tir ou shooting method.

Après avoir approximé un problème auxiliaire de contrôle optimal à temps terminal libre avec un problème linéaire-quadratique à temps terminal fixé, on résout ce dernier par la méthode de support de programmation quadratique, et ce, sans discrétisation au préalable du système dynamique linéaire. On obtient ainsi une idée assez précise de la structure des commutations avec une approximation du temps minimal et de l'état terminal nul. La deuxième étape utilise ensuite une procédure finale, genre méthode de tir, pour calculer le temps minimal d'une manière plus précise. Cette dernière consiste à résoudre un système d'équations à l'aide de la méthode de Newton, en commençant l'itération initiale par les approximations obtenues lors de la première étape.

2 Etude du problème de contrôle optimal en temps minimal

Le problème de temps minimal consiste à transférer d'une manière optimale un point matériel d'un état initial $x(0) = x^0$ vers un état final $x(t_1) = x^1 = 0$. La donnée de ce temps minimal est très importante puisqu'elle indique simplement qu'en un temps inférieur le transfert n'est pas possible.

2.1 Position du problème

Considérons le problème de contrôle optimal suivant :

$$J(u_1, t_1) = t_1 \rightarrow \min, \quad (1)$$

$$\dot{x}(t) = Ax(t) + bu_1(t), \quad x(0) = x^0, \quad (2)$$

$$x(t_1) = 0, \quad (3)$$

$$u_1(t) \in U = [-L, L], t \geq 0, \quad (4)$$

où $\dot{x}(t) = \frac{dx}{dt}$, $x(t) = (x_1(t), x_2(t), \dots, x_n(t))'$ est un vecteur de dimension n qui représente l'état du système à l'instant t , $x(0) = x^0 \neq 0$ étant l'état initial du système; $u_1(t)$ est une fonction scalaire continue par morceaux à valeurs dans U , appelée commande du système; A représente une matrice carrée d'ordre n , qui caractérise le système; b est un vecteur de dimension n et L une constante positive.

La fonctionnelle $J(u_1, t_1)$ est appelée critère de qualité de type *Mayer*, où le temps terminal t_1 est non fixé, considéré comme une variable inconnue à trouver. Le symbole $(\cdot)$ représente celui de la transposition. En outre, on suppose que le système (2) est contrôlable.

Définition 1 La commande $u_1(t)$, $t \in T = [0, t_1]$, est dite commande admissible du problème (1)-(4) si :

- (i) elle est continue par morceaux sur T et continue à droite en ses points de discontinuité :

$$\lim_{t \rightarrow t_j, t > t_j} u_1(t) = u_1(t_j + 0) = u_1(t_j), \quad j = \overline{1, s};$$

- (ii) $-L \leq u_1(t) \leq L$, $t \in T$;

- (iii) la trajectoire engendrée vérifie : $x(t_1) = 0$.

Définition 2 Une commande admissible $u_1^0(t)$, $t \in [0, t_1^0]$, est dite optimale si elle réalise le minimum du critère de qualité :

$$J(u_1^0, t_1^0) = \min_{u_1, t_1} J(u_1, t_1), \quad (5)$$

où u_1 parcourt l'ensemble de toutes les commandes admissibles et $t_1 > 0$. La trajectoire correspondante $x^0(t)$, $t \in [0, t_1^0]$ est dite trajectoire optimale.

En outre, on appelle commande suboptimale (ou ϵ -optimale), toute commande admissible $u_1^\epsilon(t)$, $t \in [0, t_1^\epsilon]$, satisfaisant l'inégalité

$$J(u_1^\epsilon, t_1^\epsilon) - J(u_1^0, t_1^0) \leq \epsilon, \quad (6)$$

où le couple (u_1^0, t_1^0) est optimal, et ϵ est un nombre positif ou nul choisi à l'avance.

Théorème 1 (Principe du Maximum [24])

Pour que l'instant t_1^0 et le contrôle admissible $u_1^0(t)$, $t \in T = [0, t_1^0]$, réalisent le minimum dans le problème (1)-(4), il est nécessaire et suffisant que les conditions suivantes soient vérifiées :

- (i) $H(x^0(t), \Psi^0(t), u_1^0(t)) = \max_{v \in U} H(x^0(t), \Psi^0(t), v)$, $t \in T = [0, t_1^0]$;
(ii) $H(x^0(t_1^0), \Psi^0(t_1^0), u_1^0(t_1^0)) = 0$,

où $x^0(t)$ est la trajectoire du système direct (2) correspondant à $u_1^0(t)$, $\Psi^0(t)$ étant une solution non identiquement nulle du système conjugué :

$$\dot{\Psi} = -A'\Psi = -\frac{\partial H}{\partial x}, \quad H(x, \Psi, u_1) = \Psi'(Ax + bu_1). \quad (7)$$

Pour résoudre numériquement ce problème, plusieurs méthodes ont été proposées, utilisant notamment le principe du maximum, la programmation dynamique, ainsi que d'autres approches [6,8,16,21,20,22]. Malheureusement, les algorithmes élaborés ne sont ni stables ni très efficaces [7,13].

Dans ce travail, on présente une nouvelle approche constructive qui consiste à considérer un problème auxiliaire pour approximer le temps minimal et l'état terminal nul. Ensuite, on applique une procédure finale, genre de méthode de tir, pour calculer plus précisément la solution optimale.

3 Etude du problème auxiliaire avec un temps terminal libre

3.1 Position du problème et lien avec le problème original

Considérons le problème auxiliaire de contrôle optimal suivant :

$$J_a(u_1, t_1) = \frac{1}{2} \|x(t_1)\|^2 + t_1 \rightarrow \min, \quad (8)$$

$$\dot{x}(t) = Ax(t) + bu_1(t), \quad x(0) = x^0, \quad (9)$$

$$u_1(t) \in U = [-L, L], t \in T = [0, t_1]. \quad (10)$$

Une commande admissible du problème (8)-(10) est une commande qui vérifie les conditions (i) et (ii) de la définition 1.

Ici le problème d'optimisation consiste à chercher un temps $\hat{t}_1 > 0$, et une commande admissible $\hat{u}_1(t)$, $t \in [0, \hat{t}_1]$, réalisant le minimum de la fonctionnelle (8). On a le lemme suivant :

Lemme 1 Soient $\hat{t}_1$ et $\hat{u}_1(t)$, $t \in [0, \hat{t}_1]$, un couple optimal du problème (8)-(10). Deux cas peuvent éventuellement se présenter :

- a) Si $\hat{x}(\hat{t}_1) = 0$, alors $\hat{t}_1$ est le temps minimal du problème (1)-(4).
- b) Si $\hat{x}(\hat{t}_1) \neq 0$, alors le couple $(\hat{u}_1, \hat{t}_1)$ n'est pas admissible dans le problème (1)-(4). On a alors :

$$t_1^0 \geq \frac{1}{2} \|\hat{x}(\hat{t}_1)\|^2 + \hat{t}_1, \quad (11)$$

où t_1^0 est le temps minimal du problème (1)-(4).

Preuve 1

- a) Supposons que $\hat{t}_1$ n'est pas le temps minimal dans le problème (1)-(4). Il existe alors un temps $\bar{t}_1 > 0$ et une commande $\bar{u}_1$ tels que $\bar{x}(\bar{t}_1) = 0$ et $\bar{t}_1 < \hat{t}_1$. Donc

$$J_a(\bar{u}_1, \bar{t}_1) = \frac{1}{2} \|\bar{x}(\bar{t}_1)\|^2 + \bar{t}_1 < \frac{1}{2} \|\hat{x}(\hat{t}_1)\|^2 + \hat{t}_1,$$

ce qui contredit l'optimalité de $(\hat{u}_1, \hat{t}_1)$ dans le problème (8)-(10). Par conséquent, $\hat{t}_1$ est le temps minimal du problème (1)-(4).

- b) Soit un couple (u_1^0, t_1^0) tel que $x^0(t_1^0) = 0$, solution du problème (1)-(4). En vertu de l'optimalité du couple $(\hat{u}_1, \hat{t}_1)$ dans le problème (8)-(10), on aura alors :

$$J_a(u_1^0, t_1^0) = \frac{1}{2} \|x^0(t_1^0)\|^2 + t_1^0 \geq \frac{1}{2} \|\hat{x}(\hat{t}_1)\|^2 + \hat{t}_1.$$

Comme $x^0(t_1^0) = 0$, on obtient donc

$$t_1^0 \geq \frac{1}{2} \|\hat{x}(\hat{t}_1)\|^2 + \hat{t}_1. \square$$

3.2 Principe du Maximum du problème de contrôle optimal avec un temps terminal libre

Soient $\hat{t}_1$ et $\hat{u}_1(t)$, $t \in [0, \hat{t}_1]$, une solution optimale du problème (8)-(10). Pour un tel temps fixé $\hat{t}_1$, la commande $\hat{u}_1(t)$, $t \in [0, \hat{t}_1]$, est une solution optimale du problème suivant :

$$J_a(u) = \frac{1}{2} \|x(\hat{t}_1)\|^2 + \hat{t}_1 \rightarrow \min_u, \quad (12)$$

$$\dot{x}(t) = Ax(t) + bu_1(t), x(0) = x^0, \quad (13)$$

$$|u_1(t)| \leq L, t \in T = [0, \hat{t}_1]. \quad (14)$$

En effet, si l'on suppose qu'il existe une autre commande admissible $\tilde{u}_1(t)$, $t \in T = [0, \hat{t}_1]$, telle que $J_a(\tilde{u}_1) < J_a(\hat{u}_1)$, alors cette dernière inégalité contredirait l'optimalité du couple $(\hat{u}_1, \hat{t}_1)$ dans le problème (8)-(10). Donc pour la commande $\hat{u}_1$, le principe du maximum suivant est vérifié pour le temps terminal fixé $\hat{t}_1$:

Théorème 2 (*Principe du Maximum [11]*)

La commande $\hat{u}_1(t)$, $t \in T = [0, \hat{t}_1]$, est optimale dans le problème (12)-(14) si et seulement si, le long de $\hat{u}_1(t)$ et des trajectoires correspondantes $\hat{x}(t)$ et $\hat{\Psi}(t)$ du système direct (13) et du système conjugué :

$$\dot{\Psi} = -A'\Psi, \quad \Psi(\hat{t}_1) = -x(\hat{t}_1), \quad (15)$$

le hamiltonien $H(x, \Psi, u_1) = \Psi'(Ax + bu_1)$ atteint son maximum :

$$H(\hat{x}(t), \hat{\Psi}(t), \hat{u}_1(t)) = \max_{v \in U} H(\hat{x}(t), \hat{\Psi}(t), v), \quad t \in T = [0, \hat{t}_1]. \quad (16)$$

Pour trouver la condition que doit satisfaire l'instant optimal $\hat{t}_1$ dans le problème (8)-(10), comparons le processus optimal engendré par la commande optimale $\hat{u}_1(t)$, $t \in T = [0, \hat{t}_1]$, à deux autres processus dont la durée de l'un est supérieure à la durée optimale $\hat{t}_1$, et l'autre inférieure à $\hat{t}_1$. En effet, considérons tout d'abord la commande :

$$\tilde{u}_1(t) = \begin{cases} \hat{u}_1(t), & t \in [0, \hat{t}_1], \\ \hat{u}_1(\hat{t}_1), & t \in [\hat{t}_1, \hat{t}_1 + \epsilon], \epsilon > 0. \end{cases} \quad (17)$$

De l'équation (13), on obtient la trajectoire correspondante à la commande $\tilde{u}_1(t)$:

$$\begin{cases} \tilde{x}(t) = \hat{x}(t), & t \in [0, \hat{t}_1] \\ \tilde{x}(\hat{t}_1 + \epsilon) = \tilde{x}(\hat{t}_1) + \epsilon \frac{d\tilde{x}(\hat{t}_1)}{dt} + o(\epsilon) = \hat{x}(\hat{t}_1) + \epsilon[A\hat{x}(\hat{t}_1) + b\hat{u}_1(\hat{t}_1)] + o(\epsilon). \end{cases}$$

Par conséquent, on aura :

$$\begin{aligned} 0 \leq \Delta J_a(\hat{u}_1, \hat{t}_1) &= J_a(\tilde{u}_1, \hat{t}_1 + \epsilon) - J_a(\hat{u}_1, \hat{t}_1) \\ &= \epsilon \frac{\partial \varphi}{\partial x}(\hat{x}(\hat{t}_1), \hat{t}_1)[A\hat{x}(\hat{t}_1) + b\hat{u}_1(\hat{t}_1)] + \epsilon \frac{\partial \varphi}{\partial t}(\hat{x}(\hat{t}_1), \hat{t}_1) + o(\epsilon), \end{aligned}$$

où $\varphi(x, t) = \frac{1}{2} \|x\|^2 + t$, avec $\frac{\partial \varphi}{\partial x} = x$ et $\frac{\partial \varphi}{\partial t} = 1$.

D'où

$$\begin{aligned} \Delta J_a(\hat{u}_1, \hat{t}_1) &= J_a(\hat{u}_1, \hat{t}_1 + \epsilon) - J_a(\hat{u}_1, \hat{t}_1) = \epsilon \hat{x}'(\hat{t}_1)[A\hat{x}(\hat{t}_1) + b\hat{u}_1(\hat{t}_1)] + \epsilon + o(\epsilon) \\ &= \epsilon[\hat{x}'(\hat{t}_1)(A\hat{x}(\hat{t}_1) + b\hat{u}_1(\hat{t}_1)) + 1 + \frac{o(\epsilon)}{\epsilon}] \geq 0. \end{aligned}$$

Pour $\epsilon > 0$ assez petit, on en déduit que

$$\hat{x}'(\hat{t}_1)[A\hat{x}(\hat{t}_1) + b\hat{u}_1(\hat{t}_1)] + 1 \geq 0.$$

Comme $\hat{\Psi}(\hat{t}_1) = -\hat{x}(\hat{t}_1)$, on peut alors écrire :

$$H(\hat{x}(\hat{t}_1), \hat{\Psi}(\hat{t}_1), \hat{u}_1(\hat{t}_1)) \leq 1. \quad (18)$$

D'autre part, considérons la commande :

$$\bar{u}_1(t) = \hat{u}_1(t), t \in [0, \hat{t}_1 - \epsilon], \epsilon > 0.$$

La trajectoire correspondante vérifie alors :

$$\bar{x}(t) = \hat{x}(t), t \in [0, \hat{t}_1 - \epsilon].$$

En particulier, on a pour $\epsilon > 0$ assez petit :

$$\bar{x}(\hat{t}_1 - \epsilon) = \hat{x}(\hat{t}_1 - \epsilon) = \hat{x}(\hat{t}_1) - \epsilon \frac{d\hat{x}(\hat{t}_1)}{dt} + o(\epsilon).$$

Par conséquent, on obtient

$$\begin{aligned} \Delta J_a(\hat{u}_1, \hat{t}_1) &= J_a(\bar{u}_1, \hat{t}_1 - \epsilon) - J_a(\hat{u}_1, \hat{t}_1) \\ &= -\epsilon \hat{x}'(\hat{t}_1)[A\hat{x}(\hat{t}_1) + b\hat{u}_1(\hat{t}_1)] - \epsilon + o(\epsilon) \\ &= \epsilon[-\hat{x}'(\hat{t}_1)(A\hat{x}(\hat{t}_1) + b\hat{u}_1(\hat{t}_1)) - 1 + \frac{o(\epsilon)}{\epsilon}] \geq 0, \end{aligned}$$

d'où l'on déduit :

$$H(\hat{x}(\hat{t}_1), \hat{\Psi}(\hat{t}_1), \hat{u}_1(\hat{t}_1)) \geq 1. \quad (19)$$

D'après les inégalités (18) et (19), on conclut que l'instant optimal $\hat{t}_1$ doit vérifier l'égalité suivante :

$$H(\hat{x}(\hat{t}_1), \hat{\Psi}(\hat{t}_1), \hat{u}_1(\hat{t}_1)) = 1. \quad \square \quad (20)$$

On a ainsi démontré le théorème suivant :

Théorème 3 Pour que l'instant $\hat{t}_1$ et le contrôle admissible $\hat{u}_1(t)$, $t \in T = [0, \hat{t}_1]$, réalisent le minimum dans le problème (8)-(10), il est nécessaire et suffisant que les conditions suivantes soient vérifiées :

- (i) $H(\hat{x}(t), \hat{\Psi}(t), \hat{u}_1(t)) = \max_{v \in U} H(\hat{x}(t), \hat{\Psi}(t), v)$, $t \in T = [0, \hat{t}_1]$;
- (ii) $H(\hat{x}(\hat{t}_1), \hat{\Psi}(\hat{t}_1), \hat{u}_1(\hat{t}_1)) = 1$,

où $\hat{x}(t)$ est la trajectoire du système direct (9) correspondant à $\hat{u}_1(t)$ et $\hat{\Psi}(t)$ étant la solution du système conjugué (15).

Remarque 1 Le cas a) du lemme 1 ne peut pas avoir lieu, car la condition terminale (15) aurait donné $H(\hat{x}(\hat{t}_1), \hat{\Psi}(\hat{t}_1), \hat{u}_1(\hat{t}_1)) = 0$, ce qui contredirait la relation (20).

4 Approximation du problème auxiliaire

4.1 Transformation du problème à temps terminal libre en un problème à temps terminal fixé

Pour utiliser les conditions d'optimalité du problème (8)-(10), nous commençons d'abord par transformer la formulation de Mayer à temps terminal libre en une formulation de Mayer à temps terminal fixé.

Dans cette transformation, il s'agit ni plus ni moins que de se débarrasser de la variable $t_1 > 0$ dans le système (8)-(10) et de se ramener à une situation de problème de commande optimale, où l'instant terminal est fixé.

On va traiter la variable inconnue $t_1 > 0$ comme une commande supplémentaire u_2 ; la connaissance de u_1 et de u_2 permet ensuite de reconstruire l'état $x(t)$, $\forall t \in T = [0, t_1]$. Deux observations vont alors guider notre démarche :

- $s = t/t_1$ est une quantité comprise entre 0 et 1, et ce, quelle que soit la valeur de t_1 .
- $u_2 = t_1$ est une constante positive arbitraire.

Cela nous conduit à tout formuler en termes de la nouvelle variable $s \in [0, 1]$.

La relation $t = st_1$ nous conduit à poser :

$$\begin{cases} \tilde{x}(s) = x(st_1), & s \in [0, 1]; \\ \tilde{u}_1(s) = u_1(st_1), & s \in [0, 1]; \\ \tilde{u}_2(s) = t_1 \equiv \text{const}, & s \in [0, 1]. \end{cases} \quad (21)$$

Après ce changement de variables, $\tilde{x}(\cdot)$ et $\tilde{u}_1(\cdot)$ ont les mêmes dimensions que $x(\cdot)$ et $u_1(\cdot)$ respectivement. De plus, on a

$$\tilde{x}(0) = x(0) = x^0 \quad \text{et} \quad \tilde{x}(1) = x(t_1).$$

Soit une nouvelle variable définie par $\tilde{x}_{n+1}(s) = st_1$. Cette variable vérifie donc

$$\dot{\tilde{x}}_{n+1}(s) = t_1 = \tilde{u}_2, \quad s \in [0, 1], \quad \tilde{x}_{n+1}(0) = 0.$$

Pour définir le nouveau problème de commande optimale, on va considérer

$$\begin{aligned} z(s) &= (\tilde{x}(s), \tilde{x}_{n+1}(s)) \in \mathbb{R}^n \times \mathbb{R}, \\ \tilde{u}(s) &= (\tilde{u}_1(s), \tilde{u}_2) \in U \times \mathbb{R}^+. \end{aligned}$$

En augmentant ainsi la dimension du vecteur d'état et celui de la commande, via le changement de variables $t = st_1$, on va reformuler le problème (8)-(10) en un problème de commande optimale équivalent, posé sur $[0, 1]$. Commençons par la fonction critère. Soit la fonction φ définie par

$$\begin{aligned} \varphi : \mathbb{R}^n \times \mathbb{R} &\rightarrow \mathbb{R} \\ z = (\tilde{x}, \tilde{x}_{n+1}) &\rightarrow \varphi(z) = \frac{1}{2} \|\tilde{x}\|^2 + \tilde{x}_{n+1}. \end{aligned}$$

On a donc

$$\varphi(z(1)) = \frac{1}{2} \|\tilde{x}(1)\|^2 + \tilde{x}_{n+1}(1) = \frac{1}{2} \|x(t_1)\|^2 + t_1.$$

Ecrivons le système dynamique en fonction de la nouvelle variable $s \in [0, 1]$:

$$\begin{cases} \dot{\tilde{x}}(s) = \frac{d}{ds}x(st_1) = \dot{x}(st_1)t_1 = [A\tilde{x}(s) + b\tilde{u}_1(s)]\tilde{u}_2, & \tilde{x}(0) = x^0, \\ \dot{\tilde{x}}_{n+1}(s) = \tilde{u}_2, & \tilde{x}_{n+1}(0) = 0. \end{cases}$$

On est ainsi conduit à un problème de *Mayer* à temps terminal fixé suivant :

$$\tilde{J}_a(\tilde{u}) = \frac{1}{2}z'(1)\tilde{D}z(1) + \tilde{c}'z(1) \rightarrow \min, \quad (22)$$

$$\dot{z}(s) = \begin{pmatrix} [A\tilde{x}(s) + b\tilde{u}_1(s)]\tilde{u}_2 \\ \tilde{u}_2 \end{pmatrix}, \quad z(0) = \begin{pmatrix} x^0 \\ 0 \end{pmatrix}, \quad (23)$$

$$\tilde{u}(s) = (\tilde{u}_1(s), \tilde{u}_2) \in U \times \mathbb{R}^+, \quad s \in [0, 1], \quad (24)$$

où

$$\tilde{D} = \begin{pmatrix} I_n & 0 \\ 0 & 0 \end{pmatrix}, \quad \tilde{c} = (0, 1),$$

représentent respectivement une matrice carrée d'ordre $(n+1)$ et un vecteur de dimension $(n+1)$, dont les n premières composantes sont nulles et la dernière vaut 1. Le vecteur $z(s) \in \mathbb{R}^{n+1}$ représente le vecteur d'état à l'instant s , $z(0) = z^0 = (x^0, 0)$ étant l'état initial du système; $\tilde{u}(s) = (\tilde{u}_1(s), \tilde{u}_2)$, $s \in [0, 1]$, est la commande du système, où $\tilde{u}_1(s)$ est une fonction scalaire continue par morceaux, à valeurs dans $U = [-L, L]$ et $\tilde{u}_2$ un paramètre de $\mathbb{R}^+$.

4.2 Linéarisation du système dynamique au voisinage du point $(z, \tilde{u}) = (z^0, \tilde{u}^0)$

Le système dynamique (23) obtenu après transformation du problème (8)-(10) est non linéaire. La commande d'un système non linéaire étant difficile à réaliser, alors une linéarisation autour de l'état initial sera établie (voir [10]), d'autant plus qu'il s'agit uniquement d'approcher dans cette étape l'état terminal nul et le temps minimal. Pour cela, considérons le système d'équations différentielles ordinaires non linéaires de la forme :

$$\frac{dz(s)}{ds} = g(z(s), \tilde{u}(s)) = \begin{pmatrix} [A\tilde{x}(s) + b\tilde{u}_1(s)]\tilde{u}_2 \\ \tilde{u}_2 \end{pmatrix}, \quad z(0) = \begin{pmatrix} x^0 \\ 0 \end{pmatrix}, \quad s \in [0, 1], \quad (25)$$

où la fonction non linéaire g est définie de $\mathbb{R}^{n+1} \times \mathbb{R}^2$ dans $\mathbb{R}^{n+1}$.

Il s'agit alors de trouver une équation différentielle ordinaire linéaire de la forme :

$$\begin{cases} \frac{dz(s)}{ds} = \tilde{A}z(s) + \tilde{B}\tilde{u}(s) + \tilde{r}, \\ z(0) = z^0 = (x^0, 0), \end{cases} \quad (26)$$

approchant l'équation non linéaire (25) aux mêmes conditions initiales, où $\tilde{A}$ représente une matrice carrée d'ordre $(n+1)$, $\tilde{B}$ une matrice de dimension

$(n+1) \times 2$ et $\tilde{r}$ un vecteur constant à déterminer.

En effet, en développant le second membre de l'équation (25) autour du point $(z, \tilde{u}) = (z^0, \tilde{u}^0)$, avec $\tilde{u}^0 = (\tilde{u}_1^0, \tilde{u}_2^0) = (0, 1)$, et en négligeant les termes d'ordre supérieurs à 1, on obtient :

$$\dot{z}_i = g_i(z^0, \tilde{u}^0) + \sum_{j=1}^{n+1} \left[\frac{\partial g_i}{\partial z_j}(z^0, \tilde{u}^0) \cdot (z_j - z_j^0) \right] + \left[\frac{\partial g_i}{\partial \tilde{u}_1}(z^0, \tilde{u}^0) \tilde{u}_1 + \frac{\partial g_i}{\partial \tilde{u}_2}(z^0, \tilde{u}^0) \cdot (\tilde{u}_2 - 1) \right], \quad (27)$$

où $z_i, i = \overline{1, n+1}$ sont les composantes de z et $\tilde{u}_1, \tilde{u}_2$ sont les composantes de la commande. Comme $g_i(z^0, \tilde{u}^0) \neq 0$, il s'ensuit de (27) que

$$\dot{z}_i = \sum_{j=1}^{n+1} \frac{\partial g_i}{\partial z_j}(z^0, \tilde{u}^0) z_j + \sum_{k=1}^2 \frac{\partial g_i}{\partial \tilde{u}_k}(z^0, \tilde{u}^0) \tilde{u}_k + \tilde{r}_i, \quad i = \overline{1, n+1}, \quad (28)$$

où $\tilde{r}_i = g_i(z^0, \tilde{u}^0) - \sum_{j=1}^{n+1} \frac{\partial g_i}{\partial z_j}(z^0, \tilde{u}^0) z_j^0 - \frac{\partial g_i}{\partial \tilde{u}_2}(z^0, \tilde{u}^0)$, représente une constante. Posons :

$$\tilde{A} = \begin{pmatrix} \frac{\partial g_1}{\partial \tilde{x}_1} & \frac{\partial g_1}{\partial \tilde{x}_2} & \cdots & \frac{\partial g_1}{\partial \tilde{x}_{n+1}} \\ \frac{\partial g_2}{\partial \tilde{x}_1} & \frac{\partial g_2}{\partial \tilde{x}_2} & \cdots & \frac{\partial g_2}{\partial \tilde{x}_{n+1}} \\ \cdots & \cdots & \cdots & \cdots \\ \frac{\partial g_{n+1}}{\partial \tilde{x}_1} & \frac{\partial g_{n+1}}{\partial \tilde{x}_2} & \cdots & \frac{\partial g_{n+1}}{\partial \tilde{x}_{n+1}} \end{pmatrix} = \begin{pmatrix} A \tilde{u}_2^0 & 0 \\ 0 & 0 \end{pmatrix} = \begin{pmatrix} A & 0 \\ 0 & 0 \end{pmatrix},$$

$$\tilde{B} = \begin{pmatrix} \frac{\partial g_1}{\partial \tilde{u}_1} & \frac{\partial g_1}{\partial \tilde{u}_2} \\ \frac{\partial g_2}{\partial \tilde{u}_1} & \frac{\partial g_2}{\partial \tilde{u}_2} \\ \cdots & \cdots \\ \frac{\partial g_{n+1}}{\partial \tilde{u}_1} & \frac{\partial g_{n+1}}{\partial \tilde{u}_2} \end{pmatrix} = \begin{pmatrix} b \tilde{u}_2^0 & b \tilde{u}_1^0 \\ 0 & 1 \end{pmatrix} = \begin{pmatrix} b & 0 \\ 0 & 1 \end{pmatrix}, \text{ et } \tilde{r} = \begin{pmatrix} -A x^0 \\ 0 \end{pmatrix},$$

où $\tilde{A}, \tilde{B}$ et $\tilde{r}$ ont pour argument le point $(z^0, \tilde{u}^0)$.

On est ainsi conduit à un problème linéaire quadratique de contrôle optimal avec un temps terminal fixé suivant :

$$\tilde{J}_a(\tilde{u}) = \frac{1}{2} z'(1) \tilde{D} z(1) + \tilde{c}' z(1) \rightarrow \min, \quad (29)$$

$$\dot{z} = \tilde{A} z + \tilde{B} \tilde{u} + \tilde{r}, \quad z(0) = (x^0, 0), \quad (30)$$

$$\tilde{u}(s) = (\tilde{u}_1(s), \tilde{u}_2(s)) \in U \times \mathbb{R}^+, \quad s \in [0, 1]. \quad (31)$$

En résolvant ce problème par la méthode de support [12], on obtient une solution optimale $\tilde{u}^*(s) = (\tilde{u}_1^*(s), \tilde{u}_2^* = t_1^*)$, avec la trajectoire optimale $\tilde{x}^*(s)$, $s \in [0, 1]$.

5 Résolution du problème originel par la procédure finale

La deuxième phase consiste à utiliser une procédure finale, genre de méthode de tir, pour calculer le temps minimal du problème (1)-(4) d'une manière plus précise.

Ici, on considère que le problème (1)-(4) est simple et que la commande optimale $\tilde{u}_1^*(s)$ possède $(n-1)$ points de commutation tels que $0 < s_1^0 < s_2^0 < \dots < s_{n-1}^0 < 1$. En vertu des relations (11) et (21), on pose alors

$$t_1 = \frac{1}{2} \|\tilde{x}^*(1)\|^2 + t_1^*, \quad \tau_i^0 = s_i^0 t_1, \quad i = \overline{1, n-1}, \quad \tau^0 = (\tau_1^0, \tau_2^0, \dots, \tau_{n-1}^0), \quad \text{et} \quad \tau_n^0 = t_1.$$

On forme la matrice supposée être de rang complet :

$$Q_0 = (q(\tau_i^0), \quad i = \overline{1, n-1}), \quad \text{avec} \quad q(t) = F(t)b, \quad t \in T = [0, t_1], \quad (32)$$

où $\dot{F} = -FA$, $F(t_1) = I_n$. On calcule le vecteur des potentiels

$\lambda^0 = (y^0, 1) \in \mathbb{R}^{n-1} \times \mathbb{R}$ tel que $Q_0' \lambda^0 = 0$, et la co-commande $E^0 = \Psi^{0'}(t)b$, où $\Psi^0(t)$, $t \in T = [0, t_1]$, est la solution du système conjugué

$$\dot{\Psi}^0 = -A' \Psi, \quad \Psi(t_1) = \lambda^0.$$

En outre, on calcule la quasi-commande

$$\omega^0(t) = L \text{sign} E^0(t), \quad t \in T,$$

et la trajectoire $\kappa^0(t)$, $t \in T$, correspondant à $\omega^0(t)$, $t \in T$. On suppose de plus que $\dot{E}(\tau_i^0) \neq 0$, $i = \overline{1, n-1}$. La procédure finale consiste alors à trouver la solution $v = (y, \tau, \tau_n) \in \mathbb{R}^{n-1} \times \mathbb{R}^{n-1} \times \mathbb{R}$, $\tau = (\tau_1, \tau_2, \dots, \tau_{n-1})$, $\tau_n = t_1^0$, de $(2n-1)$ équations :

$$\begin{cases} F_1(v) = F_1(y, \tau, \tau_n) = \kappa(\tau, \tau_n) = 0; \\ F_2(v) = F_2(y, \tau, \tau_n) = E(y, \tau, \tau) = 0, \end{cases} \quad (33)$$

où $E(y, \tau, \tau) = (E(y, \tau, \tau_i), \quad i = \overline{1, n-1})$, et $E(y, \tau, t) = (y, 1)' q(t)$, $t \in [0, \tau_n]$, et $\kappa(\tau, t)$, $t \in [0, \tau_n]$, est la trajectoire (2) correspondant à la quasi-commande $\omega(\tau, t)$, $t \in [0, \tau_n]$:

$$\omega(\tau, t) = \begin{cases} -L \text{sign} \dot{E}(\tau_1^0), & \text{si } t \in [0, \tau_1[; \\ -L \text{sign} \dot{E}(\tau_i^0), & \text{si } t \in [\tau_{i-1}, \tau_i[, \quad i = \overline{2, n-1}; \\ L \text{sign} \dot{E}(\tau_{n-1}^0), & \text{si } t \in [\tau_{n-1}, \tau_n]. \end{cases} \quad (34)$$

De (34), on déduit alors le système d'équations :

$$\begin{cases} F_1(y, \tau, \tau_n) = \kappa(\tau, \tau_n) = \kappa^0(t_1) - 2L \sum_{i=1}^{n-1} \text{sign} \dot{E}(\tau_i^0) \int_{\tau_i^0}^{\tau_i} q(t) dt = 0; \\ F_2(y, \tau) = E(y, \tau, \tau) = 0. \end{cases} \quad (35)$$

On résout le système (35) avec la méthode de Newton, en commençant par l'approximation initiale $v^0 = (y^0, \tau^0, \tau_n^0)$. La commande optimale et le temps minimal t_1^0 dans le problème (1)-(4) seront alors de la forme suivante :

$$u_1^0(t) = \omega(\tau, t), \quad t \in T^0 = [0, t_1^0], \quad \text{et} \quad t_1^0 = \tau_n. \quad (36)$$

6 Conclusion

Dans cet article, on a présenté une nouvelle méthode hybride pour résoudre le problème de contrôle optimal en temps minimal, en se focalisant sur l'étude théorique. Dans un prochain travail, on va élaborer un algorithme de résolution, le programmer sous Matlab pour faire une comparaison avec d'autres méthodes, en particulier la méthode utilisée par Gnevko [13].

Références

1. Agarwal, M. : *A Systematic Classification of Neural-Network-Based Control*. IEEE conf. on control. 75–93 (1997).
2. Alvarez, L., Horowitz, R., Li, P. : *Traffic flow control in automated highway systems*. Control Engin. Practice. 7, 1071–1078 (1999).
3. Bibi, M.O. : *Cours de Post-Graduation sur le Contrôle Optimal*. Université de Bejaia (2011).
4. Bonnans, F., Rouchon, P. : *Analyse et commande de systèmes dynamiques*. Edition de l'Ecole Polytechnique, Paris (2006).
5. Boscain, U., Piccoli, B. : *Optimal Synthesis for Control Systems on 2-D Manifolds*. Springer SMAI series. 43 (2004).
6. Eaton, J.H. : *An iterative solution to time-optimal control*. J. Math. anal. and Appl. 5 (2), 329–344 (1962).
7. Fedorenko, R.P. : *Approximate solution of optimal control problems*. Nauka, Moscow (1978).
8. Fodden, E.J., Gilbert, E.G. : *Computational aspects of the time optimal control problem*. Comput. Methods optim. Problems. Academic Press, New York (1964).
9. Fresard, M., Boncoeur, J. : *Controlling the biological invasion of a commercial fishery by a space competitor : a bioeconomic model with reference to the bay of St-Brieuc scallop fishery*. Agricultural and Resource Economics Review. 35(1), 78–97 (2006).
10. Jordan, A.J. : *Linearization of non-linear state equation*. Bulletin of the Polish Academy of Sciences : Technical Sciences. 54 (1) (2006).
11. Gabasov, R., Kirillova, F.M. : *Maximum principle in optimal control theory*. Nauka i Technika, Minsk (1974).
12. Gabasov, R., Kirillova, F.M., Kostyukova, O.I., Raketsky, V.M. : *Constructive Methods Of Optimization. Part 4 : Convex Problems*. University Press, Minsk (1984).
13. Gnevko, S. V. : *Numerical method for solving the linear time optimal control problem*. International Journal of control. 44 (1), 251–258 (1986).
14. Kohle, M. : *Techno-Economic Optimum Sizing of a Stand-Alone Solar Photovoltaic System*. IEEE Transaction on Energy Conversion. 24 (2), 511–519 (2009).
15. Kostyukova, O., Kostina, E. : *Analysis of Properties of the Solutions to Parametric Time-Optimal Problems*. Computational Optimization and Applications. 26, 285–326 (2003).
16. La salle, J.P. : *The time optimal control problem*. Ann. Math. Studies. 5 (45), (1960).

17. Lee, E.B., Markus, L. : *Foundations of optimal control theory*. John Wiley, New York (1967).
18. Louadj, K. : *Résolution de problèmes paramétrés de contrôle optimal*. Thèse de doctorat. Université de Tizi-Ouzou (2012).
19. Maurer, H. : *Second order sufficient conditions for control problems with free final time*. in Proceedings of 3rd European Control Conference. Roma, Italy. 3602–3606 (1995).
20. Moiseev, N.N. : *Numerical Methods in Optimal Systems Theory*. Nauka, Moscow (1971).
21. Neustadt, L.W. : *Synthesizing time optimal control systems*. J. Math. anal. and Appl. 1 (4), 484–492 (1960).
22. Pshenichny, B.N. : *Numerical Method for Time Optimal Linear control Systems*. Journal of Computational Mathematics and Mathematical Physics. 8 (6), 1345–1351 (1968).
23. Poggidini, L., Spadini, M. : *Bang-bang trajectories with a double switching time in the minimum time problem*. ESSAIM : COCV. 22, 688–709 (2016).
24. Pontryagin, L.S., Boltyanskii, V.G., Gamkrelidze, R.V., Mishchenko, E.F. : *Mathematical Theory of optimal processes*. Interscience, NewYork (1962).
25. Protas, B., Wesfreid, J.E. : *Drag force in the openloop control of the cylinder wake in the laminar regime*. Phys. Fluids, 14 (2), 810–826 (2002).
26. Stoer, J., Bulirsch, R. : *Introduction to numerical analysis*. Springer-Verlag, Berlin (1980).
27. Trigui, R., Jeanneret, B., Badin, F. : *Modélisation systémique de véhicules hybrides en vue de la prédiction de leurs performances énergétiques et dynamiques - Construction de la bibliothèque de modèles VEHLIB*. Recherche Transports Sécurité, (83), 129–150. ISSN 0761-8980 (2004).
28. Thiaux, Y. : *Optimisation des profils de consommation pour minimiser les coûts économique et énergétique sur cycle de vie des systèmes photovoltaïques autonomes et hybrides*. Thèse d'HDR, Ecole normale supérieure de Cachan, Paris (2010).
29. Hiriart-Urruty, J.B. : *Les Mathématiques du mieux faire. volume 2. La commande optimale pour les débutants*. Editions Ellipses, Paris (2001).
30. Zhang, R., Chen, Y. : *Control of hybrid dynamical systems for electric vehicles*. American Control Conference. Arlington, VA. 2884–2889 (2001).

A Comprehensive Taxonomy for Omics Data Integration and Analysis Methods

Amina Lemsara, Salima Ouadfel, and Mohamed Batouche

Department of Computer Science, NTIC Faculty,
University of Constantine 2-Abdelhamid Mehri, 25000 Constantine, Algeria
`amina.lemsara@univ-constantine2.dz`
`ouadfel@gmail.com`
`mohamed.batouche@univ-constantine2.dz`

Abstract. Thanks to the advances in omics technologies, high dimensional data from different molecular levels are produced making possible integrative analysis of multi-omics data related to the same biological system. Using methods that integrate multi-omics datasets can help uncover significant biological interactions and common molecular structure promoting our understanding of complex systems such as diseases. Several methods for omics data integration have been proposed in the literature. In this paper, we present a review of these methods along with a taxonomy. This latter takes into account classifications from several existing reviews and introduces a structured framework by analyzing methods from three main views that are the biological issue of interest, characteristics of the used data and the technical solution proposed for an integrative analysis.

Keywords: omics data, integrative analysis, data fusion, biologic system

1 Introduction

As a consequence of the high throughput of biological data, different types of omics data are now available. The variety of data refers especially to the molecular levels from where data are generated including genomic, transcriptomic, epigenomic and proteomic levels.

To decipher available omics data, the scientific community has adopted a new paradigm in its way to deal with data. It is to integrate datasets from different sources in order to extract new knowledge that the individual processing methods (i.e., process each dataset separately) are unable to explore. Hence, integrating data from different molecular sources can help uncover significant biological interactions and identify common structure within samples allowing us to understand more about these systems. To this end, different integrative methods have been proposed. Such methods are scattered in the literature with different purposes and formalisms. Many reviews of such methods have been introduced in the literature [1–5]. They analyze integrative solutions from different

aspects. For example, *Bersanelli et al.* in [1] analyze integrative methods focusing on the technical aspect. In [2], *Gligorijevic et al.* focus on computational methods dealing with big data in precision medicine. In [3], *Richmondson et al.* give an overall view of existing integrative studies focusing on the motivation and rationale aspects. Other reviews focus on a special biological problem such as [4] where *DeKeersmaker et al.* focus on methods for prokaryotes network reconstruction. *Wei* in [5] focuses on the analysis of cancer data. *Wang* and *Gu* in [6] cover integrative methods for cancer classification. Even though different reviews related to omics data integration have been addressed, an overall review describing integrative methods from both biological and technical aspects are needed in order to bring more clarity to community scientists.

In this paper, we aim to introduce a generic and a systematic view of existing methods related to the integrative analysis of multi-omics data. Thus, we propose a taxonomy of three main criteria that are the biological issue investigated, the used data and the proposed technique. According to this taxonomy, we analyze several related works.

The rest of the paper is organized as follows. Section 2 reveals the need for omics data integration. In section 3, a description of the proposed taxonomy is given. In section 4, we present an analysis of relevant works related to omics data integration topic. Finally, a conclusion and future works are drawn in section 5.

2 The need for omics data integration

On the genomic scale, integrative analysis refers to the situation where multiple omics datasets are analyzed simultaneously in order to gain a better understanding of biological systems. Data can be of different molecular levels (Genomic, transcriptomic, epigenomic, etc.) representing the same samples. This type of integration is called vertical integration. The Cancer Genome Atlas (TCGA) tumor data provide an easy way to achieve such integration. Another type of integration referred to as horizontal integration. It is achieved using the same omic level of data, but derived from different experimental platforms (different technologies and experimental conditions) [3].

Biologic systems such as diseases are known to be complex systems where analysis of one biological component is insufficient to draw a complete view. Cancer, for example, is assumed nowadays as a complex disease of dysregulated pathways where different molecular levels are affected, rather than simple alterations of genes [7]. Thus, several works that treated cancer in an integrative way have proven to be better in describing complex traits.

Moreover, biologic interactions between omics data have motivated the idea of integration. For example, the relation between gene expression and DNA methylation has been investigated in many studies [8, 9] that emphasize the implication of DNA methylation in the gene regulation process. The somatic mutation is assumed to affect interacted gene [10].

On the other side, the goal behind omics data integration is to bridge the gaps between the ability to generate data and the capacity to interpret them

in order to improve our understanding of biological systems. Moreover, many platforms share the same goal. Thus, integrating data across platforms may provide enriched analyses of the same issue.

3 The proposed taxonomy for omics data integration methods

Integration analysis problem has been treated in different ways replying to different biological issues. We present, in this section, the proposed taxonomy that provides a comprehensive view of omics data integration research works.

We consider three main criteria to classify research works; the biological issue investigated, the used data and the proposed technique, as shown in figure 1.

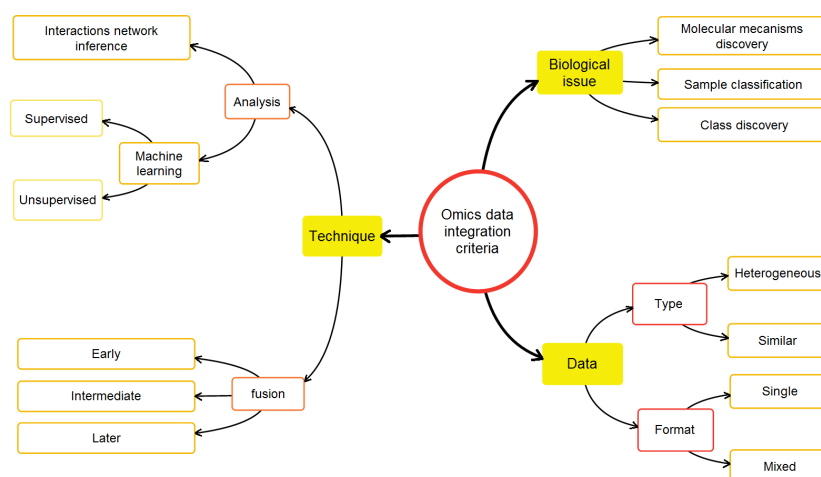


Fig. 1: The proposed taxonomy for omics data integration methods including the three main criteria; the biological issue, the used data and the proposed technique.

The first criterion refers to the biological problem treated that can be divided into three categories:

- Discovery of molecular mechanisms: E.g. interactions between proteins and Gene expression.
- Sample classification: E.g. functional classification of genes.
- Class discovery: E.g. disease subtyping.

The second criterion refers to the used data. Data are described under two sub-criteria: the type; heterogeneous or similar data and the format; single or mixed format.

- Similar data: data of one omic level, obtained from different platforms (technologies, conditions and/or different organisms).
- Heterogeneous data: data of different omic levels describing the same samples.
- Single format: data of the same format that can be numeric, categorical, binary or network.
- Mixed format: data of different formats are used.

The third criterion reflects the method formalism proposed. Each method is discussed under two axes; the fusion level and the type of analysis.

- Analysis allows making predictions on data. We distinct two main types of analysis methods, interaction network inference and machine learning.
 - Interaction network inference, this kind of methods is widely used in biology. It allows to qualitatively representing an interaction network as a graph whose nodes are the molecular components (proteins, genes, transcription factors, etc.) and arcs represent interactions between these components.
 - Machine learning methods, that allow analyzing data and extracting knowledge hidden in them in order to make predictions. They are split into two main categories; supervised learning methods that require a target variable (dependent variable) which depends, causally, on other variables (independent variable) and unsupervised learning methods whereby variables are processed in the same way without need for a target variable.
- Fusion is the stage where integration of multiple datasets is done. Basically, there are three main levels of omics data fusion [11].
 - Late fusion that involves an independent processing of each dataset according to the characteristics of data. Statistical results of various processing are then combined to provide a global result.
 - Early fusion that combines different datasets in one large input before analyzing them.
 - Intermediate fusion where data from different sources are converted into an intermediate form (e.g. kernel matrix) before combining and analyzing them.

4 Analysis of omics data integration methods

In this section, we analyze several research studies related to multi-omics datasets integration according to the taxonomy proposed above (Figure 1). Thus, the analysis involves the following three aspects; the biological issue addressed, characteristics of the used data and the technical solution proposed to tackle the issue involved.

4.1 Biological issue aspect

Modeling interactions between molecular levels improves our understanding of biological mechanisms including causal relationship. In this regard, many research works are carried out in order to infer interaction networks from multiple omics datasets. Among methods related to this topic, there are PARADIGM proposed by *Vaske et al.* in [12] and the method proposed by *Mosca and Milanesi* in [13]. PARADIGM attempts to infer molecular pathways altered in patient samples using different molecular datasets. While the method proposed in [13] aims to analyze multiple omics data as well as interactions between them.

Sample classification is another question that has been treated using multiple omics datasets in order to uncover genotype-phenotype associations. *Lanckriet et al.* [14], *Jiang et al.* [15] and *Cun and Frohlich* [16] have addressed this issue for different purposes. *Lanckriet et al.* in [14] propose a method to predict functions of yeast proteins. *Jiang et al.* in [15] and *Cun and Frohlich* in [16] attempt to identify molecular biomarkers for cancer diagnosis.

On another side, a lot of attention has been devoted to class discovery in biology and precision medicine especially disease subtyping and biomarker discovery. Among methods related to this issue there are iCluster [17], md_ module [18], SNF [19] and MVDA [20]. These methods aim to identify subtypes in cancer without any prior information in [17–19] and including clinical data in [20]. Otherwise, *Bardar et al.* in [21] have treated this issue in order to infer gene groups from their profiles. Figure 2 presents the classification of integrative methods analyzed according to the biological issue aspect.

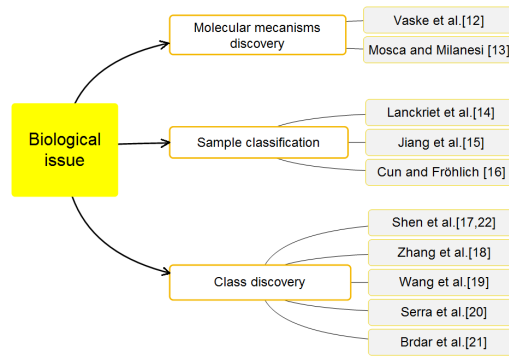


Fig. 2: Classification of omics data integration methods according to the biological issue aspect

4.2 Data aspect

As mentioned previously, based on the type of the used data, integrative methods can be divided into two categories. The first category represents methods dealing with similar data where the data of the same omic type, but across multiple platforms, are integrated as it is addressed by studies carried out in [13, 15, 21]. The second category refers to methods that aim to integrate data of different omics types representing the same samples. This is the case of PARADIGM [12], iCluster [17], md_module [18], SNF [19], MVDA [20] and methods proposed by *Mosca and Milanesi* [13], *Lanchkiet et al.* [14] and *Cun and Fröhlich* [16]. Furthermore, according to the format of data, there are methods which deal with the same format such as iCluster, md_module and methods proposed in [15, 21]. and there are which have been addressed to integrate data of different formats, that is the case of iClusterPlus [22], SNF [19], MVDA [20] and methods introduced in [13, 14, 16]. Figure 3 presents the classification of integrative methods analyzed according to the data aspect.

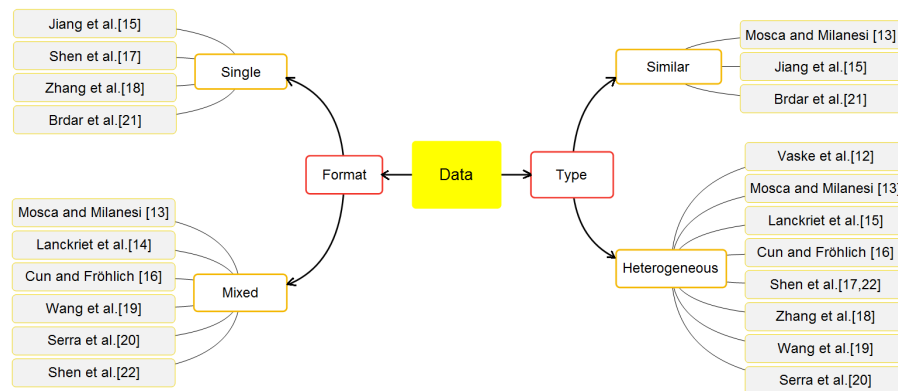


Fig. 3: Classification of omics data integration methods according to the data aspect

4.3 Technical aspect

In response to biological issues using multiple omics data, various integrative methods have been proposed. These methods are either based on statistical, machine learning or network-based computational techniques including a fusion process to combine data. Figure 4 summarizes methods classification according to the computational technique used and the fusion level.

Network inference methods are used in modeling interactions between molecular levels. It is adopted by *Mosca and Milanesi* in [13]. Their proposed method

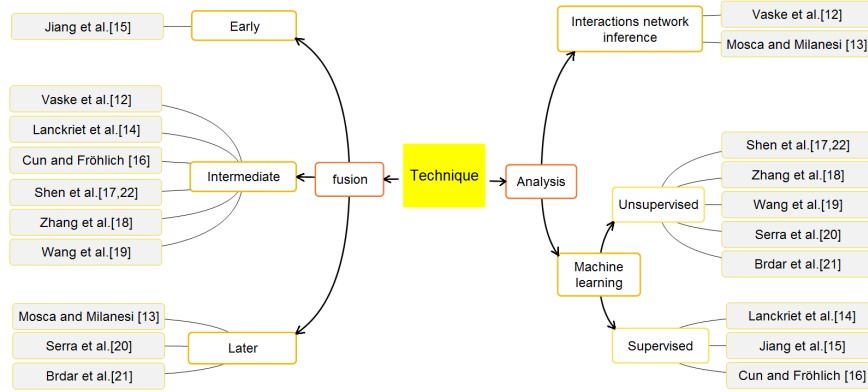


Fig. 4: Classification of omics data integration methods according to the technical aspect

uses multi-objective optimization to identify multiple optimal networks (subnetworks) from a multiple-weighted network. The different omics data are mapped as a multiple-weighted network where nodes represent omics data and edges represent the interaction between two nodes. Significant subnetworks are identified using multi-objective optimization according to several criteria such as biological component-level and topological properties. Finally, topological and biological characterization of the optimal networks generated and their quality evaluation are carried out.

In [12], to infer molecular pathways altered in patient samples, *Vaske et al.* propose PARADIGM method. The method involves two main steps; firstly modeling gene activity by a factor graph (Bayesian network), variables in the graph represent the state of different molecular entities such as mRNA expressions and protein expressions. The directed edges represent regulatory or signaling effects between two entities. Secondly, the inferred pathway activity of samples is defined by a matrix where measurements represent likelihoods of the inferred activity of each entity associated with each sample.

Otherwise, supervised machine learning methods are used to predict functions of genes and proteins in [14] and to identify biomarkers in [15,16].

To predict genes and proteins functions in [14], *Lanckriet et al.* proposed to fuse heterogeneous data (gene, protein, etc.) through a kernel function then performing a classification technique. The process consists in representing each omics dataset as a similarity matrix. Similarity is measured between pairs of samples using a kernel function specific to each dataset. Then, the similarity matrices are combined. The classification is carried out by solving semidefinite program (SDP) problem.

In [15], the proposed integrative method allows merging gene expression data provided by different studies to increase the sample size. The method involves preprocessing and normalization of data. Then, an integration process is performed by transforming data of different distributions to a similar one. Finally, Gene shaving (GS) method based on Random Forests (RF) and Fisher's Linear Discrimination (FLD) were applied to the joint data set for cancer gene selection.

The method proposed in [16], it aims to construct then analyze multi-weighted graph to discover molecular biomarkers. It integrates network information with omics data into one classifier. This is done by smoothing gene-wise t-statistics over a molecular network from the prior knowledge domain (e.g. PPI network) using a random walk kernel as a diffusion network method. Then, a permutation test conducted to select significant features on which a Support Vector Machine (SVM) classifier is trained to predict sample type.

Regarding unsupervised machine learning methods, they are adopted to discover new subtypes in [17–20] through the proposed methods iCluster, md_module, SNF and MVDA. iCluster [17] is an integrative clustering method that is based on a common latent variable model. It consists in representing the subtypes of tumors as latent variables estimated simultaneously. The clustering system based on K-means is represented as a set of Gaussian latent variable models associated with each dataset. Z the latent Gaussian component that is common across all models connects these models and allows integrative clustering of samples. iCluster is limited by the integration of only numeric omics data. To allow integration of different formats, Mo and Shen propose iClusterPlus in [22], a variant of iCluster method that considers different distributions according to the type of data processed.

md_module [18] is a method based on a joint non-negative matrix factorization jNMF. jNMF enables us to simultaneously factor two or more types of genomic data sharing the same set of samples. This method extends the NMF method [23], used for single view clustering, to integrate heterogeneous datasets represented as different views of the same samples. This is done by updating the utility function of the original NMF in the way that allows factoring multiple data matrices, simultaneously, to a common basis matrix and multiple individual matrices. The common matrix representing the shared structure among different datasets allows a fuzzy integrative clustering of samples.

SNF [19] is an alternative method to integrate heterogeneous omics data of different formats. The method involves three main steps; firstly, modeling each dataset as a similarity graph. Graph nodes represent samples and edges represent the similarity between pairs of samples, calculated using scaled exponential similarity kernel measure. Then, merging graphs associated to each dataset in a single one using a non-linear method based on the message-passing theory that updates every network, iteratively, making it more similar to the other. Then, to obtain cluster of samples, a basic clustering method can be applied such as spectral clustering.

MVDA [20] for its part, suggests a pipeline solution to identify sample clusters by integrating heterogeneous omics data of different formats such as gene

expression, RNA sequence and clinical data. The pipeline is based on consensus clustering method that performs a late fusion of data. The pipeline involves firstly preprocessing data in order to reduce the dimension of data using cluster-based correlation analysis, then select significant features using a ranked-based method. For each dataset, Samples were clustered using a basic clustering method. Then, a late integration was performed by concatenating membership matrices of separate clustering. Finally, a clustering method is used to obtain final clusters.

Unsupervised methods are also adopted to discover new gene group functions in [21] where *Brdar et al.* propose a profile-based clustering method. The method allows combining separate clustering performed on each dataset to identify an integrative clustering of multiple datasets. This method involves three main steps. Firstly, modeling each dataset as a gene network with as a measure: Mutual Information, Pearson Correlation or Euclidean Distance. Secondly, applying a clustering algorithm to each gene networks constructed. Thus, we obtain a clustering of each dataset. Finally, combining the obtained clusters into a single matrix on which a non-negative matrix factorization technique is applied in order to extract final clusters.

5 Conclusion

The emergence of new methods for omics data integration will yield a more informative result, and hence enriches analyses in biology and genomic medicine. In this context, many research works have been introduced. They can be categorized under three main criteria, the biological issue of interest, characteristics of used data and the technical solution proposed to analyze multiple omics datasets collectively. Although many methods to integrate omics data have been addressed, it is still needed for powerful and advanced strategies to provide a full analysis of biologic systems taking profit from available data including clinical data and prior knowledge. Integrative analysis in biology and genomic medicine is a recent research avenue and is being developed. Parallel computing, visualization of integrative analysis outcomes and improving the quality of omics data integration are several issues that need to be addressed in future works.

References

1. Bersanelli, M., Mosca, E., Remondini, D., Giampieri, E., Sala, C., Castellani, G., Milanesi, L.: Methods for the integration of multi-omics data: mathematical aspects. *BMC bioinformatics* 17(Suppl 2), 15 (2016)
2. Gligorijević, V., Malod-Dognin, N., Pržulj, N.: Integrative methods for analyzing big data in precision medicine. *Proteomics* 16(5), 741–758 (2016)
3. Richardson, S., Tseng, G.C., Sun, W.: Statistical methods in integrative genomics. *Annual Review of Statistics and Its Application* 3, 181–209 (2016)
4. De Keersmaecker, S.C., Thijs, I., Vanderleyden, J., Marchal, K.: Integration of omics data: how well does it work for bacteria? *Molecular microbiology* 62(5), 1239–1250 (2006)

5. Wei, Y.: Integrative analyses of cancer data: a review from a statistical perspective. *Cancer informatics* (Suppl. 2), 173 (2015)
6. Wang, D., Gu, J.: Integrative clustering methods of multi-omics data for molecule-based cancer classifications. *Quantitative Biology* 4(1), 58–67 (2016)
7. Kim, Y.A., Cho, D.Y., Przytycka, T.M.: Understanding genotype-phenotype effects in cancer via network approaches. *PLoS Comput Biol* 12(3), e1004747 (2016)
8. Su, M., Dou, X., Cheng, H., Han, J.D.J.: Integrative epigenomics. In: *Computational and Statistical Epigenomics*, pp. 127–139. Springer (2015)
9. Kormaksson, M., Booth, J.G., Figueroa, M.E., Melnick, A.: Integrative model-based clustering of microarray methylation and expression data. *The Annals of Applied Statistics* pp. 1327–1347 (2012)
10. Hofree, M., Shen, J.P., Carter, H., Gross, A., Ideker, T.: Network-based stratification of tumor mutations. *Nature methods* 10(11), 1108–1115 (2013)
11. Hamid, J.S., Hu, P., Roslin, N.M., Ling, V., Greenwood, C.M., Beyene, J.: Data integration in genetics and genomics: methods and challenges. *Human genomics and proteomics* 1(1) (2009)
12. Vaske, C.J., Benz, S.C., Sanborn, J.Z., Earl, D., Szeto, C., Zhu, J., Haussler, D., Stuart, J.M.: Inference of patient-specific pathway activities from multi-dimensional cancer genomics data using paradigm. *Bioinformatics* 26(12), i237–i245 (2010)
13. Mosca, E., Milanesi, L.: Network-based analysis of omics with multi-objective optimization. *Molecular BioSystems* 9(12), 2971–2980 (2013)
14. CKRIET, G., Deng, M., Cristianini, N., NOBLE, W.: Kernel-based data fusion and its application to protein function prediction in yeast. In: *Pacific Symposium on Biocomputing 2004: Hawaii, USA, 6-10 January 2004*. p. 300. World Scientific (2003)
15. Jiang, H., Deng, Y., Chen, H.S., Tao, L., Sha, Q., Chen, J., Tsai, C.J., Zhang, S.: Joint analysis of two microarray gene-expression data sets to select lung adenocarcinoma marker genes. *BMC bioinformatics* 5(1), 81 (2004)
16. Cun, Y., Fröhlich, H.: Network and data integration for biomarker signature discovery via network smoothed t-statistics. *PloS one* 8(9), e73074 (2013)
17. Shen, R., Olshen, A.B., Ladanyi, M.: Integrative clustering of multiple genomic data types using a joint latent variable model with application to breast and lung cancer subtype analysis. *Bioinformatics* 25(22), 2906–2912 (2009)
18. Zhang, S., Liu, C.C., Li, W., Shen, H., Laird, P.W., Zhou, X.J.: Discovery of multi-dimensional modules by integrative analysis of cancer genomic data. *Nucleic acids research* p. gks725 (2012)
19. Wang, B., Mezlini, A.M., Demir, F., Fiume, M., Tu, Z., Brudno, M., Haibe-Kains, B., Goldenberg, A.: Similarity network fusion for aggregating data types on a genomic scale. *Nature methods* 11(3), 333–337 (2014)
20. Serra, A., Fratello, M., Fortino, V., Raiconi, G., Tagliaferri, R., Greco, D.: Mvda: a multi-view genomic data integration methodology. *BMC bioinformatics* 16(1), 261 (2015)
21. Brdar, S., Crnojević, V., Zupan, B.: Integrative clustering by nonnegative matrix factorization can reveal coherent functional groups from gene profile data. *IEEE journal of biomedical and health informatics* 19(2), 698–708 (2015)
22. Mo, Q., Shen, R.: iclusterplus: integrative clustering of multiple genomic data sets (2013)
23. Lee, D.D., Seung, H.S.: Learning the parts of objects by non-negative matrix factorization. *Nature* 401(6755), 788–791 (1999)

GBAS: Graph Based approach for Arabic Stemming

Imene Boukhalifa¹ and Sihem Mostefai²

¹ New Technologies of Information and Communication Faculty,
Abdelhamid Mehri University, Constantine 2, Algeria,
boukhalifa.imene@hotmail.com,

² New Technologies of Information and Communication Faculty,
Abdelhamid Mehri University, Constantine 2, Algeria,
sihem.mostefai@univ-constantine2.dz

Abstract. In Natural Language Processing (NLP) field, Arabic is primarily one of the most challenging languages. These challenges are reflected in the difficulties brought by its rich vocabulary. Stemming as a technique of Arabic NLP is increasingly becoming a significant research domain since Arabic has a complex morphology. Stemming approaches are several: the first one is based on manually made dictionaries, the second one on the morphological analysis of the language, and the third one on statistical studies. In this study, we suggest a new graph based-approach for Arabic stemming. We represent the word by a directed weighted graph having a set of connected components. Each of these components has a specific representation. Our stemmer's efficiency is proved by the results obtained.

Keywords: Arabic, Stemming, Natural language processing, Graph

Introduction

Arabic is one of the richest languages in terms of vocabulary where a single word can have different meanings. This latter is defined not only by the morphological derivation of the word but also by its diacritics, for instance, سَلَّمَ sallama (to greet), سُلَّامٌ sullamo (a ladder), سَلَامٌ silmo (peace). Moreover, Arabic has multiple ways of expressing the same idea, in addition to its various rhetorical expressions, which leads to a certain ambiguity for understanding the meaning of sentences. These characteristics contributed to the shortage of tools and resources for this language comparing it with other languages such as English. For this aim, giving a special attention for the development of tools to process such a complex language is highly required. The domain of treating and processing the human language is called Natural Language Processing (NLP). NLP, which can also be referred to as "computational linguistics", includes many tasks such as tokenization, part-of-speech tagging, named entity recognition, semantic role labeling and stemming. Stemming as a technique of NLP has become a significant research

domain for Arabic language since it has a complex morphology. This task usually comprehends removing suffixes and prefixes. However, this process is not always obvious when the language is highly inflectional especially concerning irregular verbs such as the verb (كان) in the past tense will change to (يَكُونُ) in the future. There are several types of stemming approaches in ANLP based on: manually made dictionaries, morphological and statistical studies, and translation.

This paper presents a new approach for Arabic stemming. We attempt to represent each verb (stem) by a labeled graph while taking into account all its features. Then, we model a detected verb of a text by a graph in the same way. Finally, we compare the graph of the detected word with all the graphs of stems in order to identify the most similar one. Indeed, the corresponding stem graph is a subgraph in the graph representing the detected word. The proposed solution was tested on a corpus dedicated to a plagiarism detection shared task [1] in Arabic documents, containing 1174 documents in different domains such as science, politics, sport and more. Our stemmer has showed good performances which are quantified in the last section.

The rest of this paper is organized as follows: the second section tackles a study of related works. The third section explains our approach with the techniques used and the settings selected while the fourth one tackles the obtained results and their discussion. Finally, the paper concludes with perspectives for further investigations.

Related Work

Arabic stemming approaches in the state of the art can be divided into four main categories based on the used techniques: manually made dictionaries, morphological or statistical studies, and translation.

The use of translation is an important step in [2] and [3]. In [2], stems are found by employing Google's translator to reveal the equivalent English stems and then by the same process getting the Arabic stem. [3], however, uses two sets of texts, one is the English translation of the other. Then, the stems are found by removing suffixes and prefixes. This process is called *light stemming*. Light stemming refers to the process of stripping off a small set of prefixes and/or suffixes without trying to deal with infixes or recognize patterns and find roots [4]. [5] proposed a stemmer that removes suffixes and prefixes, then match the result with the available patterns based on some rules to find the root without

using a dictionary. [6] introduced a light stemmer which operates on a set of texts collected from France Press agency. [7] proposed a hybrid stemmer which combines both Khoja Stemmer, a light stemmer and N-grams, then select the most suitable root by calculating the *Dice measure*. The three techniques were tested on a dictionary of 9000 Arabic words.

While statistical stemmers group word variants using clustering [4], the morphological analyzer attempts to find words roots. The first attempt of applying machine learning to stemming Arabic texts was introduced by [8]. Their approach was based on a neural network combined with numerical relations between the characters. However, only three and four letters words were taken in consideration in this study.

Moreover, a rule based approach was proposed in [9]. This stemmer was divided into three modules: a module that contains rules for prefixes, the second dedicated for the suffixes and the last one concentrates on alternating letters to their original form.

Recent researches on stemming are moving towards graph modeling because of its solid theory and rich concepts. Thus, we suggest modeling the stems and the detected words by graphs whose nodes and edges take into consideration the most important characteristics of the Arabic words.

Proposed Approach

We suggest in this paper a novel approach of stemming in the Arabic language. It aims to associate a giving word to the appropriate stem (verb) by using the graph modeling. Indeed, a database of 450 Arabic verbs was created and each verb was represented by a direct labeled graph where each main property of Arabic words finds its equivalent. The proposed graph is composed of two types of nodes: gray nodes which represent the letters and green nodes that represent the main characteristics of the letter. These green nodes are labeled and each label reflects the characteristic type. The table 1 and the figure 1 below illustrate the different characteristics of letters and their corresponding labels.

Characteristic	Label
Simple Ascendent	SA
Descendent	D
High Simple Point	HSP
Low Simple Point	LSP
High Double Point	HDP
Low Double Point	LDP
High Triple Point	HTP
Loop	Loop

Table 1. The characteristics of letters.

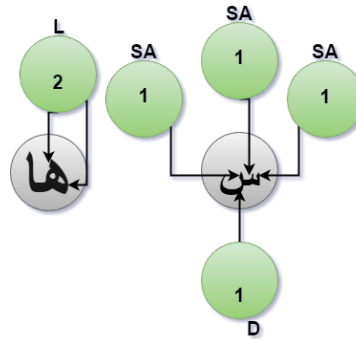


Fig. 1. The letters س and ه's characteristics.

The SA characteristic means that the letter has an ascendent part. However, HSP, LSP, HDP, LDP and HTP characteristics determine the number and the position of the points in the letter, whereas, the L characteristic represents the number of loops. The following figure 2 illustrates an example representing the word سَاعَدَ.

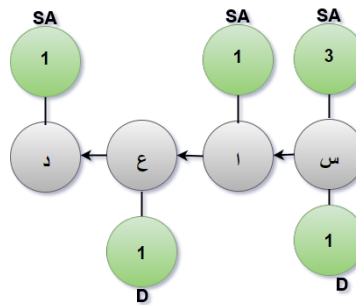


Fig. 2. The graph representation of the word سَاعَدَ.

From the figure above, it can be clearly seen that the word سَاعَدَ contains four letters, the first has three simple ascendants and one descendant while the second letter has only one simple ascendent. The third letter, however, contains

one descendent and the fourth one has only one simple ascendent. According to this modeling, any variation of the word **سَاعَدَ** by adding suffixes and prefixes will be an extension of the previous graph as seen in the figures 3 and 4. Thus, the process of finding a stem of a given word will be searching for the core graph among the graphs of the database according to a certain threshold equal to or greater than 75% of the nodes of the stem graph. This threshold was chosen according to the morphology of the Arabic language. In general, a conjugated verb does not change its form unless it is an irregular one where it changes only one letter in most cases.

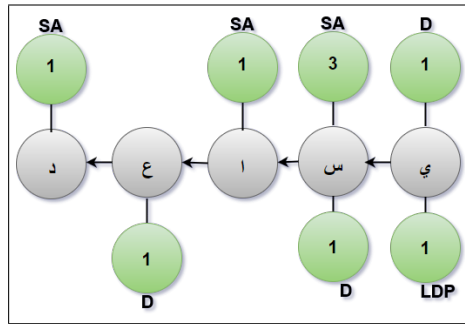


Fig. 3. The graph representation of the word يُسَاعِدُ.

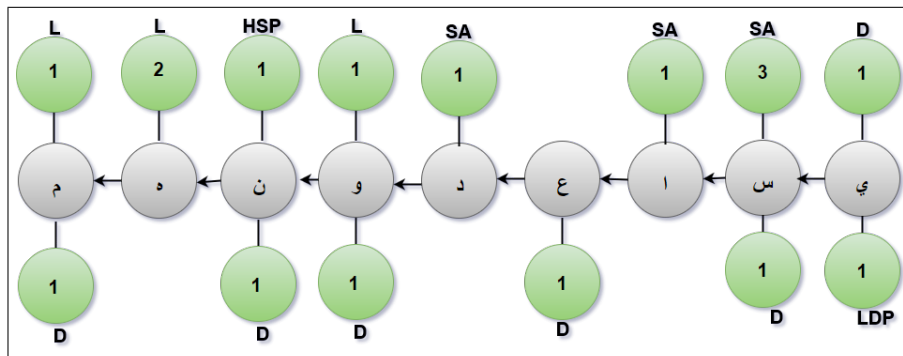


Fig. 4. The graph representation of the word يُسَاعِدُونَهُمْ.

Results and Discussion

Our Stemmer was tested on the corpus of AraPlagDet shared task [1]. This text collection contains 1174 Arabic documents that tackles different domains. To assess the performances of our stemmer, two measures were calculated: *precision* and *recall*. Recall calculates the stemmed cases among the whole cases in a given dataset, while precision assesses the correct stemmed cases among all the cases false or correct ones. Their formulas are represented in the equations 2 and 1 respectively. CSC is the Correct Stemmed Cases, FSC represents the False Stemmed Cases where irregular verbs change their form which causes a wrong stemming process, while NSC stands for the Not Stemmed Cases. Our stemmer has proved efficient. The table 2 shows the obtained results.

$$Precision = \frac{CSC}{CSC + FSC} \quad (1)$$

$$Recall = \frac{CSC}{CSC + NSC} \quad (2)$$

Measures	Dataset1	Dataset2	Dataset3
Precision	87%	82%	92%
Recall	13%	18%	8%

Table 2. Experimental results of the proposed stemmer on three datasets.

On the basis of the experimental results presented above, we discuss here the effectiveness of our proposed method for stemming Arabic documents. It can be seen that our stemmer was tested on different datasets. The first dataset tackles scientific domains, the second deals about historical topics while the third dataset contains georgaphical texts. Our stemmer showed a precision of 87%, 82% and 92%; recall, however, was 13%, 18% and 8% in scientific, historical and georgaphical topics respectively. Our approach shows remarkable results in most cases, however, it stays less effective when handling historical texts since Arabic irregular verbs (الافعال المتعّلة) change their form from one tense to another. A few examples are presented in the table 3.

verb	prefixes	suffixes	stem
يَسْتَخْرِجُونَهُمْ	يَسْت	وَهُمْ	خرج
يُودِعُ	يُ	/	ودع
مَرَّتْ	/	ت	مرَّ
يَكُونُ	ي	/	كَانَ

Table 3. Examples from the datasets.

Conclusion

This paper presented a new approach of stemming in the Arabic language. Our stemmer aims to associate each verb with its stem using a graph representation that contains the main characteristics of its letters. The selecting process is based on a database of 450 arabic stems that was created in our approach. Thus, a stem is a subgraph of the verb's representation where each main property of Arabic words finds its equivalent.

The experimental results has demonstrated that our stemmer shows good performances. This latter, however, still needs some improvements concerning irregular verbs (الافعال المعتلة). In this perspective, we aim to enhance this study by its comparaison with other techniques of Arabic stemming in the state of the art. Furthermore, a deep study of the irregular verbs must be taken in consideration as a future work.

References

1. I. Bensalem, I. Boukhalfa, P. Rosso, L. Abouenour, K. Darwish, and S. Chikhi, "Overview of the araplagdet pan@ fire2015 shared task on arabic plagiarism detection.," in *FIRE Workshops*, pp. 111–122, 2015.
2. A. Chen and F. C. Gey, "Building an arabic stemmer for information retrieval.," in *TREC*, vol. 2002, pp. 631–639, 2002.
3. M. Rogati, S. McCarley, and Y. Yang, "Unsupervised learning of arabic stemming using a parallel corpus.," in *Proceedings of the 41st Annual Meeting on Association for Computational Linguistics-Volume 1*, pp. 391–398, Association for Computational Linguistics, 2003.

4. I. A. Al-Sughaiyer and I. A. Al-Kharashi, "Arabic morphological analysis techniques: A comprehensive survey," *Journal of the American Society for Information Science and Technology*, vol. 55, no. 3, pp. 189–213, 2004.
5. R. Alshalabi, "Pattern-based stemmer for finding arabic roots," *Information Technology Journal*, vol. 4, no. 1, pp. 38–43, 2005.
6. L. S. Larkey, L. Ballesteros, and M. E. Connell, "Improving stemming for arabic information retrieval: light stemming and co-occurrence analysis," in *Proceedings of the 25th annual international ACM SIGIR conference on Research and development in information retrieval*, pp. 275–282, ACM, 2002.
7. M. Hadni, A. Lachkar, and S. A. Ouatik, "A new and efficient stemming technique for arabic text categorization," in *Multimedia Computing and Systems (ICMCS), 2012 International Conference on*, pp. 791–796, IEEE, 2012.
8. H. M. Alserhan and A. S. Ayesh, "An application of neural network for extracting arabic wordroots.," *WSEAS Transactions on Computers*, vol. 5, no. 11, pp. 2623–2627, 2006.
9. T. M. T. Sembok and B. A. Ata, "Arabic word stemming algorithms and retrieval effectiveness," in *Proceedings of the World Congress on Engineering*, vol. 3, pp. 3–5, 2013.

A hybrid direction method for solving linear programs with bounded variables

Khalil Djeloud¹, Mohand Bentobache¹ and Mohand Ouamer Bibi²

¹Laboratory of Pure and Applied Mathematics, University of Laghouat, 03000;
kdjeloud@gmail.com, mbentobache@yahoo.com

²LaMOS Research Unit, University of Bejaia, 06000 Bejaia, Algeria
mobibi.dz@gmail.com

Abstract. In this work, we suggest a new hybrid direction method for solving linear programming problems with bounded variables. The suggested algorithm uses a new hybrid direction in order to move from one feasible solution to a better one. We have proved that this direction is an ascent feasible direction. Moreover, optimality and suboptimality criteria are derived and updating formulas are constructed for the suboptimality estimate. Finally, a numerical example is given for illustration purpose.

Keywords: linear programming, adaptive method, support feasible solution, hybrid direction, suboptimality estimate, long step rule.

1 Introduction

Linear programming (LP) is a very important field in applied mathematics. Indeed, many practical problems arising in industry, finance, planning and scheduling, optimal resource allocation, etc., can be formulated as LP problems. There exists several methods for solving linear programs, we can cite: the simplex method [8], interior point methods [13,14,17], support and adaptive methods [7,10,11,15,18,19], etc.

In [2,3,4,5,6], the authors suggested algorithms using the concept of hybrid direction for solving linear programs with bounded variables. These directions are called hybrid because they take for some components the adaptive direction components and for others a quantity depending on the reduced costs ones and a given parameter.

In this work, we suggest a new hybrid direction, where the used quantity is the inverse of the reduced costs multiplied by a given non-negative parameter. We prove that the suggested direction is an ascent feasible direction. Moreover, optimality and suboptimality criteria are derived, updating formulas are constructed for the suboptimality estimate and the long step rule suggested in [11] is adapted and applied for changing the support in our method.

In Section 2, we present the problem and give some definitions. In Section 3, we define our direction, we describe and justify the steps of an iteration of the suggested method. In Section 4, we illustrate our approach by a numerical example. Finally, we conclude the paper and give some future works.

2 Problem statement and definitions

The linear programming problem with bounded variables is presented in the following standard form:

$$\begin{aligned} \max \quad & z(x) = c^T x \\ \text{subject to} \quad & Ax = b, \quad l \leq x \leq u, \end{aligned} \quad (1)$$

where $c, x \in \mathbb{R}^n$; $b \in \mathbb{R}^m$; A an $(m \times n)$ -matrix with $\text{rank} A = m < n$; $l, u \in \mathbb{R}^n$, with $\|l\| < \infty$ and $\|u\| < \infty$. The symbol $(^T)$ is the transposition operation.

- We define the following sets of indices:

$$I = \{1, 2, \dots, m\}, \quad J = \{1, 2, \dots, n\}, \quad J = J_B \cup J_N, \quad J_B \cap J_N = \emptyset, \quad |J_B| = m.$$

- The matrix A and the different vectors can be partitionned as follows:

$$x = x(J) = (x_B, x_N), \quad x_B = x(J_B) = (x_j, j \in J_B), \quad x_N = x(J_N) = (x_j, j \in J_N);$$

$$c = c(J) = (c_B, c_N), \quad c_B = c(J_B) = (c_j, j \in J_B), \quad c_N = c(J_N) = (c_j, j \in J_N);$$

$$A = A(I, J) = (A_B, A_N), \quad A_B = A(I, J_B) = (a_{ij}, i \in I, j \in J_B);$$

$$A_N = A(I, J_N) = (a_{ij}, i \in I, j \in J_N).$$

- The set J_B is called a support if $\det A_B \neq 0$.
- An n -vector x is called a feasible solution (FS) if it satisfies the constraints of problem (1).
- A FS x^0 is called optimal if $z(x) \leq z(x^0)$, for all feasible solution x of problem (1).
- A FS x^ε is called ε -optimal or suboptimal if $z(x^0) - z(x^\varepsilon) \leq \varepsilon$, where x^0 is an optimal solution for the problem (1) and ε a nonnegative number chosen beforehand.
- A pair $\{x, J_B\}$ comprising a feasible solution x and a support J_B is called a support feasible solution (SFS).
- An SFS is called nondegenerate if $l_j < x_j < u_j$, $j \in J_B$.
- We define the multipliers vector π and the reduced costs vector Δ by:

$$\pi^T = c_B^T A_B^{-1} \text{ and } \Delta^T = \pi^T A - c^T = (\Delta_B^T, \Delta_N^T),$$

where $\Delta_B^T = c_B^T A_B^{-1} A_B - c_B^T = 0$ and $\Delta_N^T = \pi^T A_N - c_N^T$.

- We define the following sets of indices:

$$J_N^+ = \{j \in J_N, \Delta_j > 0\} \text{ and } J_N^- = \{j \in J_N, \Delta_j < 0\}.$$

The quantity $\beta(x, J_B)$, called the suboptimality estimate [10], is defined by:

$$\beta = \beta(x, J_B) = \sum_{j \in J_N^+} \Delta_j (x_j - l_j) + \sum_{j \in J_N^-} \Delta_j (x_j - u_j). \quad (2)$$

Theorem 1. (Optimality criterion [10])

Let $\{x, J_B\}$ be an SFS of the problem (1). So the relations

$$\begin{cases} \Delta_j \geq 0, & \text{if } x_j = l_j, \\ \Delta_j \leq 0, & \text{if } x_j = u_j, \\ \Delta_j = 0, & \text{if } l_j < x_j < u_j, \quad j \in J_N, \end{cases}$$

are sufficient, and in the case of nondegeneracy also necessary, for the optimality of the FS x .

Theorem 2. (Suboptimality criterion [10])

Let $\{x, J_B\}$ be an SFS of the problem (1) and $\varepsilon > 0$. If $\beta(x, J_B) \leq \varepsilon$, then the SFS $\{x, J_B\}$ is ε -optimal.

Corollary 1. Let $\{x, J_B\}$ be an SFS of the problem (1). If $\beta(x, J_B) = 0$, then the SFS $\{x, J_B\}$ is optimal.

3 An iteration of the hybrid direction method

Let $\{x, J_B\}$ be an SFS for the problem (1) and $\eta > 0$. We introduce the following sets of indices:

$$\begin{aligned} J_0 &= \{j \in J_N : \Delta_j = 0\}, \quad J_1 = \{j \in J_N : 0 < \Delta_j(x_j - l_j) \leq \eta\}, \\ J_2 &= \{j \in J_N : 0 < \Delta_j(x_j - u_j) \leq \eta\}, \quad J_3 = \{j \in J_N : \Delta_j(x_j - l_j) > \eta\}, \\ J_4 &= \{j \in J_N : \Delta_j(x_j - u_j) > \eta\}, \quad J_5 = \{j \in J_N : \Delta_j > 0, x_j = l_j\}, \\ J_6 &= \{j \in J_N : \Delta_j < 0, x_j = u_j\}. \end{aligned} \quad (3)$$

Thus,

$$J_N^+ = J_1 \cup J_3 \cup J_5, \quad J_N^- = J_2 \cup J_4 \cup J_6, \quad J_N = J_0 \cup J_N^+ \cup J_N^- = \bigcup_{i=0}^6 J_i,$$

and

$$\beta(x, J_B) = \sum_{j \in J_1 \cup J_3} \Delta_j(x_j - l_j) + \sum_{j \in J_2 \cup J_4} \Delta_j(x_j - u_j). \quad (4)$$

Remark 1. Let us remark that

$$\bigcup_{1 \leq i \leq 4} J_i = \emptyset \Rightarrow \beta = 0 \Rightarrow x \text{ is optimal.}$$

Let us define the quantity $\mu(\eta, x, J_B)$ as follows:

$$\mu = \mu(\eta, x, J_B) = \sum_{j \in J_3} \Delta_j(x_j - l_j) + \sum_{j \in J_4} \Delta_j(x_j - u_j) - \eta |J_3 \cup J_4|. \quad (5)$$

Remark 2. We have $\mu \geq 0$ and $\beta \geq \mu$. Indeed,

$$\begin{aligned}\mu &= \sum_{j \in J_3} \Delta_j(x_j - l_j) + \sum_{j \in J_4} \Delta_j(x_j - u_j) - \eta|J_3 \cup J_4| \\ &= \sum_{j \in J_3} [\Delta_j(x_j - l_j) - \eta] + \sum_{j \in J_4} [\Delta_j(x_j - u_j) - \eta] \geq 0.\end{aligned}$$

Moreover, we have

$$\beta - \mu = \sum_{j \in J_1} \Delta_j(x_j - l_j) + \sum_{j \in J_2} \Delta_j(x_j - u_j) + \eta|J_3 \cup J_4| \geq 0. \quad (6)$$

If $\bigcup_{1 \leq j \leq 4} J_i = \emptyset$, i.e., x is optimal, then $\beta = \mu = 0$.

3.1 Change of the feasible solution

Let $\{x, J_B\}$ be an SFS for the problem (1) and $\eta > 0$. If $\beta(x, J_B) \leq \varepsilon$, then x is optimal. Else, we define the direction d as follows:

$$\begin{aligned}d_j &= l_j - x_j, & \text{if } j \in J_1; \\ d_j &= u_j - x_j, & \text{if } j \in J_2; \\ d_j &= -\frac{\eta}{\Delta_j}, & \text{if } j \in J_3 \cup J_4; \\ d_j &= 0, & \text{if } j \in J_0 \cup J_5 \cup J_6; \\ d_B &= -A_B^{-1} A_N d_N.\end{aligned} \quad (7)$$

In order to improve the objective function, we compute the step length θ^0 along the direction d as follows: $\theta^0 = \min\{\theta_{j_1}, 1\}$, where

$$\theta_{j_1} = \min_{j \in J_B} \theta_j, \text{ with } \theta_j = \begin{cases} \frac{u_j - x_j}{d_j}, & \text{if } d_j > 0; \\ \frac{l_j - x_j}{d_j}, & \text{if } d_j < 0; \\ \infty, & \text{if } d_j = 0. \end{cases} \quad (8)$$

Then the new feasible solution is $\bar{x} = x + \theta^0 d$. Since $Ad = 0$, the vector d is a feasible direction. Furthermore, d is an ascent direction. Indeed, the objective function increment is given by:

$$\begin{aligned}\bar{z} - z &= -\theta^0 \sum_{j \in J_N} \Delta_j d_j \\ &= -\theta^0 \left(\sum_{j \in J_1} \Delta_j d_j + \sum_{j \in J_2} \Delta_j d_j + \sum_{j \in J_3 \cup J_4} \Delta_j d_j \right) \\ &= \theta^0 \left(\sum_{j \in J_1} \Delta_j (x_j - l_j) + \sum_{j \in J_2} \Delta_j (x_j - u_j) + \eta|J_3 \cup J_4| \right) \\ &= \theta^0 (\beta - \mu) \geq 0.\end{aligned}$$

Remark 3. Note that in the simplex method, only one nonbasic component of the current solution is changed and all the others remain the same. However in our method, almost all the nonbasic components of the current solution are changed. This property can gives a better improvement for the objective function. Moreover, contrarily to the simplex method which starts by extreme points only, our method can starts with interior points.

Lemma 1. *The following updating formula holds:*

$$\bar{\beta} = \beta(\bar{x}, J_B) = (1 - \theta^0)\beta(x, J_B) + \theta^0\mu(\eta, x, J_B) \leq \beta(x, J_B).$$

Proof. We have

$$\begin{aligned} \bar{\beta} &= \beta(\bar{x}, J_B) = \sum_{j \in J_N^+} \Delta_j(\bar{x}_j - l_j) + \sum_{j \in J_N^-} \Delta_j(\bar{x}_j - u_j) \\ &= \sum_{j \in J_1 \cup J_3 \cup J_5} \Delta_j(\bar{x}_j - l_j) + \sum_{j \in J_2 \cup J_4 \cup J_6} \Delta_j(\bar{x}_j - u_j). \end{aligned}$$

However, for $j \in J_5$, $d_j = 0$ and $x_j = l_j$, so we get $\bar{x}_j = x_j + \theta^0 d_j = l_j$. Similarly, we get for $j \in J_6$, $\bar{x}_j = u_j$. Hence

$$\begin{aligned} \bar{\beta} &= \sum_{j \in J_1 \cup J_3} \Delta_j(\bar{x}_j - l_j) + \sum_{j \in J_2 \cup J_4} \Delta_j(\bar{x}_j - u_j) \\ &= \beta(x, J_B) + \theta^0 \left(\sum_{j \in J_1} \Delta_j d_j + \sum_{j \in J_2} \Delta_j d_j + \sum_{j \in J_3 \cup J_4} \Delta_j d_j \right) \\ &= \beta(x, J_B) - \theta^0 \left(\sum_{j \in J_1} \Delta_j(x_j - l_j) + \sum_{j \in J_2} \Delta_j(x_j - u_j) + \eta |J_3 \cup J_4| \right). \end{aligned}$$

Following Formula (6), we deduce that

$$\bar{\beta} = \beta - \theta^0(\beta - \mu) \leq \beta.$$

Moreover, we have

$$\bar{\beta} = (1 - \theta^0)\beta + \theta^0\mu. \quad \square$$

Remark 4. If $J_3 = J_4 = \emptyset$, then $\mu = 0$. So we find the updating formula of the suboptimality estimate [10].

Proposition 1. *(Sufficient conditions for optimality and suboptimality of $\bar{x}$)*

If $\theta^0 = 1$ and $\mu(\eta, x, J_B) \leq \varepsilon$, then the feasible solution $\bar{x}$ is ε -optimal.

If $\theta^0 = 1$ and $\mu(\eta, x, J_B) = 0$, then the feasible solution $\bar{x}$ is optimal.

Proof.

If $\theta^0 = 1$ and $\mu(\eta, x, J_B) \leq \varepsilon$, then by using Lemma 1, we have

$$\beta(\bar{x}, J_B) = \mu(\eta, x, J_B) \leq \varepsilon.$$

Following Theorem 2, we deduce the suboptimality of $\bar{x}$.

If $\theta^0 = 1$ and $\mu(\eta, x, J_B) = 0$, then we have $\beta(\bar{x}, J_B) = \mu(\eta, x, J_B) = 0$. By using Corollary 1, we deduce the optimality of $\bar{x}$. $\square$

If $\theta^0 = \theta_{j_1} < 1$ and $\beta(\bar{x}, J_B) > \varepsilon$, then we change the support J_B .

3.2 Change of the support

We define the n -vector κ and the real number α_0 as follows:

$$\kappa = x + d \text{ and } \alpha_0 = \kappa_{j_1} - \bar{x}_{j_1},$$

where j_1 is the leaving index computed with Formula (8). We compute the dual direction as follows:

$$t_{j_1} = -\text{sign}(\alpha_0), \quad t_j = 0, \quad j \neq j_1, \quad j \in J_B, \quad t_N^T = t_B^T A_B^{-1} A_N. \quad (9)$$

Let us define the following sets:

$$J_0^+ = \{j \in J_0 : t_j > 0\}, \quad J_0^- = \{j \in J_0 : t_j < 0\}, \quad (10)$$

and the quantity:

$$\alpha = -|\alpha_0| + \sum_{j \in J_0^+ \cup J_3} t_j(\kappa_j - l_j) + \sum_{j \in J_0^- \cup J_4} t_j(\kappa_j - u_j). \quad (11)$$

The new reduced costs vector and the new support are computed by:

$$\bar{\Delta} = \Delta + \sigma^0 t \text{ and } \bar{J}_B = (J_B \setminus \{j_1\}) \cup \{j_0\}, \quad (12)$$

where

$$\sigma^0 = \sigma_{j_0} = \min_{j \in J_N} \{\sigma_j\}, \text{ with } \sigma_j = \begin{cases} \frac{-\Delta_j}{t_j} & \text{if } \Delta_j t_j < 0, \\ 0 & \text{if } j \in J_0^+, \kappa_j \neq l_j, \\ 0 & \text{if } j \in J_0^-, \kappa_j \neq u_j, \\ +\infty & \text{otherwise.} \end{cases} \quad (13)$$

Remark 5. If $\sigma^0 = +\infty$, then the problem (1) is infeasible. So in the following, we assume that $\sigma^0 < +\infty$.

The suboptimality estimate corresponding to the new feasible solution and the new support is given by:

$$\bar{\beta} = \beta(\bar{x}, \bar{J}_B) = \sum_{j \in \bar{J}_N, \bar{\Delta}_j > 0} \bar{\Delta}_j(\bar{x} - l_j) + \sum_{j \in \bar{J}_N, \bar{\Delta}_j < 0} \bar{\Delta}_j(\bar{x} - u_j).$$

By using duality theory, we can prove the following result:

Proposition 2. *The suboptimality estimate $\beta(\bar{x}, \bar{J}_B)$ can be written as follows:*

$$\beta(\bar{x}, \bar{J}_B) = \beta(\bar{x}, J_B) + \sigma^0 \alpha. \quad (14)$$

If $\alpha \leq 0$, then $\beta(\bar{x}, \bar{J}_B) \leq \beta(\bar{x}, J_B)$. If $\alpha > 0$, then we update appropriately the value of the parameter η by applying the following procedure:

Procedure 1. (Procedure of updating the value of η)

- (1) Compute $\bar{J}_3 = \{j \in J_N : \Delta_j(\bar{x}_j - l_j) > \eta\}$;
- (2) Compute $\bar{J}_4 = \{j \in J_N : \Delta_j(\bar{x}_j - u_j) > \eta\}$;
- (3) If $\bar{J}_3 \cup \bar{J}_4 = \emptyset$, then $\bar{\eta} = \eta$,
 else compute $\eta_1 = \max_{j \in \bar{J}_3} \{\Delta_j(\bar{x}_j - l_j)\}$, $\eta_2 = \max_{j \in \bar{J}_4} \{\Delta_j(\bar{x}_j - u_j)\}$ and set
 $\bar{\eta} = \max\{\eta_1, \eta_2\}$.

Thanks to Proposition 2, we can apply the long step rule procedure suggested in [11] for our method. This procedure is described in the following steps:

Procedure 2. (Long step rule procedure)

- (1) Compute σ_j , $j \in J_N$ and $\sigma^0 = \min_{j \in J_N} \sigma_j$ with (13);
- (2) If $\sigma^0 = \infty$, then the problem (1) is infeasible;
- (3) Sort the indices $\{i \in J_N : \sigma_i \neq \infty\}$ by the increasing order of the numbers
 σ_i : $\sigma_{i_1} \leq \sigma_{i_2} \leq \dots \leq \sigma_{i_p}$, $i_k \in J_N$, $\sigma_{i_k} \neq \infty$, $k = 1, \dots, p$;
- (4) If $p = 1$, then $j_r = i_1$, go to step (8);
- (5) For all i_k , $k = 1, \dots, p$, compute $\Delta V_{i_k} = |t_{i_k}|(u_{i_k} - l_{i_k})$;
- (6) Compute V_{i_k} , $k = 0, \dots, p$, where

$$\begin{cases} V_{i_0} = V_0 = \alpha, \\ V_{i_k} = V_0 + \sum_{s=1}^k \Delta V_{i_s} = V_{i_{k-1}} + \Delta V_{i_k}, \quad k = 1, \dots, p; \end{cases}$$

- (7) Find the index $j_r = i_q$ such that $V_{i_{q-1}} < 0$ and $V_{i_q} \geq 0$;
- (8) Let j_1 be the index computed by (8). Set $\bar{J}_B = (J_B \setminus \{j_1\}) \cup \{j_r\}$, $\sigma^0 = \sigma_{j_r}$
 and $\bar{\Delta} = \Delta + \sigma^0 t$;
- (9) Compute $\beta(\bar{x}, \bar{J}_B) = \beta(\bar{x}, J_B) + \sum_{k=1}^q (\sigma_{i_k} - \sigma_{i_{k-1}}) V_{i_{k-1}}$ with $\sigma_{i_0} = 0$.

3.3 Scheme of the algorithm

Let $\{x, J_B\}$ be an initial SFS for the problem (1) and $\varepsilon, \eta > 0$. The scheme of the hybrid direction method for solving linear programs with bounded variables (HDMLP) is described in the following steps:

Algorithm (HDMLP).

- (1) Compute $\pi^T = c_B^T A_B^{-1}$, $\Delta_N^T = \pi^T A_N - c_N^T$;
- (2) Compute the suboptimality estimate β with (2);
- (3) If $\beta = 0$, then the algorithm stops with the optimal SFS $\{x, J_B\}$;
- (4) If $\beta \leq \varepsilon$, then the algorithm stops with ε -optimal SFS $\{x, J_B\}$;
- (5) Compute the sets J_i , $i = 1, \dots, 4$ with (3);
- (6) Compute μ with (5);
- (7) Compute the primal search direction d with (7);
- (8) Compute the primal step length θ_{j_1} with (8) and $\theta^0 = \min\{\theta_{j_1}, 1\}$;

- (9) Compute $\bar{x} = x + \theta^0 d$ and $\bar{z} = z + \theta^0(\beta - \mu)$;
- (10) If $\theta^0 = 1$, then
 - (10.1) If $\mu = 0$, then $\bar{x}$ is optimal. Stop.
 - (10.2) If $\mu \leq \varepsilon$, then $\bar{x}$ is ε -optimal. Stop.
 - (10.3) Else, compute the value $\bar{\eta}$ with Procedure 1 set $\eta = \bar{\eta}$, $x = \bar{x}$, $z = \bar{z}$, $\beta = \mu$ and go to step (5);
- (11) If $\theta^0 = \theta_{j_1} < 1$, then compute $\bar{\beta} = (1 - \theta^0)\beta + \theta^0\mu$. If $\bar{\beta} \leq \varepsilon$ then the algorithm stops with the ε -optimal SFS $\{\bar{x}, J_B\}$;
- (12) Compute $\kappa = x + d$ and $\alpha_0 = \kappa_{j_1} - \bar{x}_{j_1}$;
- (13) Compute the dual direction t with (9);
- (14) Compute J_0 with (3), J_0^+ and J_0^- with (10);
- (15) Compute α with (11), if $\alpha > 0$ then
 - (15.1) If $J_0 \neq \emptyset$ then choose an index $j_0 \in J_0^+ \cup J_0^-$, set $\bar{\Delta} = \Delta$, $\bar{J}_B = (J_B \setminus \{j_1\}) \cup \{j_0\}$, $\bar{\beta} = \beta$ and go to step (17);
 - (15.2) If $J_0 = \emptyset$, then compute the value $\bar{\eta}$ with Procedure 1, set $\eta = \bar{\eta}$, $x = \bar{x}$, $z = \bar{z}$, and go to step (2);
- (16) Compute the dual step length σ^0 , the new reduced costs vector $\bar{\Delta}$, the new support $\bar{J}_B$ and the new suboptimality estimate $\bar{\beta}$ with the long step rule (Procedure 2);
- (17) Set $x = \bar{x}$, $J_B = \bar{J}_B$, $z = \bar{z}$, $\Delta = \bar{\Delta}$, $\beta = \bar{\beta}$, go to step (3).

4 Numerical example

$$\begin{aligned} \max \quad & z(x) = c^T x, \\ \text{s.t.} \quad & Ax = b, \quad l \leq x \leq u, \end{aligned}$$

where $A = \begin{pmatrix} 3 & -1 & 1 & 1 & 0 \\ -1 & -4 & 1 & 0 & 1 \end{pmatrix}$, $l = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}$, $u = \begin{pmatrix} 1 \\ 4 \\ 5 \\ 5 \\ 19 \end{pmatrix}$, $c = (1, -3, 1, 0, 0)^T$, $b = \begin{pmatrix} 1 \\ 2 \end{pmatrix}$.

Iteration 1:

- We set $\eta = 1$, $J_B = \{4, 5\}$, $J_N = \{1, 2, 3\}$, $x = (0, 0, 0, 1, 2)^T$, $z(x) = 0$.
- The vector of multipliers π and the vector of reduced costs Δ :
 $\pi^T = c_B^T A_B^{-1} = (0, 0)$, $\Delta^T = \pi^T A - c^T = (-1, 3, -1, 0, 0)$.
- The suboptimality estimate β :

$$\beta(x, J_B) = \sum_{j \in J_N, \Delta_j > 0} \Delta_j (x_j - l_j) + \sum_{j \in J_N, \Delta_j < 0} \Delta_j (x_j - u_j) = 6.$$

- The sets J_i , $i = 1, \dots, 6$:
 $J_0 = \emptyset$, $J_1 = \emptyset$, $J_2 = \{1\}$, $J_3 = \emptyset$, $J_4 = \{3\}$, $J_5 = \{2\}$, $J_6 = \emptyset$.
- The quantity μ : $\mu = \Delta_3(x_3 - u_3) - \eta = 4$.

- The primal search direction d :
 $d_1 = u_1 - x_1 = 1$, $d_2 = 0$, $d_3 = \frac{-\eta}{\Delta_3} = 1$, $d_B = -A_B^{-1}A_N d_N = (-4, 0)^T$.
Then $d = (1, 0, 1, -4, 0)^T$.
- The primal step length θ^0 :
 $\theta_4 = \frac{l_4 - x_4}{d_4} = \frac{1}{4}$, $\theta_5 = +\infty$, then $\theta^0 = \theta_{j_1} = \theta_4 = \frac{1}{4}$. So

$$\bar{x} = x + \theta^0 d = (\frac{1}{4}, 0, \frac{1}{4}, 0, 2)^T, \quad \bar{z} = z + \theta^0(\beta - \mu) = \frac{1}{2}.$$

- The new suboptimality estimate: $\bar{\beta} = (1 - \theta^0)\beta + \theta^0\mu = \frac{11}{2}$.
- The vector κ and α_0 : $\kappa = x + d = (1, 0, 1, -3, 2)$ and $\alpha_0 = x_{j_1} - \bar{x}_{j_1} = -3$.
- The dual direction t :
 $t_4 = 1$, $t_5 = 0$, $t_N^T = t_B^T A_B^{-1} A_N = (3, -1, 1)^T$. Then $t = (3, -1, 1, 1, 0)^T$.
- The quantity α : since $J_0^+ = J_0^- = J_3 = \emptyset$, we get

$$\alpha = -|\alpha_0| + \sum_{j \in J_4} t_j(\kappa_j - u_j) = -7 < 0.$$

- The dual step length σ^0 :
 $\sigma_1 = \frac{-\Delta_1}{t_1} = \frac{1}{3}$, $\sigma_2 = \frac{-\Delta_2}{t_2} = 3$, $\sigma_3 = \frac{-\Delta_3}{t_3} = 1$. So $\sigma_1 \leq \sigma_3 \leq \sigma_2$.

$$\Delta V_1 = |t_1|(u_1 - l_1) = 3, \quad \Delta V_3 = |t_3|(u_3 - l_3) = 5, \quad \Delta V_2 = |t_2|(u_2 - l_2) = 4.$$

$$V_0 = \alpha = -7, \quad V_1 = V_0 + \Delta V_1 = -4, \quad V_3 = V_1 + \Delta V_3 = 1.$$

Since $V_1 < 0$ and $V_3 > 0$, we set $j_r = 3$, $\sigma^0 = \sigma_3 = 1$.

- The new reduced costs vector $\bar{\Delta}$ and the new support :
 $\bar{J}_B = \{3, 5\}$, $\bar{J}_N = \{1, 2, 4\}$, $\bar{\Delta} = \Delta + \sigma^0 t = (2, 2, 0, 1, 0)^T$.
- The new suboptimality estimate $\bar{\beta}$:

$$\beta(\bar{x}, \bar{J}_B) = \beta(\bar{x}, J_B) + \sigma_1 V_0 + (\sigma_3 - \sigma_1) V_1 = \frac{1}{2}.$$

Iteration 2:

- $J_B = \{3, 5\}$, $J_N = \{1, 2, 4\}$, $x = (\frac{1}{4}, 0, \frac{1}{4}, 0, 2)^T$, $z = \frac{1}{2}$, $\Delta = (2, 2, 0, 1, 0)^T$.
- We have $J_1 = \{1\}$, $J_0 = J_2 = J_3 = J_4 = J_6 = \emptyset$, $J_5 = \{2, 4\} \Rightarrow \mu = 0$.
- The primal search direction d :
 $d_1 = \frac{-1}{4}$, $d_2 = 0$, $d_4 = 0$, $d_B = (\frac{3}{4}, -1)^T$. Then $d = (\frac{-1}{4}, 0, \frac{3}{4}, 0, -1)^T$.
- The primal step length θ^0 :
 $\theta_3 = \frac{u_3 - x_3}{d_3} = \frac{19}{3}$, $\theta_5 = \frac{l_5 - x_5}{d_5} = 2 \Rightarrow \theta^0 = 1$. Therefore,

$$\bar{x} = x + d = (0, 0, 1, 0, 1)^T \text{ and } \bar{z} = z + (\beta - \mu) = 1.$$

Since $\theta^0 = 1$ and $\mu = 0$, then $\bar{x} = (0, 0, 1, 0, 1)^T$ is optimal and $z(\bar{x}) = 1$.

5 Conclusion

In this paper, we have suggested a new hybrid direction method for solving linear programming problems with bounded variables. An optimality and suboptimality criteria are given and the long step rule is used for changing the support. In future work, we will exploit the c programming language routines of the open source LP solver GLPK [16] and the initialization technique suggested in [1] in order to implement efficiently our method and compare its performance with the simplex method on netlib test problems [12].

References

1. Bentobache, M. and Bibi, M.O.: A Two-phase Support Method for Solving Linear Programs: Numerical Experiments, *Mathematical Problems in Engineering*, Vol. 2012, Article ID 482193, 28 pages doi:10.1155/2012/482193
2. Bentobache, M.: On mathematical methods of linear and quadratic programming. Ph.D. diss., University of Bejaia, 2013
3. Bentobache, M. and Bibi, M.O.: A hybrid direction algorithm with long step rule for linear programming: Numerical experiments. *Advances in Intelligent Systems and Computing* 359, 333-344 (2015)
4. Bentobache, M., Bibi, M.O.: *Numerical Methods of Linear and Quadratic Programming*. French Academic Presses, Germany, 2016 (in french)
5. Bentobache, M., Bibi, M.O.: A new method for solving linear programs. To appear in proceedings of ROADEF 2017, 22-24 Fevrier (2017)
6. Bibi, M.O., Bentobache, M.: A hybrid direction algorithm for solving linear programs. *International Journal of Computer Mathematics* 92, No. 1, 201-216 (2015)
7. Brahmi, B., Bibi, M.O.: Dual support method for solving convex quadratic programs, *Optimization* 59, No. 6, 851-872 (2010)
8. Dantzig, G.B.: *Linear Programming and Extensions*, Princeton University Press, Princeton, N.J., 1963
9. Dikin, I.: Iterative solution of problems of linear and quadratic programming. *Soviet Mathematics Doklady* 8, 674-675 (1963)
10. Gabasov, R., Kirillova, F.M.: *Methods of linear programming*. Vol. 1, 2 and 3, Edition of the Minsk University, 1977, 1978 and 1980 (in Russian)
11. Gabasov, R., Kirillova, F.M., Prischepova, S.V.: *Optimal Feedback Control*. Springer-Verlag, London, 1995
12. Gay, M.: Electronic mail distribution of linear programming test problems. *Mathematical Programming Society COAL* 13, 10-12 (1985). Data available at <http://www.netlib.org/lp/data>.
13. Khachiyan, L.G.: A polynomial algorithm for linear programming. *Soviet Mathematics Doklady* 20, 191-194 (1979)
14. Karmarkar, N.: A new polynomial-time algorithm for linear programming. *Combinatorica* 4, 373-395 (1984)
15. Kostina, E.: The long step rule in the bounded-variable dual simplex method: Numerical experiments. *Mathematical Methods of Operations Research* 55, 413-429 (2002)
16. Makhorin, A.: *GNU Linear Programming Kit, Reference Manual Version 4.9*. Draft Edition, 2006. Software available at www.gnu.org/software/glpk/glpk.html.

17. Mehrotra, S.: On the implementation of a (primal-dual) interior point method. *SIAM Journal on Optimization* 2, 575–601 (1992)
18. Radjef, S., Bibi, M.O: An effective generalization of the direct support method in quadratic convex programming. *Applied Mathematical Sciences* 6, No. 31, 1525-1540 (2012)
19. Radjef, S., Bibi, M.O: A new algorithm for linear multiobjective programming problems with bounded variables. *Arab J Math* 3, 79–92 (2014)

Modélisation de la Consommation d'Energie dans les réseaux Ad hoc via l'outil des réseaux de Petri stochastiques

LEKADIR Ouiza, ADEL-AISSANOU Karima, BOUZIT Nacera et HADIDI Nadia

Unité de recherche LaMOS (Laboratoire de Modélisation et d'Optimisation des Systèmes)

University of Bejaia 06000, Algeria.

ouizalekadir@gmail.com

ak_adel@yahoo.fr

bouzitnacera@yahoo.fr

riccyriccynina@gmail.com

Abstract. Le but de ce travail est de modéliser la décharge de batterie des nœuds dans les réseaux ad hoc sans fil. Beaucoup de travaux mettent l'accent sur la consommation d'énergie dans de tels réseaux. Même, la recherche dans les couches les plus élevées du modèle ISO, tient compte de la consommation d'énergie et son efficacité. En effet, les nœuds qui forment ce genre de réseaux sont mobiles, donc pas de recharge instantanée de batterie. Ainsi, ce type spécial de réseaux ad hoc sont des réseaux de capteurs sans fil utilisant des batteries non rechargeables. Tous les nœuds avec une batterie épuisée sont considérés comme morts et quittent le réseau. Pour prendre en compte la consommation d'énergie, dans ce travail, nous modélisons en utilisant l'outil des réseaux de Petri (RdP), la décharge de la batterie compte tenu de la fonction d'activation instantanée et de distribution de désactivation de ces nœuds.

Keywords: Réseau Ad-hoc sans fil, Consommation d'énergie, Réseaux de Petri, Evolution des performances

1 Introduction

L'évolution récente de la technologie dans le domaine de la communication sans fil et l'apparition des unités de calcul portables poussent aujourd'hui les chercheurs à faire des efforts afin de réaliser le but des réseaux : *"L'accès à l'information n'importe où et n'importe quand"*.

En 1997, un groupe de travail de l'IETF (Internet Engineering Task Force), nommé MANET (Mobile Ad hoc NETWORKS), fut mis en place afin d'assurer la normalisation de protocoles de routage basés sur IP (Internet Protocol) et adaptés aux spécificités des réseaux mobiles auto-configurables sans infrastructure.

Les réseaux MANET offrent un potentiel important dans les applications sans fil futures en particulier grâce à leur indépendance et auto-configuration et à leur faible coût de déploiement (puisque ils ne nécessitent pas d'infrastructure), à

leur adaptabilité aux changements de topologie et à la mobilité de leurs nœuds. Toutefois, ces caractéristiques présentent de nouveaux défis et problématiques. L'un des objectifs majeur pour les réseaux Ad hoc consiste à ce que les ressources des terminaux mobiles soient utilisées d'une manière optimale. Cependant, l'une des grandes contraintes de cet objectif concerne le support énergétique. En effet, la principale contrainte dans les communications sans fil est la durée de vie limitée des terminaux mobiles dont le support énergétique représente souvent une batterie à capacité limitée. Cette contrainte est beaucoup plus importante dans les réseaux Ad hoc, car chaque paquet envoyé ou reçu ainsi que chaque utilisation du terminal exploite les ressources de la batterie. Cependant, les activités liées à la communication ne sont pas les seules qui consomment de l'énergie, même en périodes de veille, un mobile doit assurer sa connectivité au réseau à travers l'envoi périodique de messages de contrôle et l'écoute du canal. Ainsi, la consommation d'énergie est un facteur primordial pour la durée de vie des nœuds, et par conséquent, celle de tout le réseau. D'ailleurs, plusieurs travaux ont démontré que l'activité réseau est très coûteuse en énergie.

La complexité de ces réseaux de communication a renforcé la nécessité de pouvoir les représenter mathématiquement. Les modèles mathématiques obtenus permettent de décrire leur évolution et de les quantifier précisément en ayant recours aux outils mathématiques tels que : les processus de Poisson, les processus de Markov, les Réseaux de Petri (RdP), ...

Dans la littérature, une attention particulière a été consacrée à la conception de solutions les moins coûteuses en énergie en s'inspirant de différents modèles et outils mathématiques tels que les RdP et les chaînes de Markov [9], [23], [20], [11] et des protocoles d'acheminement optimaux [12], [10], [24] qui réduisent la consommation d'énergie de ces réseaux.

Vu la consommation non interrompue en énergie des nœuds mobiles Ad hoc et dans le but d'analyser cette consommation, on s'intéressera à modéliser la consommation d'énergie d'un nœud dans un réseau Ad hoc en ayant recours aux Réseaux de Petri Stochastiques (RdPS). Ainsi, dans cet article, notre objectif sera de modéliser la consommation d'énergie d'un nœud Ad hoc en s'inspirant de l'idée de modélisation par un modèle de processus "ON/OFF", réalisé dans [9] ou l'analyse c'est basée sur la théorie des chaînes de Markov, cependant notre modélisation est réalisée via l'un des formalisme des réseaux de Petri (RdP) qui est les RdP stochastiques généralisés (RdPSG).

Ainsi, ce papier est organisé comme suit: Après cette section introductive, la section 2 sera dédiée à la notion de consommation d'énergie dans les réseaux Ad hoc, où sera donnée la manière avec laquelle l'énergie se consomme dans ces réseaux, les axes de recherche actuels sur la conservation de l'énergie dans ces réseaux, ..., etc. Dans la section 3, un bref aperçu sur l'outil de modélisation RdP sera donné. Dans la section 4, la problématique étudiée sera bien posée. La section 5, sera consacrée à la modélisation du problème, son analyse, sa simulation et à l'interprétation des résultats. Enfin, ce papier s'achèvera par une conclusion ou seront données quelques perspectives de travail envisagées.

2 Consommation d'énergie dans les réseaux Ad hoc

La consommation d'énergie est un critère primordial dans la conception des protocoles de routage pour les réseaux Ad hoc, les nœuds mobiles utilisent généralement les équipements de stockage d'énergie autonomes pour se procurer de l'énergie et par conséquent ont une durée de vie limitée. En réalité, dans les réseaux Ad hoc l'épuisement de l'énergie d'un nœud n'affecte pas uniquement sa capacité à recevoir ou émettre, mais également sa capacité à acheminer les données pour les autres nœuds ce qui peut diminuer les performances du réseau ou carrément isoler certains segments du réseau.

Actuellement, il existe trois axes de recherche dans le domaine de conservation de la consommation d'énergie spécifiques aux réseaux Ad hoc :

- L'économie d'énergie qui minimise la dissipation d'énergie lorsqu'un nœud mobile est inutilisé. On introduit un état de veille, plus économe en énergie que l'état actif et on tente de maximiser la durée que passent les nœuds en veille.

- Le contrôle de la puissance de transmission qui consiste à maintenir la capacité du réseau et à acheminer le trafic de données avec un coût énergétique minimal. On permet aux nœuds de déterminer la puissance de transmission minimale suffisante pour maintenir la connectivité du réseau.

- La distribution de la charge d'acheminement des données dont l'objectif principal est d'équilibrer la consommation d'énergie entre les nœuds mobiles.

Bien entendu, les trois approches ne sont pas exclusives, d'ailleurs certains protocoles de routage à basse consommation d'énergie combinent ces différentes approches.

2.1 Nœud mobile et consommation d'énergie

Un nœud mobile possède typiquement plusieurs composants matériels qui consomment de l'énergie, à savoir: le processeur, le disque, l'écran et l'interface de communication sans fil, etc. L'interface sans fil consomme jusqu'à 50% de l'énergie globale du nœud mobile [22]. Les protocoles de routage à basse consommation d'énergie proposés dans la littérature pour les réseaux Ad hoc, cherchent soit à minimiser l'énergie dissipée lors des communications actives (durant les opérations d'émission et de réception, acheminement inclus) ou celle consommée dans les périodes inactives (quand l'interface sans fil n'effectue aucune communication).

2.2 Sources de perte d'énergie

Ils existent plusieurs sources de perte d'énergie dans un réseau Ad hoc. Quelques sources sont utiles tandis que d'autres sont considérées comme des pertes qui doivent être réduites ou éliminées. Certaines pertes d'énergie lors des communications sont dues aux facteurs suivants:

- **Le mode inactif:** L'interface sans fil de communication gaspille de l'énergie sans effectuer aucune tâche utile. Une solution possible est de mettre périodiquement les nœuds en veille.

► **Les collisions:** Elles surviennent surtout dans des conditions de trafic élevé. Les paquets affectés par les collisions ne sont pas correctement reçus et l'énergie consommée lors de leurs émissions et réceptions est gaspillée.

► **Le surcoût des protocoles:** Cela fait référence aux paquets de contrôle que génèrent les différents protocoles de communication, et qui imposent une consommation d'énergie supplémentaire à ce qui est strictement nécessaire à la transmission des données.

Notre objectif étant, la modélisation de la consommation d'énergie d'un nœud Ad hoc via le formalisme des réseaux de Petri (RdP) qui est les RdP stochastiques généralisés (RdPSG), on donnera quelques notions de bases de ce formalisme dans la section suivante.

3 Réseaux de Petri

Un RdP est un outil graphique et mathématique décrivant les relations entre des conditions et des événements d'un système étudié. Formellement, un *RdP* est un graphe bipartite orienté constitué de places, de transitions et d'arcs qui relient les transitions aux places et les places aux transitions. Formellement un *RdP* généralisé est un 4-uplet, $RdP = (P, T, Pre, Post)$, où:

- P est l'ensemble fini des places $\{p_1, p_2, \dots, p_m\}$, avec $m = |P|$;
- T est l'ensemble fini des transitions $\{t_1, t_2, \dots, t_n\}$, avec $n = |T|$;
- $Pre : P \times T \rightarrow \mathbb{N}$ est l'application d'incidence avant. $Pre(p, t)$ contient la valeur entière (valuation, poids) associée à l'arc allant de la place p à la transition t .
- $Post : P \times T \rightarrow \mathbb{N}$ est l'application d'incidence arrière; $Post(p, t)$ contient la valeur entière (valuation, poids, multiplicité) associée à l'arc allant de la transition t à la place p .

Définition 1 [RdP marqué]: Un marquage M_l d'un *Rdp* est une application, $M_l : P \rightarrow \mathbb{N}$, qui associe à chaque place $p_i \in P$ du *RdP*, $M_l(p_i)$ jetons (marques). Le marquage initial, noté M_0 , donne pour chaque place le nombre initial de jetons. Ainsi, le couple $\mathfrak{S} = (Rdp, M_0)$ est dit *RdP marqué*.

3.1 Les réseaux de Petri stochastiques généralisés

► Les RdPS (Réseaux de Petri Stochastiques) ou SPN (Stochastic Petri Nets) ont été introduits pour répondre à certains problèmes d'évaluation de performance faisant intervenir des phénomènes aléatoires. Ainsi, les transitions d'un RdP contiennent des temps de franchissements aléatoires.

► Les réseaux de Petri stochastiques généralisés (RdPSG) sont une extension des RdPS autorisant deux classes de transitions:

- Des transitions instantanées à temporisation nulle (transition immédiate), qui sont franchies immédiatement dès qu'elles sont sensibilisées;
- Des transitions temporisées à qui correspondent des variables aléatoires

déterminant la durée de franchissement.

La définition formelle d'un RdPSG est donnée par la définition suivante:

Définition 2 Un RdPSG est un huit-uplet $(P, T, Pré, Post, Inh, pri, W, M_0)$ [8], où:

- P est l'ensemble des places;
- T est l'ensemble des transitions temporisées et des transitions immédiates;
- $Pr, Post, Inh: P \times T \rightarrow \mathbb{N}$ sont les fonctions d'incidence avant, d'incidence arrière et d'inhibition respectivement;
- $pri: T \rightarrow \{0, 1\}$ est la fonction de priorité qui associe à chaque transition temporisée la valeur 0 et à chaque transition immédiate la valeur 1;
- $W: T \rightarrow \mathbb{R}^+$ est une fonction qui associe à chaque transition temporisée un taux de franchissement;
- $M_0: P \rightarrow \mathbb{N}$ est le marquage initial du réseau.

Graphiquement, les places d'un RdP sont représentées par des cercles, les transitions par des rectangles (ou des traits) et les arcs par des flèches. Les applications d'incidences avant Pr et arrière $Post$ sont représentées par des arcs pondérés et orientés entre une place et une transition. Les arcs de valuation nulle ne sont pas représentés, les valuations qui sont égales à 1 sont omises, dans tous les autres cas la valuation est explicitement indiquée. Le marquage initial est donné par un nombre à l'intérieur de chaque place ou des points lorsque le nombre est suffisamment petit. Dans la représentation graphique d'un RdPSG, les transitions immédiates sont représentées par des traits et les transitions temporisées par des rectangles.

4 Position du problème

Dans un réseau Ad hoc, il n'y a pas d'administration centralisée. Chaque nœud peut rejoindre le réseau ou le quitter à tout moment. En général, les nœuds qui forment ce réseau sont mobiles, donc pas de recharge instantanée de batterie. Un nœud dont la batterie est épuisée est considéré inutilisable, mais après la mise en tension, le nœud devient de nouveau fonctionnel.

Dans notre étude, on s'intéresse à la modélisation de l'énergie consommée par un nœud en fonction de ses différents états, sachant qu'un nœud fonctionnel peut se mettre en veille, un nœud en veille peut devenir fonctionnel, un nœud éteint se place en mise en tension et un nœud chargé redevient fonctionnel.

5 Modélisation du problème par un RdPSG

Pour modéliser un nœud d'un réseau Ad hoc on va faire appel à l'un des formalismes des RdP qui est les RdPSG, ce choix du formalisme est dicté par le fait que le processus régissant la dynamique de la consommation d'énergie est un processus stochastique markovien. Le modèle qu'on a obtenu via ce formalisme est illustré dans la figure (Fig. 1.), où:

- $P = \{P_1, P_2, P_3\}$ un ensemble de trois places telles que:
 - P_1 : représente l'état fonctionnel du nœud.
 - P_2 : représente l'état veille ou déchargé (éteint).
 - P_3 : représente l'état de mise en charge.
- $T = \{t_1, t_2, t_3, t_4\}$ est un ensemble de quatre transitions, telles que:
 - t_1 : représente le passage de l'état actif à l'état veille suivant une loi exponentielle, avec un taux de franchissement μ .
 - t_2 : représente le passage de l'état veille à l'état actif suivant une loi exponentielle, avec un taux de franchissement λ .
 - t_3 : représente le passage de l'état déchargé à l'état en charge suivant une loi exponentielle, avec un taux de franchissement α .
 - t_4 : représente le passage de l'état de charge à l'état actif suivant une loi exponentielle, avec un taux de franchissement β .

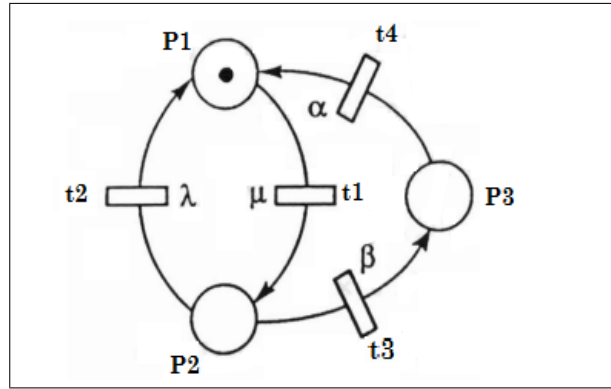


Fig. 1. Le RdPSG associé à la dynamique d'un nœud Ad hoc.

A partir de ce RdPSG, nous avons obtenu le graphe des marquages accessibles qui lui est associé, illustré dans (Fig. 2.) suivante:

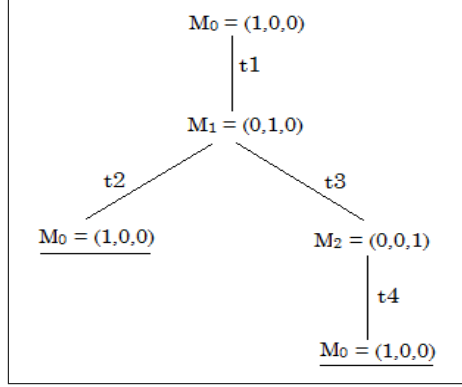


Fig. 2. Le graphe des marquages accessibles.

A partir de ce graphe des marquages on peut directement extraire la chaîne de Markov à temps continu (CMTC) régissant notre RdPSG. Les états de cette CMTC sont les marquages du graphe d'accessibilité. La CMTC qu'on a obtenue est illustrée dans (Fig. 3.):

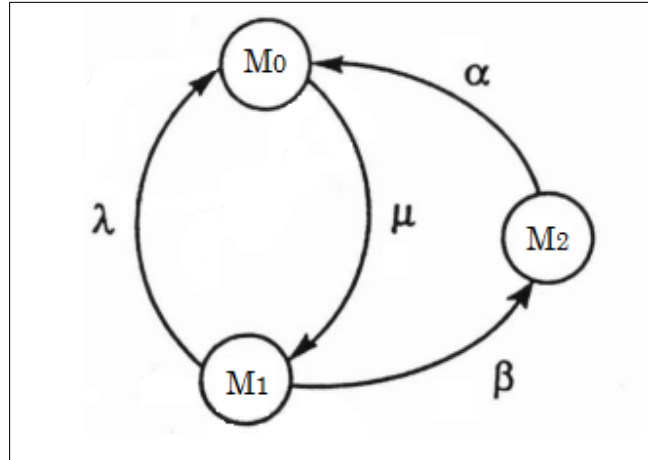


Fig. 3. La chaîne de Markov obtenue.

Le générateur infinitésimal associé à cette CMTC est défini par Q :

$$Q = \begin{pmatrix} -\mu & \mu & 0 \\ \lambda & -(\lambda + \beta) & \beta \\ \alpha & 0 & -\alpha \end{pmatrix}$$

La CMTC est ergodique, puisque c'est une chaîne finie et irréductible, donc la distribution stationnaire $\Pi = (\pi_1, \pi_2, \pi_3)$ de cette CMTC existe, on peut l'obtenir en résolvant le système d'équations linéaires suivant:

$$\begin{cases} \pi Q = 0; \\ \sum_{i=1}^3 \pi_i = 1. \end{cases}$$

On aura alors:

$$\begin{cases} \mu\pi_1 = \lambda\pi_2 + \alpha\pi_3; \\ (\lambda + \beta)\pi_2 = \mu\pi_1; \\ \alpha\pi_3 = \beta\pi_2; \\ \pi_1 + \pi_2 + \pi_3 = 1. \end{cases}$$

Après la résolution de ce système, on obtient:

$$\pi = (\pi_1, \pi_2, \pi_3) = \frac{1}{\alpha(\lambda + \mu) + \beta(\alpha + \mu)} \left(\alpha(\lambda + \beta), \quad \alpha\mu, \quad \beta\mu \right); \quad (1)$$

5.1 Calcul de la probabilité du temps durant lequel le nœud est fonctionnel

En s'inspirant du travail réalisé dans [9], on a obtenu la probabilité du temps durant lequel le nœud est fonctionnel à partir de la CMTC qu'on a obtenue via les RdPSG. En effet, d'après notre modélisation donnée dans (Fig. 4.), on a défini les états suivants :

ON: le nœud est opérationnel (état 1).

OFF: le nœud est en veille ou en charge (état 2 ou état 3).

Ainsi, l'état de déchargement de la batterie du nœud en fonction des deux états précédents est illustré dans (Fig. 4.) suivante:

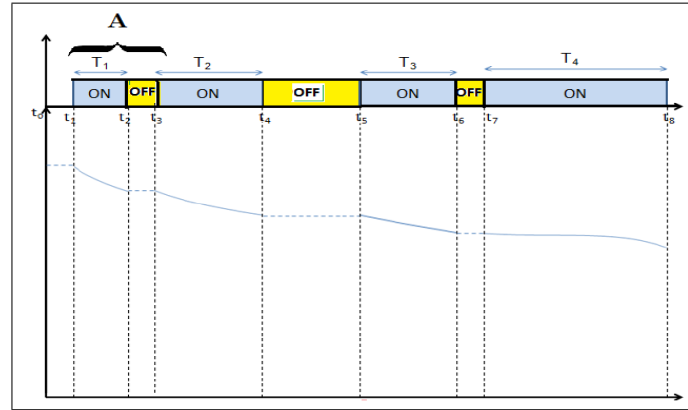


Fig. 4. État de la batterie.

Pour calculer la probabilité du temps durant lequel le nœud est fonctionnel, nous avons considéré:

$$P(t, t + T) = e^{T * a_{ij}}; \quad (2)$$

avec a_{ij} , les éléments de la Matrice A illustrée dans la figure (Fig.4), la probabilité de transition de l'état i à l'état $j \neq i$ sachant que le nœud a passé une période T à l'état initial.

On a :

$$P(t, t + T) = P(0, T) = e^{T * a_{ij}}.$$

Soient:

P_0 : la probabilité que le nœud soit à l'état opérationnel (état 1),

P_1 : la probabilité que le nœud soit à l'état veille ou en charge (état 2 ou état 3).

Ainsi, la probabilité de transition de l'état ON à l'état OFF est donnée dans la proposition suivante.

Proposition 1 La probabilité de transition de l'état ON à l'état OFF, notée P est donnée par:

$$P \simeq [(e^{-\lambda t} + e^{-\alpha t}) * (e^{-\mu T} + e^{-(\mu+\beta)T})]^n. \quad (3)$$

Preuve: On a :

$$\begin{aligned} P &= P_0(0, t_1) * P_1(t_1, t_1 + T_1) * P_0(t_1 + T_1, t_2) * P_1(t_2, t_2 + T_2) * P_0(t_2 + T_2, t_3) \\ &\quad * P_1(t_3, t_3 + T_3) * \dots * P_1(t_n, t_n + T_n) * P_0(t_n + T_n, t) \\ &= P_0(0, t_1) * P_1(0, T_1) * P_0(0, t_2 - t_1 - T_1) * P_1(0, T_2) * P_0(0, t_3 - t_2 - T_2) * P_1(0, T_3) \\ &\quad \dots * P_1(0, T_n) * P_0(0, t - t_n - T_n). \\ &= f_0(t_1) * f_1(T_1) * f_0(t_2 - t_1 - T_1) * f_1(T_2) * f_0(t_3 - t_2 - T_2) * f_1(T_3) * \dots * \\ &\quad f_1(T_n) f_0(t - t_n - T_n) \\ &= (e^{-\lambda t_1} + e^{-\alpha t_1}) * (e^{-\mu T_1} + e^{-(\mu+\beta)T_1}) * (e^{-\lambda(t_2-t_1-T_1)} + e^{-\alpha(t_2-t_1-T_1)}) * \\ &\quad * (e^{-\mu T_2} + e^{-(\mu+\beta)T_2}) * (e^{-\lambda(t_3-t_2-T_2)} + e^{-\alpha(t_3-t_2-T_2)}) * (e^{-\mu T_3} + e^{-(\mu+\beta)T_3}) * \\ &\quad \dots * (e^{-\mu T_n} + e^{-(\mu+\beta)T_n}) * (e^{-\lambda(t-t_n-T_n)} + e^{-\alpha(t-t_n-T_n)}). \end{aligned}$$

On pose:

$$A = (e^{-\lambda(t_1 + e^{-\alpha t_1})}) * (e^{-\mu T_1} + e^{-(\mu+\beta)T_1}).$$

$$B = (e^{-\lambda(t_2-t_1-T_1)} + e^{-\alpha(t_2-t_1-T_1)}) * (e^{-\mu T_2} + e^{-(\mu+\beta)T_2}).$$

$$C = (e^{-\lambda(t_3-t_2-T_2)} + e^{-\alpha(t_3-t_2-T_2)}) * (e^{-\mu T_3} + e^{-(\mu+\beta)T_3}).$$

$\vdots$

$$Z = (e^{-\mu T_n} + e^{-(\mu+\beta)T_n}) * (e^{-\lambda(t-t_n-T_n)} + e^{-\alpha(t-t_n-T_n)}).$$

On constate que $A \geq B$ et $A \geq C$, ..., $A \geq Z$, alors on peut approximer la probabilité P par A^n :

$$P \simeq [(e^{-\lambda t} + e^{-\alpha t}) * (e^{-\mu T} + e^{-(\mu+\beta)T})]^n.$$

On introduit une constante de normalisation C , telle que:

$$P(T = \theta < t) \leq C * [(e^{-\lambda t} + e^{-\alpha t}) * (e^{-\mu \theta} + e^{-(\mu+\beta)\theta})]^n.$$

La valeur de C est obtenue en calculant:

$$\int_0^t P(T = \theta < t) d\theta = 1.$$

■

Cas Particulier: Pour $n=1$, on a:

$$P(T = \theta < t) \leq C * (e^{-\lambda t} + e^{-\alpha t}) * (e^{-\mu \theta} + e^{-(\mu+\beta)\theta}).$$

$$C = \frac{1}{(e^{-\lambda t} + e^{-\alpha t}) * (\frac{1}{\mu}(1 - e^{-\mu t}) + \frac{1}{\mu+\beta}(1 - e^{-(\mu+\beta)t}))}.$$

On aura alors:

$$P(T = \theta < t) \leq \frac{(e^{-\mu \theta} + e^{-(\mu+\beta)\theta})}{\frac{1}{\mu}(1 - e^{-\mu t}) + \frac{1}{\mu+\beta}(1 - e^{-(\mu+\beta)t})}. \quad (4)$$

5.2 Simulation, résultats et interprétation

Pour valider les résultats obtenus via les RdPSG, nous avons simulé les états stationnaires de notre système en utilisant la matrice des probabilités de transition et ce en utilisant le langage de programmation Matlab.

Une comparaison des résultats obtenus par les deux approches (RdPSG et la chaîne de Markov simulée) est effectuée et elle est présentée dans (Table 1.) suivante :

Les paramètres	Les résultats de la simulation	Les résultats analytiques
$\mu=5, \lambda=4,$ $\beta=3, \alpha=2$	$\pi_1 = 0.3516; \pi_2 = 0.2581;$ $\pi_3 = 0.3903$	$\pi_1 = 0.3590; \pi_2 = 0.2564;$ $\pi_3 = 0.3846$
$\mu=1, \lambda=6,$ $\beta=2, \alpha=4$	$\pi_1 = 0.8435; \pi_2 = 0.1048;$ $\pi_3 = 0.0517.$	$\pi_1 = 0.8421; \pi_2 = 0.1053;$ $\pi_3 = 0.0526.$
$\mu=9, \lambda=3,$ $\beta=8, \alpha=1$	$\pi_1 = 0.1204; \pi_2 = 0.1045;$ $\pi_3 = 0.7751.$	$\pi_1 = 0.1196; \pi_2 = 0.0978;$ $\pi_3 = 0.7826.$

Table 1. comparaison des résultats

Les probabilités stationnaires sont interprétées comme suit:

- π_1 : est la fraction de temps pendant laquelle le nœud est fonctionnel,
- π_2 : est la fraction de temps pendant laquelle le nœud est en veille,
- π_3 : est la fraction de temps pendant laquelle le nœud est éteint.

D'après les résultats donnés dans (Table 1.), il est clair que les résultats obtenus par simulation et ceux obtenus analytiquement sont très très proches.

A partir de cette distribution stationnaire π , il est possible de calculer quelques indices de performance, tels que:

- $[(\lambda + \beta)\pi_2]^{-1} = [\mu\pi_1]^{-1}$: est le temps moyen entre deux instants consécutifs où le nœud est en veille,
- $[\alpha\pi_3]^{-1} = [\beta\pi_2]^{-1}$: est le temps moyen entre deux extinctions consécutives (ie, la durée de vie de la batterie entre deux mise en charge).

6 Conclusion

Dans ce travail, nous nous sommes intéressés à la gestion de la consommation d'énergie d'un nœud d'un réseau Ad hoc. Nous avons modélisé la dynamique du nœud par un modèle de processus "ON/OFF". L'analyse de ce modèle, nous l'avons faite en se basant sur le formalisme des RDPSPG (Réseaux de Petri Stochastiques généralisés). En effet, ce formalisme nous a permis d'extraire la chaîne de Markov (CM) associée à ce modèle. A partir de cette CM, nous avons calculé la distribution stationnaire du modèle pour obtenir enfin une majoration de la probabilité P du temps durant lequel le nœud est fonctionnel en fonction du nombre ' n ' de transitions entre l'état "ON" et l'état "OFF". Pour le cas où $n = 1$, on a même calculé cette probabilité avec un calcul exacte des constantes. Ces résultats obtenus, nous les avons validés par simulation.

Comme perspectives de travail, nous suggérons de considérer les points suivants:

- A court terme, il sera question de calculer d'autres caractéristiques telles que la probabilité du temps durant lequel le nœud est fonctionnel (son énergie est non nulle), cette caractéristique est plus intéressante puisque l'objectif de notre travail est l'étude de la consommation de l'énergie. Cependant, le calcul de cette caractéristique nécessitera une modélisation via les RdPSPG en considérant des places supplémentaires à notre modèle, qui nous permettront de pouvoir calculer les durées de temps séparant deux instants consécutifs dans un état.
- Il serait judicieux de modéliser le problème en considérant les états: veille et éteint séparément, avec l'un des formalismes des RdP.
- Il serait aussi intéressant de calculer la valeur exacte de la distribution de probabilité du temps de fonctionnement d'un nœud.

References

1. D. Bécaye. *"Effets de la mobilité sur les protocoles de routage dans les réseaux ad hoc"*. Mémoire d'ingénieur, Université Mouloud Mammeri, Tizi Ouzou, 2007.
2. A. Berrabah, H. Saidi. *"Balancement de charges dans les réseaux Ad Hoc"*. Mémoire de Master, Université Abou Bakr Belkaid, Tlemcen, 2013.
3. A. Bouzaher. *"Approche agent mobile pour l'adaptation des réseaux mobiles Ad hoc"*. Mémoire de Magister, Université Mohamed Khider, Biskra, 2011.
4. N. Boukhechem. *"Routage dans les réseaux mobiles Ad Hoc par une approche à base d'agents"*. Mémoire de Magister, Université de Constantine, 2008.
5. S. Chettibi. *"Protocole de routage avec prise en compte de la consommation d'énergie pour les réseaux mobiles Ad-hoc"*. Mémoire de Magister, Université de Constantine, 2008.
6. D. Djenouri, N. Badache. *"A Novel Approach for Selfish Nodes Detection in MANETs : Proposal and Petri Nets Based Modeling"*. Proceedings of the 8th International Conference on Telecommunications, 2005. ConTEL 2005.
7. A.C. Geniet. *"Les réseaux de Petri. Un outil de modélisation"*. Springer-Verlag Berlin Edition 2, 2006.
8. S. Hakmi. *"Evaluation des performances des systèmes prioritaires à l'aide des réseaux de Petri stochastiques généralisés"*. Mémoire de Magister, Université Abderrahmane Mira, Béjaia, 2011.

9. M. Heni, A. Bouallegue, R. Bouallegue. *"Energy consumption model in Ad hoc mobile network"*. International Journal of Computer Networks & Communications; May 2012, Vol. 4 Issue 3, p207
10. M. Heni, R. Bouallegue. *"Power control in reactive routing protocol for mobile Ad Hoc network"*. Higher School of Communication, Ariana, Tunisia, 2012.
11. A. Idrissi. *"How to minimise the energy consumption in mobile Ad Hoc networks"*. International Journal of Artificial Intelligence & Applications (IJAIA), Vol.3, No.2, March 2012.
12. Hanen Idoudi, Wafa Akkari, Abdelfatteh Belghith, Miklos Molnar. *Alternance synchrone pour la conservation d'énergie dans les réseaux ad hoc*. [Rapport de recherche] PI 1812, 2006, pp.44. <inria-00116073>
13. L. Ikhlef, S. Bouanani. *"Evaluation des performance du réseau $[M/M/1/2 \rightarrow \bullet/M/1/1]$ via les réseaux de Petri"*. Mémoire de Master, Université de Béjaia, 2011.
14. R.D. Joshi, P.P. Rege. *"Verification of energy efficient Optimised Link State Routing Protocol using Petri Net"*. Collège d'ingénierie de Pune, Inde 2011.
15. S. Khelifa, Z. Mekakia Maaza. *"Un protocole de routage ER-AODV à basse consommation d'énergie pour les réseaux mobiles Ad hoc"*. Université Mohamed Boudiaf USTO-MB Oran.
16. M.K. Molly. *"Performance analysis using stochastique Petri nets"*. 1982.
17. P. Prasad, B. Singh, A.K. Sahoo. *"Validation of Routing Protocol for Mobile Ad Hoc Networks using Colored Petri Nets"*. Institut national de Technologie Rourkela, 2009.
18. G. Pujolle. *"Minimisation de la consommation d'énergie dans les réseaux Ad hoc"*. Université Pierre et Marie Curie, 2005.
19. C. Ramchandani. *"Analysis of asynchronous concurrent systems by timed Petri nets"*. Combridge, 1974.
20. A. Shareef, Y. Zhu. *"Effective Stochastic Modeling of Energy-Constrained Wireless Sensor Networks"*. Université de Maine USA, 2012.
21. J. Sifakis. *"Etude du comportement permanent des réseaux de Petri temporisés"*. Edite par l'Institut de Programmation de Paris, 1977.
22. O. Smail. *"Routage multipath dans les réseaux Ad hoc"*. Thèse de Doctorat, Université Mohamed Boudiaf USTO-MB Oran, 2015.
23. C. Zhang, Z. Mengchu. *"A Stochastic Petri Net Approach to Modeling and Analysis of Ad Hoc Network"*. University Heights, 2003.
24. J. Zhu, C. Qiao, X. Wang. *"On Accurate Energy Consumption Models for Wireless Ad hoc Networks"*. Université de New York à Buffalo, 2006.

Nouvelle approche pour l'optimisation continue : Optimisation par Filtres Morphologiques

Khelifa Chahinez Nour El Houda¹, Belmadani Abderrahim²

²Laboratoire d'Optimisation des Réseaux Electriques

¹Département d'Informatique, Faculté des Mathématiques et Informatique
Université des sciences et de la technologie d'Oran – Mohamed Boudiaf
Bp 1505, Oran el M'Naouer, Algérie

Email : ¹khelifa.chahinez@gmail.com, ²abderrahim.belmadani@gmail.com
¹chahinez.khelifa@univ-usto.dz

Résumé. De nos jours, différentes approches pour le développement d'algorithmes d'optimisation ont été publiées. Ces approches sont soit inspirées des systèmes naturels telles que : les algorithmes génétiques, soit basées sur l'intelligence d'essaim : les colonies de fourmis, colonies d'abeilles, l'algorithme d'optimisation par Coucou. D'autres imitent certains processus définis par l'homme telles que le recuit simulé en métallurgie ou la recherche d'harmonie dans l'improvisation d'un groupe de musiciens de Jazz. Dans cet article, nous proposons une nouvelle approche inspirées des méthodes de traitement d'images. En effet, notre algorithme d'Optimisation par Filtres Morphologique (OFM) utilise les transformations morphologiques plus précisément l'érosion pour la recherche de l'optimum global dans un espace multidimensionnel. Nous avons testé notre algorithme OFM sur une série de fonctions tests de la littérature. Les résultats obtenus sont très probants ce qui montre que OFM pourrait avoir une place non négligeable dans l'ensemble des algorithmes d'optimisation.

Mots clés: optimisation, stochastique, érosion numérique, métaheuristiques, morphologie mathématique.

I. Introduction

Les métaheuristiques sont des algorithmes stochastiques avec un comportement généralement itératif conçues pour la résolution des problèmes d'optimisation. Pour la majorité, ces algorithmes sont inspirés des systèmes naturels, citons à titre d'exemple: les algorithmes génétiques qui ont été proposés par Holland en 1975[1] puis développé par De Jong [2] et Golberg[3]. Ces algorithmes qui imitent l'évolution biologique des espèces et s'inspirent de la théorie darwiniste de sélection naturelle pour améliorer la qualité des individus de la population. D'autre qui sont basés sur l'intelligence d'essaim : L'algorithme des colonies de fourmis qui a été introduit par Colomi et al[4], leurs idée est née du comportement des fourmis qui communiquent entre elles

par le dépôt de substances chimiques, appelées phéromones, sur le sol. L'algorithme de colonies d'abeilles qui a été proposée en 2005 par Karaboga [5] et qui s'inspire du comportement des abeilles mellifères qui ont un mécanisme de déplacement très efficace ce qui leur permet d'attirer l'attention d'autres abeilles de la colonie aux sources alimentaires trouvées dans le but de collecter des ressources diverses. L'algorithme d'optimisation Coucou (Cucko Search Optimization Algorithm) est une très jeune méta-heuristique proposé en 2010 par Xin-She Yang et Suash Deb [6] et qui s'inspire lui aussi du comportement de reproduction d'une espèce spéciale d'oiseaux parasites de nids appelés « Coucous ». La recherche tabou qui s'inspire de la mémoire des êtres humains et qui a été introduite principalement par Glover [7], Glover et Laguna [8] et qui fait appel à un ensemble de règles et de mécanismes généraux pour guider la recherche de manière intelligente [7]. Ainsi que le recuit simulé qui est inspiré d'un processus métallurgique, introduit par Kirkpatrick [9]. Cette méthode tire son principe de techniques utilisées par les métallurgistes qui, pour obtenir un alliage sans défaut, font alterner des cycles de réchauffage et de refroidissement lent des métaux. La recherche d'harmonie est une méta-heuristique récente proposée par Geem et ses collègues [10]. Cette approche qui s'inspire de la recherche de la meilleure harmonie dans un orchestre où chaque musicien joue une note.

II. L'algorithme proposé

Notre contribution consiste en une nouvelle approche stochastique d'optimisation. Cette approche s'inspire des transformations de la morphologie mathématique qui sont essentiellement utilisés en traitement d'images. En effet nous nous inspirons de la transformation dite EROSION à fin de parcourir l'espace de recherche et trouver l'optimum local dans le voisinage du filtre morphologique (communément appelé : élément structurant). Nous employons plusieurs filtres fonctionnant en parallèle, ceci garantissant l'intensification de la recherche dans notre algorithme. Nous avons aussi mis en place un mécanisme pour éviter le blocage au niveau des optimums locaux et d'examiner des points non visités précédemment ce qui conduit à la diversification dans notre recherche.

2.1 La morphologie mathématique

La morphologie mathématique a été développée par G. Matheron dans les années 60-70 puis par J. Serra et de son équipe [11]. C'est une théorie non linéaire qui peut être définie par l'étude des objets en fonction de leur forme, de leur taille, des relations avec leur voisinage, de leur texture, et de leurs niveaux de gris ou de leur couleur [11]. Grâce à son efficacité, la morphologie mathématique a été employée dans plusieurs domaines, citons à titre d'exemple son application aux images radar à synthèse d'ouverture (RSO) pour la différenciation des tissus urbains. Les travaux de Sheeren et al [12] qui ont proposé une approche fondée sur l'application et l'enchaînement automatiques des opérations issues de la morphologie mathématique pour l'extraction automatique des bâtiments dans les images satellitaires à très haute résolution. Sophie et son équipe [13] qui ont combinés les opérations de la morphologie mathématique et les critiques de cohérence spatiale pour la détection et la segmentation automatique de texte dans des images fixes ou issues de séquences vidéo. La morphologie mathéma-

tique a été aussi employé dans le domaine biomédicale en se basant sur ses opérations pour le filtrage et la segmentation d'image couleur ou multi-spectrales afin de développer des applications en microscopie biomédicale quantitative [14]. Elle aussi été appliqué par Risson [15] afin d'extraire les informations pertinentes sur les conditions d'éclairage dans lesquelles ont été prise les photos et les utiliser dans les divers domaines d'imagerie telle que la réalité augmenté, la post-production cinématographiques, l'indexation d'images... D'autres travaux ont combinés les transformations de la morphologie mathématique à la pré-topologie mathématique pour le traitement d'images satellitaires [16].

2.2 Opérateurs morphologiques fondamentales

Les opérateurs morphologiques cherchent à extraire à l'aide d'une comparaison ensembliste entre l'élément structurant et la surface pertinente d'une image à niveau de gris en considérant sa représentation en sous graphe.

Les deux opérateurs de base de la morphologie mathématique sont l'érosion et la dilatation. Ils sont notés respectivement $\varepsilon B(X)$ et $\delta B(X)$ et définis par les équations suivantes :

$$\varepsilon B(X) = \{z: B_z \subseteq X\} \quad (1)$$

$$\delta B(X) = \{z: B_z \cap X \neq \emptyset\} \quad (2)$$

Où :

L'érosion de X par B est le lieu des positions de l'origine z de l'élément structurant B_z quand celui-ci est inclus dans X.

Et Le dilaté de X par B est le lieu des implantations de l'origine z de l'élément structurant B_z quand celui-ci rencontre X.

Les équations (1) et (2) peuvent être généralisées à des transformations numériques en remplaçant les concepts ensemblistes par leurs équivalents fonctionnels. L'érosion numérique est définie comme suit :

$$\varepsilon B(f) = \inf\{f(y), y \in B_x\} \quad (3)$$

De même, la dilatation numérique est :

$$\delta B(f) = \sup\{f(y), y \in B_x\} \quad (4)$$

A partir de ces opérateurs, d'autres opérateurs sont définis comme des combinaisons d'érosions et de dilations. Ainsi, l'ouverture morphologique correspond à une érosion suivie d'une dilatation et la fermeture morphologique est définie comme une dilatation suivie d'une érosion:

$$\gamma(X) = (X \ominus B) \oplus B^t \quad (5)$$

$$\gamma(X) = (X \ominus B) \oplus B^t \quad (6)$$

3.2 Contribution

Comme précédemment stipulé, nous nous inspirons des transformations de la morphologie mathématique, plus précisément l'érosion. Nous considérons notre espace de recherche comme une image et faisons parcourir nos filtres comme des éléments structurants.

Un filtre appliqué à une solution de la fonction objective renvoie le minimum trouvé au niveau de ces voisins (le nombre de voisins à prendre en compte pour chaque solution est défini au préalable). Si l'un des voisins représente une meilleure solution le centre de l'élément structurant (filtre) est déplacé sur celui-ci, sinon le filtre qui a une taille définie au préalable voit sa taille réduite à fin de vérifier un voisinage plus proche. Notre critère d'arrêt est l'épuisement de la recherche dans tous l'espace. Ceci se traduit par une taille des filtres négligeable (proche de zéro).

Les étapes de l'algorithme peuvent être résumées en :(Fig 1).

1. Initialiser le nombre de filtre, nombre de voisins au niveau de chaque filtre, la dimension de l'espace de recherche.
2. Placer les filtres dans l'espace de recherche
3. calculer s'il y a un optimum local au niveau de chaque filtre.
 - Si oui : Se déplacer vers l'optimum local.
 - Si non : réduire la taille du filtre.
4. Recalcule des voisins dans le filtre:
5. Vérifier le critère d'arrêt :
 - Si oui : arrêter le processus et retourner l'optimum global.
 - Si non : répéter (3) (4) (5) jusqu'à atteindre le critère d'arrêt.

Afin de mieux expliquer notre approche, nous avons choisi de l'illustrer avec un exemple discret et un seul filtre :

Itération1. Positionnement du filtre, calcul des voisin**itération2.** Déplacement sur le Min

100	71	112	69	36	80	91	87	41	102
1024	214	186	35	78	95	110	95	75	198
123	1024	236	65	92	80	163	76	93	601
69	564	807	98	60	34	361	87	63	902
860	2564	201	50	83	65	22	86	29	86
95	2876	978	57	95	50	59	64	78	75
457	5001	264	87	56	78	847	71	94	76
1546	8796	546	861	46	769	410	87	64	58
2540	3640	756	463	34	324	475	91	71	72
2002	2010	547	159	65	61	795	83	87	64

100	71	112	69	36	80	91	87	41	102
1024	214	186	35	78	95	110	95	75	198
123	1024	236	65	92	80	163	76	93	601
69	564	807	98	60	34	361	87	63	902
860	2564	201	50	83	65	22	86	29	86
95	2876	978	57	95	50	59	64	78	75
457	5001	264	87	56	78	847	71	94	76
1546	8796	546	861	46	769	410	87	64	58
2540	3640	756	463	34	324	475	91	71	72
2002	2010	547	159	65	61	795	83	87	64

Itération 3. Réduction de la taille du filtre**Itération 4.** Réduction de la taille du filtre

100	71	112	69	36	80	91	87	41	102
1024	214	186	35	78	95	110	95	75	198
123	1024	236	65	92	80	163	76	93	601
69	564	807	98	60	34	361	87	63	902
860	2564	201	50	83	65	22	86	29	86
95	2876	978	57	95	50	59	64	78	75
457	5001	264	87	56	78	847	71	94	76
1546	8796	546	861	46	769	410	87	64	58
2540	3640	756	463	34	324	475	91	71	72
2002	2010	547	159	65	61	795	83	87	64

100	71	112	69	36	80	91	87	41	102
1024	214	186	35	78	95	110	95	75	198
123	1024	236	65	92	80	163	76	93	601
69	564	807	98	60	34	361	87	63	902
860	2564	201	50	83	65	22	86	29	86
95	2876	978	57	95	50	59	64	78	75
457	5001	264	87	56	78	847	71	94	76
1546	8796	546	861	46	769	410	87	64	58
2540	3640	756	463	34	324	475	91	71	72
2002	2010	547	159	65	61	795	83	87	64

Itération 5. Déplacement puis Arrêt

100	71	112	69	36	80	91	87	41	102
1024	214	186	35	78	95	110	95	75	198
123	1024	236	65	92	80	163	76	93	601
69	564	807	98	60	34	361	87	63	902
860	2564	201	50	83	65	22	86	29	86
95	2876	978	57	95	50	59	64	78	75
457	5001	264	87	56	78	847	71	94	76
1546	8796	546	861	46	769	410	87	64	58
2540	3640	756	463	34	324	475	91	71	72
2002	2010	547	159	65	61	795	83	87	64

Nous commençons par définir la taille de l'espace de recherche (10*10 dans ce cas). Ceci peut être réalisé grâce à la transformation : Ceci nous garantit que toutes les variables se verront localisées dans l'hyper-cube de côté R.

$$x' = \frac{x - x_{\min}}{x_{\max} - x_{\min}} R \quad (7)$$

Dans cet exemple, le nombre de voisin est fixé à 5. La première itération consiste en un choix aléatoire de la position du centre du filtre (Solution initiale générée aléatoirement) et la génération de ses voisins (la direction de voisinage varie aléatoirement à chaque itération). Le centre du filtre est la valeur 50 et les cinq voisins sont (22, 35, 860, 1546 et 846). L'étape suivante consiste à vérifier si l'un des voisins offre une meilleure solution (test : existe-t-il un optimum de la figure 1). Le voisin à la valeur 22 est donc automatiquement choisi comme nouveau centre du filtre (opération de déplacement). Il faut, ensuite, calculer ses voisins. Ce qui donne les valeurs (110, 198, 86, 58 et 410). A l'itération 3, aucun des voisins n'offre une meilleure solution que le centre du filtre ce qui conduit à une réduction de la taille de celui-ci. Les nouveaux voisins sont alors (163, 93, 29, 847 et 94). A l'itération 4, la même opération de réduction est opérée ce qui conduit aux nouvelles valeurs des voisins (361, 87, 6, 64 et 59). L'itération 5 est un déplacement du centre du filtre sur le minimum trouvé 6. N'ayant plus la possibilité de réduire la taille du filtre (ayant atteint un seuil préalablement défini), l'algorithme va s'arrêter et retourner la valeur minimale trouvée.

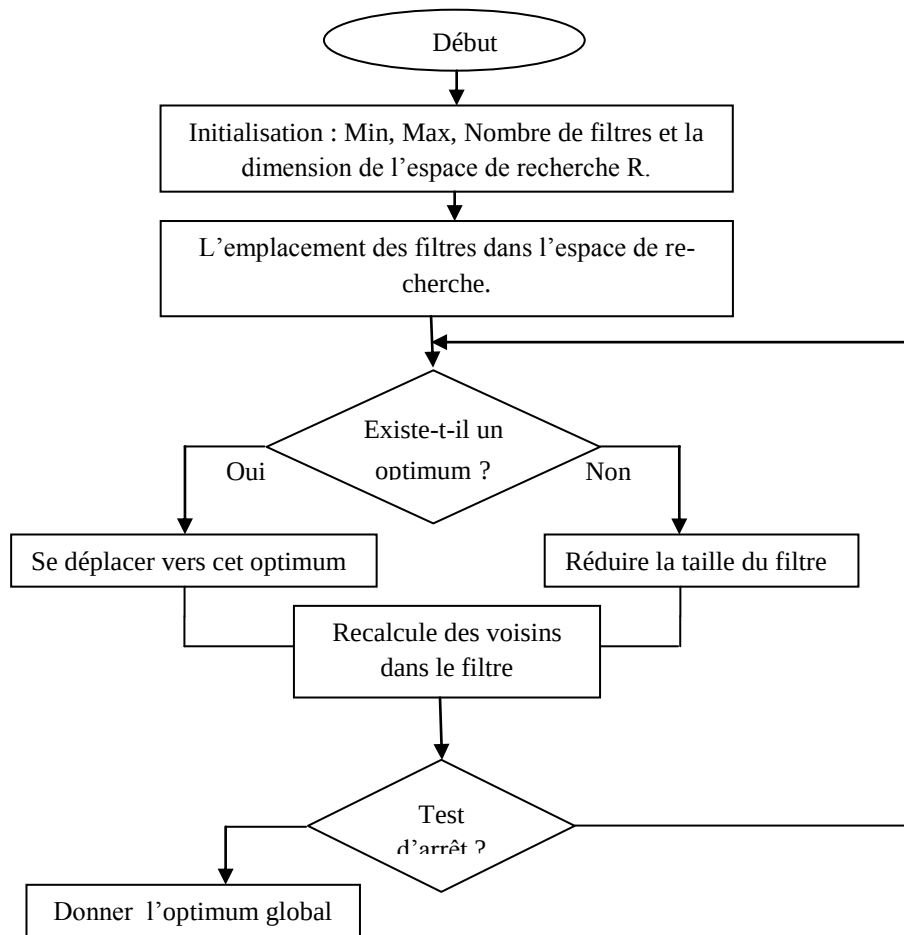


Fig. 1:organigramme général de l'algorithme OFM

1. Application et comparaison

Pour tester le bon fonctionnement de notre algorithme, nous l'avons appliqué sur les fonctions test suivantes :

3.1 La fonction Rosenbrock

$$f(x) = \sum_{i=1}^n (100 * (x_i - x_{i-1}^2)^2 + (x_{i-1} - 1)^2)$$

Pour $x_i \in [-30; 30]$ la valeur du minimum est 0

3.2 La fonction De Jong (sphérique)

$$f(x) = \sum_{i=1}^N x_i^2$$

Pour $x_i \in [-5.12; 5.12]$ la valeur du minimum est 0

3.3 Le problème de Schwefel

$$f(x) = \sum_{i=1}^N |x_i| + \prod_{i=1}^N |x_i|$$

Pour $x_i \in [-10; 10]$ la valeur du minimum est 0

1.4 La fonction Rotated hyper-ellipsoid

$$f(x) = \sum_{i=1}^N (\sum_{j=1}^i x_j)^2$$

Pour $x_i \in [-100; 100]$ la valeur du minimum est 0

1.5 La fonction de Schwefel 2

$$f(x) = \sum_{i=1}^N (x_i \sin(\sqrt{|x_i|}))$$

Pour $x_i \in [-500; 500]$ la valeur du minimum est -12569.5

3.6 La fonction Rastrigin

$$f(x) = \sum_{i=1}^N (x_i^2 - 10 \cos(2\pi x_i) + 10)$$

Pour $x_i \in [-5.12; 5.12]$ la valeur du minimum est 0

1.6 La fonction Six-Hump Camel-Back

$$f(x) = 4x^2 - 2.1x_1^4 + \frac{1}{3}x_1^6 + x_1x_2 - 4x_2^2 + 4x_2^4$$

Pour $x_i \in [-5; 5]$ la valeur du minimum est 0

3.8 La fonction de Michalewicz

$$f(x) = -\sum_{i=1}^N \sin(x_i) * \sin\left(\frac{i * x_i^{20}}{\pi}\right)$$

Pour $x_i \in [0; \pi]$ la valeur du minimum est -9.66

3.9 La fonction de Ackley's Path

$$f(x) = -20 \exp\left(-0.2 \sqrt{\frac{1}{30} \sum_{i=1}^N x_i^2}\right) - \exp\left(\frac{1}{30} \cos(2\pi x_i) + 20 + e^1\right)$$

Pour $x_i \in [-1; 1]$ la valeur du minimum est 0

3.10 La fonction de Griewank

$$f(x) = \frac{1}{4000} \sum_{i=1}^N x_i^2 - \prod_{i=1}^N \cos\left(\frac{x_i}{\sqrt{i}}\right) + 1$$

Pour $x_i \in [-600; 600]$ la valeur du minimum est 0

3.11 La fonction d'Easom

$$f(x) = -\cos x_1 * \cos x_2 * \exp\left(-((x_1 - \pi) + (x_2 - \pi))^2\right)$$

Pour $x_i \in [-100; 100]$ la valeur du minimum est -1

3.12 La fonction Golden-Price

$$f(x) = (1 + (x_1 + x_2 + 1)^2 * (19 - 14x_1 + 3x_1^2 - 14x_2 + 6x_1x_2 + 3x_2^2)) * (30 + (2x_1 - 3x_2)^2 * (18 - 32x_1 + 12x_1^2 + 48x_2 - 36x_1x_2 + 27x_2^2))$$

Pour $x_i \in [-2; 2]$ la valeur du minimum est 3

Le tableau suivant montre les résultats de 30 exécutions du programme, on respectant l'intervalle et la dimension de l'espace pour chaque fonction.

Table 1. Résultats de 30 exécutions.

Function	Best	Worst	Medium	Ecar-type	Space dim	computing time
Rosenbrock	0	0	0	0	30	9,968
De Jong	0	0	0	0	30	9,578
Schwefel's Problem	0	0	0	0	30	10,265
Rotated hyper-ellipsoid	0	0	0	0	30	18,518
Schwefel's function	-12569.5	-12567.9	-12568.71	1,124	30	17,659
Rastrigin	0	0	0	0	30	13,977
Six-Hump Camel-Back	0	0	0	0	2	1,202
Michalewicz's	-9.66	-2.27	-5,965906	5,224	10	23,681
Ackley's Path	0	6,51E-19	3,252E-19	4,59E-19	4	2,995
Griewank	0	0	0	0	30	18,128

Easom's	-1	-1	-1	0	2	13,463
Goldstein-Price's	3	3	3	0	2	1,061

Les simulations ont été réalisées à l'aide de l'environnement DELPHI, sur un Acer Intel Core (TM) i3-4005U (1,7 GHz) avec une mémoire totale de 4 Go. Pour toute les applications de notre algorithme nous avons utilisé un nombre de filtres NB-filtres=20 et une taille minimale des filtres (critère d'arrêt : Taille de chaque filtre inférieure ou égale à $\epsilon = 10^{-9}$). Nous avons pu atteindre l'optimum global pour chacune de ces fonctions test. Notons ici que la taille minimale des filtres est notre critère d'arrêt et que cette taille peut aller à 10^{-24} ou plus petite pour plus de précision.

De ce fait, les résultats concernant la fonction de MICHALEWICZ'S peuvent facilement être améliorés avec une meilleur précision. Nous avons omi de mentionner le temps de calcul qui pour une précision de 10^{-9} est négligeable. Ce temps de calcul devient estimable à mesure que la dimension de l'espace, le nombre de filtres et le nombre de voisins augmentent et à mesure que la précision devient plus fine.

2. Conclusion

Dans cet article, nous proposons une nouvelle approche d'optimisation inspirée des méthodes de traitement d'images. En effet, notre algorithme d'Optimisation par Filtres Morphologique (OFM) s'inspire des transformations morphologiques plus précisément de l'érosion pour la recherche de l'optimum global dans un espace multi-dimensionnel.

Nous avons testé notre algorithme OFM sur une série de fonctions tests de la littérature. Les résultats obtenus sont très probants ce qui montre que OFM pourrait avoir une place non négligeable dans l'ensemble des algorithmes d'optimisation.

Nous espérons adapter notre algorithme pour le traitement des problèmes d'optimisation discrets. Ensuite, nous nous tournerons vers les problèmes plus complexes de l'ingénierie.

Références

1. J.H. Holland. Genetic Algorithms and the optimal allocation of trials, SIAM Journal of Computing. Vol. 2, N° 2, pp. 88-105, 1973.
2. K.A De Jong. An analysis of the behavior of a class of genetic adaptative systems, Phdthesis, University of Michigam, Dissertation Abstract International 36(10), 5140B, University Microfilms No. 76-9381, 1975.
3. D.E. Goldberg. Genetic Algorithms in Search, Optimization and Machine Learning, ISBN: 0201157675, Addison-Wesley, 1989.
4. A. Colomi, M. Dorigo, V. Maniezzo. Distributed Optimization by Ant Colonies. Proceedings of the 1rst European Conference on Artificial Life. pp 134-142, Elsevier Publishing.
5. D. Karaboga. An idea based on honeybee swarm for numerical optimization. Technical Report TR06, Erciyes University, Engineering Faculty, Computer Engineering Department, 2005.
6. Yang, X.-S., and Deb, S. (2010), "Engineering Optimisation by Cuckoo Search", Int. J. Mathematical Modelling and Numerical Optimisation, Vol. 1, No. 4, 330–343 (2010).
7. F. Glover. Future paths for integer programming and links to artificial intelligence. Computers and Operations Research. Vol. 13, N° 5, pp 533-549, 1986.
8. F. Glover and M. Laguna. Tabu Search. Kluwer Academic Publishers, Boston. 1997.
9. S. Kirkpatrick, C. D. Gelatt, M. P. Vecchi. Optimization by simulated annealing. Journal of Science. Vol. 220, N° 4598, pp. 671-680, 1983.
10. Z. W. Geem, J. H. Kim and G. V. Loganathan. A New Heuristic Optimization Algorithm: Harmony Search. Simulation. Vol. 76, N° 2, pp. 60-68, 2001.
11. J. Serra, Image Analysis and Mathematical Morphology, Academic Press, New-York, 1982.
12. D. Sheeren, S. Lefèvre et al, La morphologie mathématique binaire pour l'extraction automatique des bâtiments dans les images THRS, Revue internationale de Géomatique. Vol 17, N° 3-4, 2007
13. S. Schüpp, Y. Chahir, and A. Elmoataz. Détection et extraction automatique de texte en vidéo: une approche par morphologie mathématique.
14. Angulo, J and Serra, J Automatic analysis of DNA microarray images using mathematical morphology. Bioinformatics, 19(5): 533-562, 2003.
15. Valéry Risson, Application de la morphologie mathématique à l'analyse des conditions d'éclairage des images couleur. Thèse de doctorat.
16. A. Belmadani, Combinaison de la morphologie mathématique et la pré-topologie mathématique pour le traitement d'images satellitaires. Thèse de Magister. USTO. Juin 2000.

Fusion de minuties pour une reconnaissance efficiente des empreintes digitales

Mohamed Hédi Ghaddab¹, Khaled Jouini², and Ouajdi Korbaa³

¹ ghaddab.mohamedhedi@gmail.com

² j.khaled@gmail.com

³ Ouajdi.Korbaa@centraliens-lille.org

Laboratoire MARS* LR17ES05

Institut Supérieur d'Informatique et des Techniques de Communication (ISITCom)
Université de Sousse, Tunisie

Résumé La reconnaissance efficiente des empreintes digitales est tributaire de la qualité des empreintes de référence (pré-stockées), auxquelles sont comparées les empreintes requêtes. Une des techniques classiques pour améliorer la qualité des empreintes de référence est de fusionner plusieurs impressions d'une même empreinte en une seule.

Une empreinte se caractérise par un ensemble de points, dits *minuties* correspondant à la terminaison ou à la bifurcation de ses crêtes. Du fait de leur haut pouvoir discriminant, les minuties sont à la base de la majorité des algorithmes de comparaison d'empreintes.

Cet article propose une nouvelle approche de fusion, intitulée FZC. FZC se distingue de la majorité des approches existantes par (i) l'application d'alignements différents sur différentes régions de l'empreinte; et (ii) une estimation de la crédibilité d'une minutie qui ne se base pas sur sa fréquence dans les différentes impressions, mais sur la qualité de la capture de la région dans laquelle la minutie se trouve. Les résultats expérimentaux obtenus confirment le bien-fondé de l'approche proposée.

Keywords: Comparaison d'empreintes; Super-modèle; Fusion de minuties.

1 Introduction

Le traitement automatisé d'empreintes digitales est au cœur de nombre d'applications sensibles, des plus classiques telles que celles ayant trait à la lutte contre la criminalité, aux plus récentes comme la sécurité d'accès aux données dans le contexte des réseaux ouverts. Le traitement automatisé des empreintes digitales suit sommairement les étapes suivantes : acquisition, extraction de caractéristiques et de descripteurs, stockage et comparaison (*i.e.* vérification) ou recherche (*i.e.* identification). Dans la suite nous nous intéressons à la comparaison automatisée d'empreintes digitales.

Une empreinte digitale se caractérise par un ensemble de crêtes et de points singuliers, dits *minuties* correspondant à la terminaison ou à la bifurcation des crêtes. Une minutie m est communément décrite par un quadruplet $m =$

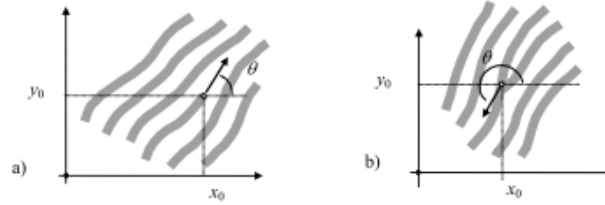


FIGURE 1. Grands types de minuties : (a) terminaison ; (b) Bifurcation. [11].

(x, y, θ, t) , où (x, y) correspondent à sa localisation spatiale, θ à l'angle formé par l'axe horizontal et la tangente à la ligne de crête au point de minutie (figure 1) et t à son type (bifurcation ou terminaison). Bien qu'il existe différentes approches pour la comparaison d'empreintes digitales, celle basée sur les points de minuties demeure largement la plus répandue [13]. La comparaison de deux empreintes basée sur leurs minuties revient sommairement à superposer les minuties (*i.e.* alignement) et à calculer un score en fonction des minuties concordantes.

Pour une même empreinte, il est extrêmement difficile que deux captures différentes aboutissent à une même représentation. Cette variabilité intra-classe s'explique par des propriétés intrinsèques et extrinsèques de l'empreinte. Les différentes techniques de capture d'empreintes ont pour propriété commune de transformer l'objet tridimensionnel qu'est l'empreinte en une image bidimensionnelle. Cette perte de dimension, conjuguée à des propriétés intrinsèques de l'empreinte comme l'élasticité de la peau et la différence entre les forces appliquées lors de l'apposition du doigt, fait qu'il soit extrêmement difficile que des captures différentes d'une même empreinte aboutissent à une localisation spatiale identique des minuties. Les propriétés extrinsèques ont trait aux erreurs introduites lors de l'extraction de minuties. Dans le cas d'images d'empreintes de mauvaise qualité ou d'empreintes obtenues par encrage du doigt, les algorithmes de détection de minuties peuvent introduire un grand nombre de fausses minuties (*i.e.* *spurious minutiae*) et ne pas être en mesure de détecter toutes les vraies minuties (*i.e.* *missing minutiae*). Les minuties déplacées, manquantes ou fausses sont la principale source de faux rejets et d'erreurs de vérification.

La qualité des empreintes de référence pré-stockées, auxquelles sont comparées les empreintes requêtes, a une incidence importante sur l'efficacité de tout système de reconnaissance d'empreintes. Une des techniques classiques pour améliorer la qualité des empreintes de référence est de fusionner plusieurs impressions d'une même empreinte en une seule. Sans perte de généralité, nous utilisons dans la suite le terme modèle (*template*) pour désigner l'ensemble de minuties extraites d'une impression d'empreinte. La fusion de plusieurs modèles d'une même empreinte en un seul modèle, dit *super-modèle* (*super template*), a deux grands objectifs : la consolidation du modèle de référence (*template consolidation*) et l'adaptation du modèle de référence aux variations que peut subir une empreinte au fil du temps suite par exemple à des coupures ou des cicatrices causées par des blessures (*template learning* ou *template adaptation*) [17].

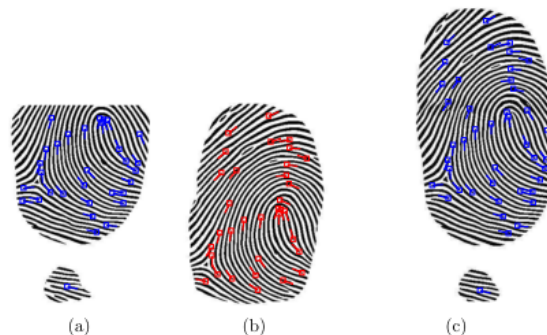


FIGURE 2. Fusion au niveau capteur : (a) Impression 1. (b) Impression 2. (c) Image composite [15]

Cet article propose une nouvelle approche pour la construction de super-modèles, intitulée FZC. Les difficultés majeures rencontrées lors de la construction d'un super-modèle sont l'estimation de la "crédibilité" d'une minutie et l'alignement des impressions à fusionner lorsque ceux-ci sont fortement affectés par les distorsions. L'approche que nous proposons part du constat que les distorsions et les déformations n'affectent pas d'une manière uniforme les différentes régions d'une empreinte. Elle se distingue de la majorité des approches existantes par l'application d'alignements différents sur différentes régions de l'empreinte. L'approche proposée se distingue également par une estimation de la crédibilité d'une minutie qui ne se base pas sur sa fréquence dans les différentes impressions, mais sur la qualité de la capture de la région dans laquelle la minutie se trouve. Les résultats expérimentaux obtenus sur les bancs de test standards, confirment le bien-fondé de l'approche proposée.

Dans la suite de l'article, la section 2 donne une brève revue des principales approches de fusion d'empreintes. La section 3 présente l'approche FZC. La section 4 étudie expérimentalement les performances de FZC. Suit la conclusion.

2 Principales approches de fusion d'impressions d'une empreinte

En biométrie le terme "fusion au niveau capteur" (*Sensor-Level*) est utilisé pour désigner la combinaison de données biométriques brutes pour former une caractéristique biométrique composite. Le terme "fusion au niveau caractéristiques" (*Feature-Level*) est communément utilisé pour désigner la combinaison des vecteurs de caractéristiques obtenus de différentes sources (différents capteurs, échantillons ou modalités biométriques). La fusion à ce niveau pré-traite les données brutes pour en extraire des descripteurs, puis combine les descripteurs extraits en un seul.

Dans le cas précis des empreintes digitales, la fusion au niveau du capteur se matérialise le plus souvent par la fusion d'images correspondant à différentes

captures d'une même empreinte. La fusion des caractéristiques se matérialise le plus souvent par la combinaison de minuties extraites de différentes impressions. Dans la suite de cette section nous présentons les principales approches de fusion d'images et de minuties d'une même empreinte⁴.

2.1 Fusion d'images

Jain et Ross [7] ont proposé une approche de "mosaïquage" (*mosaicking*) fusionnant les différentes images d'une empreinte en une seule (voir figure 2). L'approche proposée utilise un algorithme de comparaison basé sur les minuties pour trouver le bon alignement entre les images à fusionner. Les images sont ensuite rendues compatibles en normalisant les contrastes et les intensités des pixels. L'image finale est obtenue par la superposition de toutes les (sous-)images alignées, lissées et normalisées.

Plusieurs variantes de la technique de mosaïquage ont été proposées [12,3,19,4]. La plupart de ces approches diffèrent soit dans la manière d'aligner et d'harmoniser les images soit dans le domaine d'application du mosaïquage. À titre d'exemple, dans [3] les auteurs s'intéressent aux capteurs à petites surfaces et ont pour objectif d'augmenter la surface de l'empreinte couverte par l'image de référence. Les auteurs recommandent la capture de différentes parties de l'empreinte et d'utiliser le mosaïquage ensuite pour recomposer une image complète de l'empreinte. Zhang et al. [19] appliquent la technique de mosaïquage aux capteurs à balayage d'empreintes. L'approche consiste à fusionner les images obtenues par des balayages différents d'une empreinte pour améliorer la qualité de l'image de référence destinée à être stockée.

Pour des raisons de confidentialité et de maîtrise du coût du stockage, plusieurs systèmes de reconnaissance ne stockent pas les images d'empreintes et se contentent d'enregistrer les descripteurs qu'ils en extraient. Ceci constitue une limitation importante de la technique de fusion d'images. Par ailleurs, les différentes études menées pour comparer la fusion d'images et la fusion de caractéristiques, telles que celle de [15] et de [5], ont montré que bien que les deux approches améliorent la précision de la reconnaissance d'empreintes, la fusion de caractéristiques surclasse la fusion d'images.

2.2 Fusion de minuties

La fusion de minuties extraites d'impressions différentes d'une même empreinte a pour objectif la construction d'un super-modèle (*template synthesis*), plus distinctif et plus facile à comparer que les impressions prises individuellement. Yau et al. [18] partent de l'hypothèse que le module d'extraction de minuties n'introduit que peu de fausses minuties, mais qu'il est le plus souvent incapable de détecter toutes les vraies minuties. L'approche proposée aligne les

4. Il existe une troisième alternative, dite "fusion des scores", consistant à appliquer différents algorithmes de comparaison, puis à normaliser et combiner les différents scores en un seul. La fusion de scores sort du cadre de cet article.

différentes impressions d'une empreinte, puis fait l'union des minuties extraites de chaque impression. Au fur et à mesure de l'union, les attributs (localisation spatiale, angle d'orientation et type) des minuties communes à différentes impressions sont corrigés. L'approche proposée par [18] améliore donc la tolérance aux minuties déplacées et manquantes, mais reste vulnérable aux minuties parasites. L'approche ne permet pas non plus l'adaptation de la représentation de référence aux variations que peut subir une empreinte au fil du temps.

Ramoser et al. [14] proposent une approche similaire à celle de Yau et al. [18]. Elle se distingue cependant par une méthode originale d'alignement des impressions basée sur l'algorithme RANSAC⁵. Tout comme l'approche proposée dans [18], celle proposée dans [14] ne permet pas la suppression de fausses minuties et l'adaptation continue de la représentation de référence d'une empreinte.

Jiang et Ser [8] ont proposé une extension à l'approche de Yau et al. [18] permettant la prise en compte de la fiabilité d'une minutie et l'adaptation continue du modèle de référence d'une empreinte. Dans [8] les impressions sont collectées au fur et à mesure que des empreintes requêtes sont soumises au système (adaptation en ligne). Plus une minutie se répète dans les impressions collectées d'une empreinte, plus elle est jugée fiable. Chaque fois qu'une impression de référence est comparée avec une impression requête de la même source, la fiabilité associée aux minuties communes est incrémentée de un dans le modèle de référence. Au fil de la collecte des impressions, les vraies minuties finissent par obtenir les plus hauts scores de fiabilité. Lors de la comparaison d'empreintes, les minuties aux plus bas scores sont considérées comme parasites et sont ignorées.

Jiang et Ser [8] proposent également de corriger les attributs d'une minutie, au fur et à mesure de la collecte des impressions. La valeur donnée à un attribut est ainsi la moyenne pondérée des valeurs observées de cet attribut dans les différents échantillons. Comme le reporte les auteurs, ces différentes améliorations conduisent à : (i) l'élimination progressive des minuties parasites; (ii) la correction des valeurs d'attribut des minuties; et (iii) à une plus grande précision dans la comparaison d'empreintes. Cette approche n'est cependant probante que lorsque le nombre d'échantillons collectés par empreinte est significatif : 2 ou 3 impressions ne sont pas suffisantes pour aboutir à des corrections ou à des scores de fiabilité significatifs.

L'idée d'associer un score de fiabilité aux minuties en fonction de leurs nombres d'occurrences dans différents échantillons, a été reprise par Uz et al. [17] dans leur approche intitulée comparaison hiérarchique (*i.e. Hierarchical Matching*). Contrairement à Jiang et al. [8] qui ignorent les minuties à faibles scores de fiabilité lors de la comparaison d'empreintes, Uz et al. [17] considèrent toutes les minuties détectées et se servent des scores de fiabilité comme coefficients de pondération.

5. RANSAC (*RANdom SAmple Consensus*) est un algorithme itératif qui s'utilise lorsque l'ensemble de données observées est susceptible de contenir des valeurs aberrantes. RANSAC permet de générer un alignement qui ne tient compte que des valeurs pertinentes.

Hierarchical Matching examine les minuties de manière hiérarchique selon leurs scores de fiabilité. Une triangulation de Delaunay différente est construite pour chaque sous-ensemble de minuties de même fiabilité. L'objectif de cette représentation hiérarchique est de trouver la meilleure transformation permettant d'aligner une impression donnée à la représentation de référence de l'empreinte. L'alignement de l'échantillon se fait tout d'abord par rapport aux (triangles de) minuties avec le plus haut score de fiabilité (ensemble supposé ne pas contenir de minuties parasites). Cet alignement est raffiné progressivement en considérant tour à tour les (triangles de) minuties de score plus faible.

La comparaison hiérarchique souffre de la même limitation que celle de proposée dans [8] : elle n'est probante que lorsque le nombre d'échantillons collectés par empreinte est significatif. Par ailleurs, elle se base sur la triangulation de Delaunay, connue pour être vulnérable aux minuties manquantes et parasites [11].

3 Alignement de régions compatibles pour la construction de super-modèles

3.1 Vue d'ensemble

Dans cette section nous présentons notre nouvelle approche de construction de super-modèle par fusion de minuties, intitulée FZC (pour Fusion de Zones Compatibles). L'objectif de FZC est d'améliorer la qualité de l'impression de référence d'une empreinte par : (i) la restauration des minuties manquantes ; (ii) l'élimination des minuties parasites ; et (iii) l'augmentation de la surface de l'empreinte couverte par l'impression de référence.

L'approche FZC se distingue de celles existantes en deux points clés : l'estimation de l'authenticité d'une minutie et l'alignement des impressions de l'empreinte. La majorité des approches existantes se basent sur la fréquence (*i.e.* nombre d'occurrences) d'une minutie dans les différentes impressions pour juger de son authenticité. La fréquence d'une minutie peut se révéler insuffisante lorsqu'il n'est pas possible de collecter un nombre significatif d'impressions ou bien lorsque les impressions collectées sont de mauvaise qualité ou partielles. Dans notre approche, l'authenticité d'une minutie est estimée essentiellement par la qualité de la capture de la région dans laquelle se trouve la minutie. Si une minutie se trouve dans une région affectée par une forte distorsion (*e.g.* orientations ou courbures aberrantes des lignes de crêtes) ou bien dans une région obscurcie de l'image, la minutie est jugée peu fiable. Autrement, elle est jugée fiable.

Plusieurs algorithmes et outils permettent d'associer une mesure de qualité aux minuties extraites d'une empreinte [6,2,16]. Cette mesure de qualité est considérée comme donnée dans notre approche. Nous représentons donc une minutie m par un quadruplet $m = (x, y, \theta, t, Q)$, où Q correspond à la mesure de qualité normalisée fournie par l'algorithme d'extraction de minuties.

La première impression d'une empreinte est dite impression ou modèle de référence. Chaque fois qu'une nouvelle impression de la même empreinte est disponible (lors de l'enregistrement ou ultérieurement), les minuties de l'impression

en entrée sont fusionnées avec les minuties de l'impression de référence. Le résultat de cette fusion est un nouveau modèle de référence, dit super-modèle, synthétisant l'ancien modèle de référence et le modèle en entrée. Le modèle de référence et le modèle en entrée sont désignés dans la suite respectivement par *Ref* et *I*. La fusion des minuties de *Ref* et *I* se passe en deux étapes : (i) appariement des minuties et détermination des régions locales compatibles ; et (ii) alignement des minuties de chaque paire de régions compatibles et consolidation.

3.2 Régions locales compatibles

L'appariement de *Ref* et *I* donne lieu à trois groupes de minuties : (i) les minuties de *Ref* qui n'ont pas de correspondant dans *I* ; (ii) les minuties de *I* qui n'ont pas de correspondant dans *Ref* ; et (iii) l'ensemble noté $M = \{(p_i^R, q_j^I) \mid p_i^R \in Ref, q_j^I \in I\}$ de paires de minuties de *Ref* et *I* qui concordent.

Sommairement, l'idée générale derrière FZC est de trouver les paires de régions locales de *I* et *Ref* qui correspondent selon les paires de M . Ces régions sont dites *régions compatibles*. Ensuite, pour chaque paire de régions compatibles, nous cherchons l'alignement local permettant d'inclure dans le modèle final, les minuties de *I* qui n'ont pas de correspondant dans *Ref*.

Une région locale se définit par une minutie centrale et un certain nombre de minuties qui se trouvent aux-alentours. Nous considérons que deux régions de *Ref* et *I* sont compatibles (correspondent), si elles contiennent un nombre Th_n de minuties appariées dans un rayon inférieur à Th_c . Une paire Z_i de régions compatibles se définit par

$$Z_i = (center_i^R, center_i^I, M_i) \quad (1)$$

où $center_i^R$ est la minutie centrale de la région dans l'impression *Ref*, $center_i^I$ la minutie centrale de la région correspondante dans *I* et M_i l'ensemble de paires de minuties des deux régions qui concordent.

Nous utilisons l'algorithme 1 (présenté ci-après) pour déterminer les régions compatibles de *Ref* et *I*. Cet algorithme est inspiré de l'algorithme de partitionnement en K-moyennes (*i.e.* K-means). Il construit progressivement des groupes (*clusters*) de Th_n minuties se trouvant à une distance de moins de Th_c d'une minutie centrale, qui s'ajuste chaque fois qu'une minutie est ajoutée au groupe. Une minutie n'est ajoutée à un groupe donné, que si elle apparaît à la fois dans *I* et *Ref*. Les groupes construits ainsi, sont nos régions compatibles.

Les figures 3.a, 3.b et 3.c donnent une illustration de l'application de l'algorithme 1 sur deux impressions réelles d'une même empreinte. Les ellipses de la figure 3.c correspondent aux régions compatibles obtenues à partir des impressions 3.a et 3.b. Les régions compatibles contiennent des minuties de 3.a et 3.b qui se superposent.

3.3 Alignement et fusion des régions compatibles

Comme indiqué précédemment, FZC part des constats que : (i) les distorsions n'affectent pas de manière uniforme les différentes régions d'une empreinte ; (ii)

Entrées : Ensemble M de paires de minuties de Ref et I qui se superposent

Sorties : Ensemble de régions compatibles Z

```

foreach Paire  $(p_i, q_j)$  in  $M$  do
  if  $(p_i, q_j)$  n'a pas été traitée then
    créer une nouvelle région  $Z_k$ ;
     $center_k^R = p_i; center_k^I = q_j$ ;
    Ajouter  $(p_i, q_j)$  à  $M_k$ ;
     $(p_i, q_j)$  est marquée traitée;
    repeat
       $(p'_i, q'_j)$  un autre paire de  $M$  non traitée;
      if  $distance(center_k^R, p'_i) < Th_c$  ET  $distance(center_k^I, q'_j) < Th_c$ 
        then
          Ajuster les centres;
          Ajouter  $(p'_i, q'_j)$  à  $M_k$ ;
          Marquer  $(p'_i, q'_j)$  comme traitée;
        end
    until Nombre de paires dans  $M_k < Th_n$ ;
    Ajouter  $z_k$  à  $Z$ ;
  end
end

```

Algorithm 1: Algorithme de partitionnement des impressions en régions locales compatibles.

les distorsions affectent moins les régions locales d'une empreinte que l'empreinte elle-même prise à un niveau globale.

La majorité des algorithmes de comparaison basés sur les minuties finissent à l'étape de consolidation par appliquer un alignement global pour apparier les minuties [13]. Cet alignement global, aussi bon soit-il, n'est pas en mesure d'aligner toutes les minuties qui devraient l'être, compte tenu du fait que des régions différentes ont des distorsions différentes et nécessitent donc des alignements différents. Contrairement à la majorité des approches existantes, FZC procède à des alignements différents pour les régions compatibles.

Le but de l'alignement d'une paire de régions compatibles $Z_i = (z_{Ref}, z_I)$ est de trouver une transformation locale plus fine que celle globale appliquée par l'algorithme de comparaison. Les paramètres de la transformation locale (rotation et translation) sont déduits de la différence entre les localisations spatiales de la minutie $p \in z_{Ref}$ ayant la plus haute fiabilité Q et la minutie correspondante $q \in z_I$. Une fois z_{Ref} et z_I finement alignées et que les localisations spatiales des minuties de z_I sont corrigées en fonction de la transformation locale qui en résulte, les opérations suivantes ont lieu :

1. Si une minutie q_i de z_I est appariée à une minutie p_i de z_{Ref} (i.e. $(p_i, q_i) \in M_i$) et que q_i est de meilleure qualité que p_i , alors les attributs de q_i remplacent ceux de p_i dans le modèle de référence. Si p_i et q_i ont des qualités proches à un certain degré, la valeur de chaque attribut de p_i est remplacée par la moyenne pondérée par la qualité, des valeurs de l'attribut dans p_i et q_i .

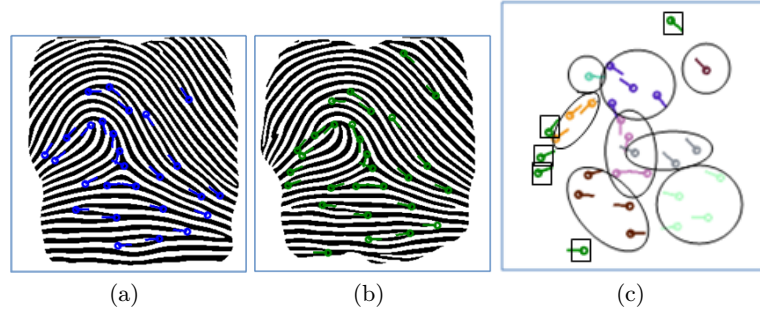


FIGURE 3. (a) Minuties extraites de l'empreinte 12_6 de la base *FVC2000 DB1_A*. (b) Minuties extraites de l'empreinte 12_3 de la même empreinte. (c) Régions compatibles et minuties sans correspondances.

2. Si une minutie q de z_I n'est appariée à aucune minutie de z_{Ref} et que la mesure de qualité qui lui est associée dépasse un certain seuil Th_Q , q est ajoutée au modèle de référence de l'empreinte, contribuant ainsi à la restauration des minuties manquantes.
3. Si une minutie p de z_{Ref} n'est appariée à aucune minutie de z_I et que la mesure de qualité qui lui est associée est en deçà de Th_Q , p est supprimée du modèle de référence.

4 Étude expérimentale

Les bancs de tests FVC sont communément utilisés pour l'évaluation des algorithmes de comparaison d'empreintes [11]. Pour valider notre approche, nous avons conduit des tests sur les quatre bases du FVC2000 [10]. Notre choix pour ces bases s'explique par le fait que les images qu'elles contiennent ont été acquises avec des scanners d'empreintes à petites surfaces de capture, sans contrôle de qualité et sans que la surface du capteur soit systématiquement nettoyée [11]. La qualité des images de la base DB1_A du FVC 2000 est particulièrement mauvaise [11].

Le banc FVC 2000 est composé de 4 bases contenant chacune 8×100 images, où 100 est le nombre d'empreintes différentes représentées dans la base et 8 le nombre d'impressions différentes par empreinte. Le protocole des bancs FVC comprend deux types de tests permettant d'estimer le FMR (*i.e. False Match Rate*) et le FNMR (*i.e. False Non-Match Rate*). La première série de tests appelée *Genuine tests* a pour objectif d'estimer le FNMR et consiste à comparer chaque impression d'une empreinte à toutes les autres impressions de la même empreinte. Le FNMR se calcule par la fraction d'impressions authentiques qui ont obtenu un score de similarité en dessous du seuil de décision. La deuxième série de tests appelée *Impostor tests* a pour objectif d'estimer le FMR et consiste à comparer la première impression d'une empreinte à la première impression de

		Nombre d'échantillons fusionnés			
		1	2	3	4
DB1_A	EER	1,572	0,649	0,622	0,707
	FMR100%	1,893	0,833	0,800	1,000
	FMR1000%	2,893	1,333	1,200	1,500
	ZeroFMR%	3,357	2,167	1,400	4,250
DB2_A	EER	1,122	0,283	0,346	0,347
	FMR100%	1,250	0,333	0,400	0,250
	FMR1000%	1,964	0,833	0,600	0,750
	ZeroFMR%	2,500	1,000	0,600	0,750
DB3_A	EER	4,228	1,831	1,985	2,020
	FMR100%	6,036	2,333	3,000	3,000
	FMR1000%	7,750	3,667	5,000	4,750
	ZeroFMR%	13,464	8,333	7,800	8,000
DB4_A	EER	2,074	0,990	1,084	1,609
	FMR100%	2,714	1,333	1,400	1,750
	FMR1000%	4,893	2,500	2,200	2,000
	ZeroFMR%	9,250	3,667	6,000	3,750

TABLE 1. Précision (en %) dans les quatres bases de données utilisées dans la compétitions FVC2000

		Nombre d'échantillons fusionnés			
		1	2	3	4
Hierarchical Matching [17]	EER	–	–	–	–
	FMR100%	–	2,810	1,750	1,230
	FMR1000%	–	4,680	3,100	2,030
	ZeroFMR%	–	8,710	5,230	3,900
FZC	EER	1,572	0,649	0,622	0,707
	FMR100%	1,893	0,833	0,800	1,000
	FMR1000%	2,893	1,333	1,200	1,500
	ZeroFMR%	3,357	2,167	1,400	4,250

TABLE 2. Étude comparative. Précision (en %) dans la base de données DB1_A utilisée dans la compétitions FVC2000

chacune des empreintes restantes. Le FMR se calcule par la fraction d'impressions imposteurs qui ont obtenu un score de similarité au-dessus du seuil de décision.

Nous avons comparé la précision de FZC à l'approche *Hierarchical Matching* proposée dans [17] (voir section 2.2). La bibliothèque libre VeriFinger [1] a été utilisée pour extraire les minuties et leur associer une mesure de qualité. Nous avons testé l'effet de la fusion de k échantillons sur la précision de la comparaison d'empreintes. Nous avons utilisé l'algorithme GMMS [9] pour le choix des k échantillons qui représentent le mieux les variations dans les impressions d'une

empreinte. Pour ne pas biaiser le calcul du FNMR, les k impressions fusionnées de chaque empreinte, sont exclues des tests *genuine*. Le nombre de tests *genuine* est donc égal à $(8 - k) \times 100$ tests.

Le tableau 1 reporte les résultats obtenus par notre approche dans les quatre bases du FVC 2000. Le tableau 2 donne les résultats obtenus par notre approche et celle de [17] sur la base DB1_A du FVC 2000. Les résultats de l'approche *Hierarchical Matching* sont ceux reportés dans [17]. Tel que le montre le tableau 2, notre approche présente une nette supériorité par rapport à [17].

Conclusion

Dans ce article nous avons proposé FZC, une nouvelle approche pour la construction d'un super-modèle à partir de différentes impressions d'une empreinte. L'objectif de FZC est d'améliorer la qualité de l'impression de référence d'une empreinte par : (i) la restauration des minuties manquantes ; (ii) l'élimination des minuties parasites ; et (iii) l'augmentation de la surface de l'empreinte que l'impression de référence couvre. FZC se distingue de la majorité des approches existantes par un alignement différent pour chaque région de l'empreinte. FZC se distingue également par une estimation de la fiabilité d'une minutie qui se base sur la qualité de la capture de la région dans laquelle elle se trouve. Ceci permet une meilleure détection des minuties parasites dans les cas où peu d'impressions sont collectées. Les résultats expérimentaux obtenus valident les idées développées dans notre approche.

Comme perspectives, nous entendons améliorer notre approche de différentes manières. Nous projetons ainsi de tester différents algorithmes de classification pour la détermination des zones compatibles. Nous projetons également d'utiliser des algorithmes d'alignement d'empreintes qui se basent sur d'autres caractéristiques que les minuties pour bénéficier à la fois des informations qu'apportent les minuties et des informations qu'apportent les autres caractéristiques. Finalement, nous entendons tester notre approche sur d'autres bancs de tests.

Remerciements Ces travaux de recherche sont effectués dans le cadre d'une thèse MOBIDOC financée par l'U.E. dans le cadre du programme PASRI.

Références

1. <http://www.neurotechnology.com/verifinger.html/>.
2. Tai Pang Chen, Xudong Jiang, and Wei-Yun Yau. Fingerprint image quality analysis. In *Image Processing, 2004. ICIP'04. 2004 International Conference on*, volume 2, pages 1253–1256. IEEE, 2004.
3. Kyoungtaek Choi, Hee-seung Choi, and Jaihie Kim. Fingerprint mosaicking by rolling and sliding. In *International Conference on Audio-and Video-Based Biometric Person Authentication*, pages 260–269. Springer, 2005.
4. Kyoungtaek Choi, Heeseung Choi, Sangyoun Lee, and Jaihie Kim. Fingerprint image mosaicking by recursive ridge mapping. *IEEE Transactions on Systems, Man, and Cybernetics, Part B (Cybernetics)*, 37(5) :1191–1203, 2007.

5. Marcos Faundez-Zanuy. Data fusion in biometrics. *IEEE Aerospace and Electronic Systems Magazine*, 20(1) :34–38, 2005.
6. Hartwig Fronthaler, Klaus Kollreider, and Joseph Bigun. Automatic image quality assessment with application in biometrics. In *2006 Conference on Computer Vision and Pattern Recognition Workshop (CVPRW'06)*, pages 30–30. IEEE, 2006.
7. Anil Jain and Arun Ross. Fingerprint mosaicking. In *Acoustics, Speech, and Signal Processing (ICASSP), 2002 IEEE International Conference on*, volume 4, pages IV–4064. IEEE, 2002.
8. Xudong Jiang and Wee Ser. Online fingerprint template improvement. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 24(8) :1121–1126, 2002.
9. Yong Li, Jianping Yin, En Zhu, Chunfeng Hu, and Hui Chen. Score based biometric template selection and update. In *2008 Second International Conference on Future Generation Communication and Networking*, volume 3, pages 35–40. IEEE, 2008.
10. Dario Maio, Davide Maltoni, Raffaele Cappelli, James L. Wayman, and Anil K. Jain. Fvc2000 : Fingerprint verification competition. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 24(3) :402–412, 2002.
11. Davide Maltoni, Dario Maio, Anil Jain, and Salil Prabhakar. *Handbook of fingerprint recognition*. Springer Science & Business Media, 2009.
12. Yiu Sang Moon, Hoi-Wo Yeung, KC Chan, and SO Chan. Template synthesis and image mosaicking for fingerprint registration : an experimental study. In *Acoustics, Speech, and Signal Processing, 2004. Proceedings.(ICASSP'04). IEEE International Conference on*, volume 5, pages V–409. IEEE, 2004.
13. Daniel Peralta, Mikel Galar, Isaac Triguero, Daniel Paternain, Salvador García, Edurne Barrenechea, José M Benítez, Humberto Bustince, and Francisco Herrera. A survey on fingerprint minutiae-based local matching for verification and identification : Taxonomy and experimental evaluation. *Information Sciences*, 315 :67–87, 2015.
14. Herbert Ramoser, B Wachmann, and Horst Bischof. Efficient alignment of fingerprint images. In *Pattern Recognition, 2002. Proceedings. 16th International Conference on*, volume 3, pages 748–751. IEEE, 2002.
15. Arun Ross, Samir Shah, and Jidnya Shah. Image versus feature mosaicing : A case study in fingerprints. In *Defense and Security Symposium*, pages 620208–620208. International Society for Optics and Photonics, 2006.
16. LinLin Shen, Alex Kot, and WaiMun Koo. Quality measures of fingerprint images. In *International Conference on Audio-and Video-Based Biometric Person Authentication*, pages 266–271. Springer, 2001.
17. Tamer Uz, George Bebis, Ali Erol, and Salil Prabhakar. Minutiae-based template synthesis and matching for fingerprint authentication. *Computer Vision and Image Understanding*, 113(9) :979–992, 2009.
18. Wei-Yun Yau, Kar-Ann Toh, Xudong Jiang, Tai-Pang Chen, and Juwei Lu. On fingerprint template synthesis. In *Proceedings of Sixth International Conference on Control, Automation, Robotics and Vision (ICARCV 2000). Singapore*, pages 5–8, 2000.
19. Yong-liang Zhang, Jie Yang, and Hong-tao Wu. A hybrid swipe fingerprint mosaicing scheme. In *International Conference on Audio-and Video-Based Biometric Person Authentication*, pages 131–140. Springer, 2005.

DSMA : Diffusion Sûre des Messages d'Alerte dans les Vanets

Hadjer Goumidi, Zibouda Aliouat, Makhoulf Aliouat

hadjer.goumidi@univ-setif.dz, zaliouat@univ-setif.dz, maliouat@univ-setif.dz

{Department of Computer Science, Ferhat Abbas Sétif 1 University,
Laboratoire des Réseaux et Systèmes Distribués}.

Résumé : Le nouveau paradigme des réseaux mobile ad hoc, appelé Vehicular Ad hoc Networks (VANET), a créé un environnement fertile pour la recherche afin de promouvoir ce nouveau concept sur lequel est fondé les systèmes de transport intelligents. Le travail rapporté dans cet article concerne la diffusion sûre des messages d'alerte dans les VANETs afin de garantir le caractère temps réel de transmission des messages d'urgence et la sûreté de leur réception. Le protocole proposé, référencé sous le nom DSMA (Diffusion Sûre des Messages d'Alerte), permet d'atteindre l'objectif fixé en offrant une solution efficace au problème soulevé. L'évaluation des performances de notre proposition a été réalisée par simulation en utilisant la combinaison de deux simulateurs Vanet-MobiSim-NS2. Les résultats ainsi obtenus en termes de taux de perte, de délai et de taux de réception des messages urgents étaient pertinents et convaincants.

Mots clés : VANETs, Diffusion sûre, messages d'alerte.

1 INTRODUCTION

Le développement technologique qu'a vu le monde d'aujourd'hui a touché tous les domaines, particulièrement le secteur de la communication qui connaît une évolution considérable par l'apparition de la technologie sans-fil qui a permis de résoudre partiellement le problème des accidents routiers. C'est dans ce contexte que rentre notre travail. Les chercheurs pensent, pour résoudre ce problème, d'exploiter cette technologie afin de permettre aux véhicules d'établir des liens entre eux, ce qui constitue les nouveaux réseaux appelés VANET (Vehicular Ad-Hoc NETwork). Une des applications prometteuse de ces réseaux consiste à permettre aux véhicules équipés de capteurs spécifiques de détecter l'environnement proche et d'avertir les conducteurs des véhicules aux alentours suffisamment tôt en cas de risques d'accidents.

Les VANET peuvent prendre en charge les systèmes de sécurité destinées à éviter les accidents de la route de deux manières : les messages de sécurité périodiques (Beacon) et les messages d'urgence, les deux partagent un seul canal de contrôle. Les messages Beacon sont des messages contenant des informations sur l'état du véhicule de l'expéditeur comme la position, la vitesse, la direction. Pour les véhicules environ-

nants dans le réseau, les messages Beacon fournissent des informations nouvelles sur le véhicule de l'expéditeur. Ces informations les aident à connaître l'état du réseau actuel et de prévoir la circulation des véhicules. Ces messages sont envoyés de manière agressive aux véhicules voisins chaque 10 seconde. Les messages d'urgence sont des messages envoyés par un véhicule détectant une situation potentiellement dangereuse sur la route. Ces informations devraient être diffusées à d'autres véhicules citoyens pour les alerter de l'occurrence d'un danger probable. Un VANET est un réseau mobile où les véhicules se déplaçant selon de grandes vitesses. Même si ces véhicules se trouvent loin du danger, ils y parviendront dans des délais très proches. Dans pareils cas, les millisecondes seront très importantes pour éviter le danger [1].

En raison de l'importance de ces messages échangés entre véhicules, il est nécessaire de s'assurer que ces derniers atteignent le plus grand nombre de véhicules dans un temps minimum. Dans ce papier, nous essayons de proposer un algorithme qui permet de répondre à ces besoins impérieux.

2 TRAVAUX CONNEXES

Dans [3] les auteurs ont proposés EMDV: Emergency Message Dissémination for Vehicular en permettant au véhicule le plus éloigné au sein de la portée de transmission de retransmettre le message d'urgence. Le choix d'un seul transitaire ne convient pas dans VANET, parce que la position est toujours en changement, et le véhicule récepteur peut devenir hors de portée lors de l'envoi ou le récepteur ne peut pas recevoir le message en raison des problèmes de canal comme le déni de service.

Dans [4] les auteurs ont proposé le protocole de routage UMB (Urban Multi hop Broadcast Protocol) qui est un protocole efficace de la norme 802.11, basé sur l'algorithme de diffusion multi-saut pour les réseaux inter véhiculaires avec support d'infrastructure. Le but du protocole est de maximiser la progression du message, éviter les problèmes du broadcast storm, du nœud caché, et les problèmes de fiabilité.

Contention-Based Broadcasting protocol (CBB), est un protocole proposé dans [5]. Il est basé sur la diffusion du message d'urgence en mode multi-saut, et selon les auteurs, il a prouvé sa supériorité sur le protocole EMDV car il définit tous les véhicules situés dans le dernier segment non vide comme des transitaires, mais le choix de plusieurs transitaires peut poser un problème de collisions si le segment a une grande densité de véhicules.

Dans [6], les auteurs proposaient EEMB : Efficient Emergency Message Broadcasting. Ce modèle a deux phases principales : La phase normale dans laquelle on transmet les messages de sécurité (Beacon) et la phase d'urgence dans laquelle on transmet les messages d'urgences. La sélection du relayer passe par trois scénarios selon la direction des véhicules.

Time Critical Emergency Message Dissemination in VANETs, ce protocole est proposé par [7]. La communication directe (V2V) est sensible aux interférences et peut être bloquée par des obstacles physiques. Si la densité du réseau est petite, le protocole propose d'utiliser la communication (V2I) pour permettre une diffusion de message d'urgence plus rapide et plus fiable en réduisant le temps d'accès au canal et

la période de contention. Mais l'utilisation de l'infrastructure est très coûteuse et ne peut pas être déployée le long de toutes les routes.

2.1 Les Protocoles implémentés

Dans [1], les auteurs ont proposé **Pso Contention Based Broadcast (PCBB)**. C'est un protocole qui permet d'augmenter le pourcentage de la réception de l'information d'urgence ; cet algorithme, basé sur la segmentation de la portée de transmission, permet au véhicule qui détecte l'événement de trouver comme transitaire, le véhicule le plus éloigné dans le dernier segment non-vidé.

Le MIN est la borne inférieure du dernier segment non-vidé, on le calcule par la fonction fitness:

$$\text{Fitness} = \frac{\text{Longueur du segment} \times \text{Progression}}{\text{Nombre de véhicules de chaque segment}} \quad (1)$$

En analysant cet algorithme on trouve un problème sérieux: le calcul de min par la fonction fitness n'est pas toujours correct. On trouve dans certains cas, min supérieur au max. Cet algorithme peut s'appliquer seulement dans les zones denses où le nombre de véhicules est grand et la distance entre ces véhicules est petite.

Dans [7], [8], les auteurs ont proposé un pourcentage de 66% tels que la probabilité de la réception de message de diffusion diminue instantanément à une distance supérieure à 66% de la portée de transmission, alors ils choisissent le nœud qui à une distance égale à 66% de la portée comme un transitaire.

Compte tenu de son importance, nous avons étudié le protocole PCBB afin de résoudre son problème et renforcer sa faisabilité. Après la sélection du transitaire, si un véhicule détecte un événement, il crée le message et le transmet dans le réseau. Lorsque les voisins reçoivent le message, ils utilisent les mêmes étapes utilisées dans le protocole PCBB.

Lorsqu'on choisit le nœud 66% non pas le nœud max (le plus éloigné) le message prend un temps de réception supérieur à celui du nœud max et cet algorithme n'utilise pas la segmentation, alors le transitaire n'a pas de nœuds potentiels, si le message n'arrive pas au transitaire, il sera simplement perdu.

3 LE PROTOCOLE PROPOSÉ

Comme premier pas de notre approche, nous allons présenter le protocole DSMA qui permet d'assurer une transmission sûre des messages d'alerte. Notre approche peut être utilisée dans toutes les zones (dense et éparses) contrairement à l'algorithme PCBB qui s'applique seulement dans les zones denses.

L'idée principale de la proposition est de transmettre le message au nœud le plus éloigné dans la portée de transmission. Un tel nœud est défini comme un transitaire. Si le transitaire ne reçoit pas le message alors on aura recours aux nœuds potentiels tel que le min de ces nœuds représentant 60% de la portée de transmission de l'émetteur. Notre proposition est basée sur une idée similaire au protocole PCBB, mais nous

n'appliquons pas l'expression fitness et nous améliorons le protocole en utilisant les paramètres : position du cas urgent, le messageID, le temps et le délai du message d'urgence.

Les deux protocoles définis précédemment souffrent de la perte des messages et à partir de ces insuffisances, on peut définir notre proposition basée sur la segmentation : On choisit le transitaire qui est le nœud le plus éloigné dans le dernier segment non vide, la borne inférieure du segment (le min) représente 60% de la portée de chaque nœud. Notre approche passe par quatre principales étapes, comme le protocole PCBB, à savoir :

3.1 Préparation de l'envoi

Chaque véhicule vérifie si sa portée de transmission est vide; si elle n'est pas vide, il doit vérifier si la RSU est dans sa portée de transmission, si oui, il la définit comme un transitaire, sinon il calcule le transitaire qui est le nœud le plus éloigné de l'émetteur et le plus proche de la RSU dans le dernier segment non-vide. La sélection du transitaire est inspirée du protocole PCBB, elle est faite après le calcul de la distance entre l'émetteur et le transitaire en utilisant l'équation (2).

$$\text{Dis} = \text{Position du Transitaire} - \text{Position de l'émetteur} \quad (2)$$

Tels que tous les véhicules du réseau connaissent la position de leurs voisins et cela à cause des messages de sécurité périodique. Puis en calculant de la distance entre le transitaire et la RSU, en utilisant l'équation (3).

$$\text{DisRSU} = \text{Position de la RSU} - \text{Position du Transitaire} \quad (3)$$

Telle que : la RSU envoie périodiquement sa localisation à tous les nœuds dans sa portée de transmission. La borne inférieure de segment (le min) représente 60% de la portée de transmission comme indiqué dans (Fig.1).

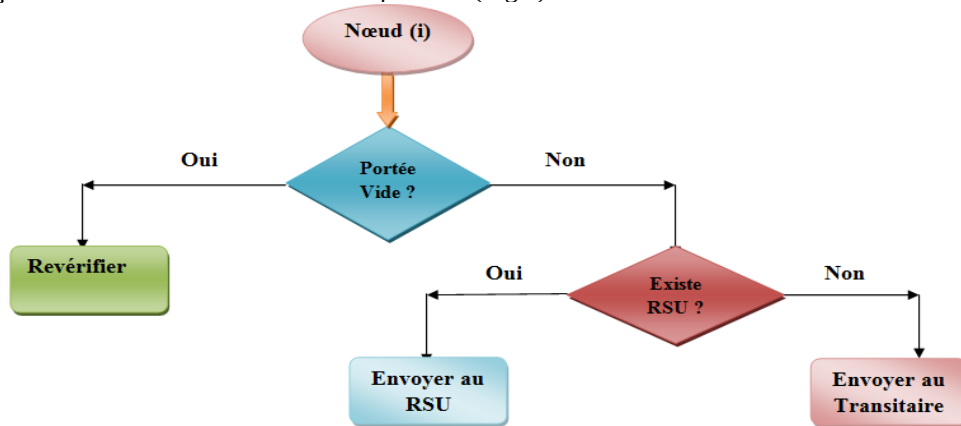


Fig. 1. Organigramme du choix du transitaire.

Lorsqu'un véhicule détecte un cas d'urgence (accident par exemple), il crée un message d'urgence « ADV_ALERT » qui contient :

Identificateur de la source. Pour ne pas retransmettre le même message, on utilise cet identificateur, qui permet à l'émetteur d'éviter de retransmettre le message.

Identificateur du message. Pour permettre de différencier les messages transmis par le même émetteur.

Identificateur de priorité. Inspiré du protocole PCBB, on utilise cet identificateur pour classer les messages selon leur importance.

ID Transitaire. Pour éviter les collisions et pour que le message doit être transmis de manière très rapide, il faut qu'on transmette le message d'urgence directement au transitaire.

Temps d'émission et Délai. Le message d'urgence nécessite un temps court, et la transmission de ce message après ce temps, occupe le canal sans aucun intérêt. L'émetteur définit son temps de transmission du message et le délai de ce message.

Position. L'absence de caractérisation de l'événement d'urgence pose le problème de charge du réseau, Pour résoudre ce problème, il faut que le véhicule, qui détecte l'événement, envoie la position de cet événement en utilisant un GPS.

Min. Il faut envoyer le min pour que chaque véhicule reçoive ce message soit connues comme véhicule potentiel ou non.

Les data. Il faut que le message d'urgence contienne des données importantes : La nature du cas urgent (accident, avalanche...), s'il y a des victimes ou des blessés. Voici le format du message « ADV_ALERT » transmis dans (Fig.2).

IDSource	IDmessage	IDpriorité	IDTransitaire	Temps émission	Délai	Position	Min	Data
----------	-----------	------------	---------------	----------------	-------	----------	-----	------

Fig. 2. Format du message d'alerte selon DSMA.

Si un véhicule détecte un événement, il vérifie s'il existe dans la table des messages d'alerte la même localisation ? Si oui, il ne crée pas de message, sinon il crée le message et il ajoute la localisation de cet événement, comme expliqué dans (fig.3). L'utilisation de la position de l'évènement peut aider les véhicules de secours (ambulances) à intervenir plus rapidement en épargnant éventuellement des vies humaines.

3.2 L'étape d'envoi

L'émetteur diffuse un message d'urgence « ADV_ALERT » pour avertir les autres véhicules à propos de tout danger potentiel.

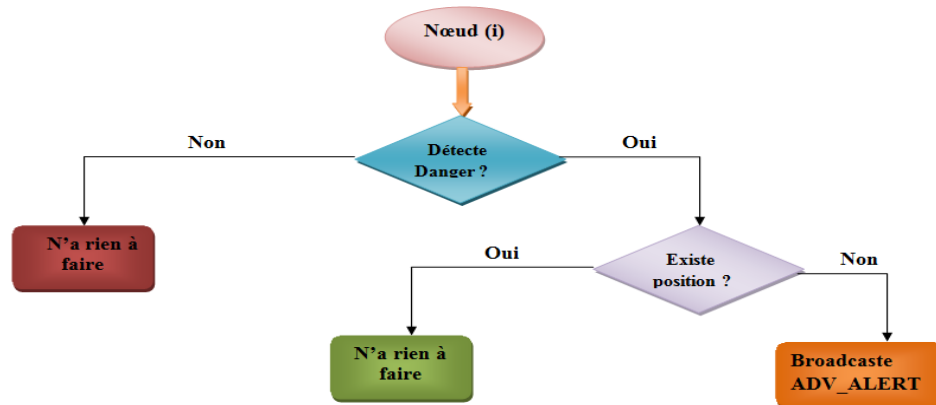


Fig. 3. Organigramme de la vérification de localisation.

3.3 L'étape de réception

Tous les véhicules existant dans la portée de transmission de l'émetteur reçoivent le message « ADV_ALERT », les véhicules du dernier segment lancent un temps de contention pour que le transitaire transmette un acquittement à l'émetteur. Puis ils détectent le canal si aucune émission (un acquittement n'été pas transmis à partir du transitaire ou d'autres véhicules, cela veut dire que le transitaire n'a pas reçu le message et le premier véhicule qui transmet l'acquittement se définit comme un transitaire. Une fois l'acquittement transmis, les autres véhicules s'abstiennent de transmettre pour éviter les collisions comme indiqué dans (Fig .4).

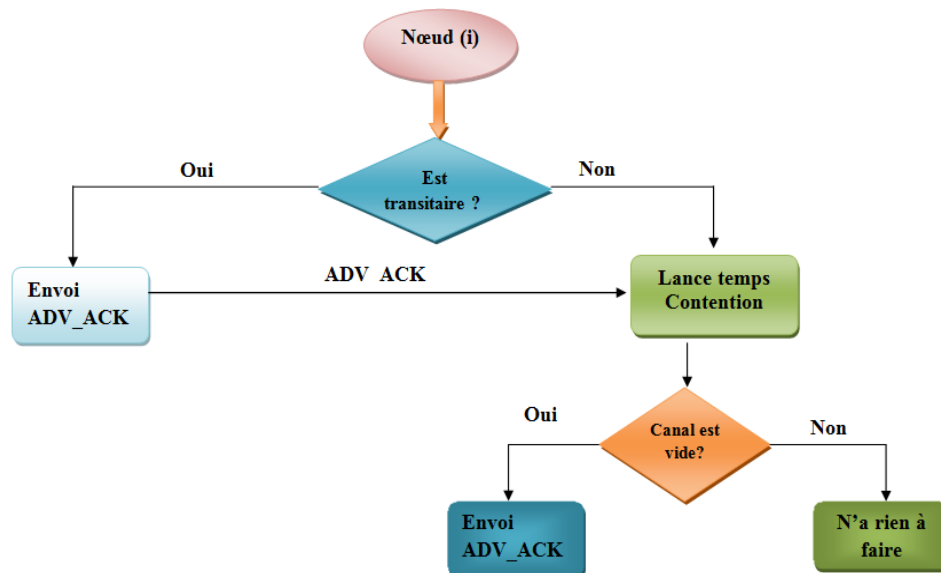


Fig. 4. Organigramme de la phase de réception du message d'alerte.

Avant l'envoi d'acquiescement, le transitaire doit vérifier s'il a enregistré le message « ADV_ALERT » dans sa table des messages d'urgences pour pouvoir envoyer l'acquiescement et retransmettre le message d'alerte à son transitaire ou non.

ID Priorité. Le transitaire vérifie l'importance du message : si tel est le cas (grand nombre), il ne l'enregistre pas sinon il vérifie les autres indices (s'ils sont acceptables) il l'enregistre dans la table selon son degré de priorité.

Temps d'émission et le Délai. Le transitaire doit calculer le délai (D) en utilisant l'équation (4).

$$D = \text{Temps réception} - \text{Temps émission} . \quad (4)$$

Si « D » est supérieur au délai défini par l'émetteur, alors il n'enregistre pas le message dans sa table de messages d'alerte et ne le transmet pas. Après chaque 30s, chaque véhicule doit faire la mise à jour de sa table, le message qui dépasse le délai est supprimé.

ID source et ID message. Sert à vérifier si le message est transmis avant ou non. Les véhicules qui reçoivent le message d'alerte doivent vérifier s'ils trouvent un message ayant le même ID source du message reçu, ils doivent vérifier son message ID : S'ils trouvent le même ID message alors pas d'enregistrement du message, Si non (pas le même ID message) enregistrer le message, cela est récapitulé dans (Fig.5).

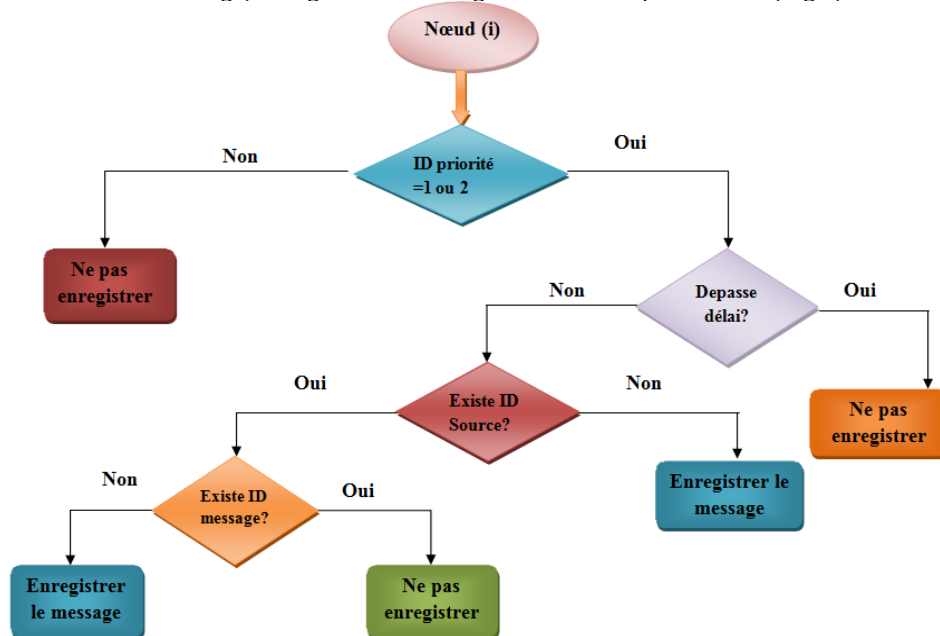


Fig. 5. Organigramme pour la vérification du message d'alerte pour l'enregistrer.

3.4 L'étape de rediffusion

Le transitaire doit rediffuser le message dans sa portée de transmission à son transitaire et suit les mêmes étapes définies précédemment.

4 SIMULATION

4.1 Paramètres de simulation

Nous avons testé un scénario de 100 nœuds où ces nœuds sont placés initialement dans des positions aléatoires et la direction de déplacement suit le modèle de mobilité IDM_LC implémenté sous VanetMobiSim. Les paramètres utilisés sont illustrés dans le tableau 1.

Table 1. Paramètres de simulation.

Paramètres	Valeurs
Topologie (m2)	1000*1000
Nombre des nœuds	100
Vitesse (m/s)	6
Couche MAC	MAC 802_11
Temps de simulation (s)	300 s

La figure 6 montre le nombre de messages transmis, perdus et reçus pour le protocole PCBB. On remarque qu'un seul message reçu, le reste est perdu (le min est supérieur au max) jusqu'à la fin de la simulation.

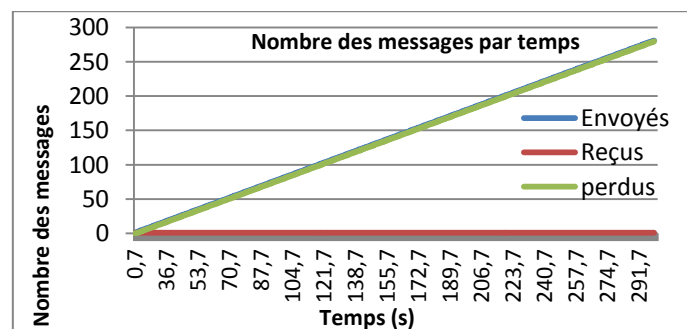


Fig. 6. Nombre de messages transmis, reçus et perdus dans le protocole PCBB.

La figure 7 montre le nombre des messages transmis, perdus ; et pour le protocole 66% on remarque que le nombre de messages perdus est grand, seulement 8 messages parmi 307 messages envoyés ont été reçus et le reste est perdu. Cette perte est due à l'absence de nœuds potentiels.

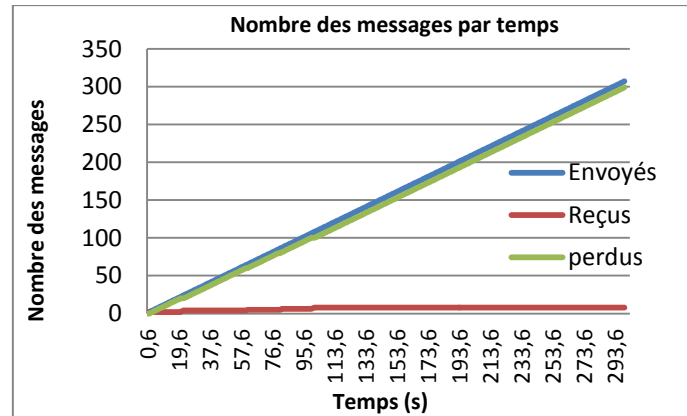


Fig. 7. Nombre de messages transmis reçus et perdus dans le protocole 66%.

La figure 8 montre le nombre de messages transmis, perdus et reçus selon DSMA.

On remarque que le nombre de messages transmis est petit (6 messages seulement) par rapport aux autres protocoles puisque chaque message transmis doit être reçu, il n'y a pas de messages perdus.

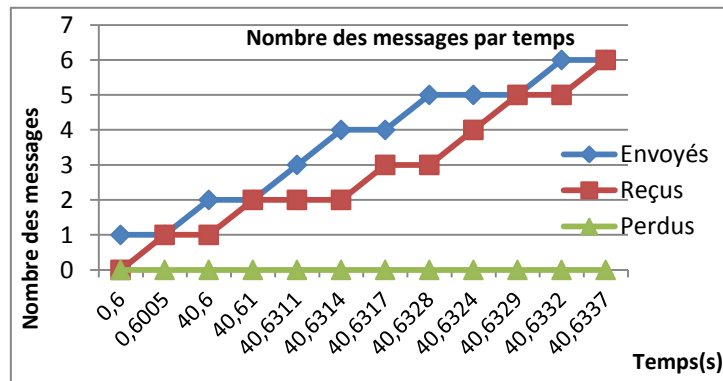


Fig. 8. Nombre de messages transmis, perdus, reçus selon DSMA.

On calcule le taux de perte des trois protocoles en utilisant la formule:

$$\text{Taux de perte} = \frac{\text{Nombre messages perdus}}{\text{nombre des messages transmit}} \times 100 . \quad (5)$$

La figure 9 montre le taux de perte selon les trois protocoles, on remarque que le taux de perte de messages de notre protocole est nul, et les deux autres protocoles ont un grand taux de perte par rapport au DSMA.

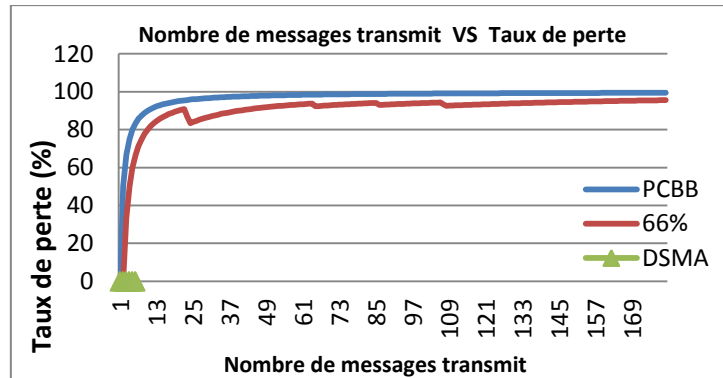


Fig. 9. Taux de perte par rapport au nombre de messages transmis selon les trois protocoles.

On calculant le délai de la réception de message d'alerte des trois protocoles, on utilise la formule (4).

La figure 10 montre le délai de réception selon les trois protocoles, on remarque que dans DSMA, le message est reçu après 40s de sa transmission, par contre dans les deux autres protocoles le message n'arrive pas au RSU.

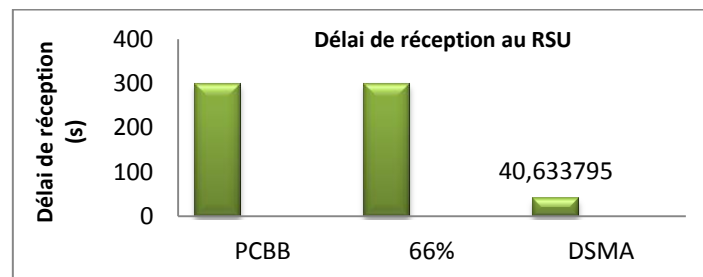


Fig. 10. Délai de réception de messages d'alerte selon les trois protocoles.

4.2 Evaluation des Performances du Protocole

Les paramètres qui ont été modifiés sont illustrés dans le tableau 2.

Table 2. Paramètres de simulation.

Paramètres	Valeurs
Topologie (m2)	1000*1000
Nombre des nœuds	50/100/200/350/500
Vitesse (m/s)	6/12/20/40/100
Couche MAC	MAC 802_11
Temps de simulation (s)	300s

La figure 11 représente le taux de réception des messages ; on remarque que le taux de réception est presque de 20% lorsque les nœuds diffusent le message entre eux. Lorsqu'une RSU le diffuse, le taux de réception est augmenté jusqu'à 80%, et dans le cas où deux RSU diffusent le message, le taux de réception devient presque 100%

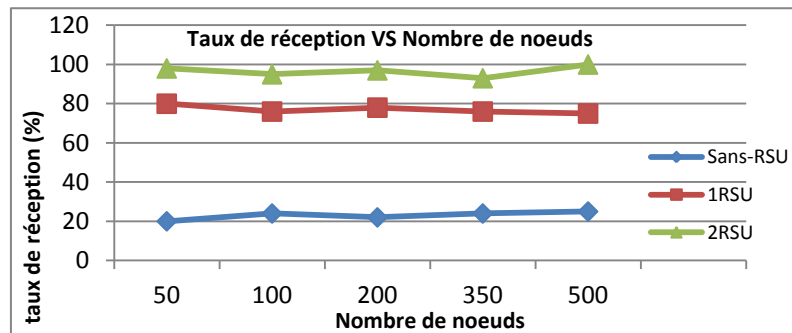


Fig. 11. Taux de réception par rapport au nombre des nœuds fonction des RSUs

La figure 12 montre le délai de réception par rapport au nombre de nœuds. Ce délai diminue avec l'augmentation du nombre de nœuds. Cela est dû aux nombres de véhicules qui transmettent le message.

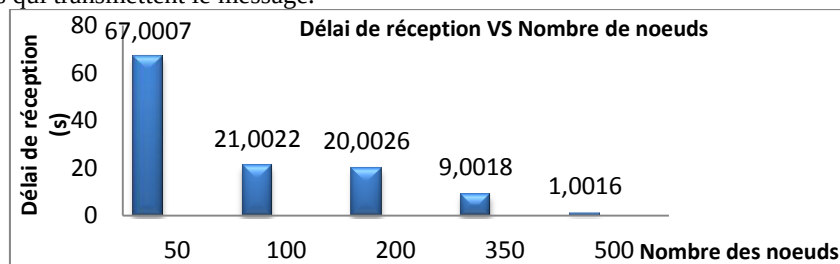


Fig. 12. Délai de réception par rapport au nombre des nœuds.

La figure 13 montre le délai de réception par rapport au changement de la vitesse, il diminue avec l'augmentation de la vitesse des nœuds. Cela est dû à la rapidité des véhicules qui transmettent le message.

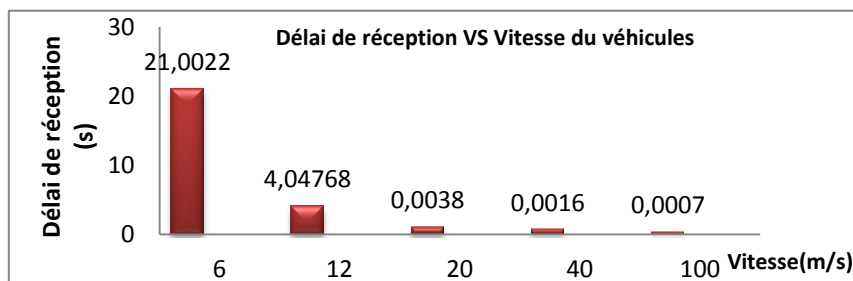


Fig. 13. Délai de réception par rapport au changement de la vitesse.

5 CONCLUSION

Dans cet article, nous avons étudié le problème de la diffusion sûre des messages d'urgence dans les VANETs et nous avons proposé un protocole référencé sur le nom DSMA (Diffusion Sûre des Messages d'Alerte). Nous avons montré par des simulations que notre protocole a un pourcentage de succès très élevé, même à haute densité et vitesse des véhicules. Nous avons montré qu'il assure un taux de réception élevé avec un délai minimal et un faible taux de perte de messages. Comme perspectives, un futur travail est souhaitable pour poursuivre l'amélioration de cet algorithme, en traitant le cas où le véhicule ne trouve pas de transporteur autre que les nœuds qui ont déjà transmis le message, celui-ci est perdu dans ce cas, la solution au problème est de connaître la direction des véhicules en utilisant le GPS. Les véhicules qui ont une direction opposée à la RSU seront exclus de la transmission du message.

6 REFERENCES

1. Ghassan, S., areq, A.: Intelligent Emergency Message Broadcasting in VANET Using PSO.In : The World of Computer Science and Information Technology Journal (WSCIT), Vol 4, Issue 7, pp. 90--100. (2014)
2. Vedha,V.,D., Seethalakshmi, Mrs.,V.: ANALYSIS OF SAFETY MEASURES AND QUALITY ROUTING IN VANETS.In: International Journal of Modern Engineering Research (IJMER) Vol.2, Issue.2, pp. 062--066. Mar-Apr (2012)
3. Korkmaz, G., Ekici, E., Özgüner, F. & Özgüner, Ü.: Urban multi-hop broadcast protocol for inter-vehicle communication systems .In : 1st ACM international workshop on Vehicular ad hoc networks, ACM.(2004)
4. Ghassan,S., Wafaa,A., Sureswaran,R. : Increasing Network Visibility Using Coded Repetition Beacon Piggybacking.In :World Applied Sciences Journal (WASJ), Vol.13,No 1, pp.100--108. (2011)
5. Jianhua He : An Efficient Emergency Message Broadcasting Scheme in Vehicular Ad Hoc Networks. Hindawi Publishing Corporation.In: International Journal of Distributed Sensor Networks. Article ID 232916, pp.11--17. September (2013)
6. Anees,A.,S., Bibin,V.,K.: Time Critical Emergency Message Dissemination in VANETs.In :International Journal of Engineering Research and General Science Vol.3, Issue.4, Part-2, July-August(2015)
7. Marc,T.,M., Daniel,J., Hannes,H. :Broadcast reception rates and effects of priority access in 802.11-based vehicular ad-hoc networks.In :First International Workshop on Vehicular Ad Hoc Networks.ACM, pp.10--18.USA 1 October(2004)
8. Mahipal, Sanjay,B.,K.: Efficient Broadcasting Protocol for Vehicular Ad-Hoc Network.In, International Journal of Computer Applications,Vol.49, No.23. July (2012)

Pareto-Optimization-Based Approach for Feature Selection in Sentiment Analysis

Fayçal Rédha Saidani ¹, Allel Hadjali ², Idir Rassoul ¹ and Djamal Belkasmi ³

¹ LARI laboratory - University of Mouloud Mammeri, Tizi-Ouzou, Algeria

² LIAS/ENSMA - University of Poitiers 1, Avenue Clement Ader, Futuroscope Cedex, France

³ DIF-FS/UMBB, Boumerdes, Algeria

faycal.saidani@ummo.dz, allel.hadjali@ensma.fr,
idir_rassoul@yahoo.fr, belkasmi.djamal@gmail.com

Abstract. This paper deals with feature selection in sentiment analysis for the purpose of polarity classification. Feature selection is a search process that aims to find a relevant feature subset from an initial set. The proposed method allows to select Pareto optimal features that are not dominated by any other feature. These non-dominated features are computed w.r.t. five criteria that rely on three metrics. The skyline calculation is carried out using an improved BNL algorithm (denoted IBNL). To demonstrate the effectiveness of the method with respect to dimensionality reduction and classification rate, some experiments are conducted on real datasets. The results show the efficiency of Pareto-optimization for the feature selection.

Keywords: Feature Selection; Polarity Classification; Skyline Queries, Feature Weighting. Sentiment Analysis

1 Introduction

With the emergence of social networks and more specifically Twitter sphere, the subjective textual data grow exponentially. Indeed, there is a very large mass of documents containing information expressing opinions, and constitute a huge source of data for various applications survey (technological, marketing, competitive, and societal). Much research at the crossroads of NLP and data mining, is addressing the problem of opinion mining.

We can distinguish two families of approaches: (i) those based on lexicons [1, 2] and (ii) those based on machine learning [3]. The former rely on a list (dictionaries, thesaurus) of subjective words either created manually or automatically. If a document includes them, it is considered as subjective document. One of the pioneering works in this family is that proposed in [4]. The authors proposed a method to estimate the orientation of a word or phrase, by comparing its similarity to seed words using Pointwise Mutual Information (PMI) and Information Retrieval. Other works focused on using adjectives or verbs as pointers of the semantic orientation [1, 5].

The machine learning based approaches use different types of classifiers, such as Support Vector Machines (SVM) [6] and Naïve Bayes (NB) [7], which are trained on a given data set of documents. In order to process the latter, text features are extracted and machine learning algorithm is trained with a labelled data set in order to construct in a supervised manner the classification model. The main strength of these approaches lies in their ability to analyse the text of any domain and produces classification models that are tailored to the problem. In addition, they can be applied to multiple languages [3].

Lexicon-based approaches suffer in most cases from the lack of specialized lexicons, unlike machine learning methods that rely mainly on data without considering human interaction. Therefore, they constitute a good alternative for our research task. However, the performance of machine learning algorithms is heavily dependent on the features (terms) used during the learning phase; therefore a preliminary task of feature selection is generally essential before any learning. The complexity of selection depends on the type of objective to be achieved.

In this paper, we investigate a novel idea that consists of leveraging skyline paradigm to optimize the feature selection process for polarity classification. This allows selecting a subset of Pareto optimal features that are not dominated by any other feature. In summary, the main contributions made in this paper are:

- A set of relevant metrics for the purpose of features selection is investigated.
- Based on these metrics (or dimensions), a feature skyline is computed, which significantly reduces the features space. By this way, irrelevant features are ruled out and then reducing the noisy in the classification process.
- A set of experiments are conducted to show the interest of our method w.r.t dimensionality reduction rate.

The paper is structured as follows. Section 2 recalls some basics notions on feature selection, skyline queries and presents the problem of interest. In Section 3, we provide some related work. In Section 4, we present our skyline-based method to features selection. Sections 5 discusses the experimental results and finally we conclude the paper.

2 Preliminaries and problem statement

This section describes the feature selection problem and some notions on skyline queries that are necessary to our proposal. Finally, we present the problem statement.

2.1 Feature Selection Problem

Feature selection is a search process that aims at finding a relevant feature subset from an initial set. The relevance of a feature subset always depends on the objectives and criteria of the problem to be solved. The problem can be defined formally as follows:

Definition1. Let X be the original set of features, with cardinality $|X| = n$, and Let $J(.)$ be an evaluation measure to be optimized defined as $J: X' \subseteq X \rightarrow \mathbb{R}$. The purpose of the feature selection is to find $X' \subset X$, such that $J(X')$ is maximum [8].

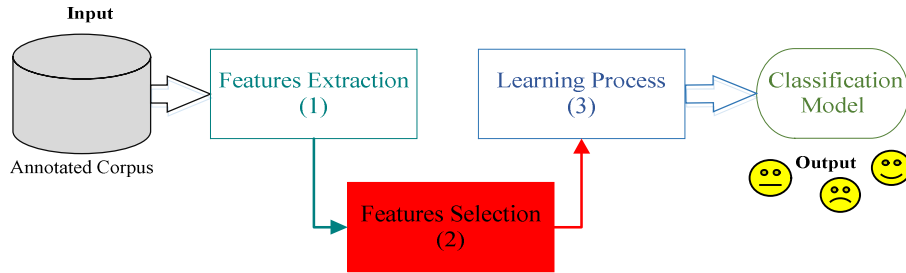


Fig. 1. Overview of sentiment classification steps.

To achieve sentiment classification using machine learning, three major steps are needed (as shown in Fig. 1). The first step is to extract an initial set of features which will serve as an input to the classifier, so, it is difficult to learn good classifiers without removing a number of irrelevant features. Thus, feature selection attempts to select the minimally sized subset in a way that the classification accuracy does not decrease. The work presented in this paper seeks to improve this step.

2.2 Skyline Queries

Skyline queries [9] are example of preference queries that can help users to make intelligent decisions in the presence of multidimensional data where different and often conflicting criteria must be taken. They rely on Pareto dominance principle which can be defined as follows:

Definition2. Let U be a set of d -dimensional points and u_i and u_j two points of U . u_i is said to dominate in Pareto sense u_j (denoted $u_i \succ u_j$) if u_i is better than or equal to u_j in all dimensions and (strictly) better than u_j in at least one dimension. Formally, we write:

$$u_i \succ u_j \Leftrightarrow (\forall k \in \{1, \dots, d\}, u_i[k] \geq u_j[k]) \wedge (\exists l \in \{1, \dots, d\}, u_i[l] > u_j[l]) \quad (1)$$

Where each tuple $u_i = (u_i[1], u_i[2], u_i[3], \dots, u_i[d])$ with $u_i[k]$ stands for the value of the tuple u_i for the attribute A_k .

Definition3. The skyline of U , denoted by S , is the set of points which are not dominated by any other point.

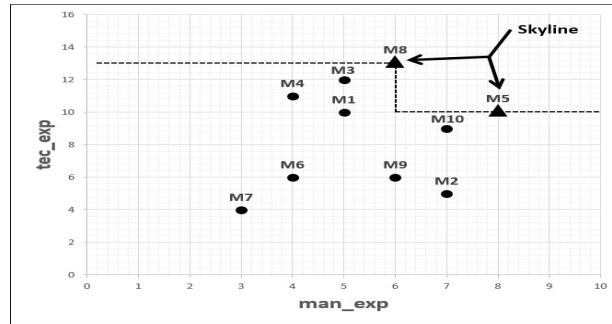
$$u \in S \Leftrightarrow \nexists u' \in U, u' \succ u \quad (2)$$

Example 1. To illustrate the Skyline, let us consider a database containing information on candidates as shown in Table 1. The list of candidates includes the following information: Code, Age, Management experience (*man_exp* in years), Technical experience (*tec_exp* in years) and distance work to home (*dist_wh* in Km).

Table 1. Liste of candidates

code	age	man_exp	tec_exp	dist_wh
M1	32	5	10	35
M2	41	7	5	19
M3	37	5	12	45
M4	36	4	11	39
M5	40	8	10	18
M6	30	4	6	27
M7	31	3	4	56
M8	36	6	13	12
M9	33	6	6	95
M10	40	7	9	20

Ideally, personnel manager is looking for a candidate with the largest management and technical experience (Max *man_exp* and Max *tec_exp*), ignoring the other pieces of information. Applying the traditional skyline on the candidate list shown in Table 1 returns the following candidates: M5, M8. As can be seen, such results are the most interesting candidates (see Fig.2).

**Fig. 2.** Skyline of candidates

2.3 Problem Statement

Given that it is not obvious with classical features selection methods to consider several and conflicting weighting metrics in order to refine the selection for sentiment analysis. We are interested here in the way of identifying a discriminant features subset using the skyline paradigm for improving the performance of the classification step in terms of accuracy.

3 Related Work

We review here the main approaches proposed in the literature about features selection for the sentiment analysis context. In approaches based on machine learning techniques [10, 11], the performance of the algorithms strongly depends on the features used in the learning task. Thus, the presence of redundant or irrelevant features can have a significant impact on the performance of the approaches.

The work done in [12] presents a new feature selection schemes that use content and syntax model to automatically learn a set of features in a review document. The features obtained yield competitive results in a maximum entropy classifier. In the same vein, in [13] explicit topic models are established to filter words and extract the training attributes of each product. Finally, several SVM classifiers are constructed to train the selected attributes to detect the corresponding implicit features for Chinese reviews.

In [14] several individual feature lists obtained by different feature selection methods are aggregated to improve sentiment classification. Let us also mention the work of [15] which proposes an optimized feature reduction that incorporates information gain and genetic algorithm as feature selection. The results show a clear improvement in terms of precision.

The main limitation of almost existing methods lies in the fact that the selected features depend on a unique optimization criterion. In order to overcome this drawback, we consider here multi-criteria decision based on the skyline paradigm in order to highlight the most discriminating features according to different criteria.

4 Presentation of the Approach

In this study, we assume that the annotated target corpus is divided into two sub-corpus: positive (*SCP*) and negative (*SCN*) corpus. The former (resp. latter) contains positive (resp. negative) documents expressing opinions in favour (resp. disfavour) of the matter of interest. By this way, one can perform a better pruning on the search space related to the features documents.

The selection approach relies on three metrics (i.e., *G_idf*, *Odds Ratio* and *POS weight*). The two first metrics are computed for both positive and negative sub-corpus in order to characterize which feature best reflects the specificity of each class in regard to the other. Each feature will be then associated with a vector of five measures (G_Idf^+ , G_Idf^- , Orr^+ , Orr^- , Pos) where *Orr* stands for Odds Ratio metric. Based on this vector, one can compute the skyline to keep only the most interesting features (in the sense of Pareto). The overall architecture of the proposed approach is illustrated in Fig. 3.

4.1 Weighting Metrics

We describe first the set of metrics used for the feature selection method.

Global tf.idf Weight.

For the first metric (i.e., Global tf.idf denoted *G_idf*), we rely on the Term Frequency-Inverse Document Frequency weight (tf-idf). The principle is to give importance to the more specific features of a document. So, to evaluate the representativeness of a feature in the corpus and not only in a document, we have to apply an aggregation function (e.g., the arithmetic mean) over the set of individual tf-idf weights of features

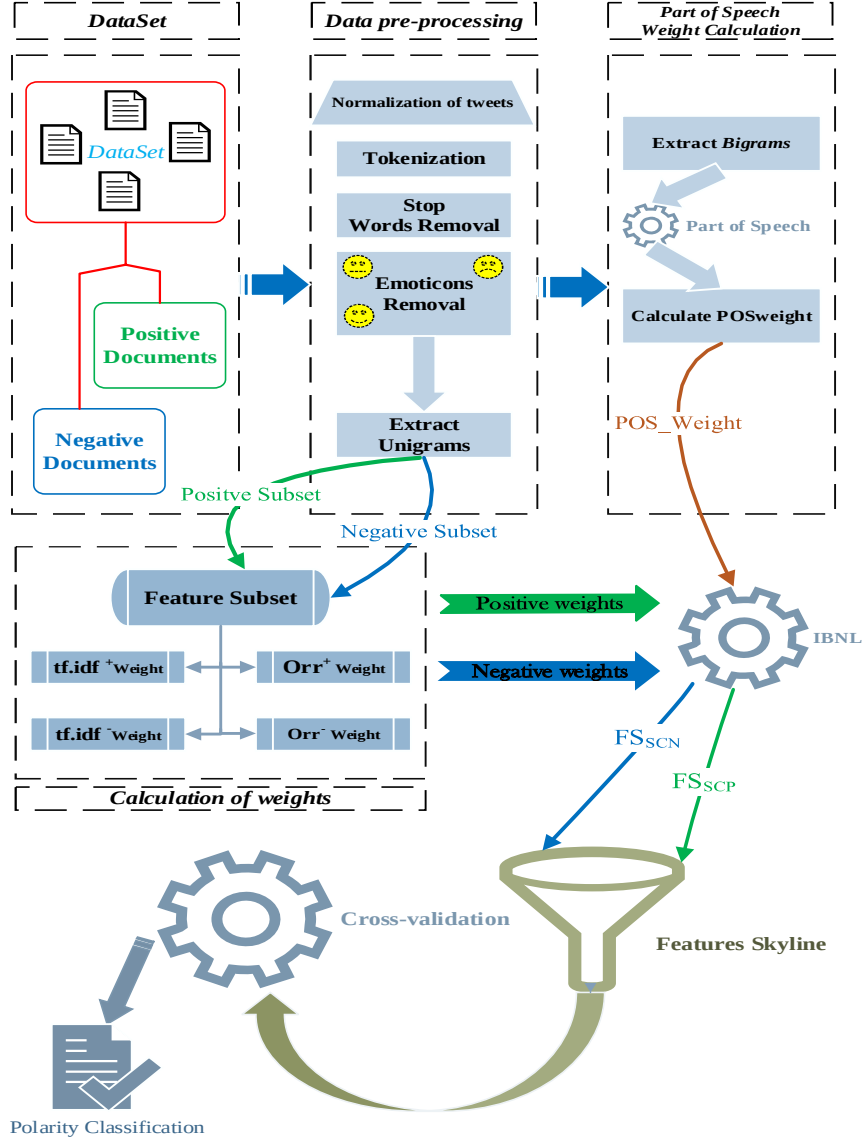


Fig. 3. Architecture of the Skyline-Based Feature Selection

computed on each document. Let C be a corpus of n documents and $f \in F$ a set of features, one can write:

$$G_idf_f = \frac{1}{n} \sum_{i=1}^n tfidf_{f,i} \quad (3)$$

In the following, G_idf^+ (resp. G_idf^-) stands for the aggregate score G_idf on SCP (resp. SCN) sub-corpus.

Odds Ratio weight.

Odds Ratio (*Orr*) gives a positive score to features that occur more often in one class than in the other, and a negative score if it occurs more in the opposite class [16]. Formally, *Orr* measure is given by (where $f \in F$):

$$Orr_f = \log \left(\frac{p(1-q)}{q(1-p)} \right) \quad (4)$$

p (resp. q) the probability of class to be predicted (resp. the opposite class) given the feature f . Following [16], a *null* value of the *Orr* parameter means that the feature belongs to the neutral polarity. As above, Orr^+ (resp. Orr^-) stands for the *Orr* weight on SCP (resp. SCN) sub-corpus.

POS Weight.

As stressed in [12], an in-depth linguistic analysis based on syntactic relations in a sentence can be useful for sentiment analysis model. For instance, collocations that contain mostly adjectives-adverbs or verbs-negation features are considered more sentiment bearing. Thus, to be more precise in our process, we have chosen to exploit this semantics by checking and accounting for linguistic collocation based on Part Of Speech (POS) using *Stanford Pos Tagger*. To this end, we introduce the following formula for measuring the POS weight (where n represents the number of documents in the target corpus):

$$POS_f = \log[1 + \sum_{i=1}^n term_frequency_f] \quad (5)$$

4.2 Feature Skyline

As mentioned above, each feature is associated with a vector of five weights expressed by $(G_Idf^+, G_Idf^-, Orr^+, Orr^-, Pos)$. The aim of this step is to compute the Feature Skyline (FS), namely, the non-dominated features (in the sense of Pareto) with respect to the above five criteria. First, a skyline FS_{SCP} is computed on the SCP corpus by maximizing G_Idf^+ , Orr^+ , Pos criteria, and minimizing G_Idf^- and Orr^- criteria. Then, a skyline FS_{SCN} is computed on the SCN corpus by maximizing G_Idf , Orr^- , Pos criteria, and minimizing the other criteria. An improved BNL algorithm (denoted IBNL) which is borrowed from our previous work [17] as illustrated in Algorithm 1, is used for calculating FS_{SCP} and FS_{SCN} . The final feature skyline FS is obtained as follows:

$$FS = (FS_{SCP} \cup FS_{SCN}) - (FS_{SCP} \cap FS_{SCN}) \quad (6)$$

As it can be seen, FS does not include the features that are common to both FS_{SCP} and FS_{SCN} since they do not convey any interesting information for the purpose of classification.

Algorithm 1. IBNLInput: A set of features $F = \{f_1, f_2, \dots, f_n\}$ Output: A Feature Skyline FS

```

1: Sort ( $F$ )
2: for  $i := 1$  to  $n - 1$  do
3:   if  $\neg f_i .dominated$  then
4:     for  $j := i + 1$  to  $n$  do
5:       status = 0;
6:       if  $\neg f_j .dominated$  then
7:         evaluate SkylineCompare( $f_i, f_j, status$ );
8:         switch status do
9:           case 1
10:             $f_i .dominated = true$ ;
11:           case 2
12:             $f_j .dominated = true$ ;
13:         if  $\neg f_i .dominated$  then
14:            $FS = FS \cup \{f_i\}$ ;
15: returns  $FS$ ;

```

SkylineCompare is a function that evaluates the dominance, in the sense of Pareto, between f_i and f_j on all skyline dimensions and returns the result in *status*. It may be equal to: 0 if $f_i = f_j$, 1 if $f_i > f_j$, 2 if $f_i < f_j$ and 3 if they are incomparable.

5 Experimental Study

5.1 Dataset and Experimental Setup

The experiments are conducted on two different data sets to evaluate the proposed approach. We focus mainly on the two following criteria: (i) dimensionality reduction rate and (ii) classification rate. In Table 2, we present the distribution of documents in datasets. The First dataset “Polarity Dataset V2.0”¹ which was extracted from the Internet Movie Database (IMDB), contains 1000 positive and 1000 negative movie reviews. The second dataset “The Twitter Sentiment Analysis Dataset” contains at the origin 1 578 627 classified tweets; each row is marked as 1 for positive sentiment and 0 for negative sentiment. A reduced version was proposed in Lightside² workbench and contains 10 662 classified tweets distributed uniformly between the two classes. A 10-fold cross-validation procedure is utilized, which means that the original data set is randomly divided into 10 equal-sized subsamples. For each time, a single subsample is used for validation, whereas the rest are kept for training. All algorithms were implemented with Java 1.8 64 bits. The Weka API [18] has been used with default settings for classifiers algorithm and all experiments were conducted on a Windows 8.1 system with Intel core i7 6500U @2.5 GHz CPU 8GB RAM.

¹ <https://www.cs.cornell.edu/people/pabo/movie-review-data/>

² <http://ankara.lti.cs.cmu.edu/side/download.html>

In order to evaluate the effectiveness of proposed feature selection method, the performances of Information Gain (IG) [19] and Pairwise Mutual Information (PMI) [20] are taken into account for comparison. Two classification algorithms are employed to measure the strength of the proposed methods. The first is Naïve Bayes (NB) [7], which is based on the probabilistic information about text features. The second algorithm is Support Vector Machine (SVM) [6], which represents the features as points in space (feature space).

Table 2. Distribution of documents in dataset

# N°	Dataset Description	#Documents	#Class	#Words
D1	Polarity dataset V2.0	2000	2	7 956
D2	The twitter sentiment analysis dataset	10662	2	20 392

Table 3. Dimensionality reduction rate with SBFS

Dataset	# Initial features	# Features after selection	Reduction rate (%)
D1	7 956	3160	60.28
D2	20 392	7 053	65.41

The well-known *precision* and *recall* are used to compute the classification effectiveness in terms of F-measure for each category. It is computed as follows:

$$F - measure = 2 \times \frac{precision \times recall}{precision + recall} \quad (7)$$

5.2 Experimental Results

The first conducted experiments focus on reduction ratio parameter. Table 3 shows the results of dimensionality reduction rate obtained by the proposed Skyline-Based Feature Selection method (called SBFS).

In order to study the impact of the reduction rate on the classification performance, we have chosen to align the reduction rate of IG and PMI methods with the rate obtained using SBFS method. Indeed, a reduced list of features can really be obtained with our method contrary to the IG and PMI which only a ranked feature list can be provided. For example in D1 dataset, a reduction rate of around 60% was obtained with our approach, so to evaluate its effectiveness we applied the same reduction rate to the top ranked features of IG and PMI methods. Fig. 4 shows the classification accuracy according to the precision rate applied with the three features selection (IG, PMI, SBFS) using SVM and NB classifiers over the two corpora D1, D2.

As shown in Fig. 4, the highest classification performance is achieved by SBFS with all datasets. The likely reason for this result is that SBFS is based on efficient multi-criteria optimization technique (i.e., Pareto optimality) applied to each polarity class.

For a better evaluation of the results, we propose to extend the experiments by considering different subsets of the initial features sets (Table 2). Different propor-

tions ranging from 20% to 80% are taken into account in the evaluation as shown in Fig. 5. The predictive performance and usefulness of each feature subset are evaluated by NB classifier in term of F-measure. We select the subset having obtained the best result and will therefore deduct its reduction rate.

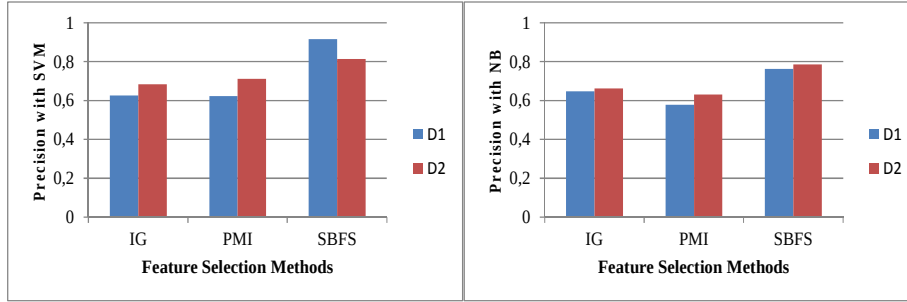


Fig. 4. Precision rate with IG, PMI, and SBFS methods

Fig. 5 shows the overall efficiency of the each method (IG, PMI) during the variation of the size in D1 and D2 datasets. As it can be seen, the best F-measure rates across different subsets are achieved at 50% and 60% of saved features respectively for D1 and D2 with IG feature selection. For PMI method, the best F-measure is achieved when 40 % and 60% of features are selected respectively for D1 and D2. These results correspond to reduction rates equivalent to 50% and 60% respectively for IG and PMI with D1 and 40% for both methods with D2. One can observe that our method SBFS is still better than IG and PMI methods.

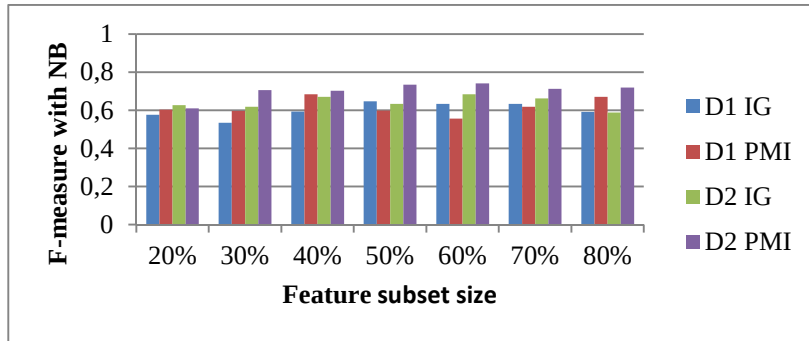


Fig. 5. Comparison of f-measure for IG and PMI on various subsets size

As shown in Fig. 6, the best classification rates across the two dataset are achieved with SBFS. However, with small dataset, it is slightly difficult for SBSF to achieve a high reduction rate; this is noticeable in negative impact on classification performance, but with the larger dataset, the reduction rate achieved by SBFS is significantly improved especially with SVM algorithm.

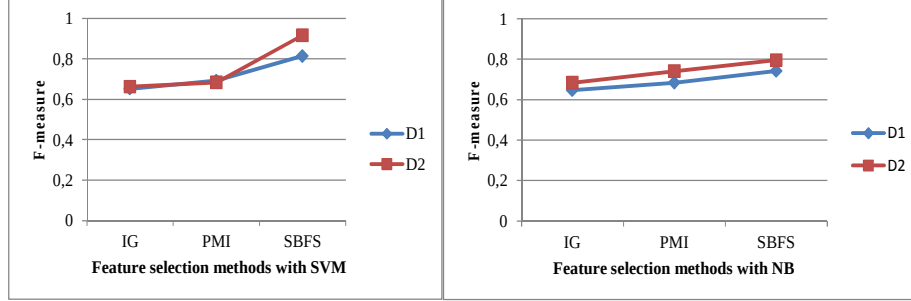


Fig. 6. Comparison of feature selection methods according to F-measure

As pointed out in [10], it is possible to obtain good classification accuracies while using less than 36% of the initial features space. This applies to SBSF as illustrated in Fig. 6 (for D2 corpus).

On the other hand, we have also conducted some experiments to analyze the required time in order to perform the learning phase that is represented in our case by the Cross-Validation including features selection step (without taking account the preprocessing step). Fig. 7 indicates that SBFS is slightly the fastest in the case of D2 with NB algorithm. But in the others cases, PMI method outperforms IG and SBFS methods. One possible explanation is that SBFS requires a double dominance check when computing the skyline, which is a time consuming task.

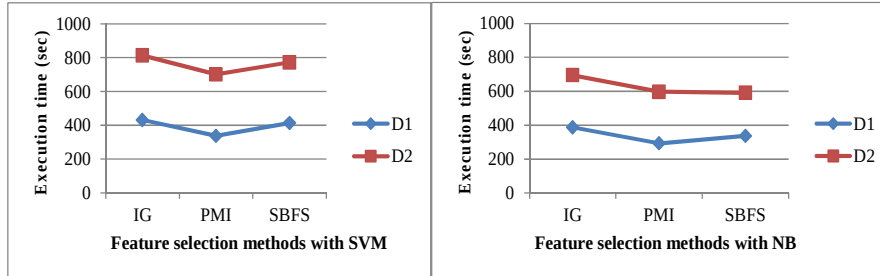


Fig. 7. Execution time (sec) using SVM and NB algorithms

6 Conclusion

In this paper, we proposed an approach that deals with the feature selection for polarity classification in sentiment opinions. The key concept of the approach is the feature skyline. Namely, the subset of features that are the most interesting (in the dominance Pareto sense) with respect to a set of relevant criteria suitably chosen. The feature skyline computed allows improving the classification performance in most cases, as shown via the conducted experiments.

As for future work, we plan first to improve the approach by using some advanced optimizations when computing the skyline. One way to improve better the

classification is to take into account more criteria (metrics). This can however lead to very large feature skylines which could be inefficient for our purpose. As second perspective consists of refining these features skyline by introducing an order relation between them thanks to a fuzzy counterpart of Pareto dominance relationship.

References

1. Maks, I., Vossen, P.: A lexicon model for deep sentiment analysis and opinion mining applications. *Decis. Support. Syst.* 53, 680-688 (2012)
2. Neviarouskaya, A., Prendinger, H., Ishizuka, M.: SentiFul: A lexicon for sentiment analysis. *IEEE Transactions on Affective Computing.* 2, 22-36(2011)
3. Boiy, E., Moens, M.F.: A machine learning approach to sentiment analysis in multilingual Web texts. *Inform retrieval.* 12, 526-558 (2009)
4. Turney, P.D.: Mining the Web for Synonyms: PMI-IR versus LSA on TOEFL. In: Raedt, L. D., Flach, P. (eds.) *ECML 2001. LNCS*, vol. 2167, pp. 491-502. Springer, Heidelberg (2001)
5. Turney, P.D.: Thumbs up or thumbs down? : semantic orientation applied to unsupervised classification of reviews. In: *40th Annual Meeting on Association for Computational Linguistics*, pp. 417-424. ACL, Stroudsburg (2002)
6. Cortes, C., Vapnik, V.: Support-vector networks. *Mach. Learn.* 20, 273-297 (1995)
7. Murphy, K.P.: Naive Bayes Classifiers. University of British Columbia, (2006)
8. Kumar, V., Minz, S.: Feature Selection. *SmartCR.* 4, 211-229 (2014)
9. Börzsönyi, S., Kossmann, D., Stocker, K.: The skyline operator. In: *17th International Conference on*, pp. 421-430. IEEE, New York (2001)
10. Abbasi, A., Chen, H., Salem, A.: Sentiment analysis in multiple languages: Feature selection for opinion classification in web forums. *ACM. T. Inform. Syst.* 26, 12 (2008)
11. O'Keefe, T., Koprinska, I.: Feature Selection and Weighting Methods in Sentiment Analysis. In: *Proceedings of the 14th Australasian document computing symposium*, pp. 67-74. (2009)
12. Duric, A., Song, F.: Feature selection for sentiment analysis based on content and syntax models. *Decis. Support. Syst.* 53, 704-711 (2012)
13. Xu, H., Zhang, F., Wang, W.: Implicit feature identification in Chinese reviews using explicit topic mining model. *Knowl-Based. Syst.* 76, 166-175 (2015)
14. Onan, A., Korukoğlu, S.: A Feature Selection Model Based on Genetic Rank Aggregation for Text Sentiment Classification. *J. Inform. Scie.* 43, 25-38 (2015)
15. Kalaivani, P., Shunmuganathan, K.L.: Feature Reduction Based on Genetic Algorithm and Hybrid Model for Opinion Mining. In *Sci. Programming-Neth.* 12, 15-26. (2015)
16. Shaw, J.: Term-relevance Computations and Perfect Retrieval Performance. *Comm. Com. Inf. Sc.* 31, 491-498 (1995)
17. Belkasmı, D., Hadjali, A.: MP2R: A Human-Centric Skyline Relaxation Approach. In: *Andreasen et al. (eds.) FQAS 2015. AISC*, vol. 400, pp. 227-241. Springer, Heidelberg (2015)
18. Hall, M., Frank, E., Holmes, G., fahringer, B.: The WEKA Data Mining Software: An Update. In: *ACM SIGKDD Explorations newsletters J.* 11, 10-18 (2009)
19. Yang, Y., Pedersen, J.O.: A comparative study on feature selection in text categorization. In *ICML.* 97, 412-420 (1997)
20. Church K.W., Hanks, P.: Word Association Norms, Mutual Information, and Lexicography. In *Comput. Linguist.* 16, 22-29 (1990)

KEY MANAGEMENT BASED AODV IN VANETS NETWORKS

Chaima BENSAID¹, BOUKLI HACENE Sofiane², FARAOUN Kamel mohamed³

¹²³ Computer science department, Djillali Liabes University at Sidi bel abbes , Sidi Bel Abbas , Algeria

Abstract. A vehicular network (VANET) is a subset of ad hoc networks where each mobile node is an intelligent vehicle equipped with communication resources (sensor). The optimal goal is that these networks will contribute to safer roads and more effective in the future by providing timely information to drivers. They are therefore vulnerable to many types of attacks. Many research studies have been proposed to secure communication in VANETs using a distributed certificate authority (DCA). We propose an approach to secure communication in VANET. Our approach has been evaluated by simulation study using the simulator network NS-2 and simulation results show that it performs.

Keywords: VANET; AODV; Ad Hoc; NS2

1. Introduction

Ad hoc networks consist of a set of self-organized of mobile nodes which cooperate using a routing protocol to facilitate the communication. They have become very popular in recent years due to their characteristics: easy deployment, lack of infrastructure, dynamic topology, mobility and minimum commissioning costs.

A VANET network is a subset of ad hoc networks where each mobile node is an intelligent vehicle equipped with communication resources (sensor). These networks have the same features as the ad hoc networks and can use the same routing protocols. However, the major problem of these networks is the security. Research studies indicate that wireless networks are more vulnerable than wired networks because of their characteristics as the dynamic topology, lack of infrastructure and the use of wireless links. A Black hole is a serious problem in VANETS, A malicious node exploits an existing vulnerability in a routing protocol such as in the Ad-hoc On Demand Distance Vector (AODV) by advertising false information as the node having the shortest path to the destination and data packets are intercepted and dropped. In this paper, we proposed a solution for detecting a Black Hole attack in VANET using as routing protocol the well-known AODV.

The rest of the paper is structured as follows. In Section 2, we describe an overview of the AODV routing protocol. In Section 3, we discuss the black hole attack. In Section 4, we present a related work, in section 5, we present our proposal. Finally, we present the simulation results and performance analysis of black hole attacks and our proposals in section 6.

2. the Routing Protocol AODV

The AODV protocol is the most widely adopted and well known reactive routing protocol in which routes are created only when they are needed [1][2]. The mobile devices or nodes in the network exchange the routing packets between them when they want to communicate with each other and maintain only these established routes. The AODV routing protocol is an adaptation of the DSDV (Destination-Sequenced Distance Vector) protocol to get dynamic link conditions [3][4]. Whenever a node wants to send the data packet to another node, it checks its routing table. If it has a fresh route to the destination node, it uses that route to send the data packet. If it does not have a route or it is not fresh enough route, then the node starts the route discovery process. So, it broadcasts Route Request message (RREQ) to its neighbors. The intermediate nodes check whether it is the destination node or it has a fresh route to go to the destination node. If it is available, the intermediate node sends back Route Reply message (RREP) to the source node. Otherwise, it forwards the RREQ message to its neighbors by using flooding approach. This process is repeated until either the destination node or a node that has a fresh enough route to the destination is found. After finishing the route discovery process, the source and the destination node can be communicate and exchange packets between them. When any node detect a link break or failure, a Route Error (RERR) message is sent to all other nodes to notify the lost of link. Hello message is used for detecting and monitoring links to neighbors.

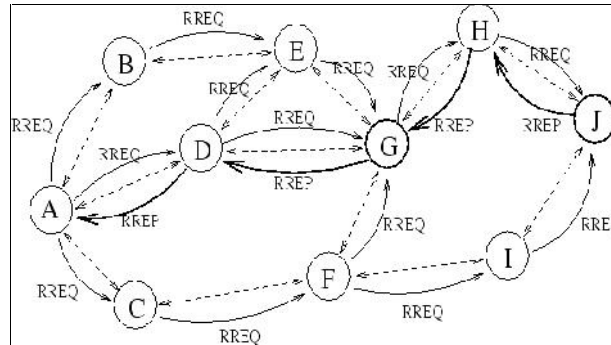


Fig. 1. RREQ and RREP packets in AODV

3. The Blackhole Attack

Blackhole attack [4] is one of the active DoS attacks possible in VANETs. In this attack, a malicious node sends a false RREP packet at the node that initiated the route discovery, in order to imposture itself as a destination node. In such a case, the source node would forward all of its data packets to the malicious node, which originally

were intended for the genuine destination. The malicious node, eventually may never forward any of the data packets to the genuine destination. As a result, therefore, the source and the destination nodes became unable to communicate with each other [5].

External node can attack by following way [6,7]:

- 1 Malicious node detects the active route and notes the destination address.
- 2 Malicious node sends a route reply packet (RREP) including the destination address field spoofed to an unknown destination address. The hop count value is set to lower values and the sequence number is set to the highest value.
- 3 Malicious node sends RREP to the nearest available node which belongs to the active route. It can also be sent directly to the data source node if route is available.
- 4 The RREP received by the nearest available node to the malicious node will relayed via the established inverse route to the data of source node.
- 5 The new information received in the route reply will allow the source node to update its routing table.
- 6 New route selected by source node for selecting data.
- 7 The malicious node will drop now all the data to which it belongs in the route.

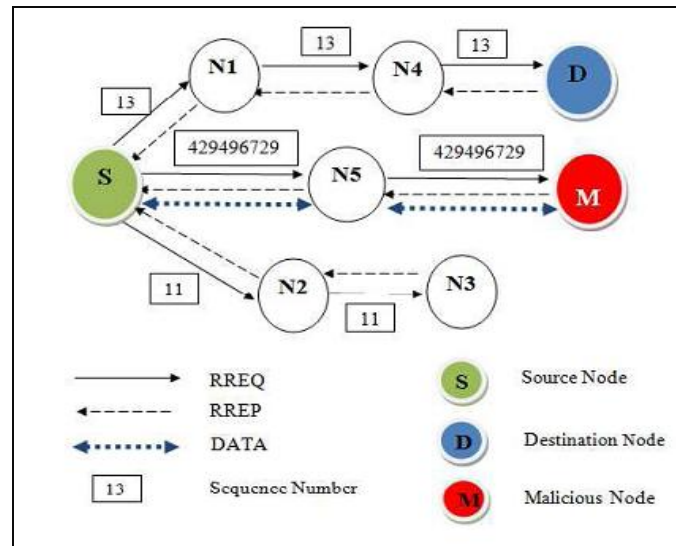


Fig. 2. The BLACKHOLE attack

The node S wants to send a data packet to the node D, and M is a malicious node. The node M responds directly to the the node D by a RREP. the node M practices a

black hole attack in the network. The attacker node can easily ignore and reject any data traffic in the network.

4. Related Works

In a related work proposed by N.R. Payal [8], the method checks the RREP destination sequence number against a threshold value which is dynamically updated [12] at every time interval. If the value is higher than the threshold, the RREP is suspected to be malicious. This method introduces ALARM packet to be sent to the neighbor nodes which contains the black list (malicious) node as a parameter. An overhead of updating threshold value at every time interval along with the generation of ALARM packet will considerably increase the routing overhead.

Zhang and Lee [9] present an intrusion detection technique for wireless ad hoc networks that uses cooperative statistical anomaly detection techniques. The use of anomaly based detection techniques results in too many number of false positives. Stamouli [10] proposes architecture for Real-Time Intrusion Detection for Ad hoc Networks (RIDAN) . The detection process relies on a state-based misuse detection system. Therefore, each node requires extra processing power and sensing capabilities. In [11], the method requires the intermediate node to send Route Confirmation Request (CREQ) to next hop towards the destination. This operation can increase the routing overhead resulting in performance degradation.

In [12], source node verifies the authenticity of node that initiates RREP by finding more than one route to the destination, so that it can recognize the safe route to destination. This method can cause the routing delay, since a node has to wait for RREP packet to arrive from more than two nodes. In

A Vani et al proposed a solution to the problem of black hole attack by comparing the sequence number with a prefixed threshold value at each time interval. If the value of received sequence number is higher than the threshold, the node is suspected to be malicious, will be added to the black list and ignore all reply messages coming through that node [13].

Subash Chandra Mandhata et al proposed a simple method that does not change the operation of the AODV protocol. It is based on two functions: collection and comparisons. The collection function consists to collect and preprocess RREP packets (Pre Process RREP). The second function compares sequence numbers of collected packets, and returns the packet with a higher sequence number if the difference is bigger. The node that sent these packets is suspected to be malicious, in this case, the source broadcasts a packet to alert neighbors containing the identity of the attacker node and all messages received from the malicious node will ignored. Each node should maintain a table of malicious nodes to isolate malicious nodes during communication [14] .

5. our proposals

The development of new communication models such as ad hoc networks make the vision of trust obsolete. Moreover, trust is not a technical problem; it is a social problem to be faced with the notion of security.

We intend to propose a new protocol based on the use of a trust model capable to secure exchanges in VANET networks, while taking into account the characteristics of these networks.

Several solutions have been proposed against the blackhole attack, we propose a method based on the study of the total behavior of each node.

We note that a blackhole node is characterized by the following characteristics:

1. A hop count = 1 and a high sequence number.
2. A number of RREQ packet sent by the blackhole = 0.
3. A number of DATA packet sent by the blackhole = 0.
4. A high number of RREP packets .

We plan to propose a new protocol based on the use of a trust model to form a secure network. Our mechanism is to provide a confidence level to each node by its traffic in the network. If a node managed to route a packet to its proper destination, its trust value increases.

Noted that the trust value V_c is a value contained in the interval $[0,1]$. So if $V_c < 0.3$, the node is added at the list of malicious nodes .otherwise, the node is selected as a trust node.

When a blackhole node receives a data packet, it removes directly, and when it receives a RREQ packet, it responds by sending a false RREP without consulting its routing table .

Based on this behavior, a legitimate node will not receive any data packet or a RREQ with a malicious node, it receives only a RREP packets.

In our model, and to assess the degree of confidence of a node, each node in the network maintains a table, in this table it saves the number of RREQ packet, RREP packet, DATA packet sent by him.

The form of the table is as follows:

	0	1			n
N_RREQ					
N_DATA					
N_RREP					

Fig. 3 the form of the cache memory

this matrix is as size of $4 \times n$ where n represents the number of nodes in the network ,Using this table, another binary matrix is generated (trust matrix),:

	0	1			n
RREQ					
DATA					
RREP					
SEQ_NUM					

Fig. 4 the form of the trust marix

This matrix is implemented as follows:

- The first line of the matrix is calculated as follows:

If(table[0][i]==0) matrix[0][i]=0;
else matrix [0][i]=1;

- The second line is calculated as follows:

If(table[1][i]==0) matrix [1][i]=0;
else matrix[1][i]=1;

- A malicious node sends a great number of RREP to intercept the data packets of its neighbors or to overload the network. For this, to fill the third line, we calculate the average of RREP packets sent by each node in the VANET network:

$$\text{average} = \sum_{i=0}^n \frac{\text{cache_memory}[2][i]}{n}$$

If(table[2][i]>average) matrix[2][i]=0;
 else matrix[2][i]=1;

- Finally the last line is filled based on the destination sequence number, it is known in AODV that a malicious node uses a very high sequence number to attract the packets of its neighbors to go through him to edit or delete them later (usually this value is 2^{32}). Why the fourth line is generated as follows:

If(table[3][i]>SEUIL) matrix[3][i]=0;
 else matrix[3][i]=1;

- Such that the SEUIL is calculated as the average between the destination sequence number stored in the RREP and the destination sequence number stored in the routing table.

Finally the trust value of each node is recorded in a trust vector; this value is generated from the trust matrix with a coefficient of each attribute among the four.

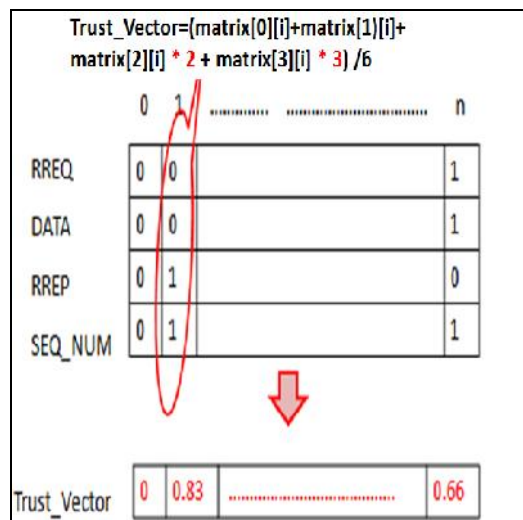


Fig.5 trust value if each node

As follows, we present the general idea of the protocol:

Step 1:

The source node S starts the route discovery phase.

Step 2:

When an intermediate node receives a RREP, it first checks if the node exists in the blacklist, if the condition is true, it deletes it directly. Otherwise it goes to Step 3.

Step 3 :

If ($V_c < 0.3$) then The node is judged Trust.

Broadcast a RREP to the source

Else Add the node to blacklist

Remove RREP

After the generation of trust value , The solution assume that all nodes have a public/private key, the network is subdivided in cluster and the cluster-head is elected using the highest trusted value and the certificate chaining is used only if there is an inter cluster communication via. Each node can obtain a certificate from cluster-head. The cluster-head issues a certificate in order to sign cluster members public-key, and the certificate is then stored in the certificate cache. The gateways issue a certificate to sign Cluster public-key and adjacent cluster gateways public certificate is stored in the gateway node.

Performance Evaluation

To evaluate the performance of our approach we used two mobility scenarios, the first is of the city of Malaga and the second is generated by the simulator vanetmobisim .

A. Malaga city scenarios:

We used three different urban VANET scenarios named U1, U2, and U3 from real areas of the downtown of Malaga, Spain [19]. These instances cover three different sizes of the same metropolitan area and they have different number of vehicles moving through the roads. Table 1 summarizes the main features of these VANET scenarios. The simulation time duration is 180 seconds with velocities between 0 and 50 km/h. [20,21].

Table1 VANET scenarios details

<i>scena rio</i>	<i>Area size</i>	<i>Numb er of vehicles</i>	<i>Number of connection</i>
U1	120 000 m ²	60	10 15
U2	240 000 m ²	60	20
U3	360 000	60	30

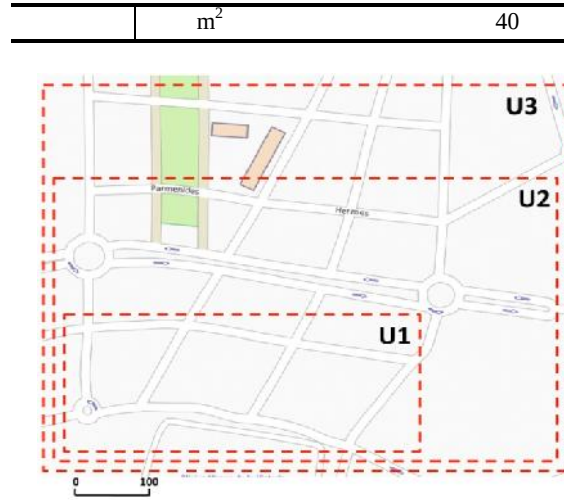


Fig. 6 Malaga urban areas

Table 2 simulation paramaters

<i>Parameters</i>	<i>value</i>
Propagation model	Nakagami
PHY layer	IEEE 802.11p
MAC layer	IEEE 802.11p
Routing layer	AODV
Transport layer	UDP
CBR packet size	1024 bytes
CBR packet rate	100kbps
Simulation time	180 s
Number of malicious node	3

B. *vanetmobisim* scenarios:

The Vehicular Ad Hoc Networks Mobility Simulator (VanetMobiSim) [22] is a set of extensions to CanuMobiSim, a framework for user mobility modeling used by the CANU (Communication in AdHoc Networks for Ubiquitous Computing) Research Group , University of Stuttgart. The framework includes a number of mobility models, as well as parsers for geographic data sources in various formats, and a visualization module. The framework is easily extensible. It is based on the concept of pluggable modules. The set of extensions provided by VanetMobiSim consists mainly on a *vehicular spatial model* using GDF-compliant data structures, and a set of *vehicular-oriented mobility models*.

We have conducted a simulation study using the famous Networks simulator ns2.35 to evaluate the performance of our Implemented approaches to:

1. AODV protocol under attack with 3 black holes with Malaga model.
2. Our approach with Malaga model .
3. AODV protocol under attack with 3 black holes with with *vanetmobisim* scenarios.
4. Our approach with *vanetmobisim* scenarios.

Table 3. simulation paramaters

<i>Parameters</i>	<i>value</i>
Propagation model	Nakagami
PHY layer	IEEE 802.11p
MAC layer	IEEE 802.11p
Areas size	1000 * 1000 m
vehicle speed	8.33 -> 13.89 m/s
CBR packet size	1024 bytes
CBR packet rate	100kbps
Vehicle number	0-15-20-30-40
Simulation time	900 s
Number of malicious node	3

The well-known metrics to evaluate routing protocol is used in our study [23]:

- **Packet Delivery Ratio (PDR):** This parameter represents the percentage of packets delivered to their destinations relative to the packet transmitted in the network.
- **The average latency of data packets (Delay):** This is the average time required to deliver data packets from the source to the destination successfully, including latency in queues, storage time in the buffer.
- **Additive costs (overhead):** The number of divided packets controls (RREQ,RREP, RERR) the number of received data packets. This criterion illustrates the amount of additives cost required for each received data packet.
- **Dropped packet:** Number of dropped packets due to either link failure or by the malicious node.
- **Certificate overhead:** the number of certificates sent in network.

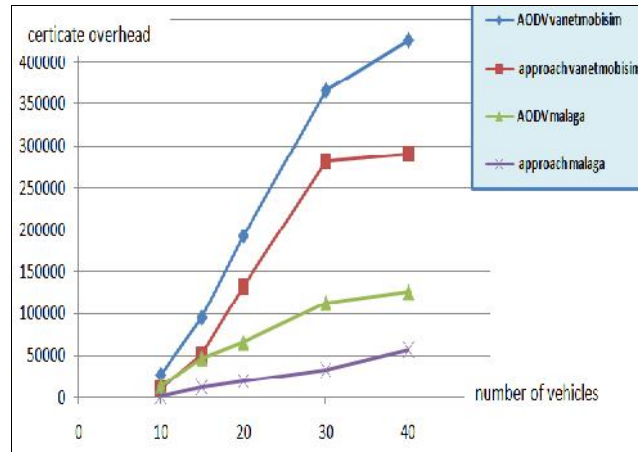


Fig. 7 certificate overhead

Fig.7 . present the evolution of the certificate overhead with the number of vehicles. We notice that the overhead is high when the number of nodes is high, especially with high mobility. This results from the fact that too many nodes request a new certificate from CA and from adjacent CA when a node migrates to surrounding cluster in the original technique. However, we depict that, our approach outperforms the original with 6.01% to 44.88% less certificates by malaga mobility model and 40.40% to 69.42% less certificates by the model generated by vanetmobisim.

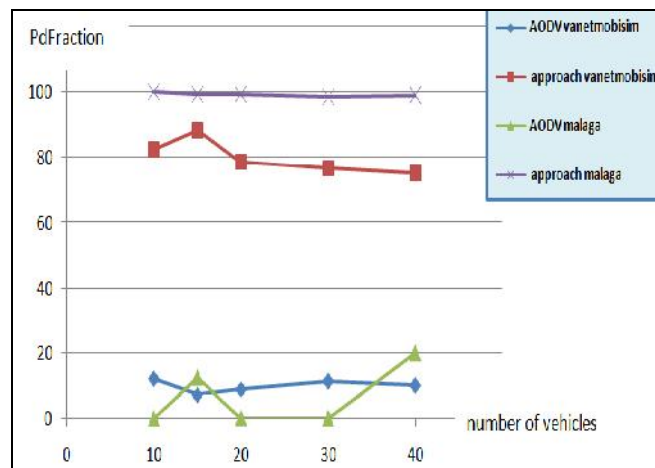


Fig. 8 PDFraction

The Figure 8 shows the evolution decreasing of PDR in protocol AODV under attack against to our approach with two mobility models. When the number of

vehicles is high, the PDR of our approach register a small degradation. This is due to frequent network topology changes with a high number of connections.

We also observe that if the number of vehicles increases the PDR increases for our solution. In the AODV protocol under attack the PDR reach 19.98% due to the high number of packets received by the fake destination (malicious node).

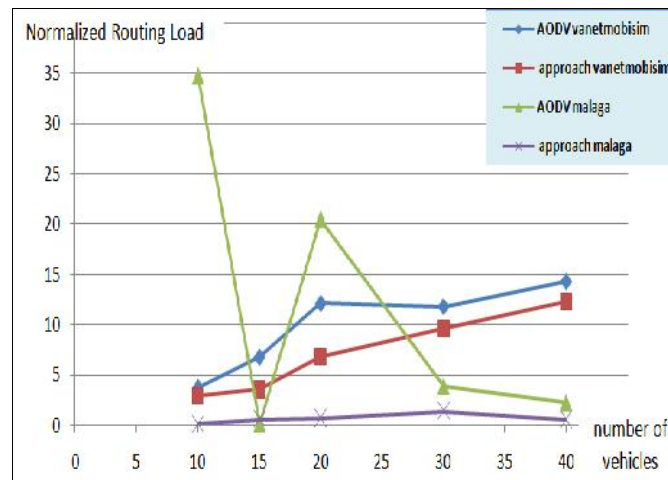


Fig.9 Normalized routing load

The simulation results in Figure 12 show that our proposal is a small routing head relative to AODV under attack. This is due to the fact that when the packet is not received in the right destination, the source node attempts to repair the paths with RRER packets.

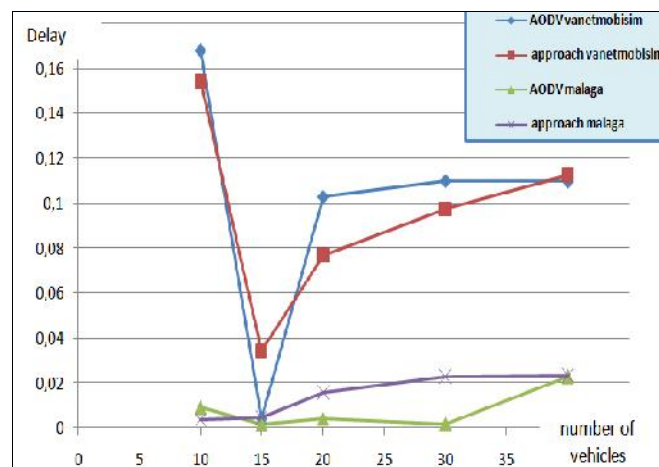


Fig. 10 Delay of communication

The Figure 10 shows the evolution of the average end to end delay, depending on number of vehicles in our approach and the AODV under attacks with two mobility models. It is found that the time required by our Proposal is higher than the AODV under attack. This can be justified by using an extra process in our solution to create the certificates and update him.

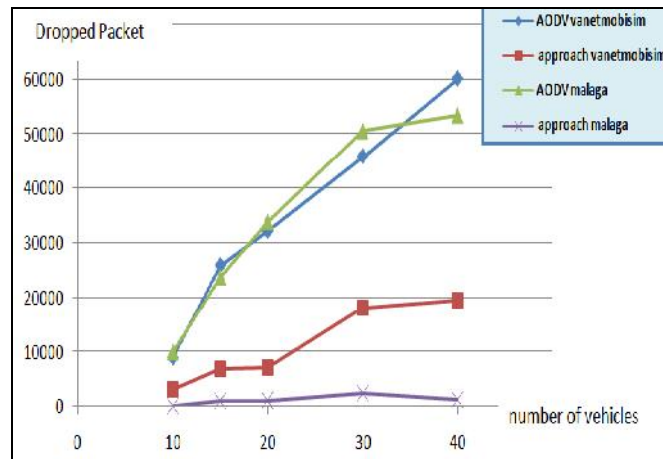


Fig. 11 Dropped packet

The figure 11 presents the number of dropped packets depending on the number of vehicles. In our proposal, with 10 sources of connection, the number of lost packets is 69 and almost no significant against to the AODV under attack 9971 packets in malaga mobility model.

Whereas the number of connections increases, it means that there is a lot of data packet sent. In this case, the malicious node intercepts a large quantity of packet so the rate of dropped packet, but in our proposal is a little minimal.

5. Conclusion

The main goal of VANET is to permit communication between vehicles in order to get safer roads and provide important information, however, this kind of network is vulnerable to many kinds of attacks because the traditional security models do not meet VANET characteristics.

In this paper, we presented the various solutions proposed by researchers on the distribution of certificates in MANET such as Partial Distribution Certification Authorities, Entire Distribution Certification Authorities Key management based on the identity and cluster based certificate chaining. We also presented our proposal to adopt a new method of distributing certificates in VANET. In Our proposal, we interested to V2V communication where the CH node acts as a virtual CA and issues certificates the cluster members and each node has a trust value generated after the study of hid behavior. The main objective of our approach is to avoid making a new

certificate request in case a node passes from a cluster to another and promotes the possibility of renewal when a certificate expires anywhere in the network. Our approach reduces certificate with 6.01% to 44.88% less certificates by malaga mobility model and 40.40% to 69.42% less certificates by the model generated by vanetmobisim.

6. References

- [1] Ochola EO & Eloff MM, "A review of black hole attack on AODV routing in MANET".
- [2] Ei Ei Khin and Thandar Phyu, "IMPACT OF BLACK HOLE ATTACK ON AODV ROUTING PROTOCOL", *International Journal of Information Technology, Modeling and Computing (IJITMC)* Vol. 2, No.2, May 2014
- [3] Charles E. Perkins & Elizabeth M. Belding-Royer, Samir R. Das (2003) Mobile Ad Hoc Networking Working Group, Internet Draft.
- [4] L. D. Zhou and J. Z. Hass "Securing ad hoc networks," *Ieee Network*, vol. 13, pp. 24-30, Nov-Dec 1999.
- [5] J. Kong. e. al, " Providing robust and ubiquitous security support for mobile ad-hoc networks," presented at the the Ninth International Conference on Network Protocols (ICNP'01), 200
- [6] CHA, J.H.J.C, and C. P. K. C. (PKI'03). "An identity-based signature from gap diffie-hellman groups," *Public Key Cryptography (PKI'03)*, 2003.
- [7] NGAI, E.C.H., and R.T.C.M.RLYU, "An authentication service against dishonest users in mobile ad hoc networks," presented at the the IEEE Aerospace Conference, 2004.
- [8] L.ESCHENAUER, V.D, and GLIGOR, "A key-management scheme for distributed sensor networks," presented at the the 9th ACM Conference on Computer and Communication Security (CCS'02), 2002
- [9] i
- [10] S.CAPKUN, L.B., J.-P., and HUBAUX, "Self-organized public-key management for mobile ad hoc networks," p. 12, 2003b.
- [11] . YI and A.R.KRAVETS, "Composite key management for ad hoc networks," presented at the First Annual International Conference on Mobile and Ubiquitous Systems: Networking and Services (MobiQuitous' 04), 2004.
- [12] S. Yi and K. R, "MOCA: mobile certificate authority for wireless ad hoc networks," in *The Second Annual PKI Research Workshop (PKI)*, 2003.
- [13] H. Y. Luo, P. Zerfos, H. J. Kong, S. W. Lu, and L. X. Zhang, "Self-securing ad hoc wireless networks," *Isc 2002: Seventh International Symposium on Computers and Communications, Proceedings*, pp. 567-574, 2002.
- [14] S.CAPKUN, L.B, and J.-P. HUBAUX, " Mobility helps peer-to-peer security " *MobileComput*, p. 8, 2006.
- [15] P. R. Zimmermann, "The official PGP user's guide" MIT Press, 1995.
- [16] M. Omar. Y. Challal. and A. Bouabdellah "Certification-based trust models in mobile ad hoc networks: A survey and taxonomy" *Journal of Network and Computer Applications* (Elsevier), Vol. 35, pp. 268-286, (2012)
- [17] G. Hahn, T. Kwon, S. Kim, and J. Song, "Cluster-Based Certificate Chain for Mobile Ad Hoc Networks," in *International Conference on Computational Science and Applications (ICCSA)* , pp. 769-778, 2006.

- [18] S. Boukli-Hacene, A. Lehireche, and A. Meddahi, "Predictive preemptive ad hoc on-demand distance vector routing," *Malaysian Journal of Computer Science*, Vol. 19, No. 2, pp. 189-195, 2006.
- [19] Jamal Toutouh, , Enrique Alba , An Efficient Routing Protocol for Green Communications in Vehicular Ad-hoc Networks,LCC, University of Malaga.2013
- [20] <http://neo.lcc.uma.es/staff/jamal/vanet/?q=node/11>
- [21] The Network Simulator Project - Ns-2. [online] Available in URL <http://www.isi.edu/nsnam/ns/>
- [22] Vanetmobisim manuel , Copyright © 2005-2006 Institut Eurécom/Politecnico di Torino
- [23] Y.-C. Hu, D.B. Johnson, and A. Perrig, "SEAD: Secure
- [24] Distance Vector Routing for Mobile Wireless Ad hoc Networks," *Proc. 4th IEEE Workshop on Mobile Computing Systems and Applications*, Callicoon, NY, June 2002, pp. 3-13.

Evolutionary Simulated Annealing Algorithm with SVM for Feature Selection and Classification

Brahim Sahmadi¹, Dalila Boughaci²

¹ LMP2M Laboratory, University Yahia Fares of Medea, Medea, Algeria
sh_brahim@yahoo.fr, sahmadi.brahim@univ-medea.dz

² LRIA/Computer Sciences Department, University of Sciences and Technology Houari
Boumediene (USTHB), Algiers, Algeria
dalila_info@yahoo.fr, dboughaci@usthb.dz

Abstract. The feature selection in data classification is a very active research field in machine learning and optimization. The combinatorial nature of the problem requires the development of specific techniques combined with several optimization methods. In this work, we propose a method based on evolutionary simulated annealing meta-heuristic combined with SVM classifier, which tries to reduce the initial size of data and to select a set of relevant variables to enhance the classification accuracy of SVM. The proposed method is evaluated on some datasets and compared with some well-known classifiers for data classification. The computational experiments show that the proposed method with optimized SVM parameters provides competitive results and finds high quality solutions.

Keywords: Evolutionary simulated annealing algorithm, support vector machine (SVM), classification, feature selection, machine learning, parameters optimization, cross-validation.

1 Introduction

Classification is one of the techniques of data mining which involves extracting a general rule or classification procedure from a set of learning examples. The term classification may refer to two distinct classes of methods: supervised classification and unsupervised classification. The unsupervised classification aims to constitute groups of examples (or groups of instances) according to the observed data, without a priori knowledge. On the other hand, the supervised classification uses a priori knowledge, on the belonging of an example to a class to construct a system which allows predicting the belonging of a new example to a class.

In a classification system, several parameters can affect its performance, notably, the quality of the features which poses problems in real applications. Some of these features may be noise-related, redundant or even unnecessary to the classification problem. Feature selection is important and necessary data pre-processing steps to increase the quality of the feature space. It aims to select a small subset of important (relevant) features from the original full feature set. It can potentially improve the performance of a learning algorithm significantly in terms of the accuracy; increase the learning speed, and simplifying the interpretation of the learnt models [1, 14]. Feature selection is used

in different tasks of learning or data mining, in the fields of image processing, pattern recognition, data analysis in bioinformatics, categorization of texts, etc. In all these domains, applications need to process data described by a very large number of attributes.

The methods used to evaluate a feature subset in the selection algorithms can be classified into three main approaches: filter methods, wrapper methods and embedded methods. Filter methods perform the evaluation independently of any classification algorithm; they are based on data and attributes [1]. Wrapper methods use the learning algorithm as an evaluation function. It therefore defines the relevance of the attributes through a prediction of the performance of the final system. Embedded methods combine the exploration process with a learning algorithm. The difference with wrapper methods is that the classifier not only serves to evaluate a candidate sub-set, but also to guide the selection mechanism.

Wrapper approaches conducts a search in the space of candidate subsets of features, and the quality of a candidate subset is evaluated by the performance of the classification algorithm trained on this subset [19]. Several wrapper methods have been proposed for feature selection, among them, the stochastic local search methods and the population based optimization metaheuristic methods, like the Genetic Algorithms (GA), the Memetic Algorithm (MA), Particle Swarm Optimization (PSO), and the Harmony Search Algorithm (HAS) [20].

In this work, we apply a wrapper method based on an Evolutionary Simulated Annealing Algorithm to the feature selection problem with the aim of aggregating an ideally minimal subset of inputs with strong discriminative power. The approach uses the SVM classifier for evaluates a feature subset.

This paper is organized as follows. Section 2 gives a general review of support vector machine methods (SVM), and evolutionary simulated annealing meta-heuristic. Section 3 describes the proposed approach for feature selection and classification. In Section 4, the obtained research results are illustrated and discussed. Finally, concluding remarks are given in Section 5.

2 Background

2.1 Support Vector Machines

Support Vector Machines (SVM) are a class of supervised learning algorithms introduced by Vladimir Vapnik [2]. The main principle of SVM is the construction of a function f called decision function that for an input vector x matches a value y , $y = f(x)$, where x is the example to classify and y is the class which corresponds to the example input. SVM are originally defined for binary classification problems, and their extension to nonlinear problems is offered introducing the kernel functions. SVM are widely used in statistical learning and has proved effective in many application areas such as image processing, speech processing, bioinformatics, natural language processing, and even data sets of very large dimensions [15].

SVM classifiers are based on two key ideas: the notion of maximum margin and the

concept of kernel function. The first key idea is the concept of maximum margin. We seek the hyperplane that separates the positive examples of negative examples, ensuring that the distance between the separation boundary and the nearest samples (margin) is maximal, they are called support vector. And as it seeks to maximize the margin, we will talk about wide margin separators [18].

Given a set of training samples $x_1, x_2, \dots, x_m$ with labels $y_1, y_2, \dots, y_m$. Let us consider a binary classification task we aim to seek the hyperplane defined with equation (1).

$$w \cdot x + b \quad (1)$$

Which separates the data into two distinct classes such that:

$$\begin{cases} w \cdot x_i + b \geq 1 - \xi_i & \text{if } y_i = 1 \\ w \cdot x_i + b \leq \xi_i - 1 & \text{if } y_i = -1 \\ \xi_i \geq 0 & \forall i \end{cases} \quad (2)$$

For no misclassifications ($\xi_i = 0$), we find the separating hyperplane which maximizes the distance between it and the closest training sample. It can be shown that this is equivalent to maximizing $2/|w|$ subject to the constraints above [3]. The second key idea in SVM is the concept of kernel function. This is transforming the data space entries in a space of larger dimension called feature space in which it is likely that there is a dividing line, in order to deal with cases where the data are not linearly separable. If no hyperplane exists (because the data is not linearly separable) we add penalty terms ξ_i to account for misclassifications. We then minimize $|w|^2/2 + C \sum_i \xi_i$ where C , the capacity, is a parameter which allows us to specify how strictly we want the classifier to fit the training data. This can be translated to the following dual problem defined in (3), which again can be solved by standard techniques [4]:

$$\begin{aligned} \text{Minimize:} \quad & \sum_i \alpha_i - \frac{1}{2} \sum_i \sum_j \alpha_i \alpha_j y_i y_j x_i \cdot x_j \\ \text{Subject to:} \quad & \begin{cases} 0 \leq \alpha_i \leq C \\ \sum_i \alpha_i y_i = 0 \end{cases} \end{aligned} \quad (3)$$

It is possible however to use a non-linear hyperplane by first mapping the sample points into a higher dimensional space using a non-linear mapping. That is, we choose a map $\phi : \mathcal{R}^n \rightarrow \mathfrak{F}$ where the dimension of $\mathfrak{F}$ is greater than n . We then seek a separating hyperplane in the higher dimensional space. This is equivalent to a non-linear separating surface in $\mathcal{R}^n$.

As shown above, the data only ever appears in our training problem in the form of dot products, so in the higher dimensional space the data appears in the form $\phi(x_i) \cdot \phi(x_j)$. If the dimensionality of $\mathfrak{F}$ is very large, this product could be difficult or expensive to compute. However, by introducing a kernel function such that $K(x_i, x_j) = (x_i) \cdot \phi(x_j)$ we can use this in place of $x_i \cdot x_j$ everywhere in the optimization problem and never need to know explicitly what ϕ is [17]. Some examples of kernel functions are:

– Linear kernel: $K(x_i, x_j) = x_i \cdot x_j$

- Polynomial kernel: $K(x_i, x_j) = (\gamma x_i \cdot x_j + r)^d, \gamma > 0$.
- RBF kernel: $K(x_i, x_j) = e^{-\frac{|x_i - x_j|^2}{2\gamma^2}}, \gamma > 0$.
- Sigmoid kernel: $K(x_i, x_j) = \tanh(\gamma x_i \cdot x_j + r)$.

Where γ , r and d are kernel parameters. The kernel parameter should be carefully chosen as it implicitly defines the structure of the high dimensional feature space $\phi(x)$ and thus controls the complexity of the final solution. Previous studies show that these two parameters play an important role on the success of SVMs [9].

2.2 Evolutionary Simulated Annealing Algorithm (ESA)

Simulated Annealing is a global research method, which was proposed by Kirkpatrick et al. [6]. It is inspired by the physical process of annealing used in metallurgy, itself based on the laws of thermodynamics set by Boltzmann. The gradual cooling of a particle system is simulated by making an analogy between the energy system and the objective function of the problem (P) on the one hand, and between the states of the system and eligible solutions for the problem in the other (see the Algorithm 1).

Algorithm 1 Pseudo code of Simulated Annealing Algorithm

Inputs : T_0 ; N ; α ; T_f ;

1. Generate randomly a current solution Sc and evaluate its energy $E(Sc)$;
 2. Set $T \leftarrow T_0$;
 3. **While** ($T > T_f$) **Do**
Begin
 - a. Set $I \leftarrow 0$;
 - b. **Repeat**
 - Select randomly a new solution Sn in the neighborhood of Sc ;
 - **If** $E(Sc) < E(Sn)$ **Then** $Sc \leftarrow Sn$;
 - Else If** $\text{random}[0, 1] < \exp((E(Sn) - E(Sc))/T)$ **Then** $Sc \leftarrow Sn$;
 - Set $I \leftarrow I+1$;
 - Until** $I = N$;
 - c. Set $T \leftarrow T * \alpha$;**End;**
 4. Final solution $\leftarrow Sc$;
-

In the simulated annealing algorithm, and in the case of a minimization problem, starting from a given configuration of the system, and subjected it to elementary modification. If this disturbance has the effect of reducing the objective function or energy ($\Delta E \leq 0$) of the system, it is accepted. Otherwise, it is accepted with probability $p = e^{-\Delta E/T}$. Applying iteratively this rule, it generates a sequence of configurations that tend toward thermodynamic equilibrium. The idea is to set the first

temperature to a high number, which gives the algorithm a very random exploration of solutions space. Then reduce the temperature in a slow manner until the temperature converges to zero.

Simulated Annealing has two drawbacks: being trapped by local minima or taking too long to find a reasonable solution. In order to overcome these drawbacks, researchers have either hybridized SA with other heuristics such as the genetic algorithms or implemented SA as parallel algorithms. The Evolutionary simulated annealing algorithm was developed to contribute to the progress in this direction. It offers an evolutionary process in which a SA algorithm is substituted for the genetic operators of crossover and mutation to evolve a population of solutions, thus it we take advantage of a diverse population provided by the GA and hill climbing provided by SA [7].

3 The Proposed Method for Feature Selection and Classification

A hybrid algorithm ESA-SVM is proposed combining features of the simulated annealing and evolutionary algorithm. It has a structure of evolutionary algorithm, which the solutions space is searched by a population of individuals. We embedded a SA into the evolutionary algorithm, which does not contain any reproduction operators (crossover, mutation) and each individual is modified independently by means of SA operator. The flowchart of ESA-SVM method is shown in Fig. 2.

3.1 Representation and Initialization

To represent the subset of selected features, we chose a binary representation of a solution in the multidimensional search space. A solution x_i is represented as an N-dimensional binary array: bit values 1 and 0 represent a selected and unselected attribute, respectively (as shown in the following figure Fig.1). The Initial population is generated randomly.

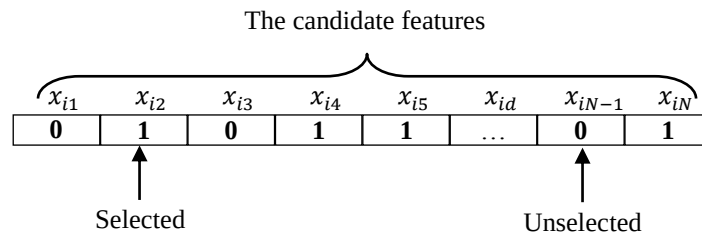


Fig.1. Binary representation of a subset of features

3.2 Objective Function

The objective function of the ESA-SVM when searching for the optimal features subset is to maximize the accuracy rate in classifying the testing dataset. This is equivalent to an optimization problem seeking a maximum solution. The classification rate *ACC* is calculated using cross-validation with 10-Folds [13]. This measure is calculated by the formula (4):

$$ACC = \left(\frac{Total\ correct}{L} \right) * 100 \quad (4)$$

Where *Total correct* is the number of examples correctly classified by the SVM classifier, and *L* is the total number of examples. The classification rate indicates whether the candidate subset permits good class discrimination.

3.3 The Evolutionary Simulated Annealing Algorithm

After initialization and parameter setting, the algorithm generates a population of individuals randomly, and for each individual operates on it with the SA operator, and evaluates whether to put it into the new population or not by a particular replacement rule. The details of ESA-SVM algorithm and SA operator are given in figure Fig. 2.

In the Simulated Annealing procedure, a new solution x'_i is randomly selected in the neighborhood of the current solution x_i . This neighborhood is defined using a move operator, which one bit change from x_i to x'_i . The algorithm always accepts x'_i as the new x_i , if x'_i is better than x_i . If x'_i is worse than x_i , there is a probability of acceptance of x'_i that depends on the current value of temperature *T*.

4 Experiments

The proposed ESA-SVM method for feature selection in data classification was implemented on a PC with an Intel Core 2 Duo CPU 2.93 GHz, 4 GB of memory and the Windows 7 operating system. The programs are coded in Java language and we have used the LIBSVM package [5] as a library for the SVMs.

4.1 Datasets

To evaluate the performance of the proposed approach, we have used 10 unbalanced datasets taken from the UCI machine learning repository [5]. Table 1 describes the characteristics of these datasets. The prediction process with the SVMs requires that the dataset must be normalized. The main advantages of such operation are to avoid attributes in greater numeric ranges dominating those in smaller numeric ranges, and to avoid numerical difficulties during the computation step. The range of each feature value can generally be linearly scaled to the range $[-1, +1]$ using formula (5), where *x*

denotes the original value, x' denotes the scaled value. Max_a is the upper bound of the feature value a , and Min_a is the lower bound of the feature value a .

$$x' = \left(\frac{x - Min_a}{Max_a - Min_a} \right) * 2 - 1 \quad (5)$$

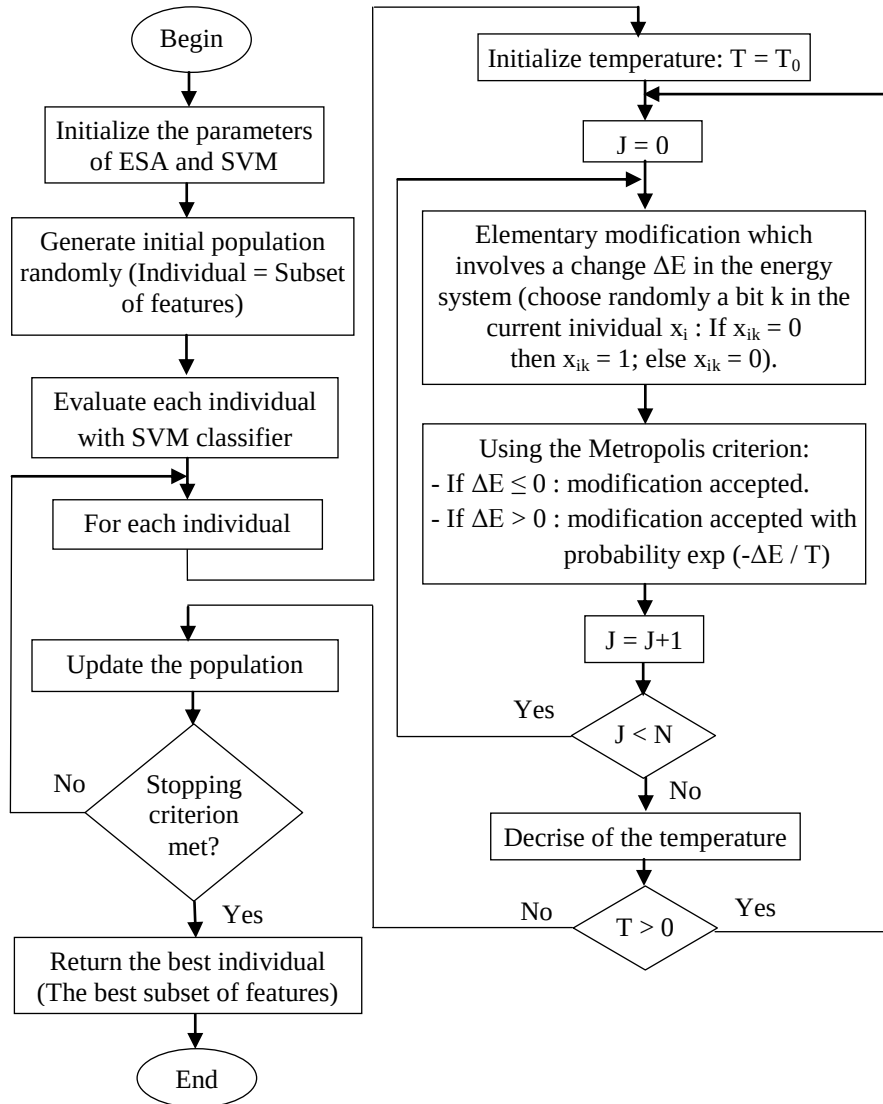


Fig.2. The flowchart of the proposed Evolutionary Simulated Annealing algorithm

Table 1. The datasets description

Dataset	Number of attributes	Number of instances	Number of classes
Breast cancer	10	683	2
Diabetes	8	768	2
Glass	9	214	6
Heart-stat	13	270	2
Ionosphere	34	351	2
Iris	4	150	3
Liver-disorders	6	345	2
Sonar	60	208	2
Vehicle	18	846	4
Vowel	10	528	11

4.2 Parameter Settings

The parameter values of the proposed algorithm are fixed by an experimental study. After a series of experiments, the different parameters are fixed empirically. The values of each parameter for the proposed method are given in Table 2.

Table 2. Parameters of ESA-SVM method

Parameters	Values
Population size	20
Maximum number of generations	50
Initial temperature	100
Number of iterations per step	10
Number of folds of cross-validation	10
Total number of runs	20

4.3 Numerical Results

Due to the non-deterministic nature of the proposed method, several executions (20) were considered for each database. The minimum value, the maximum value, and the average of the accuracy rate of the classification for each dataset are reported. The best results are in bold font. The SVM classification method is used with the RBF kernel function, and its C and Gamma parameters are optimized using an iterative search method.

Table 3 show a comparison between the results (mean of accuracy rate) obtained by the

use of SVM classifier with default parameters for RBF kernel function, the use of SVM with optimized parameters given with a grid-search method, and the results obtained with by ESA-SVM with optimized parameters. The best results are obtained with ESA-SVM algorithm for all datasets. Confirming that the feature selection and optimization of SVM parameters improves significantly the classification accuracy.

Table 3. Comparison between SVM_{default}, SVM_{Grid-search} and ESA-SVM_{Optimized}

Dataset	SVM _{default}	SVM _{Optimized}	ESA-SVM _{Optimized}
	Mean (%)	Mean (%)	Mean (%)
Breast cancer	97.04	96.70	97.47
Diabetes	77.14	77.39	78.56
Glass	58.34	71.43	75.35
Heart-stat	82.44	82.20	86.06
Ionosphere	92.17	94.66	97.01
Iris	95.63	96.50	97.93
Liver-disorders	58.41	73.59	75.17
Sonar	80.24	89.23	92.79
Vehicle	71.77	85.38	87.15
Vowel	72.83	99.31	99.81

In order to evaluate the effectiveness of the proposed method ESA-SVM, a comparison of the experimental results obtained by the method with the results of the works cited in [9, 10] is presented in the Table 4, where gives the average (Mean), the best (Max), the worst (Min) values of the classification accuracy, and the standard deviation (Sd) obtained by different methods. In [9], a hybrid search method based on both harmony search algorithm and stochastic local search, combined with a support vector machine (HAS+SVM) is given for feature selection in data classification. And the authors of [10] propose a genetic algorithm (GA) and memetic algorithm (MA) with SVM classifier for feature selection and classification.

As shown in Table 4, ESA-SVM method succeeds in finding the best results for almost the checked datasets compared to MA+SVM and HAS+SVM methods in term of classification accuracy point of view. The small standard deviations of the classifications accuracies presented show the consistency of the proposed method. The efficiency of ESA-SVM method is due to the combination of the exploration of the evolutionary and simulated annealing algorithms and parameters optimization for SVM.

Table 4. Comparison between ESA-SVM Optimized, HAS+SVM, MA+SVM and GA+SVM

Dataset	ESA-SVM Optimized				HSA+SVM				MA+SVM				GA+SVM			
	Mean (%)	Max (%)	Min (%)	Sd (%)	Mean (%)	Max (%)	Min (%)	Sd (%)	Mean (%)	Max (%)	Min (%)	Sd (%)	Mean (%)	Max (%)	Min (%)	Sd (%)
Breast cancer	97.47	97.66	97.36	0.08	97.44	97.65	97.22	0.10	97.44	97.80	97.49	0.20	97.20	97.36	96.63	0.14
Diabetes	78.56	78.78	78.39	0.11	77.79	78.13	77.47	0.17	77.71	78.26	77.34	0.21	77.53	77.99	76.82	0.27
Glass	75.35	76.64	74.30	0.68	68.80	85.04	64.48	4.08	65.00	85.98	62.62	4.03	55.28	58.41	53.74	1.39
Heart-stat	86.06	86.67	85.56	0.30	84.64	85.56	83.70	0.57	84.56	86.30	83.70	0.57	84.11	85.18	83.70	0.35
Ionosphere	97.01	97.44	96.30	0.38	93.58	95.15	92.59	0.86	93.57	96.74	92.30	1.53	93.28	93.73	92.87	0.23
Iris	97.93	98.00	97.33	0.21	--	--	--	--	96.22	97.33	96.00	0.36	95.13	96.00	94.66	0.47
Liver-disorders	75.17	75.94	74.78	0.32	72.60	73.33	71.30	0.42	71.75	73.04	70.43	0.58	63.49	64.35	62.60	0.55
Sonar	92.79	95.19	91.35	1.11	86.68	87.98	86.06	0.58	88.57	97.11	86.06	3.55	83.64	89.43	82.69	1.40
Vehicle	87.15	87.59	86.76	0.23	82.22	82.98	80.85	0.47	82.24	83.33	81.56	0.47	71.61	72.46	71.15	0.31
Vowel	99.81	100.00	99.62	0.06	90.95	91.86	90.53	0.32	93.36	94.51	93.00	0.38	76.00	78.59	59.47	4.28
Average	88.73	89.39	88.18	0.35	83.86	86.41	82.69	0.84	85.04	89.04	84.05	1.19	79.73	81.35	77.43	0.94

- The SVM with optimized parameters is used as classifier with the algorithm ESA-SVM
- -- Method did not use the dataset for test.
- The best results are in bold

5 Conclusion

In this work we proposed a wrapper method for feature selection and supervised classification based on an evolutionary simulated algorithm combined with the SVM classifier. The results obtained from tests carried out on several public databases of small and medium sizes indicate that the proposed ESA-SVM method is competitive with other meta-heuristics (Genetic algorithm, Memetic algorithm and Harmony search algorithm) for the feature selection, and experiments have shown us that the method greatly improves the learning quality and ensures the stability of the generated prediction model. It also reduces the size of the representation space by eliminating noise and redundancy.

As a continuation of this work, it would be desirable to work on the reduction of the computation time by proposing a parallel implementation of the proposed method. It is also possible to use the evolutionary simulated algorithm with other classification algorithms such as neural networks.

Acknowledgments: The authors would like to thank the developers of the Library for Support Vector Machines (LIBSVM) and the developers of Waikato Environment for Knowledge Analysis (WEKA) for the provision of the open source code. Also the authors are grateful to the reviewers of the article for their insightful comments which helped to improve the presentation of the work.

References

1. Liu H., Yu L.: Toward integrating feature selection algorithms for classification and clustering, *IEEE Trans. Knowl. Data Eng.*, vol. 17, no. 4, pp. 491–502, (2005).
2. Vapnik, V.: *The Natural of Statistical Learning theory*. Springer, New York, (1995).
3. Campbell C., Ying Y.: *Learning with support vector machines*. In: *Synthesis lectures on artificial intelligence and machine learning*, Morgan and Claypool Publishers, CA, (2011).
4. Hamel L.: *Knowledge discovery with support vector machines*. Wiley, Canada. (2009).
5. Chang CC and Lin CJ.: LIBSVM: a library for support vector machines. *ACM Transactions on Intelligent Systems and Technology*, Vol. 2, No. 3, Article 27, (2011). <http://www.csie.ntu.edu.tw/~cjlin/libsvmtools/datasets/>. Accessed 14 September 2016.
6. Kirkpatrick S., Gelatt C.D. Jr., Vecchi M.P.: Optimization by simulated annealing, *Science* 220, pp. 671–680 (1983).
7. Aydin, M.E., Fogarty: A Distributed Evolutionary Simulated Annealing Algorithm for Combinatorial Optimisation Problems. *T.C. Journal of Heuristics* 10: 269. doi:10.1023/B:HEUR.0000026896.44360.f9, (2004)
8. Eibe F., Mark A. H., and Witten I. H.: *The WEKA Workbench. Online Appendix for Data Mining: Practical Machine Learning Tools and Techniques*, Morgan Kaufmann, Fourth Edition, (2016).
9. Nekkaa M., Boughaci D.: A memetic algorithm with support vector machine for feature selection and classification. *Memetic Comp.* 7:59-73. Doi: 10.1007/s12293-015-0153-2. (7 February 2015).

10. Nekkaa M., Boughaci D.: Hybrid Harmony Search Combined with Stochastic Local Search for Feature Selection. *Neural Process Lett.* Doi: 10.1007/s11063-015-9450-5. (26 June 2015).
11. Shih-Wei Lin a,b, Zne-Jung Lee b, Shih-Chieh Chen c, Tsung-Yuan Tseng b.: Parameter determination of support vector machine and feature selection using simulated annealing approach. *Applied Soft Computing* 8, 1505–1512, (2008).
12. Grzegorz D.: Simulated Annealing Combined with Evolutionary Algorithm to Unit Commitment Problem, In *Artificial Intelligence and Hybrid Systems*. ISBN: 978-1-477554-739, iConcept Press.
13. Han J., Kamber M.: *Data mining concepts and techniques*, 2nd edn. Morgan Kaufmann, San Francisco, (2006)
14. Bonilla Huerta E.B., Duval B, Hao J.K.: A hybrid GA/SVM approach for gene selection and classification of microarray data. In: Rothlauf et al (eds) *EvoWorkshops 2006*, LNCS 3907, pp 34–44, (2006).
15. Keeman V.: *Learning and soft computing: support vector machines*. In *Neural networks and fuzzy logic models*. The MIT press, London, (2001).
16. Dash M. and Liu H.: *Feature Selection Methods for Classification, Intelligent Data Analysis*, Vol. 1, pp 131–156, (1997).
17. Keerthi, S. S., & Lin, C.-J.: Asymptotic behaviors of support vector machines with Gaussian kernel. *Neural Computation*, 15, 1667–1689, (2003).
18. Burgers C.J.C.: A tutorial on support vector machines for pattern recognition, *Data Mining Knowledge Discov.* 2, 121–167, (1998).
19. Duval B, Hao J-K.: Advances in metaheuristics for gene selection and classification of microarray data, *Briefings in bioinformatics*, Vol. 11. 1, pp 127-141, (2009).
20. El aboudi N., Benhlila I.: Review on Wrapper Feature Selection Approaches, *International Conference on Engineering & MIS (ICEMIS)*, pp. 1-5, (2016).

Magnification des images de documents dégradées issues des terminaux mobiles

Zakia Kezzoula, Soumia Faouci, Djamel Gaceb

Département Informatique (Laboratoire LIMOSE), Faculté des sciences,
Université de M'hamed Bougara Boumerdes
Route de la Gare Ferroviaire, Boumerdes 35000, Algérie

{z.kezzoula, s.faouci, d.gaceb}@univ-boumerdes.dz

Résumé. Cet article propose une nouvelle approche de super-résolution destinée aux images des documents bruitées captées en faible résolution par dispositifs mobiles. Il s'agit d'une amélioration d'une méthode non-linéaire déjà existante mais limitée par sa complexité élevée et sa faible qualité sur des images dégradées par effet de compression JPEG. Sur ce type d'images, il est nécessaire d'augmenter le périmètre d'analyse locale afin d'atteindre un meilleur rendu visuel avec une complexité réduite. Ces contraintes ont été la cible de notre contribution qui vise la linéarisation de l'approche existante par l'usage d'une approche bio-inspirée basée sur les réseaux de neurones à perceptron multicouches. Nous avons démontré que ces derniers sont capables d'apprendre le mécanisme d'une approche de super-résolution et de rendre possible son extension, vitale pour sa qualité, sans augmenter sa complexité. Il s'agit d'une nouvelle alternative d'utilisation classique des approches neuronales dans la magnification d'images.

Mots clés: Super-résolution, Interpolation d'image, sous-résolution, Images dégradées, Réseaux de neurones.

1 Introduction

La super-résolution (magnification) est une transformation permettant d'obtenir une image de haute résolution (HR) à partir d'une ou plusieurs images de basse-résolution (BR). Elle représente un domaine de recherche qui a attiré, depuis les années 90, un grand nombre de chercheurs en traitement d'image [1] [2]. La forte croissance, connue ces dernières années, peut être justifiée par un réel besoin industriel émanant de plusieurs applications: analyse et reconnaissance des images satellitaires, microscopiques, médicale, photos de tout type, vidéos, etc.). Dans le domaine de traitement automatique des documents, le passage en super-résolution devient plus que nécessaire, car la lecture optique des documents, scannés ou photographiés en faible résolution, présente des erreurs et difficultés non négligeables. Cela peut être justifié par la croissance d'utilisation des téléphones mobiles pour capturer quelques lignes de texte dans des documents de quotidien sous la présence de certains événements gênants (lieu non stable de prise de photo, ombre, poussière, etc.).

Aujourd'hui, nous témoignons la présence de plusieurs techniques de super-résolution où chacune dépend d'une application ou des spécificités liées aux types d'images. On peut les regrouper en deux grandes familles: Les techniques utilisant plusieurs images BR en entrée pour produire une image HR en sortie, les images d'entrée sont généralement issues d'une même scène prises dans différentes conditions. La deuxième famille regroupe des techniques utilisant une seule image BR en entrée pour produire une image HR en sortie. Notre contribution s'inscrit, donc, dans cette famille de techniques qui est, à son tour, divisée en trois catégories: par interpolation, transformation directe [3], [4], [5], [6] et apprentissage [7], [8], [9], [10]. D'autres méthodes de super-résolution combinent plusieurs familles de méthodes de super-résolution, comme la méthode hybride de contours-pochoirs présentée dans [12]. L'idée de cette dernière est de retrouver les contours-pochoirs d'une image BR fournie en entrée, puis appliquer ensuite sur cette image, une interpolation adaptée au type de contour retrouvé pour la transformer en image HR. Cette méthode est limitée à 12 contours-pochoirs uniquement dans un voisinage 4×4 pour rester dans des niveaux de complexité et de vitesse raisonnables. Nous avons remarqué que cette limitation a influencé considérablement la qualité de l'image magnifiée, notamment à la présence des dégradations sur des images de documents où des singularités morphologiques des caractères de texte doivent être préservées.

Nous proposons dans cet article une approche de super-résolution basée sur l'amélioration de la méthode hybride des contours-pochoirs existante. En revanche de cette dernière, notre nouvelle approche offre la possibilité d'augmenter le nombre des contours-pochoirs pour une meilleure adaptation aux dégradations. Afin d'éviter une perte significative de temps de calcul due à cette démarche, nous utilisons les réseaux de neurones (RN) et spécialement le modèle à perceptron multicouches (PMC) pour linéariser la généralisation des contours-pochoirs. Cette nouvelle technique d'utilisation d'un RN représente également une meilleure alternative aux approches neuronales existantes [11] où le RN apprend le passage direct d'une image BR vers une image HR fournie durant l'apprentissage. Les performances de l'approche hybride proposée par rapport aux approches citées ci-dessus sont confirmées par les expériences présentées à la fin de ce papier.

2 Méthodes existantes de super-résolution

Les techniques de super-résolution existantes peuvent être regroupées au départ en deux grandes familles selon le nombre d'image utilisées en entrée: La première famille regroupe les techniques qui utilisent en entrée une séquence d'images BR (souvent, une vidéo ou une rafale d'images d'une APN), pour produire une seule image HR en sortie [12]. La seconde famille regroupe des techniques qui utilisent une seule image BR en entrée, pour produire l'image HR en sortie [6], [13], [14], [15]. Les approches de la seconde famille sont à leur tour divisées en trois catégories suivantes:

2.1 Méthodes basées sur l'interpolation

Dans ces méthodes, l'interpolation d'une image repose sur le calcul des valeurs des pixels manquants dans l'image HR (à construire) à partir des pixels existants dans l'image BR (voir Fig1). En général, ces techniques sont regroupées en deux types de méthodes : a) Les méthodes classiques (ex. interpolation au plus proche voisin, bilinéaire, bicubique) sont basées sur la manipulation directe des pixels sans caractérisation statistique de contenu de l'image. Elles sont généralement rapides mais elles génèrent une sorte de floue dans l'image HR résultante. b) Les méthodes adaptatives (meilleure alternative) traitent chaque zone de l'image d'une manière différente: ex. la méthode DDT (Data Dependent Triangulation) [16] repose sur une analyse locale de la texture, la méthode NEDI (New Edge-Directed Interpolation) [17] repose sur les orientations des contours.

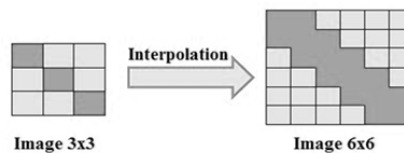


Fig. 1. Processus d'interpolation.

2.2 Méthodes basées sur la reconstruction

Ces méthodes reposent sur la reconstruction d'une image HR par transformation de l'image BR en utilisant, généralement, des opérateurs mathématiques, masques de convolution, etc. Un grand nombre de techniques de cette catégorie repose sur des transformations en ondelettes [4].

2.3 Méthodes basées sur l'apprentissage (méthodes à base d'exemples)

Ces méthodes utilisent des techniques d'apprentissage pour estimer les détails de l'image HR en sortie, et emploient souvent un dictionnaire généré d'une base de données d'images [18], [19]. Elles reposent sur deux phases: a) phase de construction du dictionnaire (création de modèle par apprentissage). b) phase de reconstruction d'une image HR en exploitant un modèle issu de la première phase. Une des variantes de ce type des méthodes est la méthodes à apprentissage par réseau de neurones [11], [20], [21], [22], [23], introduites pour régler le problème des méthodes d'apprentissage classique (coût de calcul très élevé). Ces méthodes profitent des avantages de réseau de neurones pour gérer le problème d'interpolation qui est centré sur l'optimisation de la recherche des meilleurs pixels manquants entre les pixels existants dans l'image BR.

2.4 Méthodes hybrides

Ces méthodes combinent des méthodes d'interpolation dans un concept adaptatif avec des méthodes de transformations intelligentes pour aboutir à une image HR de meilleure qualité tout en réduisant les temps de calcul d'une manière considérable. La méthode de contours pochoirs est l'une des méthodes les plus intéressantes, proposée dans ce contexte [12], elle combine une méthode d'interpolation avec la notion des contours pochoirs afin de prédire d'une manière adaptative les bords manquants dans l'image HR. Les orientations locales des contours sont détectées au départ sur l'image BR, puis utilisées en interpolation des contours pour construire les détails manquants dans l'image HR.

3 Méthode proposée : généralisation, amélioration et adaptation des contours pochoirs

Notre approche repose sur l'apprentissage du mécanisme sélectif des contours pochoirs (CP) par l'usage d'un réseau de neurones et sa généralisation permettant une analyse locale plus étendue, rapide, robuste aux dégradations et plus adaptées aux spécificités liées aux images de documents issues des caméras de dispositifs mobiles.

Il s'agit d'une amélioration de la méthode classique des CP [12] qui est limitée à 12 CP. Cette limite est l'origine de manque d'efficacité de cette approche sur des images de documents dégradées (compressions jpeg, bruit, usures). Nous avons remarqué qu'à la présence de ces artefacts sur l'image BR, il est nécessaire d'augmenter le nombre de CP et le périmètre d'analyse locale pour assurer une super-résolution de qualité (cohérence des traits et singularité des caractères préservées).

L'algorithme proposé traite chaque pixel de l'image HR en fonction de ses voisins disponibles dans l'image BR, autrement dit, il détermine pour chaque pixel, ses voisins en fonction de la fenêtre d'analyse choisie. Nous faisons appel, par la suite, à un RN au PMC (perceptron multicouches) pour déduire les pixels manquants pour la construction intégrale de l'image HR. Ce RN a été entraîné sur une base d'exemples représentative pour apprendre le mécanisme fonctionnel de la méthode de super-résolution CP et être capable de généraliser, ensuite, ses connaissances acquises sur l'ensemble de test avec des fenêtres d'analyse plus importantes dans des temps très réduits. Les entrées et les sorties du réseau doivent être normalisées par une constante empirique K . Finalement, l'image HR est reconstruite en combinant toutes les valeurs des pixels calculées par le RN. Notre algorithme repose sur les quatre étapes suivantes:

3.1 Modélisation de l'analyse locale et choix de l'architecture du RN

Dans cette étape on détermine l'architecture du RN qui dépend de la nature de l'analyse locale des pixels envisagées. Le choix de l'architecture d'un RN détermine la performance de la super-résolution ou encore sa complexité potentielle. Plusieurs paramètres doivent être définis donc, tels que le nombre d'entrées qui dépend de la

taille de la fenêtre d'analyse, le nombre de sorties, le nombre de couches cachées et le nombre de neurones par couche cachée. Il faut noter également que le plus grand problème avec les RN multicouches est de déterminer le nombre de couches cachées et le nombre de neurones pour chacune de ces couches, afin d'obtenir le meilleur rapport qualité de super-résolution/rapidité. En effet, un grand nombre de neurones augmente exagérément le temps de calcul, sans donner forcément de meilleurs résultats. Il n'existe actuellement pas de méthodes pour trouver la configuration optimale car le choix dépend fortement de la nature de l'application et de données. Pour cela, nous avons élaboré plusieurs architectures pour atteindre des meilleurs super-résolution avec les entrées/sorties de RN suivantes: a) 9 entrées et une sortie, b) 4, 9 ou 25 entrées et 5 sorties, c) 12, 16, 36 ou 64 entrées et 9 sorties. La figure suivante montre ces architectures: les cercles verts représentent les voisins de pixel dans l'image BR et les cercles bleus sont les pixels retournés par le PMC dans l'image HR.

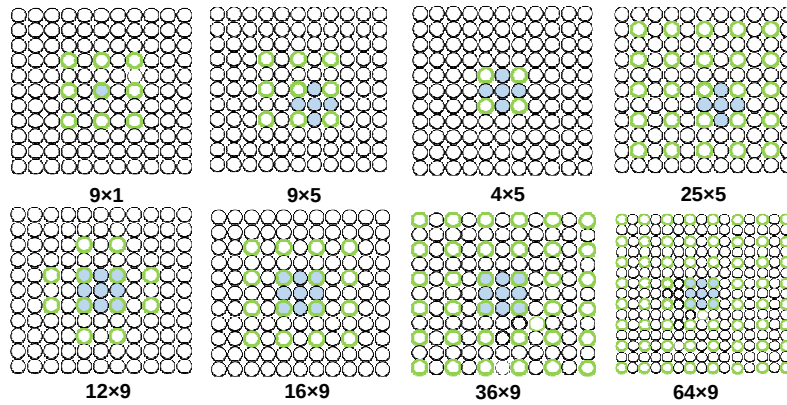


Fig. 2. Architectures du PMC.

3.2 Génération d'une base d'apprentissage représentative

Une autre difficulté est de choisir les images que l'on va faire apprendre au RN. Afin de construire une base d'apprentissage représentative, nous avons utilisé une grande variété d'images: bruitées, dégradées, floues, textuelles, graphiques, etc. Les entrées de la base ont été extraites des images résultantes de sous-échantillonnage par l'interpolation bicubique sur l'image originale, alors que les pixels de sorties désirées sont générés à partir des images originales, des images résultantes de super-échantillonnage par l'interpolation bicubique de l'image sous-échantillonnée précédemment et des images résultantes de l'interpolation par la méthode des contours-pochoirs de l'image sous-échantillonnée précédemment. Notre choix, déterminant, de la taille adéquate de la base d'apprentissage a été porté sur l'ajustement du réseau à son équilibre compromis entre sous et sur apprentissage.

3.3 Apprentissage de passage d'une résolution basse vers une haute résolution

Cette étape consiste à entraîner le réseau de neurone (PMC) en présentant des couples (entrée, sortie) au réseau. Les entrées dans ce cas représentent les voisins de pixels en cours et les sorties sont des pixels retournés par le PMC. La méthode d'apprentissage utilisée est la rétro-propagation du gradient. L'apprentissage est arrêté lorsqu'on atteint une valeur d'erreur minimale prédéfini sinon on continue l'apprentissage pour un nombre maximum d'itérations. Nous avons choisi la fonction sigmoïde comme fonction d'activation et la fonction d'erreur quadratique associée à l'algorithme de rétro-propagation comme fonction d'erreur.

3.4 Interpolation

Dans cette étape on utilise le PMC entraîné et appris précédemment pour interpoler les valeurs des pixels de l'image de HR selon l'architecture adoptée. Le PMC prend en entrée les N pixels voisins d'un pixel de l'image d'entrée (BR) et retourne comme résultats les S valeurs représentant les S pixels de l'image de sortie (HR). La figure suivante décrit la phase d'interpolation par réseau de neurones :

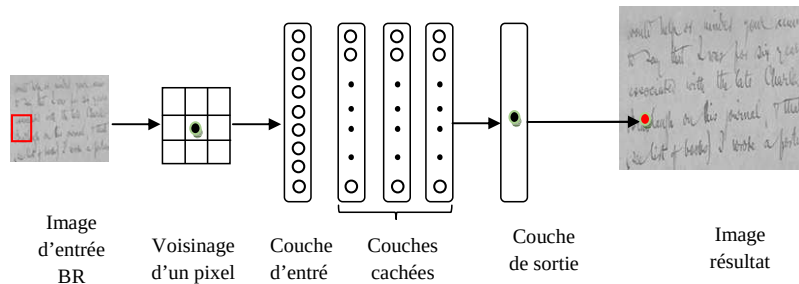


Fig. 3. Interpolation par RN.

4 Évaluation et résultats expérimentaux

Afin d'estimer la performance de notre approche, nous l'appliquons sur un ensemble de test et nous la comparons avec des méthodes bien connues, telles que les interpolations bilinéaire, bicubique et la méthode de contours-pochoirs. La comparaison repose sur la qualité d'image. Nous avons utilisé également la bibliothèque FANN (Fast Artificial Neural Network) [24] pour la mise en œuvre de notre architecture d'apprentissage et de reconnaissance RN. Pour les tests nous avons utilisé la base d'image de la compétition DIBCO de la conférence ICDAR [25].

La comparaison des méthodes citées est effectuée quantitativement en utilisant la mesure PSNR, donnée par la formule suivante :

$$PSNR = 10 \log \left(\frac{Max_I^2}{\sqrt{MSE}} \right) \quad \text{avec} \quad MSE = \frac{1}{m \times n} \sum_{i=0}^{m-1} \sum_{j=0}^{n-1} \|I(x, y) - I'(x, y)\|^2 \quad (1)$$

$I(x, y)$ est le pixel (x, y) de l'image originale, $I'(x, y)$ est le pixel (x, y) de l'image dégradée en question, m, n sont les nombres de lignes et colonnes l'image. Max_I est la valeur maximale de l'intensité lumineuse existante dans l'image originale de référence (d'une qualité de référence).

4.1 Choix de l'architecture adoptée

Afin de booster la qualité finale d'image HR résultante, nous avons appliqué une étape de post-traitement pour renforcer le contraste local des traits. Avant la comparaison avec des approches existantes, nous avons tout d'abord choisi l'architecture optimale pour notre système de RN ainsi que le meilleur type des images utilisées dans la phase d'apprentissage. Dans les tableaux qui se suivent nous considérons : *CP* (Contours-Pochoirs) ; *CP+RN-N-S* (contours-pochoirs plus RN avec une fenêtre de N pixels et S pixels en sortie) ; *ImageBicubic+RN-16-9* (RN appris à partir de l'image résultante de l'interpolation bicubique avec une fenêtre de 16 pixels et 9 pixels en sortie) ; *ImageOriginale+RN-16-9* (RN appris à partir de l'image originale avec une fenêtre de 16 pixels et 9 pixels en sortie).

Nous pouvons conclure à partir des résultats donnés dans le tableau suivant (Tab. 1) que le RN présente des meilleures performances lorsqu'il apprend sur l'image des contours-pochoirs ou sur l'image de l'interpolation bicubique par rapport à celui qui apprend directement sur l'image originale. Cela démontre qu'un apprentissage d'un processus d'une méthode de super-résolution par un RN offre une meilleure efficacité qu'un apprentissage de passage inverse fondé sur des images originales en pleine résolution. Cela est justifié par la perte d'information (ambiguïté) lors de passage vers l'image de haute résolution (retour inverse) en utilisant l'image originale.

Tab 1. Tableau représentatif de différentes versions de notre approche avec les différentes architectures d'un RN.

Méthode /Image	CP +RN-4-5	CP +RN-9-1	CP +RN-9-5	CP +RN-25-5	CP +RN12-9	CP +RN-16-9	CP +RN-36-9	ImageBicubic +RN-16-9	ImageOriginale +RN-16-9
HW1	27.38	27.38	26.74	25.24	24.71	29.03	27.09	28.51	24.86
HW2	30.52	30.52	30.63	29.73	28.85	31.71	29.75	31.46	29.15
HW3	31.32	31.32	31.86	31.86	31.24	32.74	31.21	32.57	31.37
HW4	26.02	26.02	26.05	24.59	24.5	27.14	25.15	26.29	24.58
HW5	27.6	27.6	27.72	24.77	25.27	29.19	27.03	28.43	25.40
HW6	29.08	29.08	29.18	28.14	27.93	30.21	28.05	29.84	27.74
PR1	34.03	34.03	31.82	31.83	30.46	36.33	30.73	35.71	30.50
PR2	36.00	36.00	33.71	33.89	31.36	42.09	31.90	39.50	31.74
PR3	29.19	29.19	28.56	25.76	26.76	31.66	26.04	30.41	26.41
PR4	34.37	34.37	33.38	33.22	31.46	36.90	31.99	36.58	31.57
PR5	32.77	32.77	31.42	31.74	30.28	34.58	31.36	33.92	30.43
PR7	32.00	32.00	32.01	31.83	31.36	32.84	31.37	32.43	31.47
PR8	29.97	29.97	29.12	27.51	28.46	31.18	28.02	31.23	28.55
Moyenne	30.78	23.78	30.16	29.23	28.66	32.73	29.20	32.06	28.75

Les résultats montrent que la configuration optimale pour notre approche de super-résolution est la suivante : RN appris de l'image de contours-pochoirs (apprentissage de l'approche des CP offre une meilleure qualité); Nombre d'entrées 16 (fenêtre de taille : 4×4) ; Nombre de sorties : 9.

4.2 Evaluation et comparaison avec des techniques existantes

Nous allons présenter dans cette partie les résultats des approches de super-résolution développées en les comparant visuellement ensuite quantitativement (PSNR) avec les techniques existantes.

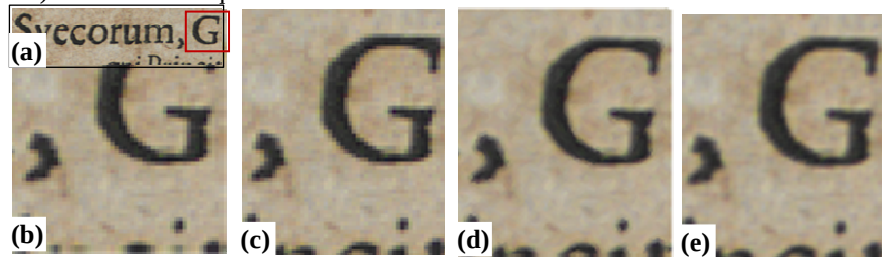


Fig. 4. Une partie d'image zoomée par facteur de 2 (a) Image originale (b) Interpolation bilinéaire (c) Interpolation bicubique (d) Contours-pochoirs (e) Notre approche.

Pour bien comparer entre les résultats des différentes approches, nous avons appliqué les différentes méthodes une deuxième fois pour un facteur de 4 sur une partie des images précédentes et voici les résultats:

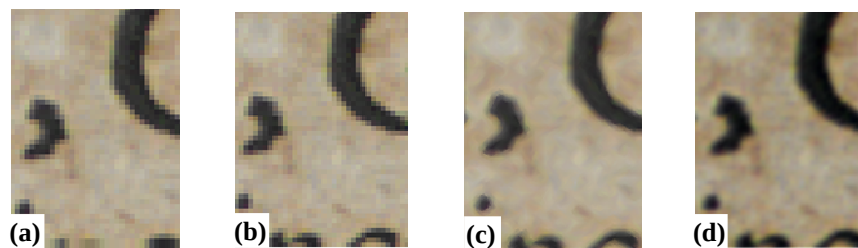


Fig. 5. Une partie d'image zoomée par facteur de 4, (a) Interpolation bilinéaire (b) Interpolation bicubique (c) Contours-pochoirs (d) Notre approche + rehaussement rapide de contraste.

Voici ci-dessous un exemple de test sur d'autres images en dehors de la base : Nous allons présenter dans cette partie un exemple de tests sur des images en dehors de la base DIBCO :

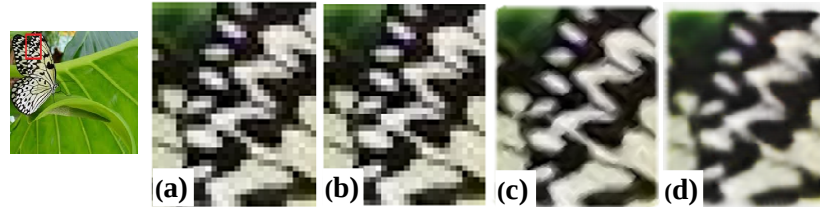


Fig. 6. Résultats de test en dehors de la base DIBCO (a) Interpolation bilinéaire (b) Interpolation bicubique (c) Contours-pochoirs (d) Notre approche.

Tab. 2. Tableau représentatif des valeurs de PSNR des images de tests (base DIBCO).

Méthodes/ Images	Interpolation bilinéaire	Interpolation bicubique	Contours pochoirs	Notre approche
HW1	23.00	23.13	28.61	29.03
HW2	27.14	27.28	31.80	33.34
HW3	29.36	29.72	32.57	32.74
HW4	22.75	26.29	27.10	27.14
HW5	22.50	22.78	28.94	29.19
HW6	26.09	26.11	30.09	30.21
PR1	28.01	28.11	35.76	36.42
PR2	28.87	29.87	39.62	42.09
PR3	24.42	24.42	31.09	31.66
PR4	29.04	29.42	36.47	38.39
PR5	28.00	28.10	34.54	35.43
PR7	30.00	30.04	32.15	32.84
PR8	26.05	26.29	30.81	31.31
Moyenne	26.55	27.04	32.27	33.06

Nous remarquons à travers les différents résultats présentés dans les figures précédentes que les méthodes d'interpolation classiques : interpolations bicubique et bilinéaire ont rendu les images floues et illisibles à cause de phénomène irréversible d'aliasing (Crénelage ou pixellisation), même si l'interpolation bicubique montre une légère supériorité que la méthode bilinéaire. Les résultats de notre approche montrent une préservation intéressante des singularités et de la cohérence des traits fins sur les bouts de caractères et de leur netteté. L'approche des contours-pochoirs présente moins de performances et a tendance de déformer les caractères de texte et les formes au passage vers des résolutions supérieures, malgré cela, elle reste meilleure que les méthodes d'interpolations classiques. Les performances de notre approche ont été renforcées par la suite en appliquant un simple rehausseur de contraste. Ceci confirme que notre approche préserve mieux les détails lors de passage vers des hautes résolutions et gère mieux l'ambiguïté lors de la prédiction des pixels inconnus grâce à sa robustesse aux différentes dégradations sur les documents et notamment sur les images bruitées (contrairement à la méthode de contours-pochoirs qui se caractérise par un manque de fiabilité au bruit).

Comme le montre les tableaux (Tab.1 et Tab.2), notre algorithme donne la moyenne de PSNR la plus élevée dans toutes les expériences. La métrique d'évaluation PSNR confirme donc la supériorité de notre approche par rapport aux méthodes: des contours-pochoirs, d'interpolation bilinéaire et bicubique. Nous pouvons en déduire donc que notre approche adaptative offre la possibilité d'obtenir une meilleure restauration d'image tout en augmentant la résolution de départ.

Le temps d'exécution de notre approche a été pris en considération également, voir (Tab.3).

Tab 3. Temps d'exécution de notre approche avec différentes architectures de RN.

Méthodes /Images	Contours-pochoirs (12 contours)	Contours-pochoirs (36 contours)	Contours-pochoirs avec RN
HW1	0.68	1.93	1.17
HW2	1.29	3.72	2.37
HW3	1.29	3.75	3.44
HW4	0.37	1.1	0.78
HW5	0.57	1.71	1.09
HW6	0.57	2.12	1.37
HW7	0.87	2.52	1.60
HW8	0.54	1.63	1.04
PR1	0.71	1.98	1.24
PR2	0.67	1.70	1.09
PR3	0.63	1.73	1.09
PR4	2.90	5.89	3.63
PR5	0.71	1.87	1.20
PR6	1.89	5.53	3.60
PR7	0.48	1.32	0.87
PR8	0.39	1.07	0.73
Moyenne	0,858	2,473	1.64

Nous constatons à travers les résultats montrés dans Tab.3 que la généralisation de la méthode de contours-pochoirs par un RN, à moindre coût temporel est possible. Celle-ci permettra de recouvrir une fenêtre d'analyse plus étendue, déterminante, pour pouvoir aboutir à une super-résolution de qualité sur des images de documents, caractérisées souvent par la présence de dégradations et des traits de caractères ou graphiques très fins. Pour de telles situations, une légère augmentation de nombre de contours pochoirs, dans l'approche classique, s'accompagne à une croissance importante, et inacceptable, des temps d'exécution avec un pourcentage de 65%. Ceci est justifié par le nombre important de chemins possibles (contours-pochoirs) pour recouvrir un périmètre d'analyse locale élevé. Notre approche, basée sur un RN, permettra d'éviter cette situation, en offrant des résultats similaires et mieux raffinés sur des dégradations que la méthode classique des CP dans des temps très réduits. Cette réduction est de 33% pour une petite fenêtre de 12 CP (voisinage de 4×4), la réduction se multiplie en utilisant des fenêtre d'analyse de plus en plus importantes.

5 Conclusion

Nous avons, dans cet article, présenté une nouvelle approche de SR adaptées aux images de documents bruitées captées en faible résolution par le biais des dispositifs mobiles. Nous avons montré à travers notre recherche bibliographique des différentes techniques de SR que certaines d'entre elles exploitent plusieurs images en entrée pour trouver l'image de sortie en résolution supérieure et d'autres utilisent une seule image de départ (BR) pour aboutir à l'image de sortie (HR). Les techniques de la seconde catégorie, plus adaptées à notre problématique, présentent trois catégories de méthodes: méthodes d'interpolation de netteté limitée, les méthodes de reconstruction et les méthodes à base d'apprentissage qui provoquent un problème de généralité nécessitant une base représentative pour aboutir à des résultats acceptables. Les méthodes hybrides gagnent en flexibilité tout en offrant le meilleur compromis entre qualité et temps de calcul. C'est pour cette raison, que notre intérêt a été porté sur une approche de cette famille. Celle-ci utilise la notion des contours pochoirs pour préserver mieux la forme des éléments de la structure physiques de document lors de passage vers l'image en haute résolution. Mais l'un de ses inconvénients majeurs dépend de la taille trop limitée de la fenêtre d'analyse locale (12 contours) qui la rend moins efficace aux images bruitées. Dans notre cas, sur des images de documents bruitées, nous avons proposé une nouvelle méthode de SR dans le but d'améliorer la méthode de contours pochoirs existante en se basant sur un RN au perceptron multi couches. Cette démarche a permis de généraliser l'approche de super-résolution pour l'usage des fenêtres d'analyse locale plus grandes sans perte en vitesse et d'adapter la méthode aux dégradations des documents. En effet, la supériorité de notre approche par rapport à l'approche existante des contours pochoirs a été confirmée par les résultats expérimentaux. Le travail que nous avons réalisé peut être complété et poursuivi sous différents aspects, notamment : Le test de notre approche sur d'autres types de RN plus évolués (profonds, RBF, etc.), l'enrichissement de la base d'apprentissage, imbriquer le rehaussement de contraste dans le RN, ceci permettra d'éliminer complètement des légères traces de flou restantes sur les bords des caractères trop dégradés.

Références

1. H. Nyquist: Certain topics in telegraph transmission theory, Transactions of the American Institute of Electrical Engineers, vol. 47, pp. 617–644 (1928).
2. R.Y. Tsai et T.S. Huang: Multiframe image restoration and registration, Advances in Computer Vision and Image Processing. (1st ed.), pp. 101–106 (1984).
3. S. Naik et N. Patel: Single image super resolution in spatial and wavelet domain, The International Journal of Multimedia and Its Applications, vol. 5, pp. 23–32(2013).
4. G. Anbarjafari, H. Demirel: Image Super Resolution Based on Interpolation of Wavelet Domain High Frequency Subbands and the Spatial Domain Input Image, vol. 32, pp. 390–394 (2010).
5. K. Zhang, X. Gao, D. Tao, X. Li: Single image super-resolution with multiscale similarity learning, IEEE Transactions on Neural Networks and Learning Systems, vol. 24, pp. 1648–1659 (2013).

6. S. Sarmadi, Z Shamsa : A new approach in single Image Super Resolution, IEEE Transactions on Image Processing, (2016).
7. C.Y. Yang, J.B. Huang, M.H. Yang: Exploiting Self-Similarities for Single Frame Super-Resolution, Asian Conference on Computer Vision (ACCV), pp. 497–510 (2010).
8. J. Yang, Z. Wang, Z. Lin, S. Cohen, T. Huang: Couple Dictionary Training for Image Super-resolution, IEEE Transactions on Image Processing, vol. 8, pp. 3467–3478 (2012).
9. M. Bevilacqua, A. Roumy, C. Guillemo, A. Morel: Single-Image Super-Resolution via Linear Mapping of Interpolated Self-Examples, IEEE Transactions on Image Processing, vol. 23, pp. 5334–5347(2014).
10. M. Bevilacqua, A. Roumy, C. Guillemot, M. Line, A. Morel: Low-Complexity Single-Image Super-Resolution based on Nonnegative Neighbor Embedding, Proceedings of the British Machine Vision Conference (BMVC), pp. 135–135 (2012).
11. Chen Change Loy: Image Image Super-Resolution Using Deep Convolutional Networks, Chinese University of Hong Kong.
12. P. Getreuer: Contour stencils for edge- adaptive image interpolation, University of California Los Angeles, Mathematics Department, U.S.A, vol. 7257 (2009).
13. C. Riedinger, M. N. Khemakhem, G. Chollet: Etude de différentes techniques de Super-Résolution de séquences vidéos. Département Traitement du Signal et des Images UMR CNRS 5141 Institut Télécom-ParisTech, PARIS, FRANCE
14. S. Naik et N. Patel: Single image super resolution in spatial and wavelet domain, The International Journal of Multimedia and Its Applications, vol. 5, pp. 23–32 (2013).
15. D. Glasner, S. Bagon, M. Irani: Super-Resolution from a Single Image, IEEE 12th International Conference on Computer Vision (ICCV), pp. 349–356 (2009).
16. J. Yang, Z. Lin, S. Cohen: Fast Image Super-Resolution Based on In-Place Example Regression, IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR), pp. 1059–1066 (2013).
17. V. Patel, K. Mistree: Review on Different Image Interpolation Techniques for Image Enhancement, International Journal of Emerging Technology and Advanced Engineering, vol. 3 (2013).
18. L. Xin, M.T. Orchard: New edge-directed interpolation, Image Processing, IEEE Transactions on Image Processing, vol. 2, pp. 1521–1527 (2001).
19. S. Fazli, M. Tahmasebi: PSO and GA based neighbor embedding super resolution, International Journal on Technical and Physical Problems of Engineering, vol. 6, pp. 17–21 (2014).
20. J. Yang, J. Wright, T. Huang, Y. Ma: Image super-resolution via sparse representation of raw image patches, Proceedings of IEEE Conference on Computer Vision and Pattern Recognition, vol. 19 (2008).
21. F. Ahmed, S.C. Gustafson, M.A. Karim: Highfidelity image interpolation using radial basis function neural networks, In Proceedings of the IEEE National Aerospace and Electronics Conference, vol. 2, pp. 588–592 (1995).
22. N. Plaziac: Image interpolation using neural networks, IEEE Transactions on Image Processing, vol. 8, pages 1647–1651(1999).
23. C. DongetC. C. LoyetK. HeImage: Super-Resolution Using Deep Convolutional Networks, IEEE Transactions on Image Processing, vol. .38, pp 295–307(2015).
24. S. Nissen: Implementation of a Fast Artificial Neural Network library (FANN). 2003.
25. I. Pratikakis, B Gatos, K. Ntirogiannis: ICDAR 2011 Document Image Binarization Contest (DIBCO 2011). In Proceedings of the 2011 International Conference on Document Analysis and Recognition ICDAR, Beijing, China, pp. 1506 –1510 (2011).

Pretrained Deep Convolutional Neural Networks for Offline Isolated Handwritten Arabic Character Recognition

Chaouki Boufenar, Adlen Kerboua, Mohamed Batouche

Computer Science Department, College of NTIC, University of Constantine 2 -
Abdelhamid Mehri,
25000 Constantine, Algeria

Abstract. Handwritten character recognition is a system widely used in the modern world and it is still an important challenge. Traditional machine-learning techniques require careful engineering and considerable domain expertise to transform raw data into a feature vector from which the classifier could classify the input pattern. To cope with this problem, the popular deep convolutional neural networks, introduced recently, have effectively replaced the hand-crafted descriptors with network features and have been shown to provide significantly better results than traditional methods. Indeed, deep learning is a powerful machine learning technique with its foundation in artificial neural networks. It is one of the fastest growing areas in machine learning, promising to reshape the future of artificial intelligence. However, the problem with deep learning is that it requires large datasets for training because of the huge number of parameters needed to be tuned by a learning algorithm. Fortunately, in practice, very few people train a deep learning network like DCNN from scratch because it is time consuming and it is relatively rare to have a dataset of sufficient size. Instead, it is common to pre-train the deep learning model on a very large dataset, and then use the pre-trained model either as an initialization or a fixed feature extractor for the task of interest. In this work, we investigate the applicability of AlexNet as pre-trained deep learning models (DCNN) on the recently proposed database for off-line isolated handwritten Arabic character, referred to as OIHACDB. The experimental results have shown that the pre-trained method competes with and even outperforms other prominent exiting methods in the field of OHACR.

1 Introduction

In real application, off-line handwritten character recognition systems are very important for the creation of electronic libraries, mail sorting, bank check, postal address and zip code recognition, to mention a few examples. This progress in this field is due to the advances in technology, development of many works and availability of international benchmarks that allowed the researchers to report in a credible way the performance of their approaches and comparing them with others that use the same databases.

Arabic characters are used by several languages such as Persian, Shahmukhi, Urdu and Jawi. However, there has been a lack of effort on handwritten Arabic character recognition (HACR) compared to Latin, Japanese and Chinese. This is due to:

- the lack of adequate support in terms of funding, and other utilities such as large benchmarks and dictionaries.
- the complexity of the Arabic script, see Table 1

This complexity comes from:

- the cursive nature of language.
- writing conditions.
- writer styles and the very similarities between some characters.

Table 1. Some Arabic isolated characters and there drawing shapes

Character name	Characters drawing shapes	Printed shape
Meem	م م م م م م م	م
Faa	ف ف ف ف ف ف ف	ف
Qaf	ق ق ق ق ق ق ق	ق
Thaa	ث ث ث ث ث ث ث	ث
Jeem	ج ج ج ج ج ج ج	ج
Taa	ت ت ت ت ت ت ت	ت
Seen	س س س س س س س	س
Sad	ص ص ص ص ص ص ص	ص
Yaa	ي ي ي ي ي ي ي	ي

The state-of-the-art research addressing the issue of HACR generally fall into the following main categories: Traditional or classical approaches based on the engineer’s expertise to extract representative features and approaches that automatically extract them from the raw image (deep learning). Fig. 1 summarizes different approaches in the field of HACR.

A traditional off-line HACR method typically employs an architecture of three main steps : pre-processing of an input handwritten character, feature extraction, and classification via machine learning methods like Support Vector Machine (SVM), Multi-Layer Perceptron (MLP), Random Forest (RF), Decision trees, etc.

Lately, a new category of machine learning methods called deep learning methods [15] has attracted a considerable amount of research and industry attention. In contrast of classical techniques, this category of algorithms deviate from the above-mentioned architecture by providing an alternative end-to-end solution to HACR without any hand-crafted feature extraction or pre-processing technique, enabling potentially high performance.

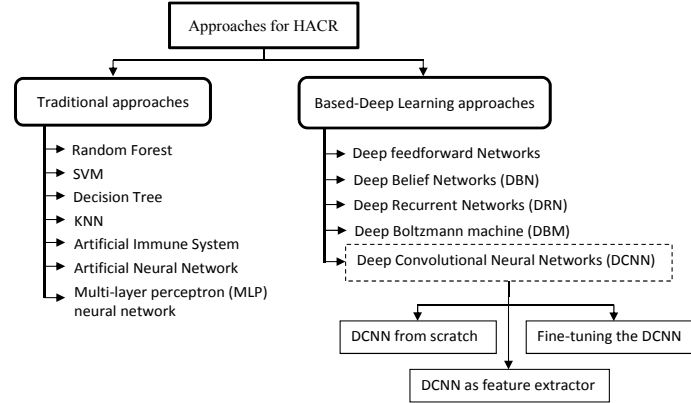
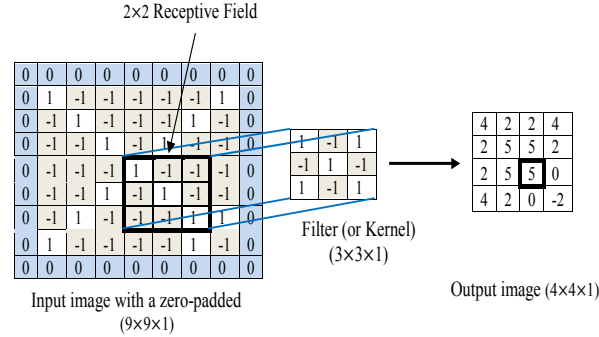


Fig. 1. HACR approaches.

One of the most prominent deep learning methods is called: Convolutional Neural Network (CNN). Yann LeCun and Yoshua Bengio introduced the concept of CNNs in 1995. A CNN is a feed-forward network with the ability of extracting topological properties that are resilient to distortions, translation, scaling, rotation and squeezing.

CNN use three mechanisms to ensure some degree of invariance in the input image: local receptive fields, weight sharing and sub-sampling [16]. Local receptive fields are connected to overlapping patches of the image. When forcing these connections to share the same weights, they become a feature map. The network is usually trained like a standard neural network by back-propagation algorithm. The basic architecture of a CNN is composed of three main layers.

- a) Convolutional layer. It consists in applying a filter to an input image. The convolution value is calculated by taking the dot product of the corresponding values in the Kernel and the channel matrices. Fig. 2 shows an example of a convolution mechanism.
- b) Pooling layer. The pooling layer is usually placed after the convolutional layer. Its primary utility lies in reducing the spatial dimensions (Width $\times$ Height) of the input volume for the next convolutional layer. The pooling layer takes a sliding window that is moved in stride across the input transforming the values into representative values. Max-pooling when the maximum value of each window is taken has been favoured over others techniques due to its better performance characteristics. Fig. 3 presents an illustration of a max-pooling operation.
- c) Fully-connected layer. The Fully Connected layer is fully connected with the output of the previous layer. Fully-connected layers are typically used in the last stages of the CNN to connect to the output layer and construct the desired number of outputs.

**Fig. 2.** The convolution operation.**Fig. 3.** The max-pooling operation.

In practice, very few people train an entire CNN from scratch (with random initialization), because it is relatively rare to have a dataset of sufficient size. Instead, it is common to pretrain a CNN on a very large dataset (e.g. ImageNet [8]), which contains 1.2 million images with 1000 categories, and then use the CNN either as an initialization or a fixed feature extractor for the task of interest. That is we called: transfer learning.

To the best of our knowledge, very little research has been conducted on the HACR with deep learning approaches. It is believed that this is due to the lack of a standard benchmark database for comparing and validating emerging approaches in HACR research field. Motivated by these ideas, we are interested in applying a pre-trained Deep Convolutional Neural Network (PDCNN) on a recently proposed database for offline HACR [5]. This work boosts the state-of-the-art in the field of HACR.

In this work, we investigate the applicability of AlexNet model [13] as a PDCNN to the HACR field. We evaluate two strategies: DCNN as feature extractor and Fine-tuning the CNN technique. We use our recently designed OIHACDB-28 database for evaluating our work.

The remaining of the paper is organized as follows. In section 2, we give a brief overview on some significant recent contributions to HACR. Section 3 presents the used architecture and section 4 describes the experiments conducted and compares their results with other leading work. Finally, section 5 presents the conclusions and some future research directions.

2 Related works

In the last two decades, many systems have been developed in the HACR field and they have achieved fairly good outcomes. Most of these systems have been based on classical methods that have insisted on hand-crafted features like Histograms, Moments, Kernels, Topological features, and so on. Numerous techniques for HACR can be reviewed in this category of approaches.

Classifiers such as feedforward Neural Networks are often used by Rashad et al in [19], Al-Jawfi in [2] and Sahlol and Suen [20] to classify handwritten Arabic character on very small databases (1000 and 750 samples, respectively) in the first two works and the CENPRMI dataset in the last.

Recently, on the same database that we use in this work, we investigated and discussed an artificial immune recognition system in the field of offline HACR [5]. We proposed a novel system composed of three main phases: pre-processing, feature extraction and classification. The system was trained and tested on an original realistic database, referred to as OIHACDB-28. Parameter tuning was performed with a grid-search algorithm and the obtained results were promising with a classification rate of 93.25% and often they outperform most well-known classifiers from Scikit Learn Library.

The major drawbacks of this category of methods is that they are a time-consuming and require an enormous effort from a domain expert engineer to design a feature extractor that transforms the raw data into an appropriate representation or feature vector from which a classifier could recognize the input pattern [15].

Deep learning algorithms are actually Deep architectures of consecutive layers. Each layer applies a nonlinear transformation on its input and provides a representation in its output [18].

The use of the deep model for the recognition of digits and various texts of different languages, such as English [7], Devanagari [17] and Chinese [6] has also attracted considerable attention.

At the best of our knowledge, a few research works using deep learning techniques are reported in HACR field. In [10] architecture based on CNN and SVM classifier was investigated to off-line Arabic handwriting recognition. The training and test sets were taken from the HACDB and IFN/ENIT databases. The performance of this model was compared with character recognition accuracies gained from state-of-the-art Arabic Optical Character Recognition, producing favourable results.

Mohamed Elleuch et al [11] investigated Deep Belief Neural Networks for handwritten Arabic word/characters recognition. Their architecture consisted of two Restricted Boltzmann Machines each one with 1000 hidden neurons. The experimental study was established on the HACDB database. The results were promising with an accuracy of 97.90%.

In [1] a new method was presented for recognizing the handwritten Farsi/Arabic digits by fusing the recognition results of a number of CNNs with gradient descent training algorithm. The experiments were conducted on extended IFH-CDB test database. The results reveal a very high accuracy classifier outper-

forming most of the previous systems. The achieved result shows 99.17% in recognition rate and was grown up to 99.98% after rejection of ten percent of hard to recognize samples.

Recently, Akm Ashiquzzaman and Abdul Kawsar Tushar [3] proposed a novel algorithm based on deep learning neural networks using appropriate activation function and regularization layer, which shows significantly improved accuracy compared to the existing Arabic numeral recognition methods. The proposed model gives 97.4% accuracy, which is the recorded highest accuracy of the dataset used in the experiment.

More recently, in [4] we investigated the applicability of Deep Convolutional Neural Network on OIHACDB-28. The model was trained from scratch under Theano framework [12] with dropout as a regularization technique to avoid overfitting to give a better generalization performance. Our results shown satisfactory recognition accuracy (97.32%) and outperform some other prominent exiting methods. However, our proposed model required a large amount of labelled training data and a great deal of expertise to ensure proper convergence.

Motivated by these last, we use in this work AlexNet as PDCNN model for Isolated HACR.

3 Used architecture

Our approach for Arabic characters recognition using a deep learning approach takes advantage of a downloadable [21] PDCNN provided in literature. These models are already trained on a large dataset such as the ImageNet database [8] with several million images. So, it is useless to waste time and resources to re-train a new model from scratch. This is a time-consuming task and it takes weeks to train with requirement of the use of large computing resources, such as an expensive GPU.

We choose to use pre-trained AlexNet model which learns a rich feature representations for a wide range of images. AlexNet is the winning solution of ImageNet Challenge 2012. This is one of the most reputed computer vision challenge and 2012 was the first time that a deep learning network was used for solving this problem.

The architecture of the AlexNet is summarized in Table 2. During training the input to our convolutional neural network is a $227 \times 227 \times 3$ images passed through a stack of different kinds of layers. We use a rectified linear units activation function (ReLU) [14], a linear activation function for the max-pooling layers, and a soft-max activation function for the output layer.

AlexNet architecture is composed of 11 layers between input and output. Lets discuss each one of them individually. Note that the output of each layer will be the input of next layer.

Before the detailed presentation of the whole process, we denote the overall structure of the algorithm, which can be divided into two principal parts: the first one deals with features extraction using convolutional layers of the CNN, while the second part concerned essentially the fine tune process of the CNN on

Table 2. The architecture used for classification on OIHACDB-28

Layer Name	Type	Function
	'data'	Input image
1	'conv1'	Convolutions 96 filters size 11×11 with stride 4 and padding 0
	'relu1'	ReLU
	'norm1'	cross channel normalization with 5 channels per element
	'pool1'	max pooling with 3×3 filters, stride 2, padding 0
2	'conv2'	convolutions with 256 filters size 5×5 with stride 1 and padding 2
	'relu2'	ReLU
	'norm2'	cross channel normalization with 5 channels per element
3	'pool2'	max pooling with 3×3 filters, stride 2, padding 0
	'conv3'	convolutions with 384 filters, size 3×3 with stride 1, padding 1
4	'relu3'	ReLU
	'conv4'	convolutions with 384 filters, size 3×3 with stride 1, padding 1
5	'relu4'	ReLU
	'conv5'	convolutions with 256 filters, size 3×3 with stride 1, padding 1
6	'relu5'	ReLU
	'pool5'	max pooling with 3×3 filters with stride 2, padding 0
7	'fc6'	fully connected layer with 4096 neurons
	'relu6'	ReLU
	'drop6'	50% dropout
8	'fc7'	fully connected layer with 4096 neurons
	'relu7'	ReLU
	'drop7'	50% dropout
9	'fc8'	fully connected layer with 28 neurons
	'prob'	softmax

our dataset. An overview of our algorithm is illustrated in Fig. 4; each step will be detailed in the remainder of this paper.

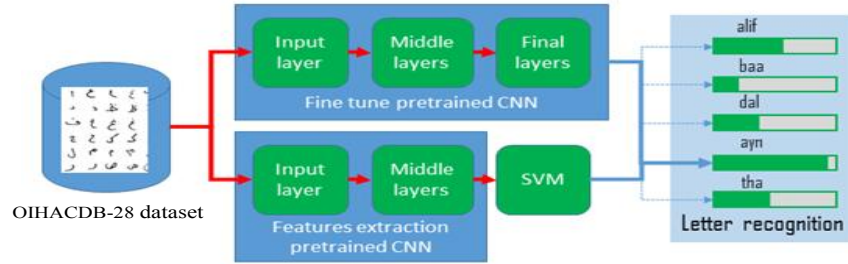


Fig. 4. Overview of our approach.

3.1 Pre-processing

Before passing the images to the DCNN, we need to perform two pre-processing operations:

- First, we segment and crop exactly the letter by computing statistics for connected regions in a binary image.
- Second, our dataset contains 5600 binary images with dimension 128×128 , as we can see in Table 2 the input layer of AlexNet initially programmed to process RGB images with three color channel. So, the input image must be $227 \times 227 \times 3$, for this we resize the images and we replicate them three times to get the right channel dimension.

3.2 Character classification

As mentioned above, we applied two deep learning methods using pre-trained CNN as fixed feature extractor and fine-tuning. We tested them with and without pre-processing, then we evaluated them in the experimental part.

a) Extract Training Features Using CNN

We extract visual features from images using pre-trained CNN, this choice avoid investing time and effort into training new one. When pre-trained CNNs are already trained using large collections images from which it can learn rich feature representations outperform hand-crafted features such as HOG, LBP, or SIFT. In our work, we use pre-trained AlexNet model. This model is trained on a subset of the ImageNet database and has learned rich feature representations for a wide range of images (1000 classes).

To extract features, we use the propriety that each layer of a CNN produces an activation, to an input image. We can easily extract features from one of the deeper layers using the activations method. We select to use the layer right before the classification layer; in Alexnet is 'fc7' layer. Which gives 4096 features by image. Notice that the first layer of the network has learned filters for capturing blob and edge features. These "primitive" features are then processed by deeper network layers, which combine the first primitives to form the higher-level ones that are better suited to recognition tasks, as they are richer in terms of image representation.

Next, we used DCNN to train a multi-class SVM classifier. A fast Stochastic Gradient Descent Solver is used for training since it helped us to speed-up the training when working with high-dimensional DCNN feature vectors (4096 neurons).

b) **Fine-tune pre-trained CNN**

We used all the layers in the pre-trained CNN except the last three. We transferred the layers to the new task by replacing those last three layers with a fully connected layer, a soft-max layer, and a classification output layer. Then, we specified the options of the new fully connected layer according to the new data initially set to 1000 classes from ImageNet. In addition, we set the fully connected layer to be of the same size as the number of classes in our dataset that is 28 classes.

We used all the layers in the pre-trained CNN except the last three ones and we transferred them to the new task by replacing those last three layers with a specified fully connected layer, soft-max layer, and classification output layer. We set the fully connected layer to be of the same size as the number of classes in our dataset which is 28 classes.

If the size of the learning images is incompatible with that of the input layer, we have to resize or crop them in the pre-processing step.

4 Experiments, comparison and discussion

4.1 Data

Large databases of HACR and words are confidential and not publicly available for non-commercial research when compared to Latin languages. Many papers in HACR have used their own small datasets. The few available databases are small and the characters are not evenly distributed over all classes. In this study, we use OIHACDB-28 which is the first version of our database OIHACDB-28 that we recently built in [5]. OIHACDB-28 contains 5600 isolated letters evenly divided into 28 classes (200 shapes per class). The dataset seems similar to MNIST dataset [9] in terms of learning complexity. The dataset can be accessed through this link: <https://www.mediafire.com/folder/74kzyuevvb8jz/Documents>.

4.2 Experiments

The algorithm is mainly written in Matlab (we used Image Processing Toolbox, Statistics and Machine Learning Toolbox, Neural Network Toolbox and Mat-

ConvNet [22]) that includes several C++ functions integrated as Mex files. The application ran on a Core i7 processor of 3.4 GHz, 16 GB of Ram. Working with CNNs can be a time-consuming task even using pre-trained model, a CUDA-capable NVIDIA GPU with computing capability of 3.0. For this, we used low cost Nvidia GPU Geforce 750m with 2 GB of Ram.

The experiment consisted of two strategies with and without pre-processing step :

- Starting with using CNN as feature extractor to classify images in our dataset. We tested the effect of the pre-processing step on the classification accuracy by training multi-class SVM classifier one-versus-all.
- Then, we evaluated the fine-tuning CNN.

We use holdout validation technique to evaluate the performance by making predictions on new data not included in training process. We randomly divide our dataset into two parts: 80% for training and the remaining 20% for testing to avoid the well-known phenomenon of over-fitting.

Table 3 illustrates the accuracies obtained with different tests. As you can see, DCNN applied as feature extractor with pre-processing achieves better score (92.94%) compared to its application without pre-processing step, when the accuracy was achieved 83.21%.

With fine-tuning strategy, when we pool-in images without modifying the pre-trained CNN, we can hold a score of 98.84% which was better than we found using the same deep learning method Depth learning with pre-processed images. The obtained results show the effect and the importance of pre-processing step in our dataset. Indeed, preprocessing step consists in making a crop to better center the images and eliminate the irrelevant pixels.

Table 3. Accuracy evaluation of the two deep learning methods using holdout validation.

Deep learning technique	Without pre-processing	With pre-processing	Average (%)
CNN as Feature extractor	83.21 %	92.95%	88.08%
Fine-tuning the CNN	98.84%	97.41%	98.12%

4.3 Comparison analysis and discussion

This section briefly touches on some comparisons between our model and recent state-of-the-art recognition systems on OIHACDB-28. For each work, Table 4 shows its publication year, the used classifier, the database used for evaluation and its accuracy. As can be seen from this table, our PDCNN using fine-tuning strategy achieves an accuracy of 98.84% with pre-processing step and outperforms our recent work that used a CNN architecture from scratch [4].

Table 4. Comparison with state of the art

Authors(Year)	Methods	Accuracy(%)
Present work (2017)	AlexNet with fine-tuning strategy	98.84
Boufenar et al (2017) [4]	DCNN from scratch	97.32
Boufenar et al (2016) [5]	AIRS	93.25

The present model achieves a better recognition rate than that obtained with random forest classifier in our previous work [5] which performed the highest accuracy (95.60%) on the same database.

5 Conclusion

Deep learning is an emerging approach within the machine learning research community. It is a hot topic in both academia and industry. In this paper, we investigated the use of AlexNet as a PDCNN model in the field of Off-line Handwritten Arabic Character Recognition (HACR). The DCNN are composed of convolution layers followed by max-pooling layers and a fully connected layer. The pre-trained models were experimented on a recently proposed database: OIHACDB-28. Two transfer learning scenarios were used: DCNN as a fixed feature extractor and fine-tuning. Experimental results have shown that pre-trained models with fine-tuning strategy provide better results with an average accuracy of 98.12% achieved on OIHACDB-28 database. As future work, we plan to:

- a) investigate other pre-trained models such as RestNet and VGGNet with data augmentation.
- b) run the proposed approach using databases other than OIHACDB-28.

References

1. Ahranjany, S.S., Razzazi, F., Ghassemian, M.H.: A very high accuracy handwritten character recognition system for farsi/arabic digits using convolutional neural networks. In: Bio-Inspired Computing: Theories and Applications (BIC-TA), 2010 IEEE Fifth International Conference on. pp. 1585–1592. IEEE (2010)
2. Al-Jawfi, R.: Handwriting arabic character recognition lenet using neural network. Int. Arab J. Inf. Technol. 6(3), 304–309 (2009)
3. Ashiquzzaman, A., Tushar, A.K.: Handwritten arabic numeral recognition using deep learning neural networks. In: Imaging, Vision & Pattern Recognition (icIVPR), 2017 IEEE International Conference on. pp. 1–4. IEEE (2017)
4. Boufenar, C., Batouche, M.: Investigation on deep learning for off-line handwritten arabic character recognition using theano research platform. In: The International Conference on Intelligent Systems and Computer Vision (ISCV), 2017 IEEE International Conference on. IEEE (2017)

5. Boufenar, C., Batouche, M., Schoenauer, M.: An artificial immune system for of-line isolated handwritten arabic character recognition. *Evolving Systems* pp. 1–17 (2016)
6. Ciregan, D., Meier, U., Schmidhuber, J.: Multi-column deep neural networks for image classification. In: *Computer Vision and Pattern Recognition (CVPR)*, 2012 IEEE Conference on. pp. 3642–3649. IEEE (2012)
7. Ciresan, D.C., Meier, U., Gambardella, L.M., Schmidhuber, J.: Convolutional neural network committees for handwritten character classification. In: *2011 International Conference on Document Analysis and Recognition*. pp. 1135–1139. IEEE (2011)
8. Deng, J., Dong, W., Socher, R., Li, L.J., Li, K., Fei-Fei, L.: ImageNet: A Large-Scale Hierarchical Image Database. In: *CVPR09* (2009)
9. Deng, L.: The mnist database of handwritten digit images for machine learning research. *IEEE Signal Processing Magazine* 29(6), 141–142 (2012)
10. Elleuch, M., Maalej, R., Kherallah, M.: A new design based-svm of the cnn classifier architecture with dropout for of-line arabic handwritten recognition. *Procedia Computer Science* 80, 1712–1723 (2016)
11. Elleuch, M., Tagougui, N., Kherallah, M.: Arabic handwritten characters recognition using deep belief neural networks. In: *Systems, Signals & Devices (SSD)*, 2015 12th International Multi-Conference on. pp. 1–5. IEEE (2015)
12. Frédéric, B., Pascal, L., Razvan, P., James, B., Ian, J.G.: Theano: new features and speed improvements. In: *Deep Learning and Unsupervised Feature Learning NIPS 2012 Workshop* (2012)
13. Krizhevsky, A., Sutskever, I., Hinton, G.E.: Imagenet classification with deep convolutional neural networks. In: *Pereira, F., Burges, C.J.C., Bottou, L., Weinberger, K.Q. (eds.) Advances in Neural Information Processing Systems 25*, pp. 1097–1105. Curran Associates, Inc. (2012), <http://papers.nips.cc/paper/4824-imagenet-classification-with-deep-convolutional-neural-networks.pdf>
14. Krizhevsky, A., Sutskever, I., Hinton, G.E.: Imagenet classification with deep convolutional neural networks. In: *Advances in neural information processing systems*. pp. 1097–1105 (2012)
15. LeCun, Y., Bengio, Y., Hinton, G.: Deep learning. *Nature* 521(7553), 436–444 (2015)
16. LeCun, Y., Bottou, L., Bengio, Y., Haffner, P.: Gradient-based learning applied to document recognition. *Proceedings of the IEEE* 86(11), 2278–2324 (1998)
17. Mehrotra, K., Jetley, S., Deshmukh, A., Belhe, S.: Unconstrained handwritten devanagari character recognition using convolutional neural networks. In: *Proceedings of the 4th International Workshop on Multilingual OCR*. p. 15. ACM (2013)
18. Najafabadi, M.M., Villanustre, F., Khoshgoftaar, T.M., Seliya, N., Wald, R., Muharemagic, E.: Deep learning applications and challenges in big data analytics. *Journal of Big Data* 2(1), 1 (2015)
19. Rashad, M., Amin, K., Hadhoud, M., Elkilani, W.: Arabic character recognition using statistical and geometric moment features. In: *Electronics, Communications and Computers (JEC-ECC)*, 2012 Japan-Egypt Conference on. pp. 68–72. IEEE (2012)
20. Sahlol, A., Suen, C.: A novel method for the recognition of isolated handwritten arabic characters. *arXiv preprint arXiv:1402.6650* (2014)
21. Team, T.M.: Cnns for matlab (2016), <http://http://www.vlfeat.org/matconvnet/>
22. Vedaldi, A., Lenc, K.: Matconvnet: Convolutional neural networks for matlab. In: *ACM International Conference on Multimedia* (2015)

Smart-Diffusion des gradients et prétraitement des images de documents en couleurs

Ryma Benabdelaziz , Djamel Gaceb , Roza Benabdelaziz
Département Informatique (Laboratoire LIMOSE), Faculté des sciences,
Université de M'hamed Bougara Boumerdes
Route de la Gare Ferroviaire, Boumerdes 35000, Algérie

{r.benabdelaziz, d.gaceb}@univ-boumerdes.dz, holarose@live.fr

Résumé. La lecture optique des documents administratifs est un domaine très exigeant en qualité et vitesse. L'extraction du premier plan, une étape systématique dans les OCR, trouve rapidement ses limites sur des images de documents dégradées. Il est nécessaire donc de prétraiter ces images avant la binarisation. Les approches linéaires comme le filtre gaussien sont incapables de préserver la netteté des traits durant le lissage. En revanche, les approches non-linéaires comme le filtre de Perona et Malik préservent mieux le tracé mais perdent en vitesse et qualité sur certains types de dégradations. Pour réduire ces limites, nous proposons, ici, une nouvelle approche de filtrage intelligent qui s'adapte localement aux dégradations présentes sur l'image. Nous exposerons les résultats obtenus sur la base d'images de la compétition DIBCO de conférence ICDAR.

Mots clés: Prétraitement des images en couleurs, binarisation, diffusion des gradients, filtrage adaptatif, évaluation de la qualité.

1 Introduction

Depuis quelques années, la technologie évolue très rapidement dans le domaine de l'acquisition d'images. Cette évolution s'accompagne d'une forte croissance de développement d'applications de vision artificielle sortant chaque jour aux niveaux industriel, académique et même chez le grand public grâce à l'usage des cameras, intégrées systématiquement dans des Smartphones. La vision artificielle, qui fait partie du domaine de l'intelligence artificielle, a pour objectif d'imiter certaines facultés de l'homme à percevoir, analyser et reconnaître le monde qui l'entoure. Pour une machine, cette tâche est très difficile du fait que l'information disponible, sous forme d'une image numérique, n'est qu'une projection ou sous-échantillonnage de la scène réelle. Cette reproduction ou numérisation entraîne souvent une perte importante d'informations lors de passage de l'analogique au numérique, de plus l'information disponible n'est pas toujours parfaite (numérisation des capteurs, déformation des objectifs, bruit, flou et présence de dégradations).

Les images de documents numérisés en couleurs contiennent environ 65 000 fois plus d'informations visuelles que les images en niveau de gris. Malgré toute cette

richesse colorimétrique dans la même image, les variations colorimétriques locales causées par différents éléments perturbateurs dégradent la phase de la segmentation, de l'analyse et de la reconnaissance optique par un OCR, étapes importantes dans tout système de lecture automatique des documents. Le prétraitement d'image joue donc un rôle déterminant dans ces systèmes afin d'améliorer leur rendement par la réduction de toute sorte de dégradations présentes dans des images de documents (bruit, flou, mauvais éclairage, fausses couleurs, ombre, etc.).

Au début de cet article, nous présentons un aperçu des méthodes de filtrage d'images existantes dans la littérature. Nous présentons, ensuite, le principe de la diffusion des gradients que nous améliorons, puis notre approche de filtrage plus adaptée aux images de documents. Nous exposons, à la fin, les résultats issus des tests effectués sur la base de la compétition DIBCO2011 de la conférence ICDAR.

2 Méthodes existantes

Vu l'importance de prétraitement par filtrage d'image (lissage), plusieurs méthodes ont été développées aujourd'hui dans différents domaines applicatifs. Ces méthodes peuvent être regroupées en catégories selon la technique utilisée (d'autres classements de méthodes dépendent du type de bruit). Le filtrage classique regroupe les filtres linéaires et non-linéaires. Les filtres linéaires (moyen, gaussien, etc.) induisent, malheureusement, une perte importante d'informations pertinentes, par apparition de flou sur les traits et les contours. Ce problème a fortement motivé l'apparition des filtres non-linéaires, comme le filtre médian, connu par sa simplicité et son efficacité en termes de suppression de points aberrants isolés (bruit impulsionnel). Sauf que le temps de calcul devient beaucoup plus élevé avec l'augmentation de la taille de la fenêtre d'analyse locale (supérieur à 5) ; problème traité dans [1]. Malgré son efficacité, des pertes de détails dans l'image ou des distorsions de certaines formes d'objets ont été remarquées [2]. Pour améliorer la qualité de filtrage, des filtres sélectifs ont été développés sous forme de filtre bilatéral (qui représente une amélioration de filtre médian permettant d'éviter ses défauts), tel que cité dans [3], est considéré comme très robuste et rapide tout en préservant les contours. Nous rencontrons également le filtrage par seuillage, introduit dans les méthodes basées sur les ondelettes qui apparaît dans plusieurs œuvres tels que [4] [5]. Selon [6], le filtrage basé sur les ondelettes est une extraction d'une structure cohérente du signal traité en considérant le bruit comme inconsistant à la base de l'ondelette choisie. Avec l'introduction du seuillage, les coefficients d'ondelettes qui sont supérieurs à un seuil donné sont considérés comme faisant partie du signal informatif de l'image [3]. Dans [7], l'aspect adaptatif est introduit dans le lissage de l'image avec l'application du filtre moyen, cette technique adaptative implique le calcul d'un masque local de convolution propre à chaque pixel de l'image. Ce dernier est appliqué ensuite séparément à chaque composant de couleur. Cette technique semble être intéressante, mais elle peut également causer un lissage léger dans des zones non bruitées. Il y a une autre catégorie de filtre, très utilisée dans des domaines plus exigeants en qualité, qui repose sur la diffusion anisotropique, dans laquelle le taux de diffusion est contrôlé par la direction des caractéristiques locales [8]. Dans le domaine adaptatif, on

cite l'approche EDP basée sur les équations aux dérivées partielles [9]. Dans le même contexte, Perona et Malik (PM) ont utilisé l'équation de la diffusion de la chaleur pour lisser l'image d'une manière adaptative [10]. Cette approche était la base de plusieurs applications en traitement d'images (traitement du bruit, de flou, fausses couleurs générées par la compression JPEG, détection des contours, segmentation, etc.). Dans cette approche, la valeur d'un pixel est modifiée par déplacement dans l'espace colorimétrique à travers des itérations en fonction de la variation du gradient. Cela offre à cette approche une meilleure précision dans les zones ambiguës et donc une meilleure préservation des contours et des traits.

Plus loin, nous trouvons également une méthode de filtrage variationnelle (ROF) [12] proposée par Rudin et al. [13], permettant une meilleure préservation des détails de l'image par minimisation de la longueur des contours. Cette technique a rapidement trouvé ses limites sur les images de documents fortement textuels en générant l'effet d'escalier, ce qui a motivé l'apparition d'une amélioration de cette méthode sous le nom (ROF2). Dans le contexte du filtrage intelligent on trouve une technique adaptative [14], une amélioration de la méthode SRAD [15], destinée uniquement aux images monochromes bruitées. Or, dans la méthode intelligente, les trois composants colorimétriques RVB sont prises en compte pour l'estimation et le filtrage de bruit. Le processus itératif, qui repose sur les équations aux dérivées partielles, permet de préserver les contours. Cette méthode est très limitée sur certaines dégradations ou sur des images texturées, et donc a tendance à éliminer des détails de l'image par le flou.

Nous pouvons conclure que ces techniques de prétraitement ne sont pas génériques à toutes les images et manquent d'efficacité sur certains types de bruit ou dégradation, qui sont spécifiques à la nature et à la complexité des images de documents qu'on souhaite traiter avec un risque de perte minimal.

Dans ce contexte, nous proposons une nouvelle approche qui agit de façon intelligente et améliore la performance des techniques existantes. Notre technique utilise d'une manière adaptée une diffusion du gradient pilotée par la qualité locale de l'image, estimée par l'analyse des variations colorimétriques locales. Ceci permet d'appliquer le traitement approprié sans déformer les contours et perdre les formes singulières des caractères et des traits. Cette préservation est essentielle pour l'étape de séparation premier/arrière-plan et donc la reconnaissance des caractères de l'image du document (voir la section 4).

3. Principe de la diffusion de gradient dans un espace 2D

L'analyse EDP existe en deux versions : isotropique et anisotropique.

1. La diffusion isotropique est un traitement qui ne présente aucune direction privilégiée, elle considère le bruit comme un signal à haute fréquence, donc les méthodes les plus utilisées sont souvent les opérations de convolution spatiale (filtre gaussien utilisé pour le processus de réduction de bruit). Ce lissage permet d'atténuer les bruits ainsi que les autres signaux à hautes fréquences (à préserver) comme les contours de tracé.

2. La diffusion anisotrope permet de distinguer entre bruit et contours à préserver [10]. Le lissage anisotrope agit, donc, d'une manière différente sur chaque partie de l'image en se basant sur la norme de gradient pour différencier les zones homogènes des contours. Ceci est réalisé à l'aide d'une fonction de diffusion anisotrope qui agira fortement dans les zones à faible gradient (les zones homogènes) et faiblement dans les zones à fort gradient (les contours et les traits de caractères). La formule suivante représente la fonction de base que Perona et Malik ont proposée pour diffuser l'image:

$$\begin{cases} \frac{\partial I(x, y, t)}{\partial t} = \text{div}(c(\|\nabla I(x, y, t)\|) \nabla I(x, y, t)) \\ I(x, y, 0) = I_o(x, y) \end{cases} \text{ avec } c(\|\nabla I\|) = \exp\left(-\left(\frac{\|\nabla I\|}{K}\right)^2\right) \quad (1)$$

où, c est une fonction de diffusion positive décroissante, k est un opérateur de divergence qui représente un seuil et peut être fixé empiriquement par estimation de bruit par l'utilisateur.

Perona et Malik proposent d'utiliser l'histogramme cumulé du gradient. La fonction de diffusion c risque d'alourdir le temps de calcul, pour le réduire, Perona et Malik ont proposé de remplacer la formule précédente de c par la suivante :

$$c(\|\nabla I\|) = \frac{\lambda^2}{(\lambda^2 + d^2)} \quad (2)$$

Où, d représente la différence entre le pixel central et ses voisins (somme des différences). Après les tests que nous avons faits, nous avons constaté que la seconde formule offre un filtrage de meilleure qualité que celui de la première formule. Bien que la technique de PM donne de très bons résultats, la diffusion reste limitée à un traitement local de 4 voisins dans une fenêtre de taille 3×3 (deux voisins horizontaux et deux voisins verticaux), ceci n'a pas été suffisant dans certains cas. Le problème avec une telle limitation de voisinage rend la méthode incapable de diffuser avec précision dans des zones de tracé qui présentent une orientation diagonale. Ceci cause une sorte de flou ou déformation dans les sommets des caractères.

Dans [11], les auteurs proposent une technique de traitement de bruit à base des tenseurs utilisant une matrice Hessienne, le modèle employé est formulé de la façon suivante:

$$\frac{dI_i}{dt} = \text{Trace}(D \cdot H_i)_{i=1 \dots n} \quad (3)$$

Avec, D est la matrice de diffusion des tenseurs et H est la matrice Hessienne de l'image, $n= 3$ pour une image RVB. Cette technique traite, d'une part, le bruit, et d'autre part, les dégradations telles que les fausses couleurs, le mauvais contraste, et un certain degré de flou dû au bruit, et cela tout en préservant les contours. Cette technique s'exécute en deux grandes étapes : a) Calcul d'une matrice Hessienne (2×2) pour chacun des pixels de l'image dans son voisinage. b) Calcul de la matrice de diffusion en se basant sur la matrice Hessienne déjà calculée à partir de ses valeurs et vecteurs propres associés. Après avoir présenté ces deux étapes, il reste l'application du model (EDP) approprié. Nous nous sommes aussi inspirés du model proposé dans [11] pour l'intégrer dans notre technique que nous détaillerons plus tard. Ce modèle

est basé sur les équations aux dérivées partielles d'ordre un et deux, que nous utilisons dans notre approche.

4 Notre approche de filtrage par diffusion guidée par la qualité

Après l'élaboration d'un grand nombre de résultats des approches de diffusion, nous constatons que certaines régions trop dégradées de l'image exigent une intense diffusion par rapport aux autres (voir la section 5). Certaines régions n'ont pas besoin d'être diffusées, car elles sont déjà homogènes et en bonne qualité, et donc, il ne sera pas très intéressant de les diffuser afin d'éviter de leurs introduire inutilement quelques perturbations colorimétriques systématiques après l'application de la diffusion. C'est dans cet esprit de principe que notre approche a été conçue. À l'inverse des approches classiques, un tel ciblage de la diffusion permettra d'augmenter la qualité de l'image résultante tout en réduisant les risques de suppressions de certains bouts de tracé et les temps de calcul. Elle consiste à appliquer localement et exclusivement la diffusion (voir section 3) sur des zones contenant des fortes fluctuations. Cette application est remise en cause à chaque pixel par une étape d'évaluation de la qualité de l'image par une appréciation rapide de l'homogénéité locale de voisinage qui dépend des zones à faibles amplitudes de gradient, très bons indicateurs de la qualité. Cette approche devrait répondre aux critères suivants: réduction efficace du bruit sans altération des zones d'image non bruitées, préservation du contraste, des contours, de la visibilité de tracé et des caractères, correction des fausses couleurs et d'effet de blocs dus à la compression JPEG.

4.1 Améliorations préliminaires apportées au concept de la diffusion

Amélioration de la méthode de PM : Afin de régler le problème de voisinage cité précédemment, nous avons augmenté la fiabilité de cette méthode à l'égard du problème, en intégrant les quatre voisins diagonaux comme supplément des quatre voisins horizontaux et verticaux (utilisés par l'approche classique). La prise en charge de toute ces orientations permettra une réduction significative des distorsions sur les caractères de texte qui accompagnent souvent l'approche classique de la diffusion. Nous avons donné également la possibilité de diffuser selon un voisinage de périmètre plus grand.

Amélioration de la méthode de Tshumperlé : La dérivée dans le domaine des images se rapporte à la différence d'intensité d'un pixel dans son voisinage, donc n'importe quelle perturbation peut influencer le calcul (sensibilité au bruit). Et afin d'éviter cela, il est préférable de lisser l'image en utilisant un filtre isotopique (filtre gaussien) avant tout calcul de dérivée. Et pour obtenir de meilleurs résultats et d'éviter l'application récurrente de la dérivation, il est préférable de calculer la dérivée seconde de la fonction gaussienne une fois pour toutes et ensuite la convoluer à chaque fois avec le voisinage traité.

4.2 Les étapes de notre approche adaptative

1. Calculer dans un rayon de voisinage (saisi par l'utilisateur) la distance colorimétrique D entre le pixel traité et ses voisins, en cumulant à chaque fois la valeur des pixels répandant à la condition suivante $\{cond1 : D < S_h\}$. S_h est un seuil pour estimer l'homogénéité locale (choisi empiriquement). D est une distance euclidienne calculée d'une façon vectorielle en cumulant la combinaison à poids égaux des sommes de différences des trois composants RVB des pixels voisins avec ceux du pixel central. Avec l'approche marginale, les différences de chacune des composants est cumulées séparément (mais celle-ci donne moins de performance).
2. Vérifier si le nombre des pixels satisfaisant *Cond.1* est majoritaire (supérieur à 50% de pixels de voisinage), le voisinage local est considéré dans ce cas comme homogène, sur laquelle il est inutile d'appliquer la diffusion. Le pixel central changera donc sa valeur par une moyenne adaptative calculée uniquement sur les pixels de couleurs similaires selon S_h (ceci renforce l'homogénéité locale de ces zones tout en évitant d'engendrer des fausses fluctuations par la diffusion et réduire d'une manière considérable les fausses couleurs et le temps de calcul).
3. Dans le cas où les variations sont importantes et la zone de voisinage est estimée comme non homogène, la diffusion (en version améliorée, telle qu'elle est décrite dans les sections 3 et 4.2) est appliquée. En revanche, la diffusion utilise les sorties de la seconde étape sur des pixels traités précédemment par une moyenne adaptative. Cet échange mutuelle entre la diffusion et le renforcement local d'homogénéité permettra de prétraiter rapidement et filtrer mieux. Il augmente également la cohérence des traits et préserve mieux la singularité des traits et des caractères tout en étant efficace sur les dégradations, notamment sur les fausses couleurs.
4. On répète les trois premières étapes jusqu'au parcours de tous les pixels de l'image.

4.3 Exploitation d'une binarisation locale pour évaluer notre approche

Notre première priorité est de réaliser un prétraitement adapté aux spécificités liées aux images de documents dégradées. La question qui se pose à ce niveau: quel est l'apport de prétraitement à l'amélioration de la qualité et la séparabilité des caractères durant la phase de la binarisation (une étape presque systématique mais déterminante aussi dans la majorité des systèmes OCRs) ?

Pour répondre à cette question, on applique sur les images que nous avons prétraitées par notre approches un seuillage automatique local par la méthode de Sauvola [17]. Cette méthodes utilise une estimation des variations locales dans un voisinage d'un pixel donné afin de déterminer son seuil qui tranche son appartenance au premier ou à l'arrière-plan. Une seule image peut avoir donc plusieurs seuils choisis d'une façon adaptative. La formule de Sauvola de calcul de seuil local représente une amélioration significative de la formule classique de Niblack par ajout d'une hypothèse sur les valeurs de niveaux de gris de texte et des pixels de fond (tous les pixels de texte ont des niveaux de gris proches de 0 et tout les pixels de l'arrière-plan ont des niveaux de gris près de 255), la formule de seuillage local devient donc :

$$S = m \times \left(1 - k \times \left(1 - \frac{e}{r} \right) \right) \quad (4)$$

k est un paramètre permettant d'atténuer les fortes fluctuations locales dus aux dégradations. Un effet d'érosion est associé au choix d'une grande valeur, alors qu'un effet de discontinuité des caractères fins est lié au choix d'une petite valeur. Donc, il est intéressant de faire un bon choix afin de fournir une bonne séparation en favorisant surtout le côté conservation de l'information pertinente et laisser les effets du bruit à l'étape de prétraitement et non pas à la valeur des paramètres (en pratique, k est fixé empiriquement à 0.5). r est la dynamique de l'écart type (fixé en pratique à 128). m et e sont respectivement la moyenne et l'écart type des couleurs dans la fenêtre d'analyse. Nous avons constaté durant notre élaboration de cette technique, que le choix d'une fenêtre d'analyse de petite taille réduit la qualité de séparation (ce qui ne correspond pas à nos objectifs), mais si nous l'augmentons, la séparation devient meilleure. Cependant, avec une taille trop grande nous constatons un effet de dilatation sur des caractères et des tâches dues à la présence de bruits sur l'image. On précise que les valeurs des paramètres de la binarisation choisis repose sur l'étude comparative des méthodes de binarisation publiée dans [8].

5 Évaluation et résultats

L'évaluation de notre approche de prétraitement et sa comparaison avec les deux méthodes de diffusions existantes (PM et Tshumperlé) est estimée à partir de la qualité des caractères issues de la binarisation d'une image de document prétraitée par chacune de ces approches. En effet, le taux d'OCR augmente sur des caractères de meilleure qualité. Pour cela, nous avons utilisé des images de la base d'images de documents de la compétition DIBCO destinée à l'origine pour l'évaluation de la binarisation[16]. Notre test a été porté sur 8 images de documents dégradés de cette base. Pour chaque image en couleurs, nous disposons d'une vérité terrain qui est une image binaire de référence (conçue par des experts en combinant les meilleures méthodes de la binarisation). Notre objectif dans cette deuxième étape est d'évaluer la performance de notre approche adaptative de prétraitement sur les images contenant du texte imprimé de différentes tailles et fortement dégradées. Ceci permettra de vérifier si l'utilisation d'une diffusion pilotée par la qualité de l'image est plus intéressants que l'utilisation classique des approches de la diffusions existantes.

La comparaison des méthodes citées ci-dessus est effectuée quantitativement en utilisant deux types de mesures : F-mesure et PSNR, données par les formules suivantes :

$$F_{\text{mesure}} = 2 \cdot \frac{R \times P}{R + P} \quad \text{ou} \quad R = \frac{TP}{TP + FN} \quad \text{et} \quad P = \frac{TP}{TP + FP} \quad (5)$$

R, P, TP, FP, FN désignent, respectivement, le rappel (sensibilité), la précision (valeur prédictive positive), le vrais positif, le faux positif et le faux négatif. Lorsque la valeur de la F-mesure est proche de 1, ceci signifie que le prétraitement a atteint sa qualité maximale et a conduit à une meilleure image binaire très similaire à la vérité

terrain (cette bonne qualité des caractères conduit à une reconnaissance et lecture précise des textes existant dans un document).

$$PSNR = 10 \log \left(\frac{Max_f^2}{\sqrt{MSE}} \right) \quad \text{avec} \quad MSE = \frac{1}{m \times n} \sum_{i=0}^{m-1} \sum_{j=0}^{n-1} \|I(x, y) - I'(x, y)\|^2 \quad (6)$$

$I(x, y)$ est le pixel (x, y) de l'image de référence, $I'(x, y)$ est le pixel (x, y) de l'image prétraitée puis binarisée, m, n sont les nombres de lignes et de colonnes de l'image. Max_f est la valeur maximale de l'intensité lumineuse existante dans l'image originale de référence (d'une qualité de référence).

5.1 Résultats d'évaluation

Nous avons comparé la qualité de prétraitement sur un ensemble E de huit images de la base DIBCO prétraitées et/ou binarisées selon les démarches suivantes:

- Binarisation des images de E sans aucun prétraitement préalable.
- Binarisation des images de E après leur prétraitement par la méthode de Perona et Malik que nous avons améliorée (*voisinage* 3×3 , $\lambda = 9$, 20 itérations).
- Binarisation des images de E après leur prétraitement par la technique de Tshumperlé que nous avons modifiée (*voisinage* 3×3 , $\lambda = 9$, 20 itérations).
- Binarisation des images de E après leur prétraitement par notre approche de smart-diffusion (*voisinage* 27×27 , $\lambda = 9$, $S_h = 56$).

Les tableaux suivants montrent les résultats obtenus des F-mesure et PSNR.

Tab. 1. F-mesure des trois approches de prétraitement.

F-Mesure	Image1	Image2	Image3	Image4	Image5	Image6	Image7	Image8
Sans Prétraitement	82.30	73.20	89.12	82.72	82.59	43.00	67.32	85.00
PM	86.97	74.95	91.90	86.13	83.19	70.06	87.21	80.00
Tshumperlé	84.60	73.78	89.99	84.32	83.49	48.07	87.44	84.58
Notre approche	90.05	77.49	91.28	93.23	83.77	76.02	87.66	84.90

Tab. 2. PSNR de trois approches de prétraitement.

PSNR	Image1	Image2	Image3	Image4	Image5	Image6	Image7	Image8
Sans prétraitement	12.12	11.07	13.97	13.78	12.72	9.15	16.96	14.3
PM	13.49	11.63	15.4	14.97	13.08	14.11	22.2	13.32
Tshumperlé	2.64	11.24	14.38	14.28	13	9.98	22.22	14.23
Notre approche	14.82	12.22	15.05	18.36	13.23	15.48	22.47	14.29

5.2 Interprétation des résultats de cette évaluation :

Cette évaluation quantitative qualifie notre approche de prétraitement comme meilleure par rapport aux approches de diffusion de Perona et Tshumperlé. Notre approche a montré une fiabilité intéressante vis-à-vis les fortes dégradations existantes proches de texte imprimé où les autres approches ont introduits des déformations irréversibles. Les valeurs de F-mesure et de PSNR de notre approche au-dessus des autres approches sont des bons indices de l'efficacité d'un traitement sélectif par rapport à un traitement itératif.

5.3 Réduction significative des temps de calcul

Pour l'évaluation des temps de calcul, nous avons effectué nos tests sur une machine HP dotée d'un processeur Intel Dual Core t4400 de 4 Go de RAM, l'unité de temps est la seconde. Le tableau suivant illustre une comparaison entre les techniques de PM et de Tshumperlé avec notre approche. Nous apercevons clairement que notre approche donne également des meilleurs résultats de temps de calcul grâce à son mode sélectif qu'elle utilise.

Tab. 3. Temps de calcul (Perona et Malik, Tshumperlé et notre méthode de prétraitement adaptatif).

Time	Image1	Image2	Image3	Image4	Image5	Image6
Perona et Malik	28.03	54.60	21.10	27.56	9.60	9.10
Tshumperlé	35.47	34.41	38.79	34.33	26.83	19.50
Notre approche	18.14	16.80	5.5	25.74	8.12	8.86

Le comparatif des temps de prétraitement (Tab.3) qualifie que notre approche est plus rapide que les deux autres méthodes grace aux mécanisme de diffusion selective pilotée par la qualité de l'image. Nous avons remarqué que sur une image d'une meilleure qualité, notre approche fait baisser d'une manière significative les temps par rapport aux autres approches, car elle diffuse moins sur une image ou une zone d'image de meilleure qualité par rapport à celle qui est trop dégradée. La diffusion sur l'intégralité de l'image nécessite plusieurs calculs itératifs (calcul des dérivées, gradients, vecteurs propres, etc.), contrairement à notre approche qui utilise, d'un côté, la diffusion sur certaines zones seulement de l'image (jugées comme dégradées ou à variations importantes), et d'un autre coté, elle calcule une simple moyenne adaptative sur d'autres régions qui sont déjà homogènes. Ce qui rend notre contribution efficace en vitesse et qualité en même temps grâce à une simple prise en compte de la qualité de l'image.

5.4 Exemples des résultats

Les figures suivantes montrent quelques résultats obtenus.



Fig.1. (a) Image originale (Base DIBCO), (b) binarisation sans prétraitement, (c) Binarisation après prétraitement par Perona et Malik (20 itérations), (d) Binarisation après prétraitement par notre approche.

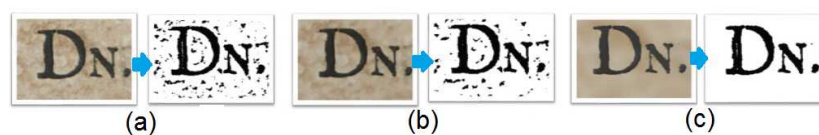


Fig. 2. (a) Binarisation sans prétraitement, (b) binarisation après prétraitement par Perona et Malik (20 itérations), (c) binarisation après prétraitement par notre approche.

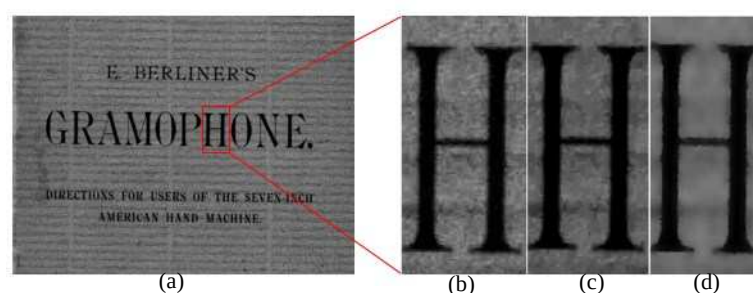


Fig. 3. (a) Image originale (base DIBCO) (b) extrait de l'image originale, (c) prétraitement avec PM originale 30 itérations, (d) prétraitement par notre technique 1 itération.

Nous avons testé également notre technique sur d'autres images de documents contemporains, externes de la base, voici un exemple de résultat :

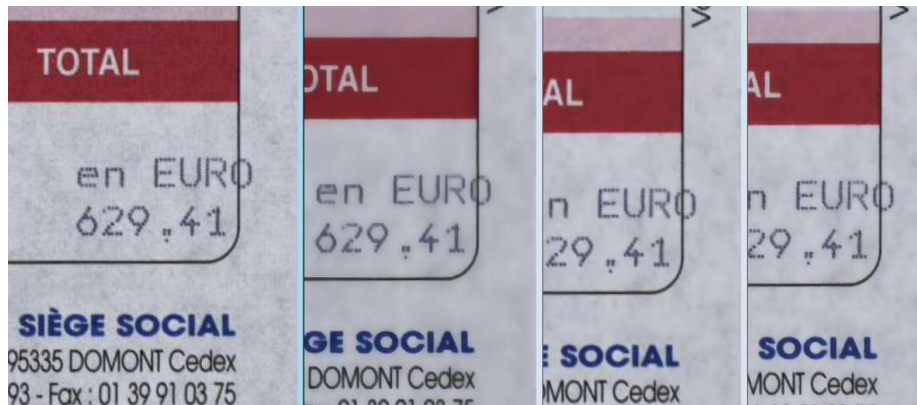


Fig. 4. Prétraitement d'une image (extraite d'une facture scannée), du gauche vers la droite, Image originale en JPEG, prétraitement PM ($t=12.35$ sec), Tshumperlé ($t=16.67$ sec), notre approche ($t=1.5$ sec).

6 Conclusion

Dans cet article, nous avons mis en évidence le problème de bruit dans le domaine de traitement d'images, qui est un phénomène très dégradant notamment sur des images qui portent de l'information textuelle comme les documents numériques. L'effet de bruit et les différentes perturbations sur les images rendent l'accès difficile au contenu porté par ces dernières. Nous avons présenté une nouvelle approche adaptative qui agit d'une manière intelligente en distinguant les zones homogènes des zones dégradées sur lesquelles la diffusion est nécessaire sans pour autant déformer les zones déjà en bonne qualité. Nous avons testé notre technique sur plusieurs images de documents dégradées et même sur des images naturelles. La comparaison des résultats de ce concept par rapport à ceux des deux approches existantes citées ci-dessus, confirme notre hypothèse sur l'importance de la prise en compte de la qualité de l'image dans toute tâche impliquée dans un système de reconnaissance, comme celle de prétraitement. Ce pilotage de prétraitement par la qualité a conduit également à une réduction significative des temps de calcul. Dans nos futures travaux, nous souhaitons évaluer cette technique en la testant sur d'autres espaces colorimétriques et nous réduisons encore les temps de calcul en introduisant une architecture parallèle.

Références

1. T. Huang, G. Yang et G. Tang : A fast two-dimensional median filtering algorithm, IEEE Transactions on Acoustics, Speech, and Signal Processing, vol. 27, n° 11, pp. 13–18 (1979).
2. J.-S. Lee : Digital image smoothing and the sigma filter, Computer Vision, Graphics, and Image Processing, vol. 24, n° 12, pp. 255–269 (1983).

3. W. Fourati et M. S. Bouhlef : Techniques de Débruitage d'Images, chez 5th Int. Conf. Sciences of Electronic, Technologies of Information and Telecommunications, (2009).
4. D. L. Donoho, I. M. Johnstone, G. Kerkycharian et D. Picard, Density estimation by wavelet thresholding, *The Annals of Statistics*, pp. 508–539 (1996).
5. D. L. Donoho, I. M. Johnstone, G. Kerkycharian et D. Picard: Universal Near Minimality of Wavelet Shrinkage, chez *Festschrift for Lucien Le Cam: Research Papers in Probability and Statistics*, D. Pollard, E. Torgersen et G. L. Yang, Éd., New, York: Springer New York, pp. 183–218 (1997).
6. A. G. Bruce et H.-Y. Gao: Waveshrink: shrinkage functions and thresholds, chez *SPIE's 1995 International Symposium on Optical Science, Engineering, and Instrumentation*, (1995).
7. N. Nikolaou et N. Papamarkos: Color Reduction for Complex Document Images,» *Int. J. Imaging Syst. Technol.*, vol. 19, n° 11, pp. 14–26 (2009).
8. D. Gaceb, F. Lebourgeois et J. Duong Adaptive Smart-Binarization Method: For Images of Business Documents, *12th International Conference on Document Analysis and Recognition*, (2013).
9. G. Aubert et P. Kornprobst: Mathematical problems in image processing: partial differential equations and the calculus of variations, vol. 147, Springer Science & Business Media, (2006).
10. P. Perona et J. Malik: Scale-Space and Edge Detection Using Anisotropic Diffusion, *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 12, n° 17, pp. 629–639, (1990).
11. D. Tschumperle et R. Deriche : Vector-valued image regularization with PDEs: A common framework for different applications,» *IEEE Transactions on Pattern Analysis and Machine Intelligence*, vol. 27, n° 14, pp. 506–517 (2005).
12. L. Piffet: Décomposition d'image par modèles variationnels: débruitage et extraction de texture (2010).
13. L. I. Rudin, S. Osher et E. Fatemi: Nonlinear total variation based noise removal algorithms, *Physica D: Nonlinear Phenomena*, vol. 60, n° 11, pp. 259–268 (1992).
14. M. Ben Abdallah, J. Malek, A. T. Azar, H. Belmabrouk, J. Esclarín Monreal, and K. Krissian : Adaptive noise-reducing anisotropic diffusion filter, *Neural Computing and Applications*, vol. 27, no. 5, pp. 1273–1300, (2016).
15. Y Yu: Acton ST Speckle reducing anisotropic diffusion. *IEEE Trans Image Process* 11(11):1260–1270 (2002).
16. I. Pratikakis, B. Gatos, et K. Ntirogiannis: ICDAR 2011 Document Image Binarization Contest (DIBCO 2011), *2011 International Conference on Document Analysis and Recognition*, (2011).
17. J. Sauvola and M. Pietikainen: Adaptive document image binarization. *Pattern Recognition*, 33:pp. 225–236(2000).

Detection Method without crowd behavior modeling by fuzzy logic

Hocine Chebi ¹, Dalila Acheli ², and Mohamed Kesraoui ³

^{1 and 3} Faculty of Hydrocarbons and Chemistry, University M'hamed Bougara Boumerdès

² Faculty of Engineering, University M'hamed Bougara Boumerdès

Laboratoire d'Automatique appliqué

Avenue de l'Indépendance 35000 Boumerdès, Algeria

¹chebi.hocine@yahoo.fr, ²dacheli2000@yahoo.fr, ³mkesraoui@umbb.dz

Abstract. This work falls within the framework of the video surveillance research axis. This work falls within the scope of video surveillance. It involves a link between automatic processing and problems related to video surveillance. The job is to analyze video streams coming from a network of surveillance cameras, deployed in an area of interest in order to detect abnormal behavior. Our approach in this paper is based on the new application and the use of fuzzy logic with an aim of avoided the modeling of abnormal crowd behavior. By introducing the notion of degree in the verification of a condition, thus allowing a condition to be in another state rather than true or false, fuzzy logic confers a very appreciable flexibility on the reasoning that uses it, taking into account inaccuracies and uncertainties. One of the work interests performed in formalized human reasoning is that the rules are spelled out in a natural language. Detecting these behaviors will increase the response speed of security services to perform accurate analysis and detection of abnormal events. We then present a comparative table followed by a study on the effects of using another technique. Despite the complexity of the sequences, our approach provides highly relevant results including in the detection areas of several crowd behaviors.

Keywords: Abnormal detection, crowd behavior, classification, segmentation, tracking, fuzzy logic, events and.

1 Introduction

The computer vision field (also called artificial vision, digital vision or cognitive vision) aims at reproducing on a computer the capacities of analysis and interpretation proper to human vision.

The general objective is to design models and systems capable of representing and interpreting the visual content of a scene. Computer vision is an exciting topic in artificial intelligence research. Although considerable progress has been made in the resolution of certain problems such as automatic video surveillance, the latter treats the framework of analysis of the behavior of people present in the scene making it

possible to predict incidents and to detect certain events. It also provides information that explains the activity of these individuals and their interests.

Therefore, the analysis of human behavior using video is rich field in various experiments [1-4]; we use the term computer vision behavior to denote the physical and apparent behavior produced as a result of movement, without psychological or cognitive interpretation. The analysis of human behavior is very useful in a large number of applications (see fig.1).



Fig. 1. Illustration of an unusual situation in a crowd scene.

We are interested in our work with crowd behavior, the automatic study of crowd behavior makes it possible to develop crowd management strategies, especially for popular events or events frequented by a large number of people (example: sports meetings, Large concerts, public events, etc). The study of the crowds makes it possible to anticipate certain abnormal behaviors in order to avoid accidents caused by disorganized movement of the crowds. Although anomalous crowd detection has received some attention in recent years, most research however, has focused on identifying and tracking abnormal behaviors (such as suspicious events, irregular behavior, uncommon behavior, unusual activity, abnormal behavior, anomaly, etc.) of a single person or a few moving objects in a crowd [5-9].

In this paper, we are led to analyze different behavior implications. The goal of the used approach is the detection of anomalies in very dense scenes based on the speed and orientation of the individuals and that of the group. The various anomalies are detected using fuzzy logic as a method of automatic classification of crowd behavior without mathematical modeling, for the management of anomalies in a group of people.

The main contribution in this article is the proposed intermediate level visual characteristic, and thus the use of fuzzy logic in the case of abnormal behavior in a crowd scene. By introducing the notion of degree in the verification of a condition, thus allowing a condition to be in another state of crowd behavior other than true or false, fuzzy logic confers a very appreciable flexibility to the reasoning that uses it. This makes it possible to take into account inaccuracies and uncertainties in the case of detection of crowd behavior. Speed is considered a characteristic of images and fuzzy rules explore the modeling of normal and abnormal behaviors. The remainder of the paper is organized as follows. In section 2, this paper briefly reviews the state of the

art on crowd motion features and anomaly detection. Section 3 represents the proposed global approach to crowd behavior detection. Section 4 is about a method of detection based on fuzzy logic. In section 5, experiment results on real world video scenes are presented. Finally, the conclusion and potential extensions of this work are described in section 6.

2 State of art

The Approaches used for the analysis of crowd behavior in video sequences generally comprises of four essential stages: detection of movement, segmentation, classification, and tracking [1, 9-14].

Whereas the temporal derivative quantifies the variation of the aspect of each pixel considered individually in the movement detection case, the optical flow is modeled by a vector field in two dimensions representing the projection on the image plane of the real movement observed (in 3D). Accordingly many methods [15-17] were proposed since the precursory article of Horn and [18] with an aim of improving the implementation of the latter. In [19], nine algorithms are studied and compared according to criteria of precision and density of the obtained vector field, but no mention is made of algorithmic complexity. The work of [20] makes it possible to fill this gap by measuring the reports/ratios (precision) / (calculations team) of these methods.

The segmentation is the stage which consists of cutting out the image in a successive way by considering the apparent movement dominating the zones not labeled then by detecting the zones not conforming to this model of movement [21]. The principle of this method rests on the taking into account of a total image in which the dominant movement and the zones in conformity with this movement forming a first area of the partition are estimated. This process of estimate and detection is reiterated on the non-conforming zones until all pixels are classified [22, 23].

The method of training for a phase of classification generates a function which corresponds to an image received in a specific labeled entry. There are several methods of training in literature, such as decision trees [24], neural networks [24], fuzzy logic, AdaBoost (Adaptive Boosting) [26], or machines with vectors of support (MVS) [27].

The methods of tracking propose to recognize and locate in the course of time the objects present in a temporal sequence of images [28, 29]. Within the framework of a human crowd, they find a particular interest in video monitoring where the follow-up of the individuals makes it possible to automatically control the comings and goings in a space. Just like recognition starting from an image, the continuation can be based on graphic properties such as colors or contours [30, 31]. Moreover, adding a temporal dimension makes it possible to suppose a continuity of the presence and position of the objects in the scene, in spite of screenings. The temporal and space consistency

of the followed characteristics can in certain cases be obtained using methods of partitioning (clustering) [32].

In our work, we propose the use of the detection movement techniques by optical flow. The latter is one of the techniques that make it possible to detect groups that move in the same direction and to extract the reasons for movement. The major advantage of this method is that it doesn't need to be modeled [33], because it consists of detecting the movement by mathematical calculation in any point of the image, which is a function of the intensity or the color of all the pixels and which is supposed to reflect the importance of the visible movement in the scene. We propose thereafter the segmentation by regrouping of the areas with an aim of providing a more precise cutting of the borders of the areas. Thereafter, we propose to use of fuzzy logic, which introduces the notion of degree in the verification of a condition, thus allowing a condition to be in a different state of crowd behavior other than true or false, by a set of fuzzy rules written in advance.

3 Easy detection with fuzzy logic

The principal stages of the suggested approach are shown in the flow chart of fig.3. In the first stage we capture the image to be processed using a camera. Thereafter, we carried out detection by optical flow; segmentation of movements, and classification. This last stage represents the new approach for the detection of abnormalities in very dense scenes while being based on the speed and orientation of the individuals and that of the group. The various anomalies are detected by a fuzzy logic approach. The next stage is the tracking of the abnormalities. Finally a test is carried out in order to extract some comprehensible information (crowd behavior detection of normal and abnormal events).

The designer of a fuzzy system must make a number of important choices. These choices are essentially based on the advice of the expert or on the statistical analysis of past data, in particular to define the membership functions and the decision matrix. Here is a synoptic overview of a fuzzy system (fig.2):

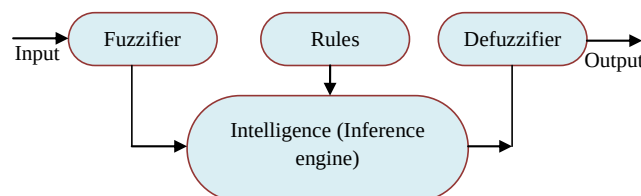


Fig. 2. Overview of a fuzzy system.

Our approach in this article relies on the new application which uses the fuzzy logic in the case of division and fusion of the crowd. The detection of these behaviors will

increase the speed of response of the security services in order to perform accurate analysis and detection of events in real time.

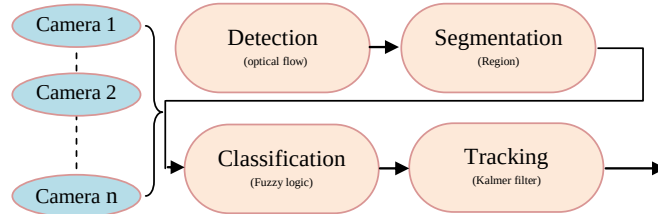


Fig. 3. General diagram represent the method of crowd behavior detection.

3.1 Motion vectors extraction

The investigated crowd activities are characterized by movement of people. The examination of the motion dynamics of the crowd is based on the so-called motion vectors obtained by the methods described below. Optical flow is applied to each pair of subsequent video frames. Two stages of the processing can be distinguished, namely feature finding and feature tracking. The first stage is based on the corner detection algorithm proposed by Shi and Tomasi [34]. For each derivative pixel (defined as the second partial derivative of the intensity of the original image) with surrounding K , the autocorrelation matrix $M(x, y)$ is calculated. Good corners occur in the points, where the smaller of two Eigen values of the matrix $M(x, y)$ is greater than a predetermined threshold value. The set of feature points (corners) for the video frame is obtained in this step.

Tracking of the dislocation of the detected corners in the subsequent frame is done using the Lukas-Kanade algorithm [35], [36]. As a result, the second set of points is produced. Connecting corresponding points in subsequent frames provides the vector. It describes the motion of the detected feature point. If the position of the feature point in consecutive frames differs then the motion of that point is considered. The obtained vectors are then assigned to classes according to magnitude and direction. In this work classification is done according to the magnitude criteria only.

The principle of computation is based on formula 1 [37], where $I(x, y, t)$ is the intensity of the image, such that (x, y) are the position points, and t is time, we have:

$$I(x, y, t) = I(x + \delta x, y + \delta y, t + \delta t) \quad (1)$$

Where δx is the spatial displacement of the point (x, y, t) after a time δt . The Taylor series gives:

$$I(x + \delta x, y + \delta y, t + \delta t) = I(x, y, t) + \frac{\partial I}{\partial x} \delta x + \frac{\partial I}{\partial y} \delta y + \frac{\partial I}{\partial t} \delta t + H.O.T \quad (2)$$

Where $(\frac{\partial I}{\partial x}, \frac{\partial I}{\partial y}, \frac{\partial I}{\partial t}) = (I_x; I_y)$ and I_t are the first-order partial derivatives of $J(x; y; t)$. (H.O.T) Denotes the terms 2 order and more, which are negligible and which are ignored in the sequel. We subtract $I(x; y; t)$ and divide by δt :

$$\begin{cases} \frac{\partial I}{\partial x} \delta x + \frac{\partial I}{\partial y} \delta y + \frac{\partial I}{\partial t} \delta t = 0 \text{ ou} \\ \frac{\partial I}{\partial x} \frac{\partial x}{\partial t} + \frac{\partial I}{\partial y} \frac{\partial y}{\partial t} + \frac{\partial I}{\partial t} \frac{\partial t}{\partial t} = 0 \\ \frac{\partial I}{\partial x} vx + \frac{\partial I}{\partial y} vy + \frac{\partial I}{\partial t} = 0 \end{cases} \quad (3)$$

As the terms are optical flow components:

$$vx = \frac{\partial x}{\partial t} \text{ and } vy = \frac{\partial y}{\partial t} \quad (4)$$

Where $(\frac{\partial I}{\partial x}, \frac{\partial I}{\partial y}, \frac{\partial I}{\partial t}) = (I_x; I_y)$ is the spatial gradient and $v = (vx, vy)$ is the velocity of the image. Equation 3 is known as the optical flow stress equation and defines a unique local constraint on the motion of the image. However, this constraint is not sufficient to compute the components of v because this constraint is a poorly formed problem. However, several methods exist for estimating the optical flow vectors based on these assumptions. So the 2D movement will have:

$$\nabla I \cdot \vec{v} = -I_t \quad (5)$$

Subsequently, the optical flow function Lucas & Kanade [37] will have:

$$f_x u + f_y v = -f_t \quad (6)$$

In matrix form with window equal 3 times 3 (3x3):

$$\begin{cases} f_{x1}u + f_{y1}v = -f_{t1} \\ f_{x2}u + f_{y2}v = -f_{t2} \\ f_{x3}u + f_{y3}v = -f_{t3} \\ \dots \\ f_{x9}u + f_{y9}v = -f_{t9} \end{cases} \quad (7)$$

Summarised of equation 7 as :

$$\begin{bmatrix} f_{x1} & f_{y1} \\ \vdots & \vdots \\ f_{x9} & f_{y9} \end{bmatrix} \begin{bmatrix} u \\ v \end{bmatrix} = \begin{bmatrix} -f_{t1} \\ \vdots \\ -f_{t9} \end{bmatrix} \quad (8)$$

Thereafter, we can write the terms of f_{xi} and f_{yi} in matrix A , which represents all the image points in Equation 9:

$$\text{With } A = \begin{bmatrix} f_{x1} & f_{y1} \\ \vdots & \vdots \\ f_{x9} & f_{y9} \end{bmatrix}$$

$$\begin{cases} A \cdot U = f_t \\ A^T A \cdot U = A^T f_t \\ U = (A^T A)^{-1} A^T f_t \end{cases} \quad (9)$$

With $(A^T A)^{-1}$ as the pseudo-inverse.

As described by [37], the algorithm finds for each point on the first image its corresponding point on the second image which minimizes the following equation 10:

$$\min \sum (f_{xi}u + f_{yi}v + f_{t1})^2 \quad (10)$$

Optical flow analysis returns a set of motion vectors in the form:

$$V_{i,t} = (x_{i,t}, y_{i,t}, m_{i,t}, \theta_{i,t}) \quad (11)$$

Where " $V_{i,t}$ " is the motion vector " i " at frame " t ", represented by the feature point at the coordinate $(x_{i,t}, y_{i,t})$, the magnitude " $m_{i,t}$ " and the orientation angle " $\theta_{i,t}$ ".

3.2 Application of fuzzy logic

Classical logic is a part of mathematics relatively well known to the public. It is on its principle that computers and most digital machines work. In classical logic, decisions are binary: either true or false. It is on this point that fuzzy logic will be distinguished from classical logic. In fuzzy logic, a decision can be both true and false at the same time, with a certain degree of belonging to each of these two beliefs.

In our work we use fuzzy logic as a method of automatic classification of crowd behavior.

For example, consider these two rules of inferences:

- ✓ If the speed magnitude is low, then behavior is normal;
- ✓ If the speed magnitude is high then behavior is abnormal.

In classical logic, an object can only be near or far. If the distance of behavior is low, then this one will be normal. In fuzzy logic, however, the object will be both normal and abnormal at the same time. Here, the object that is low will, for example, be 60% normal and 40% abnormal.

One realizes that in fuzzy logic, a fact no longer has a strict belonging to a belief, but a "fuzzy" belonging.

The objective of using fuzzy logic makes it possible to reason not on numerical variables but on linguistic variables, that is, on qualitative variables. The reasoning on these linguistic variables will allow the ability to manipulate knowledge in natural language. All that one has to return to the detection system are rules of inferences expressed in a natural language. For this work we will use the representation of the following linguistic values fig.4:

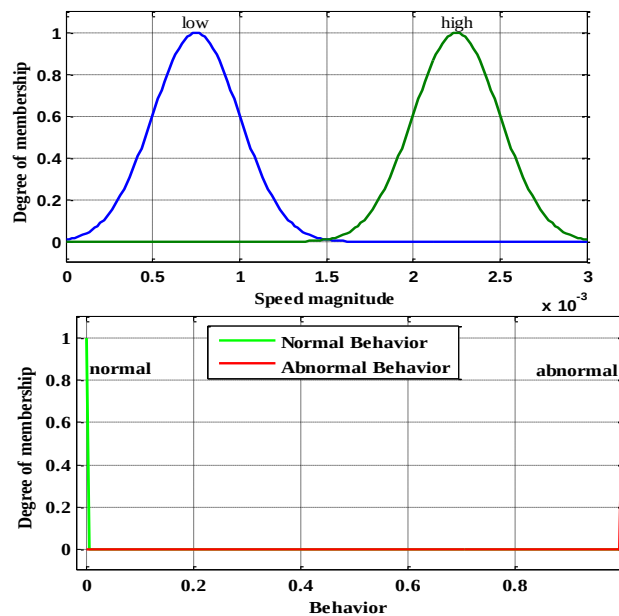


Fig.4. Image of linguistic values or fuzzy logic (a) input variable, (b) output variable.

3.3 Tracking behavior

Contrary to the algorithm of KALMAN, particles filters based on the algorithm condensation (Conditional Density Propagation), are used for non linear systems with non Gaussian models of observation. So, they are adapted well to follow a trajectory with sudden changes of direction, and for the follow-up of multiple people.

The state vector consists of address coordinates x and y of the center of gravity of anybody in the image, and the speeds according to X which we shall note v_x and according to Y that we shall note v_y in fig.4. The measures are address and coordinates on the image x and it of every point of the trajectory.

We shall thus have the following state representation:

$$\begin{cases} X(k+1) = AX(k) + BW(k) \\ Y(k) = CX(k) + V(k) \end{cases} \quad (12)$$

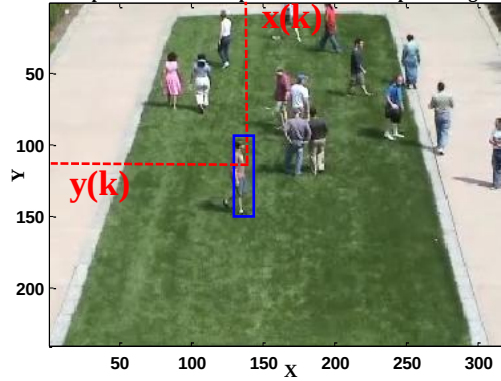


Fig.5. Illustration of tracking a person in an image.

Such as:

The state vector at instant k is $X(k) = \begin{bmatrix} x(k) \\ y(k) \\ v_x(k) \\ v_y(k) \end{bmatrix}$, and the vector of observation measures at instant k is $Y(k) = \begin{bmatrix} x(k) \\ y(k) \end{bmatrix}$.

$$A = \begin{bmatrix} 1 & 0 & T & 0 \\ 0 & 1 & 0 & T \\ 0 & 0 & 1 & 0 \\ 0 & 0 & 0 & 1 \end{bmatrix}, B = \begin{bmatrix} \frac{T^2}{2} & 0 \\ 0 & \frac{T^2}{2} \\ T & 0 \\ 0 & T \end{bmatrix}, C = \begin{bmatrix} 1 & 0 & 0 & 0 \\ 0 & 1 & 0 & 0 \end{bmatrix}$$

T : Being the period of sampling.

During this simulation, we used a noise to measure the matrix of covariance $R = \begin{bmatrix} 16 & 0 \\ 0 & 16 \end{bmatrix}$, with $x(0) = [30 \ 15 \ 15]$.

3.4 Events crowds Datasets

This section mainly proposes a new approach and a method of detecting abnormal continuous events of crowds in the case of dangerous situations, and classification by fuzzy logic with the aim of letting us show between the methods and to avoid fixing of the threshold value in types of crowd videos sequence. The videos are mainly collected from the UMN dataset [38].

Table I: DIFFERENT CLASS EVENTS OF CUHK CROWD DATASET

°	Crowd events
1	Highly mixed pedestrian walking
2	Crowd following a mainstream and well organized
3	Crowd following a mainstream but poorly organized
4	Crowd merge
5	Crowd split
6	Crowd crossing in opposite directions
7	Intervened escalator traffic
8	Smooth escalator traffic

The BEHAVE video dataset and the PETS2009 dataset [39] for performance evaluation are adopted in anomalous frame behavior detection experiments. And dataset [40] (dataset: ground truthed). The different classes events of CHHK crowd dataset are represented by table I.

4 Experimental Results

In this section, we propose the evaluation of the simulation results according to synthetic data by videos. The videos are mainly gathered from the whole of UMN data and the whole of data PETS2009, which were largely widespread for the performance evaluation, the results will be detailed in the nearest sections.

4.1 Evaluation of behavior analysis

The proposed approach is based on computing the magnitude of the motion vector. We calculate the dissimilarity of images sequence with the aim of determining the running and walking events (fig.6.a-b and fig.10). These events can be identified by using the distance measure of the vectors of optical flow.

4.2 Evaluation of tracking filter on the UMN Dataset

Having produced KALMAN filter algorithm, we obtain the following results fig.7.a. We notice after the results of tracking using the filter of KALMAN, that the hypothesis of Gaussian models (KALMAN filter) are still not being verified in practice. This favors the use of the stochastic techniques based on the simulation of random variables by approximations of Gone Up Carlo, known as particles filtering. This technique gives us the fig.7.b, where the yellow case represents the tracking of particles before the follow-up of normal behavior.

4.3 The Evaluation of Easy detection of abnormal events with fuzzy logic

The UMN dataset contains eleven videos taken from three different scenes that show behaviors of escape. The results detected from these scenes are shown in fig. 9.

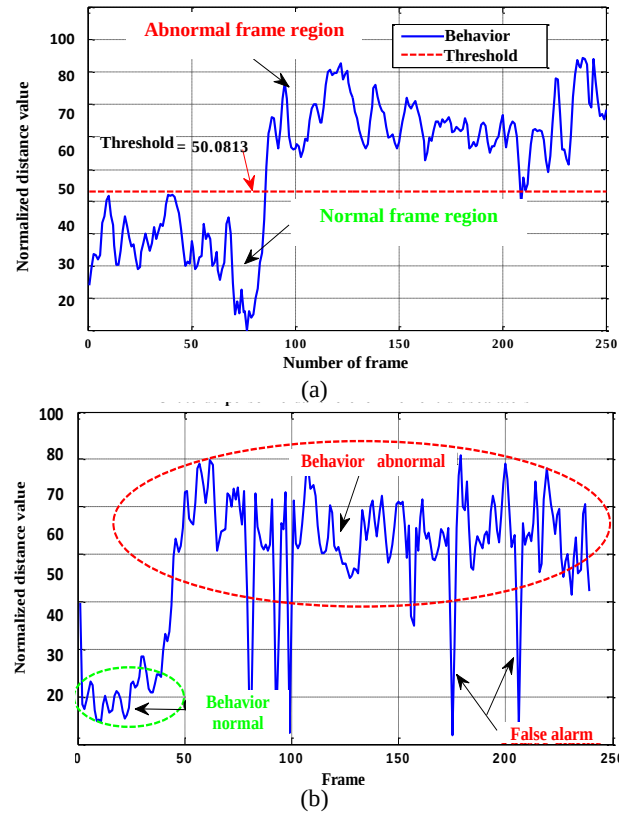


Fig.6.a. Examples of distance measure of a group of people, (a) dangerous situations, and (b) situation in an escalator.

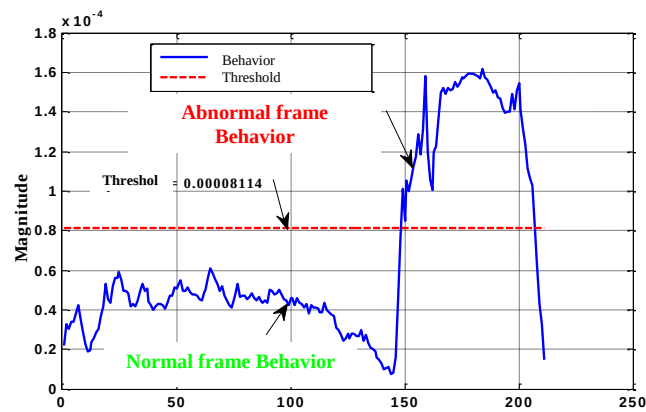


Fig.6.b. Examples of distance measure of a group of people with magnitude speed.

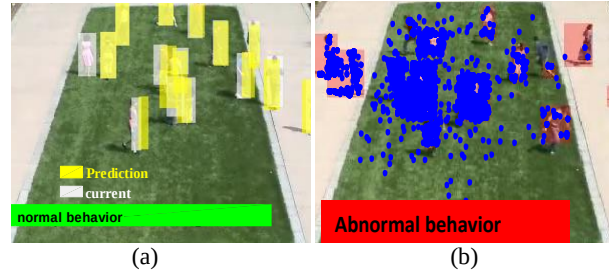


Fig.7. Illustration of the tracking results with: (a) is the KALMAN filter, and (b) is the particles filter in the case of abnormal behavior.

The normal scene is defined as a crowd of people moving with a mean velocity. And abnormal crowd behavior is shown in the second part of the video. These scenes represent the behaviors of running and walking in each video sequence.

We obtain a positive convergence in the case of abnormal behavior (see fig.10), we notice that particles are stacked on the people exhibiting abnormal behavior with the aim of the visual tracking.

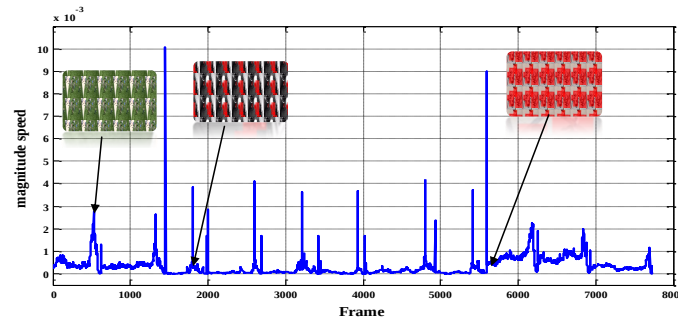


Fig.8. Evaluation the magnitude speed of the scenario UMN.

We can easily detect visual reinforcements like abnormal ones using the automatic architecture of classification. In particular, the advantage of the method suggested is obvious in outdoor and indoor scenes, the detection of normal and abnormal behavior in the various scenes is made in an automatic way one using fuzzy logic. The sufficient results of the comparison between the social force method [42, 42] and the method of the acceleration feature [43], the method suggested carries out a stable detection of behavior. Results more representative of comparison for scenarios is shown in fig. 10. We can see that the execution of the method proposed is more precise and stable in almost all the scenarios.

Fuzzy logic thus makes it possible to control complex systems that can not necessarily be modulated in an "intuitive" way. However, this method has various disadvantages. First, expressing one's knowledge in the form of natural language (and therefore qualitative) rules does not prove that the system will behave optimally. This

method cannot guarantee that the behavior is precise or optimal, or even guarantee that the rules entered by the programmer are not contradictory.

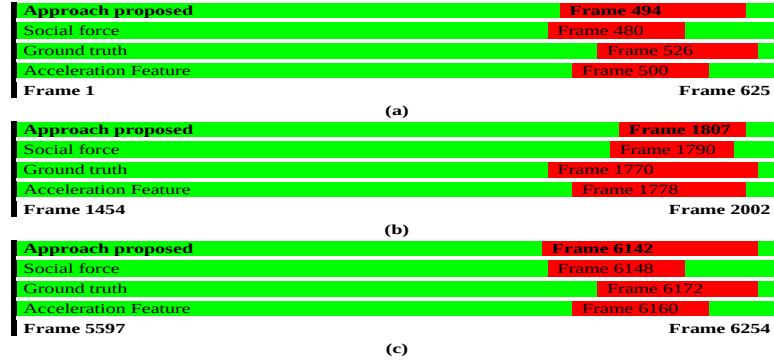


Fig.9. Crowd escapes behavior detection for several scenarios in the UMN dataset. Our performance is compared with the social force method, Ground truth and acceleration feature. (a) Result of the scenario UMN 1. (b) Result of the scenario UMN 3. (c) Result of the scenario UMN 9.

For classification, we use fuzzy logic because each image has several classes for each category of events (race or walk, division or fusion, fall in escalator, dangerous situation,...). Although this makes the system more complicated, we have nevertheless the possibility of distinguishing the events related to the speed of crowd behavior. We divided the base of videos into several sets according to dated sets used. Some simulation results and the performances of the automatic control technique thus are presented in table II. Our approach leads to better results for the three datasets used from the base of videos when the person is mobile.

TABLE II: PERFORMANCES OF USED METHODS FOR THE WHOLE CROWD BEHAVIORS DETECTION

Dataset	Crowd Behavior	Scenarios Video/images	Classification technique	Number of rules	Precision	Probability rate %
BMVA dataset	Passage	C,RT,I,IG,WT,FI,S	Fuzzy logic	02	0.990	99.07
	Brawls	A,M,RT,IG,WT,FI,S		02	0.940	94.07
UMN dataset	Full in an escalator	UMN	//	02	0.907	90.72
	Race or walk	Escalator		02	0.916	91.63
		UMN		02	0.868	86.88
PETS2009 dataset	Fusion and division	Fusion Division	//	03	0.893	89.31

4.4 Evaluation of dangerous situations with dataset BMVA

The results of the approach are represented in fig.10, 11 and 12. In our work we suppose that the number of people in an occulted group is not limited. The obtained results are encouraging when the automatic detection of anomalies is close to the real time measurement. The approach suggested shows a great robustness against false

alarm detection since the automatic detection of anomaly occurs after the real release of the anomaly.



Fig. 10. Example 1 of event detected by fuzzy logic in dangerous situations (a- normal behavior, and b- abnormal behavior).

Our approach presents some advantages as it presents a positive contribution to the detection of movement in a complex environment. However, it requires the estimation of temporal time for each image bloc and at every moment of the video sequence which makes it very greedy in computing power consumption. To show the effectiveness of the method we simulated fig.13 which shows satisfactory results in other behaviors.



Fig.11. Example 2 of event detected by fuzzy logic in situation in an escalator (a- normal behavior, and b- abnormal behavior).

4.5 Evaluation behavior fusion and division on the PETS2009 Dataset and UMN Dataset

The proposed approach is based on computing the magnitude of the motion vector which presents the optical flow in the Cartesian frame. The point $P(x,y)$ is the position of his interest point at time t where $Q(x,y)$ is the position of the same point at $t + 1$. We use the Euclidean distance to compute the distance travelled by the interest point:



Fig. 12. Detection of behavior by fuzzy logic in dataset BMVA, (a) normal behavior, (b) abnormal behavior.



Fig. 13. Detection of behavior by fuzzy logic, (a, c and d) normal behavior, (b) abnormal behavior.

$$M_i = \sqrt{(Q_{xi} - P_{xi})^2 + (Q_{yi} - P_{yi})^2} \quad (12)$$

This method is very useful when confronted with systems that are impossible, or difficult to model. Similarly, this method is very advantageous if one possesses a good level of human expertise. Indeed, it is necessary to provide to the fuzzy system a whole set of rules expressed in natural language to allow reasoning and draw conclusions. The greater the human expertise of a system, the greater the ability to add inference rules to the system of detection.

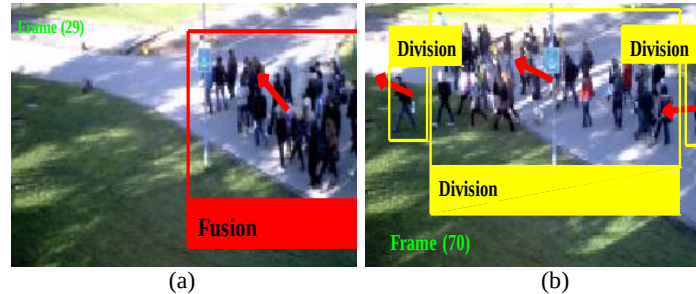


Fig. 14. Representation of the events, (a) Fusion (b) Division.

To detect the events of fusion and division we calculate the circular variance relating to the orientations of displacement of the groups in each image. Division occurs in a larger zone with a considerable divergence of the groups. But in the case of fusion occurs in a zone restricted for a short length of time. The results evolution is represented in Fig.14.

5 Conclusion

In conclusion, it can be said that fuzzy logic has the advantage of being intuitive and able to operate a large number of different systems with strong human expertise. Nevertheless, one should bear in mind that in fuzzy logic it is impossible to predict detection performance. If the settings are fine, the performance will be fine, but if there is a lack of precision in the settings, the performance will be imperfect also. In spite of this inconvenience we can say that the results show that our system proves to be very robust and precise for the sets of videos used in this application. Ultimately in the future, we want to implement the application on a real-time map and to use a hybrid method to fuzzy logic.

ACKNOWLEDGMENT

I express my sincere gratitude to Pr ACHELI DALILA and Mohamed KESRAOUI for giving me the opportunity to take part in his team, and I also thank all the people who encouraged me to finish this work.

References

1. H. Chebi, D. Acheli, and M. Kesraoui, "Dynamic detection of abnormalities in video analysis of crowd behavior with DBSCAN and neural networks", *Advances in Science, Technology and Engineering Systems Journal*, vol. 1, no. 5, pp. 56-63. (2016).
2. C. Chen, Y. Shao, and X. Bi. "Detection of Anomalous Crowd Behavior Based on the Acceleration Feature". *IEEE Sensors Journal*, 15(12), 7252-7261. (2015).

3. T. Li, H. Chang, M. Wang, B. Ni, R. Hong, and S. Yan. "Crowded scene analysis: A survey". *IEEE Transactions on Circuits and Systems for Video Technology*, 25(3), 367-386. (2015).
4. G. J. Burghouts, R. den Hollander, K. Schutte, J. W. Marck, S. Landsmeer, and E. den Breejen. "Increasing the security at vital infrastructures: automated detection of deviant behaviors". In *SPIE Defense, Security, and Sensing* (pp. 80190C-80190C). International Society for Optics and Photonics. (2011, May).
5. H. Ullah, L. Tenuti, and N. Conci. "Gaussian mixtures for anomaly detection in crowded scenes". In *IS&T/SPIE Electronic Imaging* (p. 866303). International Society for Optics and Photonics. (2013, March).
6. Z. Zhong, W. Ye, and Y. Xu. "Detecting human abnormal behaviors in crowd". *International Journal of Information Acquisition*, 4(04), 281-290. (2007).
7. A. Johansson, D. Helbing, H. Z. Al-Abideen, and S. Al-Bosta. "From crowd dynamics to crowd safety: a video-based analysis". *Advances in Complex Systems*, 11(04), 497-527. (2008).
8. Q. Wang, Q. Ma, C. H. Luo, H. Y. Liu, and C. L. Zhang. "Hybrid Histogram of Oriented Optical Flow for Abnormal Behavior Detection in Crowd Scenes". *International Journal of Pattern Recognition and Artificial Intelligence*, 30(02), 1655007. (2016).
9. T. Ko, "A survey on behavior analysis in video surveillance for homeland security applications". *Applied Imagery Pattern Recognition Workshop. AIPR '08. 37th IEEE*, vol., no., pp.1,8, 15-17 Oct. 2008.
10. M. Szczodrak, J. Kotus, K. Kopaczewski, K. Lopatka, A. Czyzewski, and H. Krawczyk. "Behavior Analysis and Dynamic Crowd Management in Video Surveillance System," *Database and Expert Systems Applications (DEXA)*, 2011 22nd International Workshop on, vol., no., pp.371,375, Aug. 29 2011-Sept. 2 2011.
11. A. El Maadi, and M.S. Djouadi, "Suspicious motion patterns detection and tracking in crowded scenes," *Safety, Security, and Rescue Robotics (SSRR)*, 2013 IEEE International Symposium on, vol., no., pp.1,6, 21-26 Oct. 2013.
12. J. Li, and G. Wang. "A shadow detection method based on improved Gaussian Mixture Model". In *Electronics Information and Emergency Communication (ICEIEC)*, 2013 IEEE 4th International Conference on (pp. 62-65). IEEE. (2013, November).
13. H. Qian, X. Wu, Y. Ou, and Y. Xu. "Hybrid algorithm for segmentation and tracking in surveillance". In *Robotics and Biomimetics, 2008. ROBIO 2008. IEEE International Conference on* (pp. 395-400). IEEE. (2009, February).
14. S. Ali, and M. Shah. "A Lagrangian Particle Dynamics Approach for Crowd Flow Segmentation and Stability Analysis," *Computer Vision and Pattern Recognition, 2007. CVPR '07. IEEE Conference on*, vol., no., pp.1,6, 17-22 June 2007.
15. J. Shi, and C. Tomasi. "Good features to track," *Computer Vision and Pattern Recognition, 1994. Proceedings CVPR'94., 1994 IEEE Computer Society Conference on*, vol., no., pp.593,600, 21-23 Jun 1994.
16. B. Atcheson, W. Heidrich, and I. Ihrke. "An evaluation of optical flow algorithms for background oriented schlieren imaging". *Experiments in fluids*, 46(3), 467-476. (2009).
17. P.J. Burt, and E.H. Adelson, "The Laplacian Pyramid as a Compact Image Code," *Communications, IEEE Transactions on*, vol.31, no.4, pp.532,540, Apr 1983.
18. B.K. Horn, and B.G. Schunck. "Determining optical flow." 1981 Technical Symposium East. International Society for Optics and Photonics, 1981.
19. J.L. Barron, J.F. David, and S.B. Steven. "Performance of optical flow techniques." *International journal of computer vision*, 43-77, 12:1:1994.

20. H. Liu, T. H. Hong, M. Herman, T. Camus, and R. Chellappa. "Accuracy vs efficiency trade-offs in optical flow algorithms". *Computer vision and image understanding*, 72(3), 271-286. (1998).
21. C. Demonceaux. "Etude du mouvement dans les séquences d'images par analyse d'ondelettes et modélisation markovienne hiérarchique". Application à la détection d'obstacles dans un milieu routier. Diss. Université de Picardie Jules Verne, 2004.
22. M. Irani, B. Rousso, and S. Peleg. "Detecting and tracking multiple moving objects using temporal integration". In *Proc. European Conference on Computer Vision*, pages 282–287, Santa Margherita Ligure, Italy, Mai 1992.
23. S. Ayer, P. Schroeter, and J. Bigun. "Segmentation of moving objects by robust motion parameter estimation over multiple frames". In *Proc. European Conference on Computer Vision*, pages 316–327, Stockholm, Suède, Mai 1994.
24. L. Grewe, and A. C. Kak. "Interactive learning of a multiple-attribute hash table classifier for fast object recognition". *Computer Vision and Image Understanding*, 61(3), 387-416. (1995).
25. H. A. Rowley, S. Baluja, and T. Kanade. "Neural network-based face detection". *IEEE Transactions on pattern analysis and machine intelligence*, 20(1), 23-38. (1998).
26. P. Viola, M. J. Jones, and D. Snow. "Detecting pedestrians using patterns of motion and appearance". *International Journal of Computer Vision*, 63(2), 153-161. (2005).
27. C.P. Papageorgiou, M. Oren, and T. Poggio. "A general framework for object detection". In *Computer vision, 1998. sixth international conference on* (pp. 555-562). IEEE. (1998, January).
28. P. F. Gabriel, J. G. Verly, J. H. Piater, and A. Genon. "The state of the art in multiple object tracking under occlusion in video sequences". In *Advanced Concepts for Intelligent Vision Systems* (pp. 166-173). (2003, September).
29. R. Serajeh, K. Faez, and A. E. Ghahnavieh. "Robust multiple human tracking using particle swarm optimization and the Kalman filter on full occlusion conditions". In *Pattern Recognition and Image Analysis (PRIA), 2013 First Iranian Conference on* (pp. 1-4). IEEE. (2013, March).
30. T. Mathes, and J. H. Piater. "Robust non-rigid object tracking using point distribution manifolds". In *Joint Pattern Recognition Symposium* (pp. 515-524). Springer Berlin Heidelberg. (2006, September).
31. T. Yang, Q. Pan, J. Li, and S. Z. Li. "Real-time multiple objects tracking with occlusion handling in dynamic scenes". In *2005 IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR'05)* (Vol. 1, pp. 970-975). IEEE. (2005, June).
32. V. Rabaud, and S. Belongie. "Counting Crowded Moving Objects". *Computer Vision and Pattern Recognition, 2006 IEEE Computer Society Conference on*, vol.1, no., pp.705,711, 17-22 June 2006.
33. S. Wang, and Z. Miao. "Anomaly detection in crowd scene". In *IEEE 10th INTERNATIONAL CONFERENCE ON SIGNAL PROCESSING PROCEEDINGS* (pp. 1220-1223). IEEE. (2010, October).
34. P.J. Burt, E.H. Adelson, "The Laplacian Pyramid as a Compact Image Code," *Communications, IEEE Transactions on*, vol.31, no.4, pp.532,540, Apr 1983.
35. K. B. Horn and G.S. Brian. "Determining optical flow." 1981 Technical Symposium East. International Society for Optics and Photonics, 1981.
36. J.L. Barron, J. F. David, and S.B. Steven. "Performance of optical flow techniques." *International journal of computer vision*, 43-77, 12:1:1994.
37. S. Birchfield. "An implementation of the Kanade-Lucas-Tomasi feature tracker". 1998.

38. UMN, Minneapolis, MN, USA. (2006). Unusual Crowd Activity Dataset of University of Minnesota. [Online]. Available: <http://mha.cs.umn.edu/movies/crowd-activity-all.avi>.
39. PETS, Vellore, India. (2009). Multisensor Sequences Containing Different Crowd Activities. [Online]. Available: <http://www.cvg.rdg.ac.uk/pets2009/a.html>.
40. Blunsden S. J, Fisher R. B. The BEHAVE video dataset: ground truthed video for multi-person behavior classification. *Annals of the BMVA*, Vol 2010(4), pp 1-12, <http://groups.inf.ed.ac.uk/vision/BEHAVEDATA/INTERACTIONS/>.
41. S. Wu, H.-S. Wong, and Z. Yu, A Bayesian model for crowd escape behavior detection. *IEEE Trans. Circuits Syst. Video Technol.*, vol. 24, no. 1, pp. 85–98, Jan. 2014.
42. R. Mehran, A. Oyama, and M. Shah, Abnormal crowd behavior detection using social force model. In *Proc. IEEE Int. Conf. Comput. Vis. Pattern Recognit. (CVPR)*, Miami, FL, USA, Jun. 2009, pp. 935–942.
43. C. Chen, Y. Shao, and X. Bi. Detection of Anomalous Crowd Behavior Based on the Acceleration Feature. *IEEE Sensors Journal*, 15(12), 7252-7261. 2015.

Suivi multi-cibles

Houssam Eddine Rouabhia¹, Brahim Farou¹, Hamid Seridi¹, Herman Akdag²

¹ LabSTIC, Université 8 mai 1945, Guelma, Algérie

²LIASD, Université Paris 8, 93526 Saint-Denis, France

rouabhia.h@gmail.com

farou@ymail.com

seridihamid@yahoo.fr

herman.akdag@ai.univ-paris8.fr

Résumé. Le suivi des objets en mouvement est une étape essentielle pour construire un système intelligent de prise de décision précis et rapide. Cependant, il y a plusieurs conditions rendent le suivi d'objet difficile. Ces conditions incluent le mouvement d'objet non-rigide, les variations d'apparence de cibles dues à des changements d'illumination, et l'encombrement d'arrière-plan. Dans cet article, nous présentons une méthode rapide et précise pour le suivi de plusieurs objets dans les applications de vidéosurveillance. Il est basé sur la technique de soustraction d'arrière-plan pour réduire l'espace de recherche et détecter les régions candidates, un descripteur incluant des caractéristiques de couleur et de forme globale en plus le filtre de Kalman pour traiter l'occlusion et rendre la méthode précise pour identifier chaque objet. Les expériences effectuées sur la base de données CAVIAR ont montré l'efficacité du système proposé.

Mots clés: suivi des objets, vidéos surveillances, GMM, CENTRIST, filtre de Kalman

1 Introduction

Le suivi des objets en mouvement est un problème difficile régi par de nombreuses circonstances d'ordre techniques et environnementales qui doivent être réconciliées en un seul algorithme. Parmi les problèmes rencontrés dans ce domaine, nous citons : les variations d'éclairage, les variations liées à l'emplacement de la caméra, la position de l'objet par rapport à la caméra, les changements d'apparence de l'objet et les occlusions.

Le suivi se trouve dans de nombreuses applications regroupées aux alentours d'un axe de recherche appelé vision par machine. Récemment, c'est la surveillance et l'interaction homme-machine qui ont amplement bénéficié des avancées technologiques en termes de matériels et logiciels intelligents. On trouve beaucoup de méthodes dans l'état de l'art, ces derniers se différencient entre eux, mais généralement elles se sont divisées en trois groupes.

Dans le premier groupe, les trackers effectuent une correspondance de la représentation des extraits cibles d'une image source et d'un modèle construit précédemment sur la base d'une image de référence. L'objectif est d'établir la correspondance entre la cible et le modèle. Un degré de similarité entre eux est calculé pour prendre une décision finale et étiqueter cette cible [1]. Dans [2] le modèle de l'objet est représenté par un ensemble de patches locaux, où chaque patch est suivi de manière indépendante et les patches votent pour le centre de l'objet. Ce modèle est utilisé pour manipuler les occlusions partielles. Dans [3] Plusieurs trackers sont échantillonnés et combinés pour gérer des scénarios difficiles. Cependant les approches basées sur ces modèles ne fonctionnent que dans le fond avec une texture simple. En outre, la complexité de l'algorithme, liée à la recherche de l'objet dans toute l'image, empêche l'utilisation de ces méthodes dans une application en temps réel [4] [5].

Les trackers du deuxième groupe traitent le problème de suivi comme une tâche de détection de l'objet au fil du temps en utilisant un classificateur adéquat pour distinguer les pixels cibles. La première étape est la détection d'objets. Il est utilisé pour localiser la position de chaque cible. Ensuite, chaque cible détectée peut être classifiée et étiquetée en fonction du résultat du classificateur. Une combinaison de segmentation et d'apprentissage en ligne est proposée dans [6]; Cependant, la procédure d'apprentissage basée sur la forêt de Hough est très coûteuse et ne peut pas être utilisée en temps réel. P-N algorithm est utilisé pour exploiter la structure sous-jacente des échantillons positifs et négatifs pour apprendre un classificateurs efficaces pour le suivi des objets [7]. Les méthodes basées sur les classificateurs utilisent un apprentissage hors ligne avant même le début du suivi. Cela paralysera le système si l'objet change constamment sa posture. La construction des classificateurs nécessite de vastes échantillons positifs et négatifs. D'autre part, la nécessité de rechercher des objets dans une zone de grande portée entraîne un coût de calcul élevé. Dans la littérature, il ya aussi quelques méthodes qui essaient de profiter de l'apprentissage en ligne [8] [9]. Cependant, ils ont l'inconvénient de converger vers l'apparence actuel de l'objet.

La troisième catégorie utilise le filtre prédictif pour estimer l'état de l'objet, compte tenu de toutes les mesures jusqu'à l'heure actuelle. Il se compose de deux phases, la première est de prédire la position de l'objet en utilisant les informations précédentes. La seconde consiste à vérifier la validité de la position prédite. Un grand nombre de techniques de filtrage ont utilisé le filtre de Kalman (KF) pour le suivi des objets en raison de son efficacité [10]. Le filtre de Kalman est également utilisé pour estimer les emplacements des objets pour construire un modèle de relation mutuelle entre deux objets pour le suivi de cibles multiples [11]. Pour une estimation rapide et robuste de l'état des cibles, la texture locale, la répartition spatiale des couleurs et les gradients orientés vers les bords avec un filtre mixte Kalman / H_∞ sont utilisés pour former un modèle robuste [12]. Cependant, les hypothèses de linéarité et de gaussianité du KF ne peuvent pas gérer des scénarios complexes, ces filtres nécessitent aussi une correction périodique des valeurs estimées pour éviter de perdre la trace de l'objet suivi. Ces corrections sont effectuées soit manuellement, soit par combinaison avec des classificateurs.

Étant donné que chaque méthode présente des avantages et des inconvénients, nous proposons une méthode impliquant les avantages de chaque méthode. Où nous avons

utilisé un classificateur avec un apprentissage en ligne. Ce classificateur est également utilisé pour collaborer avec le filtre de Kalman pour éviter de perdre l'objet en cas d'occlusion.

Dans cet article, nous visons le suivi de plusieurs objets dans des vidéos capturées à partir de caméras fixes. Dans la deuxième section, nous présentons un aperçu du système proposé et décrivons en détails chaque phase. Les résultats expérimentaux sont présentés dans la troisième section. Enfin, nous concluons par une conclusion.

2 Notre système

Pour aboutir à un système stable et performant, nous avons divisé le système proposé en quatre modules. Le premier consiste dans la sélection des régions d'intérêt (ROI) pour extraire les régions candidates. Le deuxième permet de coder chaque ROI pour trouver la meilleure correspondance en se basant sur la description de l'apparence. En parallèle, un module d'estimation de mouvement est appliqué afin de prédire les futures positions possibles de chaque objet et de choisir l'objet le plus proche. Le dernier module permet d'associer une étiquette pour chaque objet en se basant sur les informations transmises par les modules précédents.

L'initialisation du système est effectuée une seule fois lors de l'acquisition de la première image de la vidéo. Cette dernière est utilisée pour la création des modèles des objets qui seront utilisés dans la phase de suivi.

Chaque modèle est défini par son futur emplacement et un vecteur centroïde représentant la moyenne de ces vecteurs. L'extraction des caractéristiques est effectuée de façon identique à celle utilisée dans la phase de suivi. Chaque vecteur inclut la couleur et les caractéristiques de la forme globales de l'objet.

2.1 Détection des régions en mouvement :

La plupart des méthodes de suivi utilisent toute l'image comme un espace de recherche. Ce choix augmente considérablement la complexité de ces approches et génère beaucoup de faux positifs au niveau de l'étape de décision. Pour cette raison, nous avons proposé comme première étape de détecter les régions qui peuvent contenir des mouvements afin de réduire l'espace de recherche et par conséquent réduire le temps nécessaire pour le suivi des cibles.

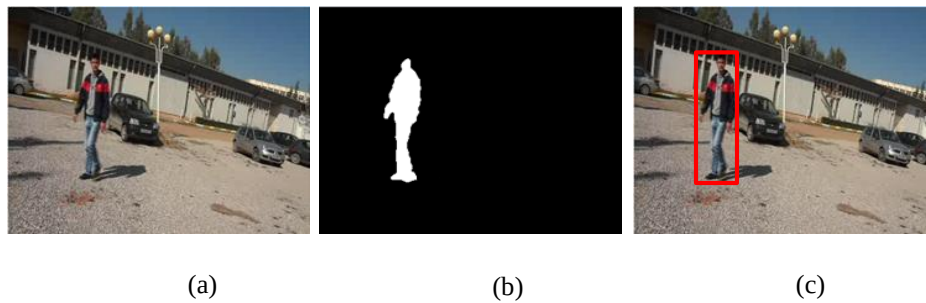
Plusieurs méthodes ont été proposées dans l'état de l'art pour résoudre le problème de la détection des objets en mouvement [différence de frame] [GMM][KDE]. Cependant, de nombreux travaux comparatifs ont montré que le modèle des mixtures de gaussiennes présente l'avantage d'être le moins gourmand en espace mémoire et en temps de calcul et donne de très bons résultats dans des environnements à faibles variations de luminosité [13].

Pour cette raison et afin d'alléger les étapes ultérieures, nous avons choisi d'utiliser les GMM pour détecter les objets en mouvement en temps réel.

Le principe des GMM est d'utiliser l'historique des valeurs d'un pixel pour estimer l'intervalle des prochaines valeurs du même pixel délimité par l'allure et l'étalement

de la gaussienne. La création du modèle consiste à attribuer pour chaque pixel un ensemble de K gaussiennes ou chacune d'elle permettra de modéliser le comportement du pixel dans un laps de temps qui dépend fortement de la valeur de K . Le résultat du modèle de mélange gaussien est illustré dans la fig. 1.

Fig. 1. Le résultat du modèle de mélange gaussien



2.2 Descripteur d'apparence.

L'apparence est une caractéristique extrêmement importante dans la détection des objets. En effet, il suffit d'avoir une idée sur la forme de l'objet pour pouvoir le classer correctement. Pour une meilleure performance, nous avons utilisé un descripteur pour chaque ROI contenant deux types de caractéristiques sous forme d'un histogramme : les couleurs et la forme globale. Cette combinaison nous a permis de séparer entre les objets similaires dans leur forme en utilisant la caractéristique couleur qui joue un peu le rôle d'un descripteur secours en cas de conflit entre l'apparence des objets. Il existe plusieurs espaces de couleur dans le domaine du traitement d'image. Nous avons utilisé dans notre approche l'espace de couleur RGB pour sa simplicité et pour éviter des transformations inutiles vu que c'est le format native d'une image en entrée. Pour la représentation de la forme globale des objets, nous avons choisi d'utiliser le descripteur CENTRIST [14]. Ce choix est justifié par la capacité de CENTRIST à donner de bons résultats en gardant une vitesse acceptable par rapport aux autres caractéristiques, sachant que le suivi des objets nécessite dans la plupart des cas des réponses en temps réel. En plus, CENTRIST n'a besoin d'aucun prétraitement sur l'image ni une étape de normalisation du vecteur de caractéristique.

L'image est convertie en histogramme de transformation du recensement en appliquant l'opérateur de Sobel à l'image d'entrée ; le gradient de Sobel de chaque pixel est remplacé par sa valeur de transformation du recensement (CT) pour générer l'image CT. La valeur CT d'un pixel est obtenue en comparant la valeur d'intensité d'un pixel avec ses huit voisinages comme suivant l'équation (1):

$$CT_i = \sum_{N=0}^7 [2^N \times S(P_i - P_N)] , S(P_i - P_N) = \begin{cases} 1, & P_i - P_N \geq 0 \\ 0, & P_i - P_N < 0 \end{cases} \quad (1)$$

Lorsque P_N ($N = 0, 1 \dots 7$) est le N ième voisin du pixel P_i , qui est le pixel central dans une fenêtre de balayage de 3×3 .

Correspondance d'apparence.

Cette étape permet de mesurer la similarité entre l'objet en cours de traitement et les modèles créés précédemment en se basant sur les vecteurs de caractéristiques. Pour cela, nous avons utilisé la distance de Bhattacharyya réputé d'être parmi les meilleures mesures de similarité dans le traitement d'image et de la vidéo [15]. Nous allons calculer le coefficient de Bhattacharyya entre le vecteur de caractéristique de la région candidate et le vecteur centroïde de chaque modèle d'objet. L'utilisation du centroïde nous permet de gagner du temps. En effet, si un modèle n'est pas proche de la description de l'objet d'intérêt, le système va calculer plusieurs fois la similarité de façon inutile vu que dès le départ le modèle ne décrit pas l'objet. Donc le centroïde permet de représenter le modèle et éviter des calculs superflus.

Si le résultat de la similarité est inférieur à un seuil $T1$ alors le modèle de l'objet et est directement rejeté. Dans le cas contraire, les modèles ayant obtenus un score supérieur à $T1$ seront directement pris en considération dans la phase de décision. Ce processus est efficace pour déterminer si ce modèle représente réellement la région candidate. L'apparence ayant obtenu le score le plus élevé sera considérée par le système en tant que modèle de l'objet.

Update.

La mise à jour permet de recalculer le vecteur centroïde pour chaque modèle. Dans le cas où une région candidate est classée comme un objet existant dans l'unité d'étiquetage, le vecteur de caractéristiques représentant cet objet sera pris en considération dans le calcul de vecteur centroïde de ce modèle.

Dans le cas où le système n'a trouvé aucun modèle qui permet de représenter l'objet, la région candidate est classée comme un nouvel objet et un nouveau modèle de l'objet est créé pour ce candidat.

2.3 Estimation de mouvement

Cette phase permet de réduire l'espace de recherche et de simplifier la tâche pour le module de classification. En effet, en se basant sur les informations relatives aux positions futures de l'objet, les mauvaises classifications peuvent être rapidement détectées et corrigées. En plus, l'estimation de la zone susceptible de contenir l'objet en mouvement permettra de réduire considérablement le temps de recherche et garantit

un degré de fiabilité non négligeable vu que la prédiction de la position de l'objet en mouvement a une très faible chance de ne pas être correcte.

Le processus d'estimation de la position future de chaque cible s'effectue en associant un filtre de Kalman à chaque ROI. Le filtre de Kalman est un estimateur optimal utilisant l'information des mesures et des états précédents. Le centre de l'objet est trouvé en premier, puis nous utilisons un filtre de Kalman pour prédire la position de celui-ci dans la trame suivante. Un filtre de Kalman est utilisé pour estimer l'état d'un système linéaire où l'état est supposé suivre une distribution de type gaussienne. Pour le modèle de mouvement d'un objet en mouvement (qui contient une sorte de bruit dynamique), et quelques observations bruyantes sur sa position, le filtre de Kalman fournit une estimation optimale de sa position à chaque instant. L'optimum est garanti dans le cas où tous les bruits sont de type gaussien. Dans ce cas, le filtre va permettre de minimiser l'erreur quadratique moyenne des paramètres estimés (par exemple position, vitesse). Le filtre de Kalman est un processus en ligne, les nouvelles observations sont traitées dès qu'elles arrivent. Le filtre de Kalman se compose de deux étapes, la prédiction pour la mise à jour du temps et la correction pour la mise à jour de la mesure. L'estimation d'état est obtenue en pondérant la somme de la prévision d'état et de la mesure de sortie. Dans notre application, le filtre de Kalman est utilisé pour estimer des états de l'objet en mouvement.

Dans le cas où la région candidate est classée comme un objet existant dans l'unité d'étiquetage, sa position sera utilisée pour corriger le filtre de Kalman. Dans le cas contraire, la région candidate est classée comme un nouvel objet, et un nouveau filtre de Kalman est associé à cet objet.

2.4 Étiquetage

À cette étape, une décision finale est prise si le ROI est un objet existant ou non. Cette décision est prise en utilisant les deux distances calculées à l'étape 2 et à l'étape 3 où:

L'ensemble des modèles d'objets $M = \{ M_J, J \in 1 \dots n \}$ et l'ensemble des ROI $= \{ ROI_I, I \in 1 \dots m \}$. La distance euclidienne calculée à l'étape 2 est définie comme $D_E = \{ D_{I,J}(ROI_I, M_J), I \in 1 \dots m, J \in 1 \dots n \}$. La distance de Bhattacharyya calculée à l'étape 3 est définie comme $D_B = \{ D_{I,J}(ROI_I, M_J), I \in 1 \dots m, J \in 1 \dots n \}$

n Est le nombre total de modèles d'objets créés avant. m Est le nombre total de ROI existant dans le frame actuel. Pour chaque modèle $ROI_I \in ROI$ Nous recherchons des modèles qui sont proches de ce ROI où :

$R_I = \{ M_S / M_S \in M, D_{I,S} \in D_E \text{ and } D_{I,S} < T \}$ Où T est un seuil.

Les ROI qui ne contiennent aucun modèle d'objet considéré comme de nouveaux objets. Le reste est étiqueté comme:

$$\text{Étiquette} = \underset{J}{\operatorname{argmin}} (D_{I,J}(ROI_I, M_J) / D_{I,J} \in D_B \text{ and } M_J \in R_J)$$

3 Expérimentation

La méthode proposée est une méthode automatique de suivi de n'importe quel objet en mouvement. Elle n'a besoin d'aucune segmentation de couleur et elle permet de suivre plusieurs objets en même temps sans aucun critère ni condition préalable sur l'objet. Nous avons testé le système proposé dans de nombreuses situations afin de montrer l'efficacité de notre approche et cela, quel que soit l'environnement d'application. Nous allons dans ce qui suit présenter des scénarios représentant les cas les plus courants qu'un système de suivi peut affronter.

3.1 Multi cible mais sans aucune similitude entre eux

L'absence de toute similitude entre les objets facilite le processus de suivi où l'utilisation de la forme, la couleur et l'emplacement donne de bons résultats et ne conduit à aucun problème.

3.2 Multi cible avec la même couleur et une forme différente ou avec la même forme et une couleur différente

L'existence d'une quelconque similitude entre deux objets dans la forme ou dans la couleur n'affecte pas réellement les résultats du suivi. En effet, les objets similaires dans la forme peuvent être facilement séparables par la caractéristique couleur et vice versa. Cependant, lorsque deux objets sont presque identiques dans leurs formes et ont une couleur similaire cela peut engendrer des erreurs de suivi et mettre le système dans une situation incohérente. La plupart des systèmes de suivi échouent dans ce cas de figure et nécessitent dans la plupart des cas une réinitialisation par l'opérateur humain. Pour remédier à ce problème, notre système utilise les informations produites par le filtre de Kalman pour faire la distinction vu que ce dernier est immunisé par rapport à ce type de problème et offre une précision beaucoup plus grande dans la distinction entre les deux objets.

3.3 Multi objet similaire

En cas de similitude totale entre deux objets se trouvant dans des positions éloignées, la distinction entre eux est basée sur la prédiction du filtre Kalman attribué à chaque objet.

Pour plus d'évaluation, nous utilisons la base de données CAVIAR pour évaluer notre méthode et comparer les résultats avec Zhang et al [16] and Xing et al [17]. La base de données CAVIAR comprend 26 séquences vidéo d'une passerelle dans un centre commercial prises par une caméra unique avec une taille d'image de 384×288 et un taux d'images de 25fps.

Nous évaluons la performance de suivi en fonction des cinq paramètres suivants proposés dans [18]:

- La plupart des trajectoires suivies (MT), le nombre de trajectoires qui sont suivies avec succès pour plus de 80% (plus cette valeur est grande, mieux un tracker est)
- La plupart des trajectoires perdues (ML), le nombre de trajectoires qui sont suivies pour moins de 20% (plus cette valeur est petite, mieux un tracker est)
- Des trajectoires partiellement suivies (PT), le nombre de trajectoires suivies entre 20% et 80%
- Fragmentation (FRMT), le nombre de fois qu'une trajectoire est interrompue (plus cette valeur est petite, mieux un tracker est)
- (IDS), le nombre de fois où deux trajectoires changent d'ID (plus cette valeur est petite, mieux un tracker est)

Les résultats sur CAVIAR sont montrés dans le Tableau 1, où GT est le nombre de trajectoires

Afin d'évaluer les performances de notre système et bien situer notre approche par rapport à l'état de l'art, nous avons choisi quatre méthodes qui utilisent un suivi multicibles sur la même base de données. Ce choix a été guidé d'une part par la robustesse de ces méthodes et d'autre part par l'absence de méthodes de suivi multicible en général et en particulier dans la base de données CAVIAR.

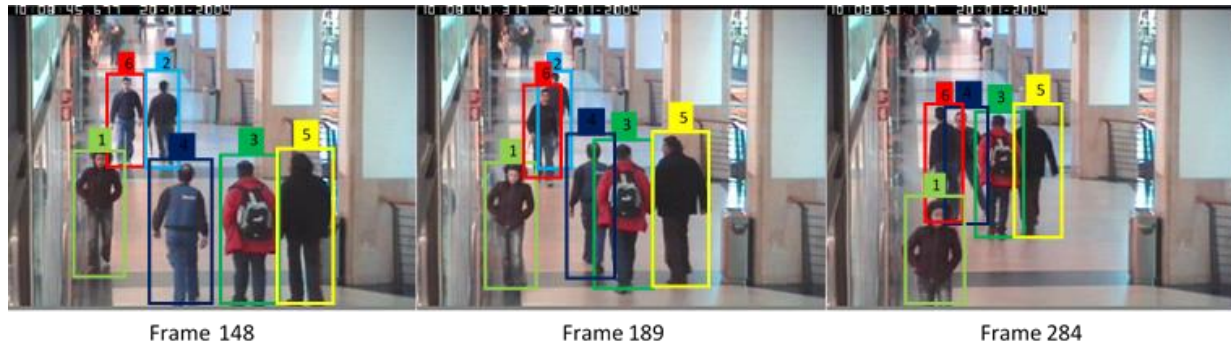
Table 1. Comparaison des résultats de suivi sur la base de données CAVIAR.

Method	GT	MT	PT	ML	FRMT	IDS
Zhang et al [16] Alg.1	140	104	29	7	58	7
Zhang et al [16] Alg.2	140	120	15	5	20	15
Xing et al [17] Étape locale	140	112	12	6	33	16
Xing et al [17] Deux étapes	140	118	17	5	24	14
Notre méthode	140	121	14	5	16	8

Le tableau 1 montre clairement, à travers le critère MT, que notre système dépasse les autres méthodes dans la capacité de suivre les cibles le plus longtemps possible avec un taux de réussite de 86,42 %. Notre système a également réussi à suivre de façon partielle 14 cibles exprimées par le critère PT. On remarque aussi que notre système présente le taux d'interruption le faible et il a eu une perte d'uniquement cinq cibles c.-à-d. 3,57 % de l'ensemble des objets suivi. L'autre point fort de notre système est la capacité de différencier entre les objets même en cas de conflit, et cela avec un faux étiquetage, de 8 objets sur la totalité des objets suivis, exprimé par le critère IDS.

Cela est dû à l'existence d'un consensus entre les unités de la méthode, qui dépend du résultat du filtre de Kalman et de l'apparence de l'objet pour la recherche d'occlusion entre les objets, ce qui réduit la perte de l'objet et l'interruption de leurs trajectoires.

Fig. 2. Résultat de suivi sur la base de données CAVIAR



L'apprentissage en ligne de l'apparence d'objets fortifie l'identification de l'objet comme le montre la Figure 2. Dans les cas d'une similarité, le filtre de Kalman nous offre par ces capacités de prédiction de distinguer entre les objets similaires. L'utilisation de l'apparence et de la position de prédiction a donné une bonne précision dans l'identification des objets. L'utilisation de ce résultat final pour corriger le filtre de Kalman et mettre à jour l'apparence des objets a permis au système de distinguer entre les objets suivis, ce qui a fortement contribué dans la baisse du nombre de fois où deux trajectoires changent leurs étiquètes.

Concernant la vitesse de traitement, notre système opère avec une fréquence d'environ 10 fps. Ce résultat peut être amélioré en utilisant une approche parallèle. Qui peut donner une application en temps réel utilisé pour le système de surveillance visuelle en ligne.

4 Conclusion

Dans cet article, nous avons étudié le problème du suivi d'un objet inconnu dans les vidéos surveillances. Nous proposons une nouvelle méthode qui décompose la tâche en quatre composantes, détection, identification, apprentissage et étiquetage. Nous avons proposé comme première étape de détecter des régions qui peuvent contenir des mouvements afin de réduire l'espace de recherche et donc de réduire le temps nécessaire pour suivre les cibles. Nous avons utilisé un descripteur pour chaque ROI contenant deux types de caractéristiques sous la forme d'un histogramme de couleur et de forme globale. Cette combinaison nous a permis de séparer les objets similaires. Parallèlement, un module d'estimation de mouvement est appliqué afin de prédire les futures positions possibles de chaque objet pour choisir l'objet le plus proche qui présente chaque candidat. Ces deux modules fonctionnent en parallèle pour améliorer la capacité d'identifier les objets. Car chaque module complète et corrige l'autre. L'apprentissage

en ligne de l'apparition d'objets renforce l'identification de l'objet aussi. La compatibilité entre ces modules permet de résoudre le problème du chevauchement entre les objets d'une manière marquée, ce qui a permis de diminuer le nombre de fois où deux trajectoires modifient leur étiquetage et augmentent la capacité de suivre les cibles le plus longtemps possible .

Bibliographie

1. Oron, S., Bar-Hillel, A., Levi, D., Avidan, S.: Locally Orderless Tracking. In: 2012 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 1940–1947 (2012).
2. Adam, A., Rivlin, E., Shimshoni, I.: Robust Fragments-based Tracking using the Integral Histogram. In: 2006 IEEE Computer Society Conference on Computer Vision and Pattern Recognition (CVPR'06). pp. 798–805 (2006).
3. Kwon, J., Lee, K.M.: Tracking by Sampling Trackers. In: 2011 International Conference on Computer Vision. pp. 1195–1202 (2011).
4. Kwon, J., Lee, K.M., Park, F.C.: Visual tracking via geometric particle filtering on the affine group with optimal importance functions. In: IEEE Conference on Computer Vision and Pattern Recognition, 2009. CVPR 2009. pp. 991–998 (2009).
5. Čehovin, L., Kristan, M., Leonardis, A.: An adaptive coupled-layer visual model for robust visual tracking. In: 2011 International Conference on Computer Vision. pp. 1363–1370 (2011).
6. Godec, M., Roth, P.M., Bischof, H.: Hough-based tracking of non-rigid objects. In: 2011 International Conference on Computer Vision. pp. 81–88 (2011).
7. Kalal, Z., Matas, J., Mikolajczyk, K.: P-N learning: Bootstrapping binary classifiers by structural constraints. In: 2010 IEEE Conference on Computer Vision and Pattern Recognition (CVPR). pp. 49–56 (2010).
8. Hare, S., Golodetz, S., Saffari, A., Vineet, V., Cheng, M.M., Hicks, S.L., Torr, P.H.S.: Struck: Structured Output Tracking with Kernels. *IEEE Trans. Pattern Anal. Mach. Intell.* 38, 2096–2109 (2016).
9. Liu, F., Shen, C., Reid, I., van den Hengel, A.: Online unsupervised feature learning for visual tracking. *Image Vis. Comput.* 51, 84–94 (2016).
10. Kalman, R.E.: A New Approach to Linear Filtering and Prediction Problems. *J. Basic Eng.* 82, 35–45 (1960).
11. Duan, G., Ai, H., Cao, S., Lao, S.: Group Tracking: Exploring Mutual Relations for Multiple Object Tracking. In: Fitzgibbon, A., Lazebnik, S., Perona, P., Sato, Y., and Schmid, C. (eds.) *Computer Vision – ECCV 2012*. pp. 129–143. Springer Berlin Heidelberg (2012).
12. Wang, H., Nguang, S.: Multi-target Video Tracking Based on Improved Data Association and Mixed Kalman/H; Filtering. *IEEE Sens. J.* PP, 1–1 (2016).

13. Farou, B., Seridi, H., Akdag, H.: Improved Parameters Updating Algorithm for the Detection of Moving Objects. In: IFIP International Conference on Computer Science and its Applications_x000D_. pp. 527–537. Springer (2015).
14. Wu, J., Rehg, J.M.: CENTRIST: A Visual Descriptor for Scene Categorization. *IEEE Trans. Pattern Anal. Mach. Intell.* 33, 1489–1501 (2011).
15. Farou, B., Seridi, H., Akdag, H.: IAJIT - Improved Gaussian Mixture Model with Background Spotter for the Extraction of Moving Objects. *J. IAJIT*. 13, (2016)
16. Zhang, L., Li, Y., Nevatia, R.: Global data association for multi-object tracking using network flows. In: IEEE Conference on Computer Vision and Pattern Recognition, 2008. CVPR 2008. pp. 1–8 (2008).
17. Xing, J., Ai, H., Lao, S.: Multi-object tracking through occlusions by local tracklets filtering and global tracklets association with detection responses. In: IEEE Conference on Computer Vision and Pattern Recognition, 2009. CVPR 2009. pp. 1200–1207 (2009).
18. Wu, B., Nevatia, R.: Detection and Tracking of Multiple, Partially Occluded Humans by Bayesian Combination of Edgelet based Part Detectors. *Int. J. Comput. Vis.* 75, 247–266 (2007).

b -chromatic number of K_m cartesian product with some particular graphs

B.FERDJALLAH¹ and F.AFFIF CHAOUICHE² and A.BERRACHEDI³

Faculty of Mathematics, USTHB, BP 32 El-Alia,
Bab-Ezzouar 16111, Algiers, Algeria.

¹bayaferdja@hotmail.com

²f_affif@yahoo.fr

³aberrachedi@usthb.dz

Abstract. The b -chromatic number of a graph $G = (V, E)$, $\varphi(G)$, is the maximum number k of colors that can be used to color the vertices of G , such that we obtain a proper coloring and each color i , $1 \leq i \leq k$ has at least one representant vertex v_i adjacent to a vertex of every color j , $1 \leq j \neq i \leq k$. We denote by S_n the complete bipartite graph $K_{1,n}$ and H_n the graph wheel; constituted by an n -cycle with another vertex adjacent to all the vertices of the cycle. In this paper, we first determine the value of $\varphi(K_m \square S_n)$, then $\varphi(K_m \square H_n)$; under some conditions on the two integers m and n .

Key Words: Chromatic number, b -coloring, b -chromatic number.

1 Introduction

All considered graphs are finite, undirected, without loops and multiple edges and connected. A graph G is denoted by (V, E) , where V and E are the vertex set and the edge set respectively. uv is an edge between u and v . $d(u, v)$ is the length of the shortest induced (u, v) -path. The order of G is the cardinality of V . By $\Delta(G)$ (resp. $\delta(G)$), we mean the maximum (resp. minimum) degree of G . A chromatic number of a graph G , is the minimum number of colors that one can use to color all the vertices of G , in such a way that *two* adjacent vertices have *two* different colors, it is denoted by $\chi(G)$.

For a given m , K_m denotes the complete graph of order m . The cartesian product of *two* graphs G and G' is the graph $G \square G'$; where $V(G \square G') = V(G) \times V(G')$ and (u, u') is adjacent to (v, v') if and only if either $u = v$ and $u'v' \in E(G')$ or $u' = v'$ and $uv \in E(G)$.

In this work, we represent the cartesian product of *two* graphs G_1 and G_2 of orders n_1 and n_2 respectively, by an $n_1 \times n_2$ -array, where each column induces a copy of G_1 and each row induces a copy of G_2 . So the entry (i, j) is corresponding to the vertex (u_i, u_j) with $u_i \in V(G_1)$ and $u_j \in V(G_2)$. The dominating system

is given by the vertices with $*$.

The b -chromatic number of a graph G is defined as the maximum number k of colors that can be used to color the vertices of G , such that *two* adjacent vertices have *two* different colors and each color i , with $1 \leq i \leq k$ has at least *one* representant vertex u_i adjacent to vertex of every color j , $1 \leq j \neq i \leq k$. u_i is called a dominating vertex for the color i . The set of all dominating vertices is a dominating system of G . If the dominating system is stable, we said that G admit a stable dominating system.

The b -chromatic number $\varphi(G)$ was introduced by R.W Irving and D.F.Manlove in [1]. B.Omoomi and R.Javadi [3] studied the cartesian product of complete graph with cycle, path and with the same complete graph.

An elementary result regarding this invariant is :

$$\chi(G) \leq \varphi(G) \leq \Delta(G) + 1.$$

In this paper, we consider the cartesian product of a complete graph K_m with a wheel H_n , which is constituted by an n -cycle with another vertex adjacent to all the vertices of the cycle.

2 The b -chromatic number of the graph $K_m \square K_{1,2,2}$.

$K_{1,2,2}$ is a tripartite complete graph. It is isomorphic to the wheel H_4 . It is proved in [2] that for 2 graphs G_1 and G_2 ,

$$\max(\varphi(G_1), \varphi(G_2)) \leq \varphi(G_1 \square G_2) \leq \Delta(G_1 \square G_2) + 1.$$

$\varphi(H_4) = 3$ and $\Delta(K_m \square H_4) = m + 3$. Then, for $m \geq 3$, we have :

$$m \leq \varphi(K_m \square H_4) \leq m + 4. \quad (1)$$

More, we have the following lemma.

Lemma 1. *For m a positive integer no null, we have :*

$$\varphi(K_m \square H_4) \leq m + 3. \quad (2)$$

Proof. Let's suppose that $\varphi(K_m \square H_4) > m + 3$. Then, we must have at least $m + 4$ vertices of degree greater or equal then $m + 3$.

But in $K_m \square H_4$, there are only m vertices of degree $m + 3$. So the result follow.

It is proved by B.Omoomi and R.Javadi in [3], that if C is a set of φ colors assigned to the graph $K_m \square G$, where $\varphi > m$; then the column corresponding to the vertex v of G in $K_m \square G$, contain at most $d_G(v)$ dominating vertices. It

follows that if G is a graph of n vertices and e edges; $d = (d_1, d_2, \dots, d_n)$ is the degree sequence of the graph G , then :

$$\varphi(K_m \square G) \leq \sum_{i=1}^n d_i = 2e. \quad (3)$$

We proved the following theorem.

Theorem 1. *Let m and n be two strictly positive integers. We have :*

$$\varphi(K_m \square H_4) = \begin{cases} m+3 & \text{if } m=2 \\ (\geq) m+2 & \text{if } 3 \leq m \leq 14 \\ (\geq) m+1 & \text{if } m=15 \\ m & \text{if } m \geq 16 \end{cases}$$

Proof. • We prove that if $m \geq 16$, then $\varphi(K_m \square H_4) = m$.

Suppose that $\varphi(K_m \square H_4) > m$.

By (3), we have necessarily $\varphi(K_m \square H_4) \leq \sum_{i=1}^5 d_i = 16$; where $(d_1, d_2, \dots, d_5)$ is the degree sequence of the graph H_4 . Thus the contradiction.

- If $m = 2$, we give in The following figure a 4-dominating coloring of the graph $K_2 \square H_4$.

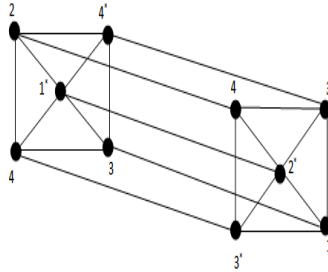


Fig. 1.

The dominating system is : $\{1^*, 2^*, 3^*, 4^*\}$.

The dominant coloration of the graph $K_2 \square H_4$ in the form of a table is represented by the following figure:

1*	3	4*	2	4
2*	1*	3	4	3*

Fig. 2.

- If $3 \leq m \leq 14$, $K_m \square H_4$ admit a $(m+2)$ -dominating coloring. It can be presented in a similar way as the previous array. For example, for $m = 7$ (Figure 3), we have The dominating system is : $\{1^*, 2^*, 3^*, 4^*, 5^*, 6^*, 7^*, 8^*, 9^*\}$.

1*	8	9	4	2
2*	5*	8	9*	3
3*	6*	2	8*	9
4*	9	1	5	8
5	1	7*	6	7
6	7	4	7	1
7	4	3	1	2

Fig. 3.

- For $m = 15$, the graph admit the following 16-dominating coloring.

1*	16*	2	14	15
2*	15*	16*	15	14
3*	14*	13*	16*	2
4*	1	12*	3	16
5	3	11*	10*	6*
6	4	1	9*	1
7	5	6	8*	3
8	7	5	1	4
9	8	4	2	11
10	9	7	4	7*
11	10	8	7	5*
12	11	9	5	8
13	12	10	11	9
14	6	14	12	13
15	13	15	13	12

Fig. 4.

We now consider the graph wheel H_n formed by connecting a single vertex to all vertices of a cycle of length n , in this section we are interested in the search for the b -chromatic number of $K_m \square H_n$.

3 b -chromatic number of the graph $K_m \square H_n$

The graph $K_m \square H_n$ has maximum degree $m + n - 1$. As $\chi(K_m \square H_n) = m$, then

$$m \leq \varphi(K_m \square H_n) \leq m + n.$$

But this graph has just m vertices of degree $m + n - 1$ and all the other ones have degree $m + 2$. So, we deduce the following corollary which is the general case of the lemma 1.

Corollary 1. *Let m and n be two positive integers, with $n > 4$. We have :*

$$m \leq \varphi(K_m \square H_n) \leq m + 3.$$

Proof. Suppose that $\varphi(K_m \square H_n) > m + 3$. Then, there are at least $\varphi(K_m \square H_n)$ vertices of degree greater or equal to $m + 3$.

But $K_m \square H_n$ contain only m vertices of degree $m + n - 1$.

Thus $m \leq \varphi(K_m \square H_n) \leq m + 3$.

Let $d = (d_1, d_2, \dots, d_n)$ be the degree sequence of the graph H_n . Since the last corollary :

$$\varphi(K_m \square H_n) \leq \sum_{i=1}^n d_i = 4n.$$

We deduce that : $\varphi(K_m \square H_n) \leq \min(m + 3, 4n)$.

To get the b -chromatic number of $K_m \square H_n$, we have the following theorem.

Theorem 2. *Let m and n be two strictly positive integers. We have*

$$\varphi(K_m \square H_n) = \begin{cases} m & \text{if } m \geq 4n \\ m + 1 & \text{if } m = 4n - 1 \\ m + 2 & \text{if } m = 4n - 2 \\ m + 3 & \text{if } m = 4n - 3 \text{ and } n \geq 6 \end{cases}$$

Proof. • $m \geq 4n$.

We have : $m \leq \varphi(K_m \square H_n) \leq m + 3$.

But we have only m vertices of degree greater than m . Thus $\varphi(K_m \square H_n) = m$.

• $m = 4n - 1$.

If $n \geq 8$; $K_m \square H_n$ admit an $(m + 1)$ -coloring given by the following algorithm.

We precise that $m - n - r + 6 + k$ is given modulo m and $m - n - r + 6 + k = m$ if $n + r - 6 - k = m$.

$$c(i, j) = \begin{cases} i & \text{if } (i, j) = (i, 1), 1 \leq i \leq m \\ m+1 & \text{if } (i, j) = (i, i+1), 1 \leq i \leq n \\ 2 & \text{if } (i, j) = (i, i), 3 \leq i \leq n \text{ and} \\ & (i, j) = (1, n+1) \\ 2(1+k) & \text{if } (i, j) = (r+2k, r+k), 3 \leq r \leq n+k-1, \\ & 1 \leq k \leq \frac{n}{2}+2 \\ n+k+3r-2 & \text{if } (i, j) = (2r+k+2, r+2), 1 \leq r \leq n-1, \\ & 0 \leq k \leq 3 \\ 2(r-1) & \text{if } (i, j) = (2r, r-1), 3 \leq r \leq n+2 \\ n+k+3r-3 & \text{if } (i, j) = (2r+k+2, r+2), 1 \leq r \leq n-1, \\ & 5 \leq k \leq m-n-3r+3 \\ 2k+1 & \text{if } (i, j) = (m-n-r+6+k, r+2), \\ & 1 \leq r \leq n-1, 0 \leq k \leq \frac{n}{2} + \frac{3}{2}(r-1) \\ 2k+1 & \text{if } (i, j) = (m-2n+7+k, n+1), \\ & 0 \leq k \leq 2n-7 \text{ or } 2n-5 \leq k \leq 2n-3 \\ 4n-11 & \text{if } (i, j) = (n+1, n+1) \\ m & \text{if } (i, j) = (2, 2) \\ m-1 & \text{if } (i, j) = (3, 2) \\ 1 & \text{if } (i, j) = (4, 2) \\ 3 & \text{if } (i, j) = (5, 2) \\ 5 & \text{if } (i, j) = (7, 2) \\ i-1 & \text{if } (i, j) = (i, 2), i = 8, \dots, m-1 \\ 6 & \text{if } (i, j) = (m, 2) \\ 2j-2+2(\frac{j+1}{2}-k) & \text{if } (i, j) = (j-\frac{j+1}{2}+k, j), 4 \leq j \leq n-4, \\ & 0 \leq k \leq \frac{j+1}{2}-2. \end{cases}$$

It is easy to check that the algorithm gives a proper coloration. Indeed, each column is a copy of the complete graph; in any line, the color of the first vertex can not be repeated and the colors of the second vertex and the last one must be different.

This is a dominating coloring and the set $\{(i, 1), 1 \leq i \leq n; (2r+k+2, r+2), 1 \leq r \leq n-1 \text{ and } 0 \leq k \leq 2; (1, 2); (2, 2); (3, 2)\}$ is a stable dominating system.

If $3 \leq n \leq 7$, it is easy to see that an $(m+1)$ -dominating coloring exist. For example for $n = 3$, we have the following coloring :

1*	12	11	4
2*	4*	12	1
3*	5*	2	12*
4	6*	1	2
5	7	3	11*
6	8	4	5
7	9	5	6
8	11	6	7
9	10	7*	8
10	1	8*	9
11	3	9*	10*

Fig. 5.

As we have only $4n - 1$ vertices of degree greater than $4n - 1$, then $\varphi(K_m \square H_n) = 4n - 1$.

- $m = 4n - 2$: $K_m \square H_n$ admit an $(m + 2)$ -coloring.

If $n \geq 7$, we have the following coloring algorithm. We precise that i is modulo m for $n + 3r + 3 \leq i \leq n + 4r + 2$, $n + 4r + 3 \leq i \leq m + r - 1$ also for $i \geq m + 1$.

$$c(i, j) = \begin{cases} i & \text{if } (i, j) = (i, 1), 1 \leq i \leq m \\ m + 1 & \text{if } (i, j) = (i, i + 1), i \leq n \\ m + 2 & \text{if } (i, j) = (i, i + 2), 1 \leq i \leq n - 1 \text{ and} \\ & (i, j) = (n, 2) \\ n + 3r - k & \text{if } (i, j) = (n + 3r + k - 7, r + 1), 1 \leq r \leq n - 1 \\ & \text{and } 0 \leq k \leq 3 \\ n + k + 3r - 2 & \text{if } (i, j) = (n + 3r + k, r + 2), 1 \leq r \leq n - 2 \\ & \text{and } 0 \leq k \leq 1 \\ n + 3r + k - 8 & \text{if } (i, j) = (n + 3r + k - 7, r + 2), 0 \leq k \leq 2 \\ & \text{and } 1 \leq r \leq n - 1 \\ n + 3r - 5 & \text{if } (i, j) = (n + 3r + 2, r + 2), 1 \leq r \leq n - 2 \end{cases}$$

$$c(i, j) = \begin{cases} i - r - 1 & \text{if } (i, j) = (i, r + 2), 0 \leq r \leq n - 3 \text{ and} \\ & r + 2 \leq i \leq n + 3r - 8 \\ i - r - 11 & \text{if } (i, j) = (i, r + 2), 2 \leq r \leq n - 2 \text{ and} \\ & n + 3r + 3 \leq i \leq n + 4r + 2 \\ i - r + 1 & \text{if } (i, j) = (i, r + 2), 2 \leq r \leq n - 2 \text{ and} \\ & n + 4r + 3 \leq i \leq m + r - 1 \\ i + 1 & \text{if } (i, j) = (i, 3), \text{ and } n + 6 \leq i \leq m - 1 \\ n & \text{if } (i, j) = (n - 1, 2) \\ n - 1 & \text{if } (i, j) = (n + 1, 2) \\ i + 2 & \text{if } (i, j) = (i, 2) \text{ and } n + 2 \leq i \leq m - 2 \\ n - 5 & \text{if } (i, j) = (m - 1, 2) \\ n - 2 & \text{if } (i, j) = (m, 2) \\ n - 6 & \text{if } (i, j) = (m, 3) \\ i - n - 1 & \text{if } (i, j) = (i, n) \text{ and } n + 2 \leq i \leq 2n - 1 \\ i - n + 1 & \text{if } (i, j) = (i, n), 2n \leq i \leq 4n - 14 \\ n - 1 & \text{if } (i, j) = (n, n) \\ n - 3 & \text{if } (i, j) = (n - 4, n + 1) \\ 1 & \text{if } (i, j) = (n - 3, n + 1) \\ n - 4 & \text{if } (i, j) = (n - 2, n + 1) \\ n - 2 & \text{if } (i, j) = (n + 2, n + 1) \\ i + 1 & \text{if } (i, j) = (i, n + 1), 1 \leq i \leq n - 6 \\ 4(n - 3) & \text{if } (i, j) = (n - 5, n + 1) \\ m & \text{if } (i, j) = (n + 1, n + 1) \\ i - 2 & \text{if } (i, j) = (i, n + 1), n + 3 \leq i \leq 4n - 11 \\ i - 1 & \text{if } (i, j) = (i, n + 1), 4n - 7 \leq i \leq m \end{cases}$$

Also here, the graph can not have a b -coloring system with more than $m + 2$ vertices. Thus the equality.

If $3 \leq n \leq 6$, an $(m + 2)$ -dominating coloring is given in Figure 6 for the case $n = 5$.

- $m = 4n - 3$

In this case, $K_m \square H_n$ admit an $(m + 3)$ -coloring if $n \geq 6$.

For $6 \leq n \leq 9$, an $(m + 3)$ -coloring can be found easily. For $n = 6$, such coloring is given in Figure 7

.

If $n \geq 10$, we give the following algorithm, presented in two parts. Note that i is calculated modulo m , for $m - 2 \leq i \leq n - 3$, or $2n + 2r + 8 \leq i \leq m + n - r - 5$ or $i \geq m + 1$.

1*	19*	2	5	4	20
2*	20*	19	18	5	1
3*	5	20	19	18	4
4	18	1	20	19	2
5*	3	18	4	20	19
6	7	17*	8	1	5
7	8	16*	6	2	6
8	6	15*	7	3	7
9	11	10	14*	11	8
10	9	11	13*	9	18
11	10	9	12*	10	12
12	13	14	1	14	11*
13	14	12	2	12	10*
14	12	13	3	13	9*
15	16	3	16	8*	17
16	17	4	17	7*	15
17	15	5	15	6*	16
18	4	8	11	15	3

Fig. 6.

1*	22*	2	23	16	24	3
2*	23*	3	24	17	22	1
3*	24	1	22	18	21	2
4*	5	22	1	23	1	24
5*	6	23	2	24	2	22
6*	4	24	3	22	3	23
7	9	8	21*	9	4	15
8	7	9	20*	7	5	14
9	8	7	19*	8	6	13
10	11	18*	12	15*	11	4
11	12	17*	10	14*	12	5
12	10	16*	11	13*	10	6
13	14	15	4	1	14	7
14	15	13	5	2	15	8
15	13	14	6	3	13	9
16	17	4	13	4	18	12*
17	18	5	14	5	16	11*
18	16	6	15	6	17	10*
20	21	19	17	21	8*	19
21	19	20	18	19	7*	20

Fig. 7.

$$c(i, j) = \begin{cases} i & \text{if } (i, j) = (i, 1), 1 \leq i \leq m \\ m+i & \text{if } (i, j) = (i, 2), 1 \leq i \leq 3 \\ m+1 & \text{if } (i, j) = (k+1, n-k), 1 \leq k \leq n-4, \\ & (n-1, n+1) \text{ and } (n, n) \\ m+2 & \text{if } (i, j) = (k+2, n-k), 1 \leq k \leq n-5, \\ & (n-2, n), (n, n+1) \text{ and } (1, 4) \\ m+3 & \text{if } (i, j) = (k+3, n-k), 1 \leq k \leq n-6, \\ & (n-2, n+1), (n-1, n), (1, 5) \text{ and } (2, 4) \\ n+k & \text{if } (i, j) = (k+3, 2), 1 \leq k \leq n-3 \\ n & \text{if } (i, j) = (n-2, 2) \\ n-2 & \text{if } (i, j) = (n-1, 2) \\ n-1 & \text{if } (i, j) = (n, 2) \\ i-n+6 & \text{if } (i, j) = (i, 2), n+1 \leq i \leq 2n-9 \\ i+3 & \text{if } (i, j) = (i, 2), 2n-8 \leq i \leq m-9 \\ i+6 & \text{if } (i, j) = (i, 2), m-8 \leq i \leq m-6 \\ m-4 & \text{if } (i, j) = (m-5, 2) \\ m-3 & \text{if } (i, j) = (m-4, 2) \\ m-5 & \text{if } (i, j) = (m-3, 2) \\ k & \text{if } (i, j) = (m-3+k, 2), 1 \leq k \leq 3 \\ m+10-n-k & \text{if } (i, j) = (k, 3), 1 \leq k \leq 3 \\ m+3-n+i & \text{if } (i, j) = (i, 3), 7 \leq i \leq n \\ k & \text{if } (i, j) = (n+3+k, 3), 1 \leq k \leq 3 \\ i-3 & \text{if } (i, j) = (i, 3), n+7 \leq i \leq m+9-n \\ n-m-3+i & \text{if } (i, j) = (i, 3), m+10-n \leq i \leq m \\ m+3-i & \text{if } (i, j) = (i, 4), 3 \leq i \leq n-4 \\ i-n-6 & \text{if } (i, j) = (i, 4), n+7 \leq i \leq 2n+3 \\ i+6-n & \text{if } (i, j) = (i, 4), 2n+4 \leq i \leq m \\ k & \text{if } (i, j) = (k+1, 5), n+7 \leq k \leq 2n+3 \\ m-k & \text{if } (i, j) = (n-2+k, 5), 0 \leq k \leq 2 \\ i-15 & \text{if } (i, j) = (i, 5), n+10 \leq i \leq n+15 \\ i-3 & \text{if } (i, j) = (i, 5), n+16 \leq i \leq m \\ i-n & \text{if } (i, j) = (i, n), n+1 \leq i \leq m-12 \\ i-9-n & \text{if } (i, j) = (i, n), m-2 \leq i \leq n-3 \\ k+6 & \text{if } (i, j) = (k, n+1), 1 \leq k \leq 3 \\ 6 & \text{if } (i, j) = (4, n+1) \\ 4 & \text{if } (i, j) = (5, n+1) \\ 5 & \text{if } (i, j) = (6, n+1) \end{cases}$$

$$c(i, j) = \begin{cases} k+9 & \text{if } (i, j) = (6+k, n+1), 1 \leq k \leq n-9 \\ m-2+k & \text{if } (i, j) = (m-14+k, n+1), 0 \leq k \leq 2 \\ k & \text{if } (i, j) = (m-12+k, n+1), 1 \leq k \leq 3 \\ n+6-k & \text{if } (i, j) = (m-5+k, n+1), 0 \leq k \leq 5 \\ i+6 & \text{if } (i, j) = (i, n+1), n+1 \leq i \leq m-15 \\ n+3r-k+6 & \text{if } (i, j) = (n+3r+k-2, r+3), 1 \leq r \leq n-3 \\ & \text{and } 0 \leq k \leq 2 \\ n+3r+k-5 & \text{if } (i, j) = (n+3r+k-4, r+1), 1 \leq r \leq n+3r+k \\ & \text{and } 0 \leq k \leq 1 \\ n+3r-5 & \text{if } (i, j) = (n+3r-3, r+3), 1 \leq r \leq n-2 \\ i-n+r+2 & \text{if } (i, j) = (i, r+5), n-1-r \leq i \leq n+3r \\ & \text{and } 1 \leq r \leq n-6 \\ i+r-n-7 & \text{if } (i, j) = (i, r+5), n+3r+10 \leq i \leq 2n+2r+7 \\ & \text{and } 1 \leq r \leq n-6 \\ i-n+r+5 & \text{if } (i, j) = (i, r+5), 2n+2r+8 \leq i \leq m+n-r-5 \\ & \text{and } 1 \leq r \leq n-6. \end{cases},$$

A dominating system is $\{(i, 1), 1 \leq i \leq n; (n+3r+k-2, r+3), 1 \leq r \leq n-3, 0 \leq k \leq 2; (4, 2); (5, 2); (6, 2); (n-2); (n-1, 3); (n, 3); (m-5, n+1)\}$.

References

1. R.W.Irving, D.F.Manlove, The b -chromatic number of a graph. Discret Applied Math. **91** (1999) 127-141.
2. M.Kouider, M.Mahéo, Some bounds for the b -chromatic number of a graph. Discrete Math. **256** (2002) 267-277.
3. B.Omoomi, R.Javadi, On the b -coloring of cartesian product of graphs. Ars Comb. 107(2012), 521-536.

Gradient Local Binary Patterns for Keyword Spotting in Handwritten Documents

Mohamed Lamine Bouibed, Hassiba Nemmour, and Youcef Chibani

Laboratoire d'Ingénierie des Systèmes Intelligents et Communicants (LISIC),
Faculty of Electronic and Computer Sciences, University of Sciences and Technology
Houari Boumediene (USTHB), Algiers 16111, Algeria
{mbouibed, hnnemmour, ychibani}@usthb.dz

Abstract. In this work, we propose a new descriptor that is called Gradient Local Binary Patterns (GLBP) for automatic keyword spotting in handwritten documents. GLBP is a gradient feature that improves the Histogram of Oriented Gradients (HOG) by calculating the gradient information at transitions of the Local Binary Pattern code. For the matching step, we use the Euclidian Distance and the Cosine Similarity. We evaluate the performance of GLBP comparatively to HOG on a benchmark database of 300 query words extracted from 100 documents written in four languages. The results obtained highlight the effectiveness of the proposed descriptor.

Keywords: Cosine Similarity, Euclidian Distance, GLBP, HOG, Matching, Word Spotting.

1 Introduction

The large amount of information contained in document databases, which store various kinds of historical books and manuscripts, makes their indexing a key component for information retrieval. Firstly, optical character recognition has been employed for indexing documents. However, this approach is useless if documents are degraded or noised. Thereby, currently researchers focus on developing document retrieval systems that are based on the analysis of some words describing the content of the researched document. This approach that is called Keyword spotting aims to identify locations of the query-word in a document image, without performing an explicit recognition stage [1]. Specifically, in keyword spotting problems a model is provided as a query, and the goal is to retrieve all the occurrences in a word images database that are close to this model in terms of a specific similarity measure [2]. Keyword spotting systems can be developed according to segmentation-based or segmentation-free approaches. The first approach proceeds by segmenting the text documents into words and then, compare each of them with the query word. The second approach searches zones in the documents which have similar characteristics as the query word [3]. Mainly, a keyword spotting system is composed of two steps that are the feature generation step and the matching step. For the matching step, training-based classifiers are rarely utilized since they are applicable for a limited word lexicon. Thereby, similarity measure techniques such as Dynamic Time Warping and Euclidian distance are extensively used to process

any number of queries [4]. Furthermore, the literature reports a wide range of feature generation schemes including, shape, texture and gradient features which were proposed for keyword spotting [5]. Specifically, in [6, 7] authors demonstrated that handwritten text characterization can be more successful when calculating features on local regions of the text images. In this respect, Fernandez-Mota et al., [8] report interesting outcomes using Local Binary Patterns (LBP) calculated on a Quad-tree partitioning of word images. On the other hand, Rusinol et al. [9] extracted the Scale Invariant Feature Transform (SIFT) descriptors in multi-scale squared regions. A Mean Average Precision (MAP) of 30.42% was obtained on George Washington which is a historical benchmark dataset. Furthermore, in [10] authors report a more successful performance of SIFT on a less constrained dataset. Precisely, the word spotting was performed on a French dataset that contains 630 documents and a MAP of 70% is reported. In [11], Kovalchuk adopted the combination of Histogram of Oriented Gradients (HOG) and Local Binary Patterns (LBP) to improve the word spotting on Georges Washington dataset. The MAP obtained for word spotting on segmented words reached 50.1%. Note that despite the diversity of the proposed approaches and retrieval protocols, the keyword spotting is a very complicated tasks and MAP values reported on benchmark datasets, still need improvement.

Presently, we propose a new feature generation scheme that aims to improve the gradient characterization. The proposed feature is the Gradient Local Binary Patterns (GLBP). It employs the uniform Local Binary Patterns, to calculate gradient magnitudes. Although the generation of features is the basis on which, we build our word spotting system, matching methods are also important. Here we employ the Euclidian distance and the cosine similarity, which are commonly used to compare documents in text mining [12, 13]. The experimental analysis is carried on a benchmark dataset which was presented in the International Conference on Frontiers on Handwriting Recognition (ICFHR 2014). The rest of this paper is structured as follows: Section 2 introduces the proposed feature. Results are presented and discussed in Section 3. Finally, main conclusions of this work are summarized in Section 4.

2 Proposed system for Word spotting

A key-word spotting system is composed of the feature generation and the matching modules. Similarly, to most of handwriting recognition systems, features are locally extracted by applying a uniform grid over word images. Precisely, the feature generation is separately performed for each cell. Then, resulting features are concatenated to constitute the word features (See figure1).

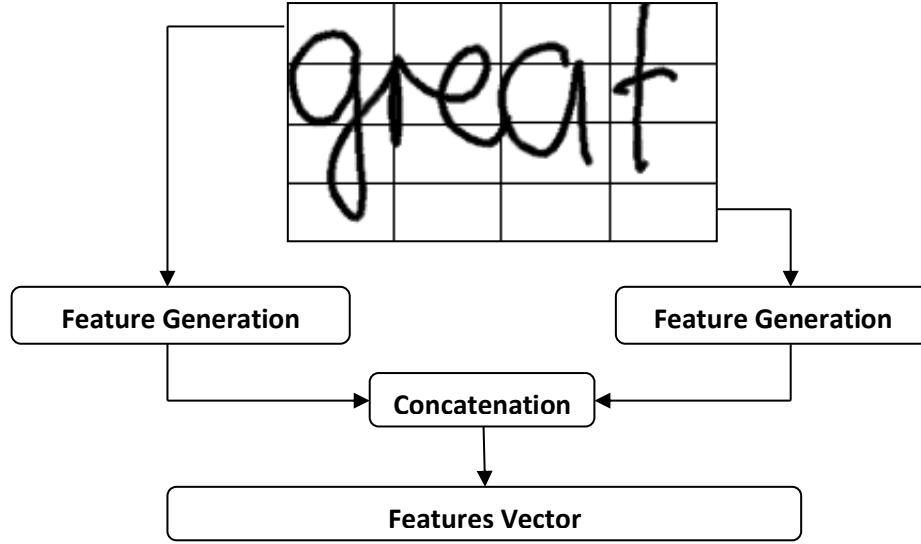


Fig. 1. Local calculation of word features.

2.1 Feature generation

For feature generation, we propose the Gradient Local Binary Patterns (GLBP) which is an improved form of the Histogram of Oriented Gradient (HOG) features. To understand, the contribution of GLBP, we briefly review the concept of the HOG descriptor.

2.1.1 Histogram of Oriented Gradient

The Histogram of Oriented Gradients (HOG) was introduced by Navneet Dalal and Bill Triggs in computer vision for human detection [14]. It has been successfully used in various handwriting recognition applications such as signature verification [15], writer gender prediction [16] and keyword spotting [11]. Roughly, HOG aims to describe the local form of an object through the distribution of the gradient intensity, which highlights the contour direction. For a given cell HOG is calculated according to the following steps:

1. For each central pixel $I(x, y)$ calculate the horizontal and vertical gradients, G_x and G_y .

$$G_x(x, y) = I(x + 1, y) - I(x - 1, y) \quad (1)$$

$$G_y(x, y) = I(x, y + 1) - I(x, y - 1) \quad (2)$$

2. Calculate the gradient magnitude and phase by using equations (3) and (4), respectively.

$$\|G(x, y)\| = \sqrt{G_x^2(x, y) + G_y^2(x, y)} \quad (3)$$

$$\theta(x, y) = \tan^{-1} \left(\frac{G_y(x, y)}{G_x(x, y)} \right) \quad (4)$$

Figure 2 illustrates the calculation of the gradient magnitude and phase for a given pixel.

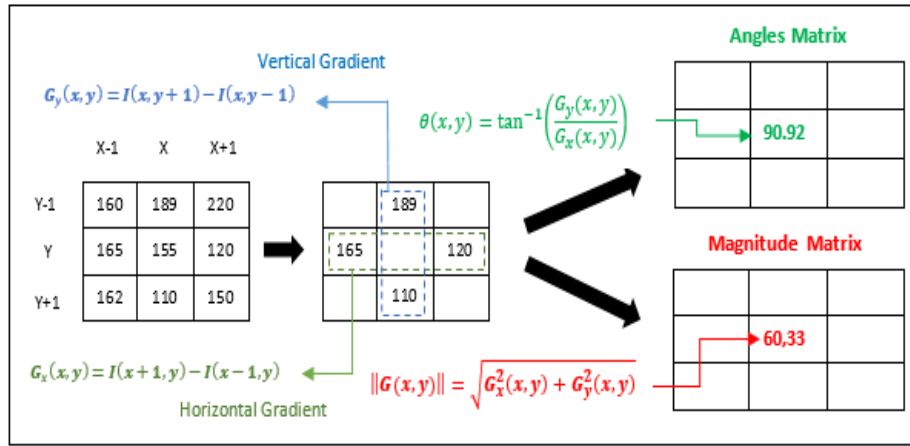


Fig. 2. Calculation of the gradient magnitude and phase for a pixel.

3. Calculation of the histogram: accumulate magnitude values for each range of orientations that are taken in the interval $[0^\circ, 360^\circ]$. Dalal and Triggs have shown that the optimal number of orientations equals 9 [14].

2.1.2 Gradient Local Binary Pattern

In many applications, such as handwriting recognition, the calculation of the HOG by considering orthogonal neighbors doesn't allow an optimal characterization of the edge, which may include some other neighbors (See figure 3). To overcome this limitation, in [15] authors proposed the Gradient Local Binary Patterns which consists of computing the gradient information at the edge position that is not necessarily situated on orthogonal neighbors. Specifically, GLBP employs the first order neighborhood of a given pixel, where edges correspond to transitions from 1 to 0 or from 0 to 1 of the uniform LBP code.

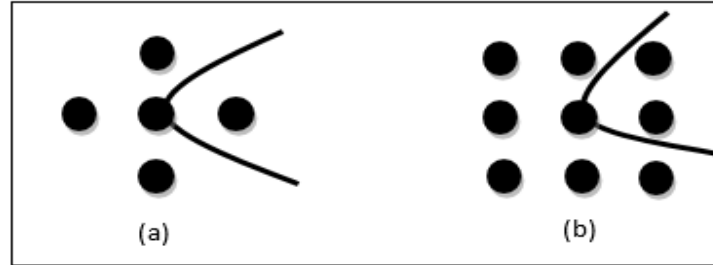


Fig. 3. Neighborhood of gradient calculation (a) HOG, (b) GLBP.

For each pixel, the GLBP feature is obtained according to the following steps:

1. As shown in figure 4, LBP code is obtained by comparing the grey level value of a central pixel to its neighborhood the code takes 1 if the neighboring pixel has higher grey level value. Otherwise it takes 0.
2. Compute width and angle values from the uniform patterns as follow (see figure 4):
 - The width value corresponds to the number of “1” in LBP code.
 - The angle value corresponds to the Freeman direction of the middle pixel within the “1” area in LBP code.
3. Compute the gradient at the 1 to 0 (or 0 to 1) transitions in uniform patterns.
4. Width and angle values define the position within the GLBP matrix, which is filled by the gradient information.

The size of the GLBP matrix is defined by all possible angle and width values. Specifically, there are eight possible Freeman directions for angle values while the number of “1” in uniform patterns can go from 1 to 7 (see figure 4). This yields 7×8 GLBP matrix, in which gradient features are accumulated. Finally, the L2 normalization is applied to scale features in the range $[0, 1]$, [15].

2.2 Similarity measures

A word spotting process calculates the similarity between the query and all reference words. The output corresponds to a vector in which reference words are ranked from the most similar to the less similar one. Presently, the matching process is based on two different measures, that are described in what follows:

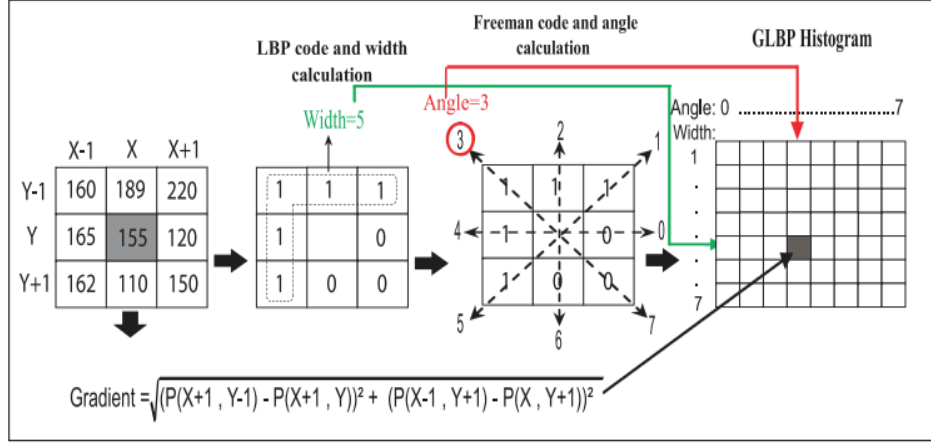


Fig. 4. Summary of the GLBP calculation for a central pixel from [18].

- **Euclidian distance**

For two feature vectors A and B, Euclidian distance corresponds to the square root of the sum of the subtraction of each element of the two vectors as follow:

$$Euc(A, B) = (\sum_{i=1}^n |a_i - b_i|^p)^{\frac{1}{p}} \quad \text{With } p=2 \quad (5)$$

Since it is a distance, it is significant when its value is low.

- **Cosine similarity**

The cosine similarity of two vectors A and B. θ is computed as:

$$Similarity(A, B) = \cos \theta = \frac{A \cdot B}{\|A\| \cdot \|B\|} \quad (6)$$

Contrary to the Euclidean distance, cosine similarity is significant when its value is high.

3 Results

The database employed in this work has been introduced in the International Conference on Frontiers in Handwriting Recognition 2014. It is composed of 100 documents written in four languages that are English, French, German and Greek¹ (see figure 5). There are 25 documents per language and 300 query words proposed to perform the keyword spotting task [19]. The measure employed for performance evaluation of the submitted word spotting algorithms is the Precision at Top k Retrieved words (P@5), which is defined as follows [19].

¹ <http://vc.ee.duth.gr/h-kws2014/#Datasets>

$$P@k = \frac{| \{ \text{relevant words} \} \cap \{ \text{retrieved words} \} |}{| \{ \text{retrieved words} \} |} \quad (7)$$

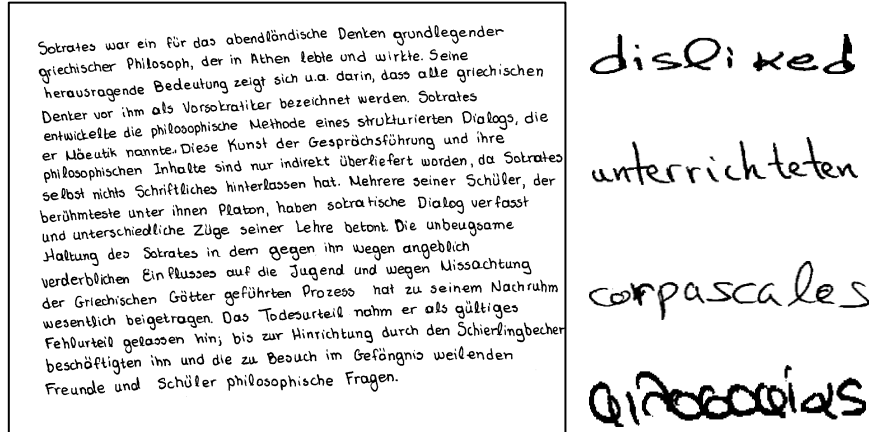


Fig. 5. Text sample from the dataset and some query words

In a first experiment, we tried to find the appropriate number of cells when partitioning words into uniform grids. So, we used a set of 11 word queries each of which is represented by 4 or 5 repetitions. For each grid configuration, feature vector of a given query word is compared to those of all words contained in all documents by using Euclidian distance. Then, among all resulting distances we keep the smallest F distances, where F is the appearance frequency of the query word. Based on the ground truth, we compute the $P@F$ criterion. Finally, we calculate the matching accuracy for each pair (n_x, n_y) and choose the configuration that gives the best accuracy. Figures 6 and 7 depict the mean $P@F$ obtained for both HOG and GLBP.

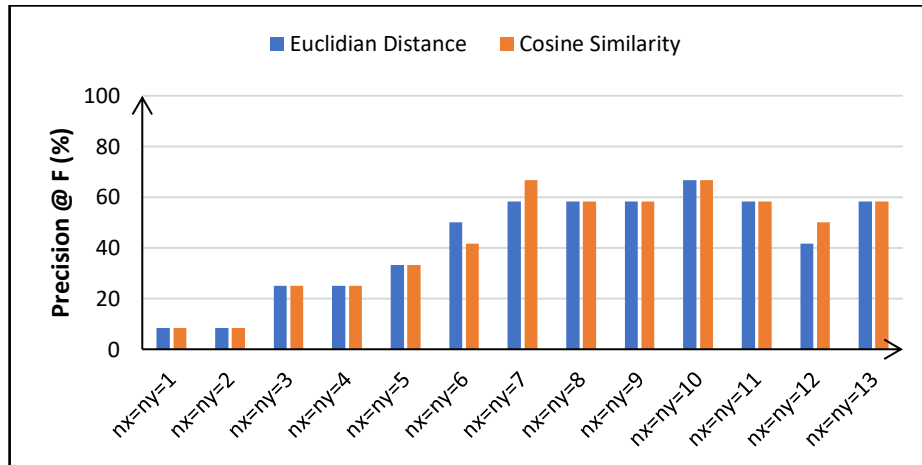


Fig. 6. Precision versus the number of cell for HOG using Cosine Similarity.

It should be noted that for the GLBP the two similarity measures provide their the best accuracy in the same configuration (nx , ny).

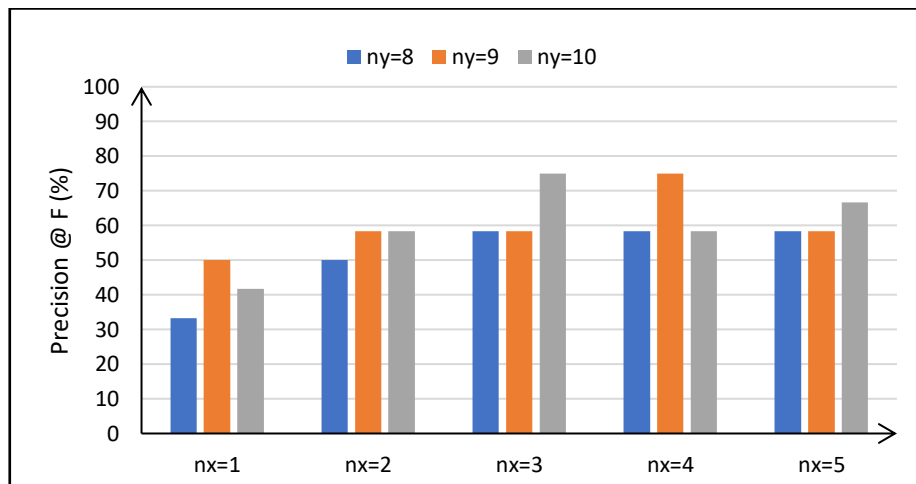


Fig. 7. Precision versus the number of cells for GLBP using Euclidian Distance and Cosine Similarity.

As an optimal configuration, we retain $nx = 3$, $ny = 10$ for GLBP and $nx = ny = 10$ for HOG.

Once the grid configuration fixed, we performed the keyword spotting over the whole dataset. The results obtained for both Euclidian distance and Cosine similarity are reported in tables 1 and 2, respectively. Since most of the query word appear approximately 5 times in the database, we evaluated the precision from rank 1 to rank 5. It is easy to see that with Euclidian distance, HOG and GLBP provide quite similar precisions. Nevertheless, with the Cosine similarity higher scores are obtained. Moreover, the GLBP outperforms HOG with at least 2% in the global precision.

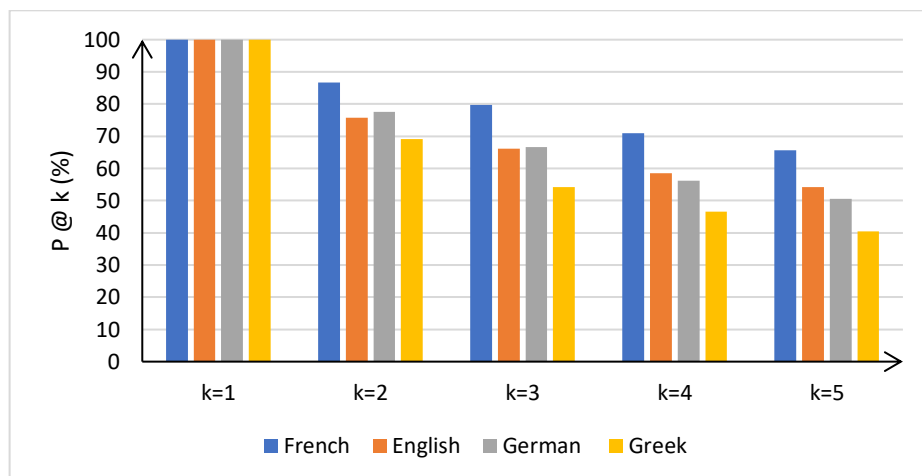
Table. 1. Precision @ k using Euclidian Distance (%).

	P@1	P@2	P@3	P@4	P@5
HOG	100	68,33	54,44	45,08	39,33
GLBP	100	69,67	53,56	45,33	39,07

Table. 2. Precision@ k using Cosine Similarity (%).

	P@1	P@2	P@3	P@4	P@5
HOG	100	68,50	53,11	44,92	39,47
GLBP	100	70,50	57,78	48,50	42,67

For a better inspection of the word spotting system behavior and its sensitivity towards language properties, keyword spotting experiments were replicated separately for each language. For example, a word written in French will be sought in a Corpus only containing documents written in that language. The results obtained by using the GLBP features associated with the Cosine similarity are illustrated in figure 8.

**Fig.8.** Precision (%) @ k using GLBP associated with Cosine Similarity for each language.

From this figure, we remark that accuracies vary from a language to another. Precisely, the best scores are obtained for French and English documents where the P@5 is about 65,71% and 54,19% respectively. The Greek language is the most critical since the P@5 is about 40,47%. This outcome was expected thanks to the high similarity in word samples of this language.

4 Conclusion

Keyword spotting appears to be an attractive alternative to obvious recognition-based retrieval approaches which are conventionally used for document retrieval. In this paper, we proposed a new system for keyword spotting that is based on GLBP features and the Cosine Similarity. Experimental analysis was performed on a benchmark database.

Results show that local features present mixed results that depend on the keyword query. Nevertheless, the GLBP associated with the Cosine Similarity provides satisfactory results. Moreover, the result reveal that the keyword spotting is influenced by the language properties. Particularly, substantial troubles in the precision scores were observed for Greek language which, has a high likeness between characters. Thereby, further experiments using more robust descriptors are necessary to cope with these confusions.

References

1. K. Zagoris, I. Pratikakis and B. Gatos, "Segmentation-based Historical Handwritten Word Spotting using Document-Specific Local Features" in International Conference on Frontiers in Handwriting Recognition, ICFHR.2014.10, pp 9-14. IEEE 2014.
2. Fernandez,D, (2010), Handwritten word spotting in old manuscript images using Shape Descriptors, Master thesis in Computer Vision and artificial intelligence, Autonomous University of Barcelona, 38 pages.
3. Blessy Varghese et al, "A Survey on Various Word Spotting Techniques for Content Based Document Image Retrieval", (IJCSIT) International Journal of Computer Science and Information Technologies, Vol. 6 (3), 2015, 2682-2686 Lehner, W. (eds.) Euro-Par 2006. LNCS, vol. 4128, pp. 1148--1158. Springer, Heidelberg (2006).
4. S. Yao, Y. Wen and Y. Lu, "HoG based Two-Directional Dynamic Time Warping for Handwritten Word Spotting", In 13th International Conference on Document Analysis and Recognition (ICDAR)2015. Pp 161-165. IEEE 2015
5. J. A. Rodriguez, F. Perronnin, "Local gradient histogram features for word spotting in unconstrained handwritten documents", In International Conference on Frontiers in Handwriting Recognition (ICFHR). 2008
6. Bouadjenek, N., Nemmour, H., Chibani, Y., 2015. Robust soft-biometrics prediction from off-line handwriting analysis, Applied Soft Computing Journal, doi: 10.1016/j.asoc.2015.10.021
7. Bertolini, D., Oliveira, L., Justino, E., Sabourin, R., 2010. Reducing forgeries in writer-independent off-line signature verification through ensemble of classifiers, Pattern Recognition Journal, Vol. 43, pp.387-396.
8. Fernandez-Mota, D., Almazan, J., Cirera, N., Fornes, A., Lladós, J.: BH2M: the barcelona historical, handwritten marriages database. In: 2014 22nd International Conference on Pattern Recognition (ICPR), pp. 256–261. IEEE (2014)
9. M. Rusinol, D. Aldavert, R. Toledo, and I. Lladós, "Browsing heterogeneous document collections by a segmentation-free word spotting method," in International Conference on Document Analysis and Recognition, pp. 63-67, 2011
10. J. Rodriguez and F. Perronnin, "Local gradient histogram features for word spotting in unconstrained handwritten documents," in International Conference on Frontiers in Handwriting Recognition, pp. 7-12, 2008.

11. A. Kovalchuk, L. Wolf, and N. Dershowitz, "A simple and fast word spotting method," in International Conference on Frontiers in Handwriting Recognition, pp. 3-8, 2014.
12. S. Tollari, (2006), Image indexing and retrieval by combining textual and visual informations, Université du Sud Toulon-Var.
13. W. Dong, Z. Wang, M. Charikar and K. Li, "Efficiently Matching Sets of Features with Random Histograms" of the 16th ACM international conference on Multimedia pp 179-188, 2008.
14. Dalal, N. and Triggs, B., Finding People in Images and Videos, PhD thesis, French National Institute for Research in Computer Science and Control (INRIA), July 2006
15. M. B. Yilmaz, B. Yanikoglu, C. Tirkaz and A. Kholmatov, "Offline signature verification using classifier combination of HOG and LBP features," 2011 International Joint Conference on Biometrics (IJCB), Washington, DC, 2011, pp. 1-7.
16. N. Bouadjenek, H. Nemmour and Y. Chibani, "Histogram of Oriented Gradients for writer's gender, handedness and age prediction," 2015 International Symposium on Innovations in Intelligent Systems and Applications (INISTA), Madrid, 2015, pp. 1-5
17. N. Jiang, J. Xu, W. Yu, and S. Goto, " Gradient local binary patterns for human detection," in Circuits and Systems (ISCAS), 2013 IEEE International Symposium on, pp. 978-981, May 2013.
18. N. Bouadjenek, H. Nemmour, Y. Chibani, "Age, Gender and Handedness Prediction from Handwriting using Gradient Features", 2015 13th International Conference on Document Analysis and Recognition (ICDAR).
19. Pratikakis, I., et al. (2014). "ICFHR 2014 Competition on Handwritten Keyword Spotting (H-KWS 2014)." 814-819.

A comparative analysis of context-management approaches for the Internet of Things

Farida Retima ¹(0000-0002-9530-7667), Saber Benharzallah ²(0000-0002-3870-438), Laid Kahloul ³(0000-0002-9739-7715), Okba Kazar ⁴(0000-0003-0522 4954)
^{1,2,3,4}Smart Computer Sciences Laboratory, Biskra University 07000, Algeria

²Batna 2 University Computer sciences department, Algeria.

¹retimafarida@yahoo.fr, ²sbharz@yahoo.fr,
³laid.k.b@gmail.com, ⁴kazarokba@yahoo.fr.

Abstract. The Internet of Things (IoT) has gained much attention during the last decade. A novel aspect like context management is fundamental requirement for development of such systems. In literature, there are different approaches enabling the context management for internet of things (IoT). This paper studies these approaches, which are well known according to our defined criteria such as heterogeneity, mobility, the influence of the physical world, scalability, security, privacy, quality of context, autonomous deployment of entities, characterization multi scales, and interoperability.

Keywords: Context management, Internet of things, Context-awareness, Context manager, middleware.

1 Introduction

The increasing number of mobile objects used in our daily life has given birth for a number of new paradigms such as the Internet of Things (IoT), which relies on a wide range of hardware, network infrastructure, communication protocols and applications. This concept IoT aims to allow products or daily objects to interoperate with a large number of systems using Internet. The services offered by this paradigm must be smart and adapted or sensitive to context changes. The context management aims to link applications which have heterogeneous needs with heterogeneous capabilities of sensors, thus it must ensure compatibility between the different actors. This concept is a generally treated through ambient environments. In the Internet of things, context management becomes more complex. This is due to the heterogeneity of the data, the authorization of the mobility of architectural elements, and the provision of information on the context of the user, who exposes the privacy of individuals near to the objects connected to the Internet [1]. Therefore, a context management approach to the Internet of Things should be developed, where it must take into account the requirements of such an environment. The rest of this paper is organized as follows: In Section 2, we introduce a software entity

responsible for the management and presentation of context information to applications; this software entity is called a context manager. In section 3, we describe some technical and non-technical issues which need to be studied to enable the Internet of things to reach its full potential. In Section 4, we will do a comparison of these approaches according to our defined criteria with their sub criteria well illustrated in Table 1. Finally, section 6 concludes this paper.

2 Context manager in ambient environment

A context manager is a software entity capable of calculating high-level information from various sources of context data. Its features include the acquisition of raw data, the processing of such data (fusion, aggregation, interpretation, inference) and then the presentation of the highest level of information to context-aware applications. These applications are typically viewed as end consumers of context information, as entities that provide the raw data are considered as original producers. The component which, in a context manager, is responsible for the context data processing is a context processor that is both a consumer and a producer, intermediate level, of context data [2]. The context information management infrastructure aims to link applications that have heterogeneous needs with sensors that have heterogeneous capabilities; it must ensure compatibility between these different actors [3].

3 The difficulties emerged with the internet of things

When talking about the Internet of Things we distinguish numerous challenges currently posed in the literature that must be studied, among these issues we are interested to the following ones: heterogeneity, mobility, and influence of the physical world, scalability, security, privacy, quality of context, autonomous deployment of entities, characterization multi scales, and interoperability.

3.1 Heterogeneity of the Internet of Things

The Internet of Things is composed of heterogeneous objects that are composed of several levels: hardware, software and communication protocols. The equipment used is very varied, functioning with various operating systems, adopting different wireless technologies. Within the same network, the workstations with high computing capabilities and large storage can coexist with other devices with limited resources [4]. Heterogeneity criterion can be divided into two sub criteria:

- Semantic heterogeneity of objects and devices of an ambient system.
- Technological heterogeneity [5].

3.2 Influence of the physical world on the internet of things

Objects connected to the Internet are much more than simple connectivity or a message transported. The Internet of things is open, means that anyone can add and removed objects [6]. The detection capability of smartphones and user mobility [7].

influence on the physical world on the Internet of things that require for developers to take into account a mechanism for data representation and processing. We can divide this criterion into three sub criteria:

- Specific techniques of detection, correction or error attenuation: for flow data (flow measurements or events) produced by sensors.
- The capacity of mobile objects to adapt to changes
- The capacity of mobile objects to handle their energy limits [8].

3.3 Security

The IoT is extremely vulnerable to attacks for several reasons. First, its components often spend most of the time without supervision; and consequently, it is easy to attack them physically, second, most of the communication is wireless, which makes listening extremely simple. Finally, most of the IoT components are characterized by low capacity in terms of energy and computing resources (which is especially the case for passive components) and so, they cannot implement complex Systems for security support. Specifically, the main security problems are authentication and data integrity [9]. We can divide this criterion into two sub criteria:

- Secure communication for peers
- Secure topology: that handles authentication of the new peer, the access authorizations to the network, and protection of routing information exchanged in the network [9].

3.4 Protection of privacy

The growing number of interactions between entities becoming closer to the user pushes to reconsider personal data security. Indeed, users transmit authentication data to many systems, often through unsecured communication channels. And because applications or services are personalized for users, the transmitted data is more sensitive than ever, in the sense that they can lead to a disclosure of their privacy [10], so to promote the acceptability by users of new applications, it is essential to provide a mechanisms to ensure respect for the privacy of users and the protection of the data handled [11]. We can divide this criterion into two sub criteria:

- The hardware layer: should ensure confidentiality during temporary collection and storage in the device.
- The protocol layer: ensure that communication is well protected.
- The application layer: monitoring and control who can see or use the context [12].

3.5 Mobility

A coherent and stable vision of the network topology is an important factor for the management of IoT applications. However, mobility related to everything makes this coherency difficult and inefficient [13]. because the smart objects in IoT are not static and change their physical location this leads to rapid changes in network

topology which in its turn affects the routing efficiency because there are many paths that become unavailable. Update routing information after the mobility of objects can cause major general costs especially when there are many nodes that have changed their physical locations simultaneously [14].

3.6 Quality of context

The QoC enables to qualify the context information using criteria such as accuracy, completeness, freshness, allows a context aware application to adapt its reactions according to QoC delivered by the Manager [15]. We can divide this criterion into three sub criteria:

- QoC of the context information acquisition phase.
- QoC of the context information of the processing phase.
- QoC of the delivery phase of contextual information to context-aware applications [16].

3.7 Scalability

IoT is composed of a growing number of smart objects these objects communicate together. The proposed middleware should maintain its stability and should have the ability to effectively manage the increasing number of objects and perform their functions with the required efficiency, although the number of connected devices is increasing and varies from place to other [14]. The scalability criterion facing in various sub criteria, including:

- Naming and addressing: because of the large scale and rapid growth of IoT devices.
- Data communication and networks: due to the high level of interconnection between many entities.
- The management of information and knowledge.
- Service management: because of the massive number of services that might be available and the need to manage heterogeneous resources [17].

3.8 Multiscale characterisation

A characterization of scenarios multi-scale aims to target specifically the design requirements of context manager, including the requirements for multi-scale. The multi-scale feature of a system can be considered according to several views as equipment, network, data ... etc. For each view, we select one or more dimensions, and for each dimension we propose a set of scales relevant to the multi-scale characterization of the studied system and therefore each system can lead to special multi-scale characterization [18].

3.9 Standalone deployment of entities

The software deployment is a complex process which aims the provision and the operational maintenance of the software. The deployment involves two categories of sites: a producer site that hosts software components and installation procedures, and a consumer site that is the target of the deployment, on which a software

element is to be executed. The deployment of a system can be performed on several consumer sites. All of these sites are deployment domains requiring to write the distribution of components, which is called a deployment plan [19].

3.10 Interoperability

It is the ability of two or more systems or components to exchange information and to use the information that has been exchanged. Two possible types of interoperability can be achieved. On the one hand, inner components can interact with each other and share information. On the other hand, different context aware middleware can communicate with each other and make use of exchanged information [20].

3.11 Context life cycle

- Context Acquisition: the aim of context acquisition is to obtain a maximum amount of data, such that the possibilities for applications to be intelligent could be maximized due to richer context information [20]. We discuss three factors: Acquisition technique, data source support (any sensors, physical sensors, virtual sensors, or logic sensors), context discovery (There are many types of context that can be used to enrich sensor data)[12]
- Context modeling: referred to as context representation. There are several popular context modeling techniques used in context-aware computing. We discuss in this phase the factor modeling techniques [12].
- Context reasoning: a process of giving high-level context deductions from a set of contexts. We discuss two factors: reasoning techniques, quality of context [12].
- Context Distribution: Context distribution is responsible for disseminating useful context information to corresponding applications. We discuss in this phase the factor distribution techniques [20].

4 The comparison of the different approaches proposed for context management in IoT

In the literature, different approaches exist to allow the context management in the Internet of Things. The study of different approaches and their characteristics according to different issues is presented in this section. The considered criteria include: Heterogeneity, influence of the physical world, the security, protection of privacy, mobility, scalability, Multi-scale characterization, Standalone deployment of entities, Interoperability, Quality of context.

Table1.Summary table of evaluation criteria

		Context management approaches for <u>IoT</u>								
The comparison criteria	Sub criteria	INCOME Approach [21]	SITAC Approach [22]	FIWARE Approach [23]	CA4IOT Approach [24]	SAMURAI Approach [25]	HYDRA Approach [26]	AURA Approach [27]	SAI Approach [28]	WISemId Approach [29]
Quality of context	QoC of acquisition phase	✓	-	-	-	✓	-	-	-	-
	QoC of transformation Phase	✓	-	-	-	✓	-	-	-	-
	QoC of delivery phase	✓	-	-	-	✓	-	-	-	-
Influence of the physical world	detection, correction, error attenuation techniques	✓	-	-	-	-	-	??	-	-
	objects mobiles adaptation	✓	✓	✓	✓	-	✓	✓	✓	✓
	energy limitation treatment	-	✓	-	-	-	-	-	-	✓
Heterogeneity of objects	semantic heterogeneity	✓	✓	✓	✓	✓	✓	✓	✓	✓
	technological heterogeneity	✓	✓	✓	✓	✓	✓	✓	✓	✓
Security	Peers communication	✓	✓	✓	-	✓	✓	??	✓	??
	Topology management	✓	✓	✓	-	✓	✓	??	✓	??
Privacy Protection	Hardware level	✓	??	✓	-	-	✓	✓	✓	??
	Protocol level	✓	??	✓	-	-	✓	✓	✓	??
	Application level	✓	??	✓	-	-	✓	✓	✓	??
Standalone deployment of entities		✓	-	-	-	✓	-	-	-	-

Context modeling	<u>Modeling Techniques</u>	Ontology based(context information) <u>QoCIM(QoS)</u>	??	Ontology	Ontology	<u>ML Context model</u>	Ontology	??	XML data model	WIOP message format
Context reasoning	Reasoning Techniques	Ontology based	??	ontology	Ontology and statistical	??	Ontology	??	??	??
	Quality of context	Conflict resolution & Validation	??	-	-	??	-	??	-	-
Context Distribution	Distribution Techniques	Publish and Subscribe	??	Publish and Subscribe	Query	Publish and Subscribe	Query	Query	Publish and Subscribe	Query
Interoperability	Internal	✓	✓	✓	✓	✓	✓	✓	✓	✓
	External	✓	✓	✓	✓	✓	✓	✓	✓	✓
mobility of objects		✓	✓	✓	✓	✓	✓	✓	✓	✓
Design Method		Model-Driven Engineering(<u>MuSCa</u> , <u>QoCIM</u>)+ Distributed Event-Based System(DEBS)	Semantic interoperability	Based on Generic Enablers	learning by understanding, maintaining context data in knowledge bases; follows a layered architecture	Multi-Tenant event-based streaming architecture	SOA-based Middleware + <u>semantic-based Model-Driven Architecture</u>	architecture-specific connectors	Based on distributed grid architecture, Follows service oriented architecture (SOA)	<u>Integration</u> <u>WSNs</u> and the Internet at service level(<u>WSN</u> and Internet nodes)
Tools of implementation		EMF(Eclipse Modeling Framework)	Social-media based <u>platform</u>	Open stack	<u>OpenIoT</u> software Project	<u>Espen</u> (Java POJOs), <u>GeoSPARQL</u> , <u>RESTful</u> APIs	<u>SourceForge</u> + <u>OSGi platform</u>	<u>Wrapped</u> <u>Microservices</u> Word and Emacs, Media Player and <u>Xanadu</u> Code or <u>AFS</u> (distributed file systems)	JOTM(an open and standalone implementation of the JTA (Java Transaction API), <u>ActiveMQ</u>	<u>WISemantic</u> Interface <u>Definition Language</u> , <u>WISemantic</u> Inter-ORB Protocol (WIOP), <u>programming languages</u> (Java, <u>mesC</u>).

✓ Problem solved

- Unsolved problem

?? Not cited in detail on the project

5 Discussion

This paper has presented a review on the latest prominent solutions for context management of IoT which are based on decoupling between the participants of the IoT, namely the context information owners and consumers of such context information. Currently, in most proposed solutions, the functionalities of middleware are achieved by distributing the tasks in a layered/distributed architecture among these solutions are INCOME approach, SITAC approach, FIWARE approach, CAIoT approach, SAMURAI approach, HYDRA approach, AURA approach, SAI approach, WISEMid approach. Only four middleware architectures take security and privacy into account by applying different strategies e.g. INCOME, FIWARE, HYDRA, SAI projects. The criterion of scalability is focused on the stability of functioning with the large number of entities, in this light; cloud-oriented techniques are an excellent solution [20] to solve the scalability problem from different aspects all the approaches that we studied are achieved this issue. It becomes possible to guarantee good applications functioning taking account management of the QoC and so more trust between participants of IoT as the case of INCOME and SAMURAI projects which achieved this issue by proposing software architecture to enforce QoC contracts between producers and consumers in the IoT. Each middleware solution focuses on different aspects in IoT such as interoperability, Multiscale characterization, Standalone deployment of entities and many more. In conclusion, the major goal is to design a solution to help users to automating the task of selecting the sensors according to the tasks at hand, such as CAIoT and WISEMid approaches which are security, privacy and QoC are missing in this two middleware. Several technical challenges have been detected in the development of future middleware architectures and more efforts should be taken into account to drive the current middleware proposals towards better ones with high level of context awareness by combining the architecture, models, and techniques into existing IoT middleware solutions, which intend to fulfill the demands in IoT paradigm. A generic domain-focused middleware solution should take into account the protecting privacy of information, ensuring the information security during processing or probably creating private clouds dedicated to public sectors, QoC resolution and realization a middleware that could adequately understand any change of current environment based on all available context information.

From this comparison of the presented approaches that has exposed the major advantages and disadvantages of each work, we deduce that a Model-Driven Approach is more suitable for the development of context-aware applications and context management system as such it facilitates developers work at design time, by providing specialized modeling languages, metamodels and code generation tools [30]. MDE (Model Driven Engineering) provides a high level of abstraction and it make available for use tools and grammars allowing the construction of context models which may be used to model context information and can be automatically transformed models to particular target platforms[31]. The modeling language used in the majority of cases is the Unified Modeling Language (UML) [32, 33]. As an

illustration of this model-driven approach, there is a scale awareness framework, called MuSCa; this framework provides a multiscale taxonomy to describe at design time the multiscale nature of complex distributed systems which includes a characterization process that can be represented in the form of a MuSCa model conforming to the MuSCa meta-model. The concepts of viewpoints, dimensions and scales constitute the core of a dedicated metamodel. For each viewpoint (Geography, User, Device and Network), a restricted view of the system is considered. Then, for this view, one or several dimensions are chosen. To identify scales for a given dimension, each dimension is associated with a numeric or a semantic measure. Depending on the type of the measure, the resulting scales match, respectively, orders of magnitude for numeric measures, or sets of elements that share common semantic characteristics for semantic measures [30].

6 Conclusion

In this paper, we have demonstrated that the new challenges emerged with IoT in terms of large-scale, heterogeneous objects, the scalability, interoperability, security, and privacy and other problems already mentioned previously require to design a software solution for multi-scale context management, from production context information until used by an application. This paper has given an overview of some approaches that exist in the literature such as the INCOME, SITAC, FIWARE, CA4IoT, SAMURAI, HYDRA, AURA, SAI, WISemid projects and make study in different well-focused criteria so it is our belief that the focus of the future work for context management in IoT should be use the Model Driven Engineering approach and proposing a model to represent the quality of context information that comes from the multiple kinds of sources information for fulfilling an QoC aware application request. Any change in the quality of context information provided by sensor can be easily made at the model level and propagated automatically to the implementation. Furthermore, the time and effort of context management development can be reduced. The goal of MDE is to ensure interactions with context management and applications.

References

1. Sam, R., Sébastien, L., Claire, L., Chantal, T: Vers une définition d'un système réparti multi-échelle. In: 8èmes journées francophones Mobilité et Ubiquité, pp.178--183, France (2012).
2. Jean-Paul, A., Sophie, C., Denis, C.: Gestion de contexte multi échelle pour l'Internet des objets. In: Journées francophones Mobilité et Ubiquité, 2014 (France).
3. Jérôme, P.: Une infrastructure de gestion de l'information de contexte pour l'intelligence ambiante. Thèse doctoral à l'Université Joseph Fourier - Grenoble 1, 2009.
4. Gaëtan, R.: Modélisation du contexte et Adaptation ,2010.
5. Peter, W.: Learning IoTs, Packt Publishing Ltd, Mumbai (2015).
6. Song, G., Xiaofei, L., Fangming, L., Yanmin, Z. : Collaborative Computing: Networking, Applications, and Worksharing. In : 11th International Conference CollaborateCom, Chine (2015).doi: 10.1007/978-3-319-28910-6.

7. Benjamin, B: Systeme de gestion de flux pour l'Internet des objets intelligents. Thèse doctoral, Université deVersailles-Saint Quentin, 2015 (France).
8. Luigi, A., Antonio, I., Giacomo, M. :The IoTs:A survey. J. Computer Networks. 54, 278-2805 (2010).doi: 10.1016/j.comnet.2010.05.010
9. Soma, B.: Role of middleware for IoTs: a study. J. Computer Science & Engineering Survey (IJCSSES),2, 94--105 (2011).doi: 10.5121/ijcses.2011.2307
10. Ahmed Malek, N.: L'intelligence ambiante et les systèmes de transport intelligents. Mémoire du diplôme magister en informatique, Université Badji Mokhtar , 2014(Annaba Algeria).
11. Zied, A.: Gestion de la qualite de contexte pour l'intelligence ambient. Thèse doctoral de Télécom, Université d'Evry-Val -D'Essonne Sud Paris, 2012.
12. Charith, P., ArkadyZ., Peter, C., Dimitrios, G.: Context Aware Computing for The IoTs:A Survey , J. IEEE Communications Surveys & Tutorials. 16, 414 – 454 (2014).
13. Jun young, K.: Secure and Efficient Management Architecture for the IoTs. In: proceedings of the 13th ACM Conference on Embedded Networked Sensor Systems, School of Computer Science and Engineering, UNSW, pp. 499--500, Australia (2015).doi: 10.1145/2809695.2822522
14. Fersi, G.: Middleware for IoTs: a study. In: IEEE International Conference Distributed Computing in Sensor Systems (DCOSS), Tunisia (2015).doi: 10.1109/DCOSS.2015.43.
15. Jean-Paul, A.: Projet INCOME : Infrastructure de gestion de COntexte Multi-Echelle pour l'Internet des Objets. In : Conférence Francophone sur les Architectures Logicielles (CAL), pp. 1--2, Toulouse (2014).
16. Pierrick, M.: Gestion de bout en bout de la Qualité de Contexte pour l'Internet des Objets: le cadriciel QoCIM", Thèse doctoral; l'Université Toulouse, 2015.
17. Daniele, M., Sabrina, S., fransisco, P.: IoTs: Vision, applications and research challenges. J. Ad Hoc Networks journal. 10, 1497-1516 (2010).doi: 10.1016/j.adhoc.2012.02.016
18. Sam, R.: Modèles, méthodes et outils pour les systèmes répartis multiéchelle. Thèse doctorat, Télécom SudParis école doctorale S&I avec l'Université d'Évry-Val-d'Essonne, 2015.
19. Raja, B.: Déploiement de systèmes répartis multi-échelles processus, langage et outils intergiciels. Thèse Doctoral de l'université de Toulouse, 2015.
20. Xin, L.,Martina, E., José, M.: Context Aware Middleware Architectures: Survey and Challenges. J. Sensors journal.15, 20570--20607(2015).doi: 10.3390/s150820570
21. Samer, Ma.M.: Models and algorithms for managing quality of context and respect for privacy in the Internet of Things. Thèse de doctorat de l'université Paris-Saclay, Telecom SudParis (2015).
22. Celestino, M., Manuel, O., Joaquim, B., Tipu, Ra., Jonathan, R.: Social Networks and Internet of Things, an Overview of the SITAC Project. In: Mumtaz, S et al.(Eds) WICON 2014. LNICST, vol.146, pp. 191—196. Springer, Portugal (2015).doi: 10.1007/978-3-319-18802-7_27
23. Martínez, R., Darío.: Internet of things virtualization and identity management via fiware generic enablers. Projet final, E.T.S.I. Télécommunications (UPM). Madrid (2015).
24. Charith, P., Arkady, Z., Peter Ch., Dimitrios, G.: CA4IOT: Context Awareness for Internet of Things. In : Proceedings of the IEEE International Conference on Green Computing and Communications, Conference on Internet of Things, and Conference on Cyber,Physical and Social Computing (iThings/CPSCoM/GreenCom), pp. 775—782. IEEE Computer Society Washington, USA (2012): doi: 10.1109/GreenCom.2012.128
25. Davy, P., Yolande, B.: SAMURAI: A Streaming Multi-tenant Context-Management Architecture for Intelligent and Scalable Internet of Things Applications. In : IE '14

- Proceedings of the 2014 International Conference on Intelligent Environments, pp. 226-233. IEEE Computer Society Washington, USA (2014).doi: 10.1109/IE.2014.43
26. Tomas, S., Jan, H.: Semantic middleware in Applications of the Internet of Things. In: Horváth.T et al. (Eds.) Znalosti 2013, pp. 24—32. IEEE, Slovakia (2010).doi: 10.1109/IOT.2010.5678448
 27. Joio, P.S., David, G.: Aura: An Architectural Framework for User Mobility in Ubiquitous Computing Environments. In : WICSA 3 Proceedings of the IFIP 17th World Computer Congress - TC2 Stream / 3rd IEEE/IFIP Conference on Software Architecture: System Design, Development and Maintenance,pp. 29-43.Spring US, New York(2002).doi: 10.1007/978-0-387-35607-5_2
 28. David, P., Federica, P., Dino, G., Agostino, L.: A Scalable Grid and Service-Oriented Middleware for Distributed Heterogeneous Data and System Integration in Context-Awareness Oriented Domains. In : Daniel.G(Eds), pp. 109-118. Springer, New York (2010).doi: 10.1007/978-1-4419-1674-7_11
 29. Shirin, Z.: Middleware for internet of things. Master Thesis, Faculty of electrical engineering, Mathematics and computer science software engineering. Pays-Bas (2013).
 30. Sam, R., Sébastien, L., Chantal, T., Claire, L., Thierry, D.: MuSCa: A Multiscale Characterization Framework for Complex Distributed Systems. In: Proceedings of the 2014 Federated Conference on Computer Science and Information Systems, pp. 1657--1665, Toulouse (2014). doi: 10.15439/2014F131
 31. Chantal, T., Zakia, K.: Building Context-Awareness Models for Mobile Applications. J. Digital Information Management. 8, 78 –87 (2010).
 32. Georgia, M. K., George, N. P., Nikolaos, D. T., Iakovos, S. V.: Context-aware service engineering: A survey. J. Systems and Software. 82, 1285—1297 (2009).doi: 10.1016/j.jss.2009.02.026.
 33. Vijay, B.: Survey of context information fusion for ubiquitous Internet-of-Things (IoT) systems. In: van den, B, Egon. (Eds.), Open Comput. Sci. 2016. vol. 6, pp. 2299--1093. De Gruyter Open, India (2016).doi: 10.1515/comp-2016-0003.

Avec le soutien de

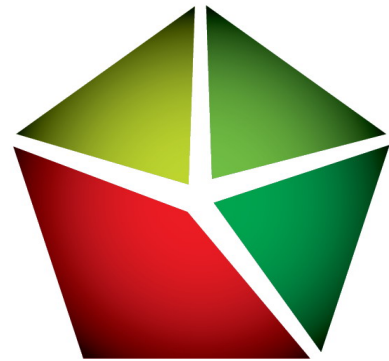
Université Akli Mohand Oulhadj (Bouira)

&

**Ministère de l'Enseignement Supérieur et de la Recherche
Scientifique**



Sonelgaz



Agence thématique de
recherche en sciences
et technologie



Office des publications
universitaires (OPU)



SARL OULED ELHADJ RABEH
ENTREPRISE DES GRANDS TRAVAUX
PUBLICS ET HYDRAULIQUES ET BATIMENTS