

APIA & CNIA & IC & JFPDA & RJCIA

IA pour l'éducation IA & Santé TAL & IA Ethique & IA
Robotique & IA France@IJCAI2018

PFIA 2018

**11^e Plate-forme
Intelligence Artificielle**

2 au 6 juillet 2018 - Nancy

Campus Sciences - Université de Lorraine
Vandœuvre-lès-Nancy

CNIA & RJCIA 2018

ACTES

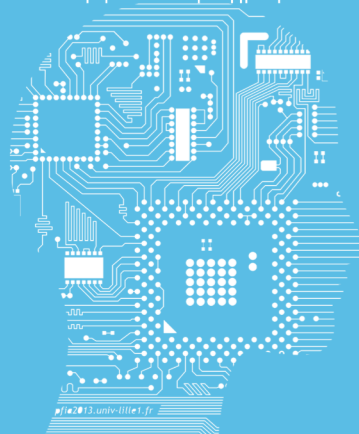
&

de la
**Conférence Nationale
en
Intelligence Artificielle**

Président du comité de programme
Jérôme Euzenat

et des
**Rencontres
des Jeunes Chercheurs
en Intelligence Artificielle**

Président du comité de programme
François Schwarzentruher



PFIA2018.LORIA.FR



Actes de la

Conférence Nationale d'Intelligence Artificielle

et des

Rencontres Jeunes Chercheurs en Intelligence Artificielle

se tenant à Nancy sur la

Plate-Forme Intelligence Artificielle

des 4 au 6 juin

2018

coordonnés par Jérôme Euzenat & François Schwarzentruher
(présidents des comités de programme)

organisés sous la direction d'Armelle Brun & Davy Monticolo
(présidents du comité d'organisation)

Introduction

Pour l'année 2018, la Conférence Nationale d'Intelligence Artificielle et les Rencontres des Jeunes Chercheurs en Intelligence Artificielle sont organisées conjointement. Elles ont lieu du 4 au 6 juillet 2018 à Nancy sur la Plate-Forme Intelligence Artificielle (PFIA). Lors de la même semaine, la plate-forme accueille la 4e Conférence Nationale sur les Applications pratiques de l'Intelligence Artificielle (APIA), les 29es journées francophones d'Ingénierie des Connaissances (IC) et les 13es Journées Francophones sur la Planification, la Décision et l'Apprentissage pour la conduite de systèmes (JFPDA).

La Conférence Nationale en Intelligence Artificielle (CNIA) s'adresse à l'ensemble de la recherche en Intelligence Artificielle (IA). Elle est l'occasion de témoigner des dernières avancées en IA et de présenter ses résultats les plus récents dans tous les aspects de l'IA. CNIA est le prolongement spécifiquement IA de la conférence originelle Reconnaissance des Formes et Intelligence Artificielle (RFIA). CNIA 2018 fait suite à l'édition 2016, organisée à Clermont-Ferrand dans le cadre de RFIA. En sus d'être une conférence d'Intelligence Artificielle, on peut assigner deux objectifs spécifiques à CNIA :

- Le rôle de CNIA est de renforcer les connexions entre les différentes sous-disciplines. CNIA n'a pas pour but de nier la diversité d'objectifs, de méthodes et de pratiques déployée dans les diverses disciplines spécialisées, pas plus que de replacer les forums spécifiques à chacune. Elle souhaite être un point de rencontre pour la communauté IA permettant de rapprocher les différentes disciplines qui la composent et d'établir des passerelles entre elles ;
- Le congrès RFIA était l'occasion d'approfondir les liens de l'Intelligence Artificielle avec la Reconnaissance des Formes. CNIA, au-delà de renforcer les liens interne à l'intelligence artificielle doit permettre d'explorer les liens vers des communautés plus diversifiées.

Les seizièmes Rencontres des Jeunes Chercheurs en Intelligence Artificielle (RJCIA) sont destinées aux jeunes chercheur-se-s en Intelligence Artificielle : doctorant-e-s ou titulaires d'un doctorat depuis moins d'un an. L'objectif de cette manifestation est double :

- permettre aux jeunes chercheurs préparant une thèse en Intelligence Artificielle, ou l'ayant soutenue depuis peu, de se rencontrer, de présenter leurs travaux, et ainsi de former des contacts avec d'autres jeunes chercheurs et d'élargir leurs perspectives en échangeant avec des spécialistes d'autres domaines de l'Intelligence Artificielle ;
- former les jeunes chercheurs à la préparation d'un article, à sa révision pour tenir compte des observations du comité de programme, et à sa présentation devant un auditoire de spécialistes, leur permettant ainsi d'obtenir des retours de chercheurs de leur domaine ou de domaines connexes.

Toute contribution relevant de l'Intelligence Artificielle est la bienvenue. Toutefois, les contributions relevant de l'Ingénierie des Connaissances auront vocation à être publiées prioritairement par la conférence associée (IC). La liste indicative des thématiques ci-dessous n'est pas exhaustive :

- Apprentissage Automatique,
- Extraction et Gestion des Connaissances,
- Interaction avec l'Humain : environnements informatiques pour l'apprentissage humain (EIAH), interface homme-machine (IHM),
- Reconnaissance des Formes, Vision,
- Représentation et Raisonnement,
- Robotique, Automatique,

- Satisfaisabilité et Contraintes,
- Systèmes Multi-Agents et Agents Autonomes,
- Traitement Automatique des Langues.

Bien que les deux conférences, CNIA et RJCIA, soient dans l'esprit distinctes, nous avons organisé conjointement un unique comité de programme et un unique programme. Nous avons reçu 25 propositions d'articles dont 18 émanant en partie de jeunes chercheurs. Nous avons sélectionné 8 articles longs et 2 articles courts pour CNIA et 6 articles pour RJCIA. À cela, s'ajoutent trois démonstrations.

Reflet de l'air du temps, l'apprentissage automatique est très présent dans l'ensemble des articles acceptés. Il est complété par des thèmes de traitement du langage naturel, de représentation de connaissance, de web sémantique, de satisfaction de contraintes et de décision. On ne va pas tirer de conclusions statistiques de si petits nombres.

Ce programme est complété sur la plate-forme de conférences invitées d'Aldo Gangemi (Université Paris Nord et université de Bologne), Zhongzhi Shi (Académie des sciences de Chine, Beijing), Moshe Vardi (Rice university, Houston), Daniela Rus (Massachusetts Institute of Technology, Cambridge), et Nicolas Guarino (CNR, Trento). Nous avons aussi proposé à quelques auteurs d'articles d'autres conférences de présenter leur travail dans le programme. C'est le cas de la 18e conférence internationale sur l'extraction et la gestion des connaissances (EGC), des 14es Journées Francophones de Programmation par Contraintes (JFPC), des 12es Journées d'Intelligence Artificielle Fondamentale (JIAF), et de l'International Joint Conference on Artificial Intelligence (IJCAI).

Nous n'aurions pas pu organiser CNIA, ni RJCIA tout seuls. Nous nous devons de remercier l'Association Française pour l'Intelligence Artificielle (AFIA) qui porte ces conférences, et son président Yves Demazeau qui sait nous rappeler à nos devoirs avec la distance nécessaire, et le Laboratoire Lorrain de Recherche en Informatique et ses Applications (LORIA) qui organise la plate-forme, en particulier Armelle Brun et Davy Monticolo les co-président-e-s du comité d'organisation.

Nous sommes aussi reconnaissants à Sandra Bringuay (co-présidente du comité de programme d'APIA) et Sylvie Ranwez (présidente du comité de programme d'IC) avec qui la coordination a été presque idéale. Les membres de notre comité de programme commun, dont les noms figurent ci-après, ont pleinement accompli leur travail d'évaluation des soumissions. Finalement, nous remercions les auteurs sans qui les conférences ne seraient pas possibles.

Isabelle Tellier, auteur, membre de notre comité de programme, nous a quitté début juin. Nous souhaitons lui rendre hommage.

Jérôme Euzenat
INRIA & Univ. Grenoble Alpes
Président du comité de programme
CNIA

Francois Schwarzentruher
École normale supérieure de Rennes
Président du comité de programme
RJCIA

Comité de programme

- Carole Adam (LIG CNRS UMR 5217 - UJF)
- Emmanuel Adam (Univ Lille Nord de France)
- Audrey Baneyx (Sciences Po, médialab)
- Isabelle Bichindaritz (State University of New York at Oswego)
- Meghyn Bienvenu (CNRS, University of Montpellier, INRIA)
- Isabelle Bloch (ENST - CNRS UMR 5141 LTCI)
- Olivier Boissier (Mines Saint-Etienne, Institut Henri Fayol, Laboratoire Hubert Curien UMR CNRS 5516)
- Grégory Bonnet (Université de Caen Normandie)
- Sylvain Bouveret (LIG - Grenoble INP, Université Grenoble-Alpes)
- Nicolas Béchet (IRISA)
- Elena Cabrio (Université Côte d’Azur, CNRS, Inria, I3S, France)
- François Charpillet (INRIA-Loria)
- Mohamed Chetouani (ISIR)
- Laurence Cholvy (ONERA-Toulouse)
- Rémy Courdier (LIM - Université de la Réunion)
- Bruno Cremilleux (Universite de Caen)
- Jérôme David (Univ. Grenoble Alpes & INRIA)
- Nicolas Delestre (LITIS, Normandie Université, INSA de Rouen Normandie)
- Tiago de Lima (University of Artois and CNRS)
- Yves Demazeau (CNRS - LIG)
- Cyril de Runz (CRESTIC)
- Gaël Dias (Normandie University)
- Catherine Faron Zucker (Université Nice Sophia Antipolis)
- Jérémy Fix (SUPELEC)
- Catherine Garbay (CNRS - LIG)
- Serge Garlatti (ENST Bretagne, GET)
- Thomas Guyet (AGROCAMPUS OUEST, UMR6074 IRISA, F-35042 Rennes)
- Salima Hassas (Universit Claude Bernard-Lyon1)
- Abir Beatrice Karami (LIRIS-CNRS)
- Sébastien Konieczny (CRIL - CNRS)
- Philippe Laborie (IBM)
- Nicolas Lachiche (University of Strasbourg)
- Jean Marie Lagniez (CRIL)
- Arnaud Lallouet (Huawei Technologies Ltd)
- Robin Lamarche-Perrin (Laboratoire d’Informatique de Paris 6)
- Philippe Lamarre (LIRIS)
- Jérôme Lang (CNRS, LAMSADE, Université Paris-Dauphine)
- Fabien Lauer (Université de Lorraine, LORIA)

- Florence Le Ber (icube)
- Christophe Lecoutre (CRIL, Univ. Artois)
- Marie Lefevre (LIRIS - Université Lyon 1)
- Jean-Guy Mailly (LIPADE, Université Paris Descartes, France)
- Pierre Marquis (CRIL, U. Artois & CNRS)
- Philippe Mathieu (University of Lille 1)
- Nicolas Maudet (Université Pierre et Marie Curie)
- Engelbert Mephu Nguifo (Université Clermont Auvergne - LIMOS)
- Fabien Michel (LIRMM - Université de Montpellier)
- Frederic Migeon (IRIT)
- Davy Monticolo (LORIA)
- Maxime Morge (Université de Lille)
- Abdel-Ilah Mouaddib (universit de Caen)
- Philippe Muller (IRIT, Toulouse University)
- Amedeo Napoli (LORIA Nancy (CNRS - Inria - Université de Lorraine) France)
- Alexandre Niveau (GREYC)
- Antoine Nongaillard (Université de Lille)
- Wassila Ouerdane (LGI-CentraleSupélec)
- Odile Papini (LSIS UMR CNRS 7296)
- Alexandre Pauchet (LITIS - INSA Rouen - Normandy University)
- Damien Pellier (Laboratoire d’Informatique de Grenoble)
- Frédéric Pennerath (Supélec)
- Henri Prade (IRIT - CNRS)
- Sylvie Ranwez (LGI2P / Ecole des mines d’Alès)
- Chedy Raïssi (INRIA)
- Regis Riveret (CSIRO)
- Mathieu Roche (Cirad, TETIS)
- Marie-Christine Rousset (University of Grenoble Alpes)
- Catherine Roussey (Irstea Clermont-Ferrand Center)
- Nicolas Sabouret (LIMSI-CNRS)
- Thomas Schiex (INRA (Institut National de la Recherche Agronomique))
- Karima Sedki (CRIL (université d’artois), France)
- Christine Solnon (LIRIS CNRS UMR 5205 / INSA Lyon)
- Isabelle Tellier (Lattice)
- Alexandre Termier (Université de Rennes 1)
- Konstantin Todorov (LIRMM / University of Montpellier)
- Charlotte Truchet (LINA, UMR 6241, Université de Nantes)
- Tim Van de Cruys (IRIT & CNRS)
- Laurent Vercouter (LITIS lab, INSA de Rouen)
- Srdjan Vesic (CRIL, CNRS – Univ. Artois)
- Serena Villata (CNRS)
- Christel Vrain (LIFO - university of Orléans)
- Bruno Zanuttini (GREYC, Normandie Univ. ; UNICAEN, CNRS, ENSICAEN)
- Antoine Zimmermann (École des Mines de Saint-Étienne)
- Pierre Zweigenbaum (LIMSI, CNRS, Université Paris-Saclay)

Table des matières

Introduction	i
Comité de programme	iii
RJCIA	1
Découverte des u-shapelets basée sur la corrélation pour le clustering de séries temporelles incertaines	1
<i>Vanel Steve Siyou Fotso, Engelbert Mephu Nguifo, Philippe Vaslin</i>	
ABClass : Une approche d'apprentissage multi-instancés pour les séquences	10
<i>Manel Zoghalmi, Sabeur Aridhi, Mondher Maddouri, Engelbert Mephu Nguifo</i>	
Méthode non paramétrique pour l'analyse et la classification des données fonctionnelles	19
<i>Papa Alioune Meissa Mbaye, Anne-Françoise Yao, Chafik Samir</i>	
Expression verbale des tendances à l'action en langue française	27
<i>Alya Yacoubi, Nicolas Sabouret</i>	
Classification d'images en apprenant sur des échantillons positifs et non labélisés avec un réseau antagoniste génératif	35
<i>Florent Chiaroni, Mohamed-Cherif Rahal, Frédéric Dufaux, Nicolas Hueber</i>	
Inférence bayésienne pour l'estimation de déformations larges par champs gaussien : application au recalage d'images multi-modales	43
<i>Thomas Deregnaucourt, Chafik Samir, Anne-Françoise Yao</i>	
CNIA	50
Mécanisme de négociation multilatérale pour la prise de décision collective	50
<i>Ndeye Arame Diago, Samir Aknine, Onn Shehory, Mbaye Sène</i>	
Choisir un encodage CNF de contraintes de cardinalité performant pour SAT	58
<i>Thomas Delacroix</i>	
Apprentissage par analogies grâce à des outils de la théorie des catégories	65
<i>Sarah Amar, Omar Bentahar, Lionel Cordesses, Thomas Ehrmann, Aude Laurent, Julien Page, Kevin Poulet</i>	
Modélisation thématique à l'aide des plongements lexicaux issus de Word2Vec	73
<i>Svitlana Galeshchuk, Bruno Chaves Ferreira</i>	
Réseaux de neurones récurrents multi-tâches pour l'analyse automatique d'arguments	78
<i>Jean-Christophe Mensonides, Sébastien Harispe, Jacky Montmain, Véronique Thireau</i>	

Découverte de cardinalité maximale contextuelle dans les bases de connaissances	86
<i>El Arby Sidi Aly, Mohamed Lamine Diakité, Arnaud Giacometti, Béatrice Markhoff, Arnaud Soulet</i>	
Trois approches pour classifier les données du web des données	94
<i>Justine Reynaud, Yannick Toussaint, Amedeo Napoli</i>	
Diagnostic automatique de l'état dépressif	102
<i>Stéphane Cholet, Hélène Paugam-Moisy</i>	
Articles courts (CNIA)	110
Analyse d'opinions multi-aspects pour la recommandation fine de restaurants	110
<i>Isabelle Tellier, Hamid Hammouche, Didier Cholvy, Jean-Baptiste Tanguy</i>	
A de novo robust clustering approach for amplicon-based sequence data	113
<i>Alexandre Bazin, Didier Debroas, Engelbert Mephu Nguifo</i>	
Démonstrations	116
Démonstration du diagnostic automatique de l'état dépressif	116
<i>Stéphane Cholet, Hélène Paugam-Moisy</i>	
OntoRev : un moteur de révision d'ontologies OWL	118
<i>Thinh Dong, Chan Le Duc, Myriam Lamolle</i>	
ABClass : A multiple instance learning approach for sequence data	123
<i>Manel Zoghلامي, Sabeur Aridhi, Mondher MadDouri, Engelbert Mephu Nguifo</i>	

Découverte des u-shapelets basée sur la corrélation pour le clustering de séries temporelles incertaines.

Vanel Steve Siyou Fotso¹

Engelbert Mephu Nguifo¹

Philippe Vaslin¹

¹ University Clermont Auvergne, CNRS, LIMOS, F-63000 Clermont-Ferrand, France

{siyou, mephu, vaslin}@isima.fr

Abstract

An u-shapelet is a sub-sequence of a time series used for clustering a time series dataset. The purpose of this paper is to discover u-shapelets on uncertain time series. To achieve this goal, we propose a dissimilarity score called FOTS whose computation is based on the eigenvector decomposition and the comparison of the autocorrelation matrices of the time series. This score is robust to the presence of uncertainty; it is not very sensitive to transient changes; it allows capturing complex relationships between time series such as oscillations and trends, and it is also well adapted to the comparison of short time series. The FOTS score is used with the Scalable Unsupervised Shapelet Discovery algorithm for the clustering of 17 datasets, and it has shown a substantial improvement in the quality of the clustering with respect to the Rand Index. This work defines a novel framework for the clustering of uncertain time series.

Mots Clef

Clustering, UShapelet, Correlation, time series.

Résumé

Un u-shapelet est une sous-séquence d'une série temporelle utilisée comme propriété pour séparer un groupe de séries temporelles en deux sous-groupes. Un sous-groupe de séries temporelles contenant le shapelet et un sous-groupe de série temporelles ne contenant pas le shapelet. Le but de cet article est de découvrir des u-shapelets sur des séries temporelles incertaines. Pour atteindre cet objectif, nous proposons un score de dissimilarité appelé FOTS dont le calcul est basé sur la comparaison des vecteurs propres des matrices d'autocorrélation des séries temporelles. Ce score est robuste à la présence d'incertitude; il n'est pas très sensible aux changements transitoires; il permet de capturer des relations complexes entre des séries temporelles telles que les oscillations et les tendances, et il est également bien adapté à la comparaison de séries temporelles courtes. Le score FOTS est utilisé avec l'algorithme Scalable Unsupervised Shapelet Discovery pour la classification non supervisée de 17 ensembles

de données, et il a montré une amélioration substantielle de la qualité de la classification non supervisée par rapport au Rand Index. Ce travail définit un nouveau cadre pour la classification non supervisée de séries temporelles incertaines.

Mots Clef

Classification non supervisée, UShapelet, Correlation, Séries temporelles.

1 Introduction

Toutes les mesures effectuées par un système mécanique ont une incertitude. En effet, le principe d'incertitude met en évidence les limites de la capacité des systèmes mécaniques à effectuer des mesures sur un système sans les perturber [1]. Ainsi, les séries temporelles des instruments de mesure sont incertaines. Ces séries temporelles produites par des capteurs constituent une vaste proportion des séries temporelles utilisées en science, que ce soit en médecine avec des Électrocardiogramme (enregistrement de l'activité électrique du cœur), en physique avec des mesures enregistrées par des télescopes, en informatique avec l'Internet des objets, etc. Ignorer l'incertitude des données au cours de leur analyse peut conduire à des conclusions approximatives ou inexactes, d'où la nécessité de mettre en œuvre des techniques de gestion des données incertaines. Plusieurs études récentes ont porté sur le traitement de l'incertitude dans l'exploration de données. Deux approches principales permettent de prendre en compte l'incertitude dans les tâches de data mining : soit elle est prise en compte lors de la comparaison en utilisant les fonctions de distance appropriées [2–7], soit son impact est réduit par les transformations effectuées sur les données. Cette dernière stratégie est utilisée nativement par l'algorithme u-shapelet.

1.1 État de l'art sur les u-shapelets

Considérons un ensemble de données composé de 6 séries temporelles correspondant aux appels d'oiseaux : 3 séries temporelles correspondant à **Moucherolle à côtés olive** (séries temporelles vertes) et 3 séries temporelles correspondant aux appels de **Moineau à couronne blanche** (séries

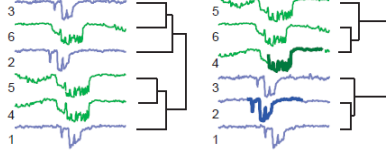


FIGURE 1 – Exemple de classification de séries temporelles utilisant d’une part la distance euclidienne (gauche), d’autre part des sous-séquences caractéristiques appelées u-shapelet (droite). [8].

temporelles bleues). Lorsque ces séries temporelles sont classées en utilisant la distance euclidienne comme mesure de dissimilarité (Fig. 1 gauche), les groupes obtenus ne sont pas homogènes ; en d’autres termes, l’algorithme n’identifie pas l’oiseau à partir de ses cris. Cependant, si nous recherchons des sous-séquences caractéristiques (u-shapelets) pour classer les séries temporelles, nous obtenons des groupes plus homogènes (Fig. 1 à droite).

Une fois cette observation faite, la question naturelle est de savoir comment trouver des sous-séquences qui caractérisent un groupe de séries temporelles, c’est-à-dire des sous-séquences qui ne sont observées que dans un sous-groupe de séries temporelles. L’algorithme de découverte d’u-shapelet répond à cette question et procède comme suit : l’algorithme prend la longueur du motif comme paramètre. Sur chaque série temporelle de la base, on fait glisser une fenêtre de la même longueur que le motif, chaque nouvelle sous-séquence obtenue par ce processus est un u-shapelet candidat.

Parmi les u-shapelets candidats, nous considérons comme un u-shapelet la sous-séquence capable de diviser l’ensemble de séries temporelles en deux sous-ensembles D_A et D_B de sorte que D_A contient toutes les séries temporelles qui possèdent le u-shapelet et D_B toutes celles qui ne contiennent pas l’u-shapelet. Deux autres contraintes sont prises en compte dans la découverte de motifs :

La première est la capacité de l’u-shapelet à construire des sous-ensembles bien séparés. La deuxième est la capacité de l’u-shapelet à construire des sous-ensembles qui ne sont pas déséquilibrés. C’est-à-dire que la taille de D_A doit être au plus k fois plus grande que celle de D_B et vice versa.

Definition 1 Deux jeux de données D_A et D_B sont dit **r-équilibré** si et seulement si $\frac{1}{r} < \frac{|D_A|}{|D_B|} < (1 - \frac{1}{r})$, $r > 1$

Definition 2 Un **u-shapelet** est une sous-séquence qui a une longueur inférieure ou égale à la longueur de la plus petite série temporelle du jeu de données, et qui permet de diviser le jeu de données en deux sous-groupes **r-équilibrés** D_A et D_B ; où D_A est le groupe de séries temporelles qui contiennent un motif **similaire** au u-shapelet et D_B est le groupe de séries temporelles qui ne contiennent pas l’u-shapelet.

La similarité entre une série temporelle et un u-shapelet est évaluée à l’aide d’une fonction de distance.

Definition 3 La distance de sous-séquence $sdist(S, T)$ entre une série temporelle T et une sous-séquence S est le minimum des distances entre la sous-séquence S et toutes les sous-séquences de T de longueur égale à celle de S .

Cette définition ouvre la question de la mesure de distance à utiliser pour $sdist$. En général, la distance euclidienne omniprésente (ED) est utilisée, mais elle ne l’est pas appropriée pour les séries temporelles incertaines [5]. Dans la section suivante, nous présentons une fonction de dissimilarité plus adaptée à l’incertitude.

Le calcul de la $sdist$ entre un u-shapelet candidat et toutes les séries temporelles d’un jeu de données est appelé **orderline**.

Definition 4 Un **orderline** est un vecteur de distance entre un u-shapelet et toutes les séries temporelles d’un jeu de données.

Le calcul d’un orderline est coûteux en temps. Un orderline pour un seul u-shapelet candidat a une complexité en temps égale à $O(NM \log(M))$ où N est le nombre de séries temporelles du jeu de données, M est la longueur moyenne des séries temporelles. L’algorithme force brute pour la découverte de u-shapelet requiert K calculs d’orderline, où K est le nombre de sous-séquences candidate. La stratégie utilisée par [8] à filtrer les K sous-séquences candidates en considérant seulement celles permettant de construire deux groupes **r-équilibrés**. Cette sélection est faite efficacement grâce à un algorithme de hachage.

L’évaluation de la qualité des u-shapelets est basée sur leur pouvoir de séparation qui est calculé comme suit :

$$gap = \mu_B - \sigma_B - (\mu_A - \sigma_A), \quad (1)$$

Où μ_A (resp. μ_B) représente la moyenne($sdist(S, D_A)$) (resp. moyenne($sdist(S, D_B)$)), et σ_A (resp. σ_B) représente l’écart-type de $sdist(S, D_A)$ (resp. écart-type de $sdist(S, D_B)$).

Si D_A ou D_B est constitué d’un seul élément (ou d’un nombre insuffisant de séries temporelles pour constituer un groupe), le gap prend la valeur zéro. Ceci assure qu’un gap élevé pour un u-shapelet correspond à une séparation réelle.

1.2 U-shapelets pour la classification non supervisée de séries temporelles incertaines

La classification non supervisée basée sur des u-shapelets est une approche introduite par [9] qui a suggéré de regrouper des séries temporelles à partir des propriétés locales de leurs sous-séquences plutôt qu’utiliser les caractéristiques globales de la série temporelle [10]. Dans ce but, le clustering basé sur les u-shapelets cherche d’abord un ensemble de sous-séquences caractéristiques des différentes catégories de séries temporelles et classe une série temporelle en fonction de la présence ou de l’absence de ces sous-séquences caractéristiques.

La classification non supervisée de séries temporelles avec des u-shapelets présente plusieurs avantages. Premièrement, la classification non supervisée basée sur les u-shapelets est définie pour les ensembles de données dans lesquels les séries temporelles ont des longueurs différentes, ce qui n'est pas le cas pour la plupart des techniques décrites dans la littérature. En effet, dans de nombreux cas, l'hypothèse de longueur égale est implicite, et le découpage à longueur égale est effectué en exploitant des compétences humaines coûteuses [8]. Deuxièmement, la classification non supervisée basée sur les u-shapelets est beaucoup plus expressive en ce qui concerne le pouvoir de représentation. En effet, une série temporelle n'est associée à un groupe que si elle contient l'u-shapelet caractéristique de ce groupe. Ainsi, une série temporelle pourrait n'être associée à aucun groupe.

De plus, il est très approprié d'utiliser la classification non supervisée basée sur les u-shapelets avec des séries temporelles incertaines parce que la stratégie de comparaison d'un u-shapelet et d'une série temporelle ignore les données non pertinentes de la série temporelle et ainsi réduire les effets négatifs de la présence d'incertitudes dans celle-ci. Malgré cet avantage, il est hautement souhaitable de prendre en compte l'impact négatif de l'incertitude lors de la découverte des u-shapelets.

1.3 Incertitude et découverte des u-shapelets

Les mesures traditionnelles de similarité comme la distance euclidienne (ED) ou la distorsion temporelle dynamique (DTW) ne fonctionnent pas toujours bien pour les séries temporelles incertaines. En effet, ces distances agrègent l'incertitude de chaque point de la série temporelle et amplifient ainsi l'impact négatif de l'incertitude. Cependant, ED joue un rôle fondamental dans la découverte des u-shapelets, car elle est utilisée pour calculer l'écart entre deux groupes formé par l'u-shapelet candidat. La découverte de u-shapelet sur des séries temporelles incertaines pourrait donc conduire à la sélection d'un mauvais candidat u-shapelet ou à l'assignation d'une série temporelle au mauvais groupe.

Dans cette étude, notre but n'est pas de définir un algorithme pour la découverte d'u-shapelets incertains, mais plutôt d'utiliser une fonction de dissimilarité robuste à l'incertitude pour améliorer la qualité des u-shapelets découverts et donc la qualité de clustering des séries temporelles incertaines.

1.4 Contributions

- Nous faisons un état de l'art sur les mesures de dissimilarité incertaines et nous les évaluons pour leur pertinence pour la comparaison de séries temporelles incertaines de petite taille.
- Nous introduisons une fonction de dissimilarité nommée corrélation frobenius pour la découverte d'u-shapelets sur les séries temporelles incertaines (FOTS); qui possède des propriétés intéressantes pour la comparaison de séries temporelles incertaines

de petite taille et ne fait aucune hypothèse sur la distribution de probabilité de l'incertitude.

- Nous mettons le code source à la disposition de la communauté scientifique pour permettre une extension de ce travail.

2 Définitions et travaux connexes

2.1 Définitions

Une série temporelle incertaine (UTS) $X = \langle X_1, \dots, X_n \rangle$ est une séquence de variable aléatoire où X_i est une variable aléatoire modélisant une valeur réelle inconnue à l'instant i . Deux modèles sont principalement utilisés pour représenter les séries temporelles incertaines : le modèle ensembliste, et le modèle basé sur la fonction de densité de probabilité de l'incertitude.

Dans le modèle ensembliste, chaque élément X_i ($1 \leq i \leq n$) de l'UTS $X = \langle X_1, \dots, X_n \rangle$ est représenté par un ensemble $\{X_{i,1}, \dots, X_{i,N_i}\}$ de valeurs observées et N_i représente le nombre d'observation à l'instant i .

Dans le modèle basé sur la distribution de probabilité, chaque élément X_i , ($1 \leq i \leq n$) de l'UTS $X = \langle X_1, \dots, X_n \rangle$ est représenté par une variable aléatoire $X_i = x_i + X_{e_i}$, où x_i est la valeur exacte qui est inconnue, et X_{e_i} est une variable aléatoire représentant l'erreur. C'est ce modèle que nous considérons tout au long de notre travail.

Plusieurs mesures de similarité ont été proposées pour les séries temporelles incertaines. Elles peuvent être regroupées en deux catégories principales : les mesures de similarité traditionnelles et les mesures de similarité incertaines.

- les mesures de similarité traditionnelles telle que la distance euclidienne sont celles conventionnellement utilisées avec les séries temporelles. Elles utilisent une seule valeur incertaine à chaque instant comme approximation de la valeur réelle inconnue.
- les mesures de similarité incertaines utilisent des informations statistiques additionnelles qui mesurent l'incertitude associée à chaque approximation de la valeur réelle. C'est le cas notamment de DUST, PROUD, MUNICH [11]. [12] démontre que les performances des mesures de similarité incertaines associées au pré-traitement sont meilleures que les performances des mesures de similarité traditionnelles sur des jeux de données contenant de l'incertitude.

2.2 État de l'art sur les mesures de similarité incertaines

Les mesures de similarité incertaines peuvent être regroupées en deux grandes catégories : les mesures de similarité déterministes et les mesures de similarité probabilistes.

Mesure de similarité déterministe. Tout comme les mesures de similarité traditionnelles, les mesures de similarité déterministes renvoient un nombre réel représentant la dis-

tance entre deux séries temporelles incertaines. **DUST** est un exemple de mesure déterministe de similarité.

DUST [13] Etant donné deux séries temporelles incertaines $X = \langle X_1, \dots, X_n \rangle$ et $Y = \langle Y_1, \dots, Y_n \rangle$, la distance entre deux valeurs X_i, Y_i est définie comme étant la distance entre leurs valeurs réelles inconnues $r(X_i), r(Y_i)$: $dist(X_i, Y_i) = |r(X_i) - r(Y_i)|$. Cette distance est utilisée pour mesurer la similarité de deux valeurs incertaines.

$\varphi(|X_i - Y_i|)$ est la probabilité que les valeurs réelles à l'instant i soient égales, connaissant les valeurs réelles à l'instant i .

$$\varphi(|X_i - Y_i|) = Pr(dist(0, |X_i - Y_i|) = 0). \quad (2)$$

cette fonction de similarité est par la suite utilisée par la fonction de dissimilarité *dust* :

$$dust(X_i, Y_i) = \sqrt{-\log(\varphi(|X_i - Y_i|)) + \log(\varphi(0))}. \quad (3)$$

La distance entre les séries temporelles incertaines $X = \langle X_1, \dots, X_n \rangle$ et $Y = \langle Y_1, \dots, Y_n \rangle$ calculée à partir de *DUST* est alors définie comme suit :

$$DUST(X, Y) = \sqrt{\sum_{i=1}^n dust(X_i, Y_i)^2}. \quad (4)$$

Le problème avec les distances déterministes incertaines comme *DUST* est que leur expression varie en fonction de la distribution de probabilité de l'incertitude, et la distribution de probabilité de l'incertitude n'est pas toujours disponible pour les jeux de données de séries temporelles.

Mesure de similarité probabiliste. Les mesures des similarités probabilistes n'exigent pas la connaissance de la distribution des probabilités d'incertitude. De plus, elles fournissent aux utilisateurs plus d'informations sur la fiabilité du résultat. Il existe plusieurs fonctions de similarité probabiliste, comme **MUNICH**, **PROUD**, **PROUDS** ou **Corrélation Locale**.

MUNICH [14]. Cette fonction de distance convient aux séries temporelles incertaines représentées par le modèle ensembliste. La probabilité que la distance entre deux séries temporelles incertaines X et Y soit inférieure à un seuil ε est égale au nombre de distances entre X et Y , qui sont inférieures à ε , sur le nombre possible de distances :

$$Pr(distance(X, Y)) \leq \varepsilon = \frac{|\{d \in dists(X, Y) | d \leq \varepsilon\}|}{|dists(X, Y)|} \quad (5)$$

Le calcul de cette fonction de distance est très coûteuse en temps.

PROUD [15] Soient $X = \langle X_1, \dots, X_n \rangle$ et $Y = \langle Y_1, \dots, Y_n \rangle$ deux UTS modélisées par des séquences de variables aléatoires, la distance **PROUD** entre X et Y est $d(X, Y) = \sum_{i=1}^n (X_i - Y_i)^2$. D'après le théorème central limite [16], la distribution cumulée approche asymptotiquement une loi normale

$$d(X, Y) \propto N\left(\sum_i E[(X_i - Y_i)^2], \sum_i Var[(X_i - Y_i)^2]\right) \quad (6)$$

Une conséquence de cette caractéristique de la distance **PROUD** est que le tableau de la loi réduite centrée normale peut être utilisé pour calculer la probabilité que la distance **PROUD** normalisée soit inférieure à un seuil :

$$Pr(d(X, Y)_{norm} \leq \epsilon). \quad (7)$$

Un inconvénient majeur de **PROUD** est son inadéquation pour la comparaison de séries temporelles de petites longueurs comme les u-shapelets. En effet, le calcul de la probabilité que la distance **PROUD** soit inférieure à une valeur est basé sur l'hypothèse qu'elle suit **asymptotiquement** une distribution normale. Ainsi, cette probabilité sera d'autant plus précise que les séries temporelles comparées sont longues (plus de 30 points de données).

PROUDS [12] est une version améliorée de **PROUD** qui suppose que les variables aléatoires qui constituent la série temporelle sont indépendantes et identiquement distribuées.

Définition 5 La forme normalisée d'une série temporelle $X = \langle X_1, \dots, X_n \rangle$ est définie par $\hat{X} = \langle \hat{X}_1, \dots, \hat{X}_n \rangle$ tel que à chaque instant i ($1 \leq i \leq n$), nous avons :

$$\hat{X}_i = \frac{X_i - \bar{X}}{S_X}, \quad \bar{X} = \sum_{i=1}^n \frac{X_i}{n}, \quad S_X = \sqrt{\sum_{i=1}^n \frac{(X_i - \bar{X})^2}{(n-1)}}. \quad (8)$$

PROUDS définit la distance entre deux séries temporelles normalisées $\hat{X} = \langle \hat{X}_1, \dots, \hat{X}_n \rangle$ and $\hat{Y} = \langle \hat{Y}_1, \dots, \hat{Y}_n \rangle$ (Définition 5) comme suit :

$$Eucl(\hat{X}, \hat{Y}) = 2(n-1) + 2 \sum_{i=1}^n \hat{X}_i \hat{Y}_i \quad (9)$$

Pour les mêmes raisons que **PROUD**, **PROUDS** ne conviennent pas à la comparaison de séries temporelles courtes. Un autre inconvénient de **PROUDS** est qu'il suppose que les variables aléatoires sont indépendantes : cette hypothèse est forte et particulièrement inappropriée pour des séries temporelles courtes comme les u-shapelets. Une

hypothèse plus réaliste avec les séries temporelles serait de considérer que les variables aléatoires constituant les séries temporelles sont M-dépendantes. Les variables aléatoires d'une série temporelle sont dites M-dépendantes si $X_i, X_{i+1}, \dots, X_{i+M}$ sont dépendantes (corrélées) et les variables X_i et X_{i+M+1} sont indépendantes. Cependant, supposer que les variables aléatoires sont M-dépendantes complexifie l'écriture de PROUDS et rend son utilisation plus difficile car elle requiert dès lors d'affecter une valeur au paramètre M.

Corrélation incertaine [7] : Les techniques d'analyse de corrélation sont utiles pour la sélection de caractéristiques dans des séries temporelles incertaines. Ces informations permettent d'identifier les éléments redondants. La même stratégie peut être utile pour la découverte de u-shapelet. La corrélation incertaine est définie comme suit :

Définition 6 (Correlation sur les séries temporelles incertaines) Etant données deux UTS $X = \langle X_1, \dots, X_n \rangle$ et $Y = \langle Y_1, \dots, Y_n \rangle$, leur corrélation est définie par

$$Corr(X, Y) = \sum_{i=1}^n \hat{X}_i \hat{Y}_i / (n - 1), \quad (10)$$

Où $\hat{X}_i$ et $\hat{Y}_i$ sont les formes normales de X_i et Y_i (Définition 5) respectivement. X_i et Y_i sont des variables aléatoires indépendantes et continues.

Si nous connaissons la distribution de probabilité des variables aléatoires, il est possible de déterminer la fonction de densité de probabilité associée à la corrélation, qui servira par la suite à calculer la probabilité que la corrélation entre deux séries temporelles soit supérieure à un seuil donné. La corrélation incertaine a cependant quelques inconvénients :

- Il est trop sensible aux changements transitoires, ce qui conduit souvent à des scores très fluctuants ;
- Il ne peut pas capturer les relations complexes dans les séries temporelles ;
- Il faut connaître la fonction de distribution de probabilité de l'incertitude ou faire une hypothèse sur l'indépendance des variables aléatoires contenues dans les séries temporelles.

En raison de tous ces inconvénients, la corrélation incertaine ne peut pas être utilisée en l'état pour la découverte d'u-shapelet. Le paragraphe suivant présente une généralisation du coefficient de corrélation qui n'est pas une fonction de similarité incertaine mais qui reste intéressante pour la découverte des u-shapelets.

Corrélation locale [17] : La corrélation locale est une généralisation de la corrélation. Elle calcule un score de corrélation évolutif dans le temps qui suit une similarité locale sur des séries temporelles multivariées basées sur une matrice d'auto-corrélation locale. La matrice d'auto-corrélation **permet de capturer des relations complexes** dans des séries temporelles comme l'oscillation clé (par

exemple, sinusoïdale) ainsi que les tendances apériodiques (par exemple, à la hausse ou à la baisse) qui sont présentes. L'utilisation de matrices d'auto-corrélation qui sont calculées sur la base de fenêtres se chevauchant permet de **réduire la sensibilité aux changements transitoires** dans les séries temporelles.

Définition 7 (Auto-covariance, fenêtre glissante) Étant donnée une série temporelle X , un ensemble de fenêtre glissante w , l'estimateur de la matrice d'autocorrélation locale $\hat{\Gamma}_t$ utilisant une fenêtre glissante est définie à l'instant $t \in \mathbb{N}$ tel que (Eq.11) :

$$\hat{\Gamma}_t(X, w, m) = \sum_{\tau=t-m+1}^t x_{\tau,w} \otimes x_{\tau,w}. \quad (11)$$

Où $x_{\tau,w}$ est une sous-séquence de la série temporelle de longueur w et commençant à τ , $x \otimes y = xy^T$ est le produit extérieur de x et y . L'ensemble d'échantillons de m fenêtres est centré autour du temps t . Nous fixons généralement le nombre de fenêtres à $m = w$.

Étant donnée les estimations $\hat{\Gamma}_t(X)$ et $\hat{\Gamma}_t(Y)$ pour les deux séries temporelles, la prochaine étape est de savoir comment les comparer et extraire un score de corrélation. Cet objectif est atteint en utilisant la décomposition spectrale ; les vecteurs propres des matrices d'auto-corrélations capturent les principales tendances apériodiques et oscillatoires, même sur des **séries temporelles courtes**. Ainsi, les sous-espaces couverts par les premiers (k) vecteurs propres sont utilisés pour caractériser localement le comportement de chaque série. La définition 8 formalise cette notion :

Définition 8 (LoCo score) Étant donnée deux séries temporelles X et Y leur score LoCo est défini par

$$\ell_t(X, Y) = \frac{1}{2} (\|U_X^T u_Y\| + \|U_Y^T u_X\|) \quad (12)$$

Où U_X et U_Y sont les k premiers vecteurs propres des matrices d'auto-corrélation locales $\hat{\Gamma}_t(X)$ et $\hat{\Gamma}_t(Y)$ respectivement, et u_X et u_Y sont les vecteurs propres ayant la plus large valeur propre.

Intuitivement, deux séries temporelles X et Y seront considérées comme étant proches lorsque l'angle formé par l'espace portant les informations de la série temporelle X et de la série temporelle Y est nul. En d'autres termes, X et Y seront proches lorsque la valeur de $\cos(\alpha)$ sera 1. La seule hypothèse faite pour le calcul de la similitude LoCo est que la moyenne des points de la série temporelle est nulle. Cette hypothèse peut facilement être vérifiée, il suffit pour cela de normaliser les séries temporelles en cours de comparaison. La fonction de similarité LoCo a de nombreuses propriétés intéressantes et ne nécessite pas :

- de connaître la distribution de probabilité de l'incertitude,

- de supposer l'indépendance des variables aléatoires ou de faire une hypothèse sur la longueur de l'u-shapelet.

Elle est donc intéressante pour la découverte de motifs, mais nous avons encore besoin d'une dissimilarité pour pouvoir découvrir des u-shapelets. Dans le paragraphe suivant, nous allons définir une fonction de dissimilarité qui a les mêmes propriétés que LoCo et c'est-à-dire, qui est robuste à la présence d'incertitude.

3 Notre approche

3.1 Fonction de dissimilarité

La fonction de similarité LoCo définie sur deux séries temporelles multivariées X et Y correspond approximativement à la valeur absolue du cosinus de l'angle formé par les espaces propres de X et Y ($|\cos(\alpha)|$). Une idée simple serait d'utiliser la valeur $\sin(\alpha)$ ou α comme fonction de dissimilarité mais cette approche ne fonctionne pas si bien ; le sinus et l'angle ne sont pas assez discriminants pour la comparaison de vecteurs propres à des fins de clustering. Nous proposons donc la mesure de dissimilarité suivante (Définition. 9).

Définition 9 (*FOTS : Frobenius cOrrelation for uncertain Time series u-Shapelet discovery*)
Étant données deux séries X et Y , leur score FOTS est défini par

$$FOTS(X, Y) = \|U_X - U_Y\|_F = \sqrt{\sum_{i=1}^m \sum_{j=1}^k (U_X - U_Y)_{ij}^2} \quad (13)$$

où $\|\cdot\|_F$ est la norme de Frobenius, m est la longueur de la série temporelle et k est le nombre de vecteurs propres.

Comme le calcul FOTS est basé sur la comparaison des k -premiers vecteurs propres des matrices d'autocorrélation de la série temporelle, il a les mêmes propriétés souhaitables de la fonction de similarité LoCo, c'est-à-dire :

- Il **permet de capturer des relations complexes** dans des séries temporelles comme les tendances oscillatoires clés (par exemple, sinusoïdales) ainsi que les tendances apériodiques (par exemple, à la hausse ou à la baisse) qui sont présentes ;
- Il permet de **réduire la sensibilité aux changements transitoires** dans les séries temporelles ;
- Il est approprié pour le **comparaison de séries temporelles courtes**.

De plus, la fonction de dissimilarité FOTS est **robust à la présence d'incertitude** due à la décomposition spectrale des matrices d'autocorrélation des séries temporelles. La robustesse du FOTS à l'incertitude est confirmée par le théorème suivant :

Théorème 1 (*Hoffman Wielandt*) [18] Si X et $X + E$ sont matrices $n \times n$ symétriques, alors :

$$\sum_{i=1}^n (\lambda_i(X + E) - \lambda_i(X))^2 \leq \|E\|_F^2. \quad (14)$$

où $\lambda_i(X)$ est la plus grande valeur propre de X , et $\|E\|_F^2$ est le carré de la norme Frobenius de E .

La section suivante explique comment le FOTS est intégré dans l'algorithme de découverte d'u-shapelets.

3.2 Algorithme des u-shapelets avec score FOTS

Dans cette section, nous ne définissons pas un nouvel algorithme SUShapelet, mais nous expliquons comment nous utilisons l'algorithme SUShapelet avec le score FOTS (FOTS-SUSh) pour faire face à l'incertitude.

Le gap est un critère essentiel pour la sélection des u-shapelets. Il est sujet à l'incertitude parce que son calcul est basé sur la distance euclidienne. Pour y remédier, nous proposons d'utiliser le score de FOTS au lieu d'une simple distance euclidienne lors du calcul du gap. L'algorithme 1 explique comment calculer l'orderline en utilisant le score de FOTS. L'algorithme 2 calcul l'orderline et trie les séries temporelles en fonction de leur proximité au u-shapelet candidat (ligne 2 et 3). Un u-shapelet est considéré comme étant présent dans une série temporelle si sa distance à la série temporelle est inférieure ou égale à un seuil. Ainsi, l'algorithme de sélection de seuil construit un cluster D_A dont la taille varie entre lb et ub (ligne 5). L'algorithme cherche alors parmi les seuils sélectionnés celui ayant un gap maximal (ligne 6 et 11).

Définition 10 La fonction de dissimilarité $sd_f(S, T)$ entre une série temporelle T et une sous-séquence S est le minimum des valeurs du score de FOTS entre la sous-séquence S et toutes les sous-séquences possible de T de longueur égales à S .

Algorithme 1 : ComputeOrderline

Input : u-shapeletCandidate : s,

Jeu de données : D

Output : Distance entre l'u-shapelet candidat et toutes les séries temporelles du jeu de données

```

1 function ComputeOrderline(s, D)
2    $dis \leftarrow \{\}$   $s \leftarrow zNorm(s)$ 
3   forall  $i \in \{1, 2, \dots, |D|\}$  do
4      $ts \leftarrow D(i, :)$ 
5      $dis(i) \leftarrow sd_f(s, ts)$ 
6   return  $dis/|s|$ 
```

Algorithme 2 : ComputeGap

Input : u-shapeletCandidate : s ,
Jeu de données : D ,
 lb, ub : lower/upper bound : nombre minimum et maximum de séries temporelles par groupe
Output : gap : distance entre deux groupes

```
1 function ComputeGap( $s, D, lb, ub$ )
2    $dis \leftarrow ComputeOrderline(s, D)$ 
3    $dis \leftarrow sort(dis)$   $gap \leftarrow 0$ 
4   for  $i \leftarrow lb$  to  $ub$  do
5      $D_A \leftarrow dis \leq dis(i), D_B \leftarrow dis > dis(i)$ 
6      $m_A \leftarrow mean(D_A), m_B \leftarrow mean(D_B)$ 
7      $s_A \leftarrow std(D_A), s_B \leftarrow std(D_B)$ 
8      $currGap \leftarrow m_B - s_B - (m_A + s_A)$ 
9     if  $currGap > gap$  then
10      |  $gap \leftarrow currGap$ 
11 return  $gap$ 
```

4 Evaluation expérimentale

4.1 Classification non supervisée avec les u-shapelets

Il existe de nombreuses façons de regrouper les séries temporelles décrites par des u-shapelets. Dans cette expérience, l'algorithme divise itérativement le jeu de données à partir de chaque u-shapelet découvert : chaque u-shapelet divise l'ensemble de données en deux groupes D_A et D_B . Les séries temporelles qui appartiennent à D_A sont celles contenant le u-shapelet et sont ensuite supprimées du jeu de données. Une nouvelle recherche de u-shapelet se poursuit avec le reste des données jusqu'à ce qu'il n'y ait plus de séries temporelles dans le jeu de données ou jusqu'à ce que l'algorithme ne soit plus capable de trouver d'u-shapelet. En guise de critère d'arrêt, nous considérons la baisse du gap. L'algorithme s'arrête lorsque le gap de l'u-shapelet nouvellement trouvé devient la moitié du gap du premier u-shapelet découvert. Cette approche est une mise en oeuvre directe de la définition d'u-shapelet.

Choisir la longueur N des u-shapelet : Le choix de la longueur de l'u-shapelet est guidé par la connaissance du domaine d'application duquel provient les séries temporelles. Dans le cadre de ces expériences, nous avons testé tous les nombres entre 4 et la moitié de la longueur des séries temporelles. Nous considérons comme longueur de l'u-shapelet celle permettant de mieux regrouper les séries temporelles.

Choisir la longueur w de la fenêtre : L'utilisation de fenêtres qui se chevauchent pour le calcul de la matrice d'auto-corrélation permet de capturer les oscillations présentes dans la série temporelle. Au cours de ces expériences, nous considérons que la taille de la fenêtre est égale à la moitié de la longueur de la forme en U.

Choisir le nombre k de vecteurs propres : Un choix pratique est de fixer k à une petite valeur ; nous utilisons $k = 4$ tout au long de ces expériences. En effet, les tendances aperiodiques clés sont capturées par un seul vecteurs propres, tandis que les principales tendances oscillatoires se manifestent dans une paire de vecteurs propres.

5 Métriques d'évaluation

Différentes mesures de la qualité de classification non supervisée de séries temporelles ont été proposées, notamment le score Jaccard, l'indice Rand, l'indice Folkes et l'indice Mallow, etc. Cependant, dans notre cas, nous avons des étiquettes de classe pour les jeux de données, nous pouvons donc utiliser cette information externe pour évaluer la véritable qualité de la classification non supervisée en utilisant l'indice Rand. De plus, l'indice Rand semble être la mesure de la qualité des groupes couramment utilisée [8–10].

5.1 Comparaison avec u-shapelet

De même que [11], nous avons testé notre méthode sur 17 jeux de données du monde réel provenant des archives UCR [19] représentant un large éventail de domaines d'application. Les ensembles de d'apprentissage et de test ont été réunis pour obtenir des jeux de données plus importants. Le tableau 1 présente des informations détaillées sur les ensembles de données testés.

Data-set	Size of dataset	Length	No. of Classes	Type
50words	905	270	50	IMAGE
Adiac	781	176	37	IMAGE
Beef	60	470	5	SPECTRO
Car	120	577	4	SENSOR
CBF	930	128	3	SIMULATED
Coffee	56	286	2	SPECTRO
ECG200	200	96	2	ECG
FaceFour	112	350	4	IMAGE
FISH	350	463	7	IMAGE
Gun_Point	200	150	2	MOTION
Lighting2	121	637	2	SENSOR
Lighting7	143	319	7	SENSOR
OliveOil	60	570	4	SPECTRO
OSULeaf	442	427	6	IMAGE
SwedishLeaf	1125	128	15	IMAGE
synthetic_control	600	60	6	SIMULATED
FaceAll	2250	131	14	IMAGE

TABLE 1 – Jeux de données

Le tableau 2 présente une comparaison entre les deux algorithmes.

5.2 Comparaison avec k-shape et USLM

k-Shape et USLM sont deux algorithmes de clustering basés sur les u-shapelets pour les séries temporelles présentées dans [10]. Dans cette section, nous comparons l'indice Rand index obtenu par FOTS-UShapelet et celui obtenu par

Datasets	RI_SUSH	RI_FOTS
50words	0.811	0.877
Adiac	0.796	0.905
Beef	0.897	0.910
Car	0.708	0.723
CBF	0.578	0.909
Coffee	0.782	0.896
ECG200	0.717	0.866
FaceFour	0.859	0.910
FISH	0.775	0.899
Gun_Point	0.710	0.894
Lighting2	0.794	0.911
Lighting7	0.757	0.910
OliveOil	0.714	0.910
OSULeaf	0.847	0.905
SwedishLeaf	0.305	0.909
synthetic_control	0.723	0.899
FaceAll	0.907	0.908

TABLE 2 – Comparison du Rand Index de SUSH (RI_SUSH) et de FOTS-SUSH (RI_FOTS). Le meilleur Rand Index est en gras

k-Shape et USLM sur 5 jeux de données¹ (Tableau 3). Les résultats de k-Shape et USLM ont été précédemment rapportés dans [10]. Cette comparaison montre qu'en général, FOTS-UShapelet donne de meilleurs résultats que k-Shape et USLM.

TABLE 3 – Comparaison entre k-Shape, USLM et FOTS-UShapelet

Rand Index	k-Shape	USLM	FOTS-UShapelet
CBF	0.74	1	0.909
ECG200	0.70	0.76	0.866
Fac.F.	0.64	0.79	0.910
Lig2	0.65	0.80	0.911
Lig.7	0.74	0.79	0.910
OSU L.	0.66	0.82	0.905

5.3 Discussion

L'utilisation du score FOTS associé à l'algorithme de SU-Shapelet fait qu'il est possible de découvrir d'autres u-shapelets que ceux trouvés par la distance Euclidienne. Le FOTS-SUSH améliore les résultats de la classification des séries temporelles parce que le score FOTS prend en compte les propriétés intrinsèques de la série temporelle et est robuste à la présence d'incertitude. Cette amélioration est particulièrement significative lorsque le score FOTS est utilisé pour la classification non supervisée de séries temporelles contenant plusieurs petites oscillations. En effet, ces oscillations ne sont pas capturées par la distance euclidienne mais par le score FOTS dont le calcul est basé sur

1. Nous considérons 5 jeux de données car se sont les jeux de données pour lesquels nous avons les résultats des algorithmes k-shape et USLM.

la matrice d'autocorrélation. Cette observation est illustrée par le résultat obtenu sur jeu de données SwedishLeaf.

Analyse de la complexité en temps. ED peut être calculé en $\mathcal{O}(n)$ et le score FOTS est calculé en $\mathcal{O}(n^\omega)$, $\leq \omega \leq 3$ en raison de la complexité temporelle des décompositions des vecteurs propres [20]. Le calcul du score FOTS est alors plus coûteux que celui de ED. Cependant, son utilisation reste pertinente pour la recherche d' u-shapelets, car ils sont souvent de petite taille.

6 Conclusion et perspective

Le but de ce travail était de découvrir des u-shapelets sur des séries temporelles incertaines. Pour répondre à cette question, nous avons proposé un score de dissimilarité (FOTS) adapté à la comparaison de séries temporelles courtes, dont le calcul est basé sur la comparaison du vecteurs propres des matrices d'autocorrélation des séries temporelles. Ce score est robuste à la présence d'incertitude, il n'est pas très sensible aux changements transitoires, et il permet de capturer des relations complexes entre des séries temporelles telles que les oscillations et les tendances. Le score FOTS a été utilisé avec l'algorithme Scalable Un-supervised Shapelet Discovery pour la classification non supervisée de 17 jeux de données de la littérature et a montré une amélioration de la qualité du regroupement évalué à l'aide de l'indice Rand. En combinant les avantages de l'algorithme des u-shapelets, qui réduit les effets néfastes de l'incertitude, et les avantages du score FOTS, qui est robuste à la présence de l'incertitude, ce travail définit un cadre original pour la classification non supervisée de séries temporelles incertaines. Dans la perspective de ce travail, nous prévoyons d'utiliser le score FOTS pour la classification non supervisée floue de séries temporelles incertaines.

Remerciements

Nous remercions cordialement le Ministère Français de l'enseignement supérieur et de la recherche qui a financé ce travail et nous remercions les reviewers pour leurs commentaires et leurs suggestions qui ont aidé à améliorer la qualité de ce travail.

Références

- [1] G. B. Folland and A. Sitaram, "The uncertainty principle : a mathematical survey," *Journal of Fourier analysis and applications*, vol. 3, no. 3, pp. 207–238, 1997.
- [2] N. B. Rizvandi, J. Taheri, R. Moraveji, and A. Y. Zomaya, "A study on using uncertain time series matching algorithms for MapReduce applications," *Concurrency and Computation : Practice and Experience*, vol. 25, no. 12, pp. 1699–1718, aug 2013.
- [3] J. Hwang, Y. Kozawa, T. Amagasa, and H. Kitagawa, "GPU Acceleration of Similarity Search for Uncertain Time Series," in *2014 17th International*

Conference on Network-Based Information Systems. IEEE, sep 2014, pp. 627–632.

- [4] K. Rehfeld and J. Kurths, “Similarity estimators for irregular and age-uncertain time series,” *Climate of the Past*, vol. 10, no. 1, pp. 107–122, 2014.
- [5] M. Orang and N. Shiri, “An experimental evaluation of similarity measures for uncertain time series,” in *Proceedings of the 18th International Database Engineering & Applications Symposium on - IDEAS '14*. New York, New York, USA : ACM Press, 2014, pp. 261–264.
- [6] W. Wang, G. Liu, and D. Liu, “Chebyshev Similarity Match between Uncertain Time Series,” *Mathematical Problems in Engineering*, vol. 2015, pp. 1–13, 2015.
- [7] M. Orang and N. Shiri, “Correlation analysis techniques for uncertain time series,” *Knowledge and Information Systems*, vol. 50, no. 1, pp. 79–116, jan 2017.
- [8] L. Ulanova, N. Begum, and E. Keogh, “Scalable clustering of time series with u-shapelets,” in *Proceedings of the 2015 SIAM International Conference on Data Mining*. SIAM, 2015, pp. 900–908.
- [9] J. Zakaria, A. Mueen, and E. Keogh, “Clustering time series using unsupervised-shapelets,” in *Data Mining (ICDM), 2012 IEEE 12th International Conference on*. IEEE, 2012, pp. 785–794.
- [10] Q. Zhang, J. Wu, H. Yang, Y. Tian, and C. Zhang, “Unsupervised feature learning from time series,” in *IJCAI*, 2016, pp. 2322–2328.
- [11] M. Dallachiesa, B. Nushi, K. Mirylenka, and T. Palpanas, “Uncertain time-series similarity : return to the basics,” *Proceedings of the VLDB Endowment*, vol. 5, no. 11, pp. 1662–1673, 2012.
- [12] M. Orang and N. Shiri, “Improving performance of similarity measures for uncertain time series using preprocessing techniques,” in *Proceedings of the 27th International Conference on Scientific and Statistical Database Management - SSDBM '15*. New York, New York, USA : ACM Press, 2015, pp. 1–12.
- [13] K. Murthy and S. R. Sarangi, “Generalized notion of similarities between uncertain time series,” Mar. 26 2013, uS Patent 8,407,221.
- [14] J. Aßfalg, H.-P. Kriegel, P. Kröger, and M. Renz, “Probabilistic similarity search for uncertain time series,” in *SSDBM*. Springer, 2009, pp. 435–443.
- [15] M.-Y. Yeh, K.-L. Wu, P. S. Yu, and M.-S. Chen, “Proud : a probabilistic approach to processing similarity queries over uncertain data streams,” in *Proceedings of the 12th International Conference on Extending Database Technology : Advances in Database Technology*. ACM, 2009, pp. 684–695.
- [16] J. Hoffmann-Jørgensen and G. Pisier, “The law of large numbers and the central limit theorem in banach spaces,” *The Annals of Probability*, pp. 587–599, 1976.
- [17] S. Papadimitriou, J. Sun, and S. Y. Philip, “Local correlation tracking in time series,” in *Data Mining, 2006. ICDM'06. Sixth International Conference on*. IEEE, 2006, pp. 456–465.
- [18] R. Bhatia and T. Bhattacharyya, “A generalization of the Hoffman-Wielandt theorem,” *Linear Algebra and its Applications*, vol. 179, pp. 11–17, jan 1993.
- [19] Y. Chen, E. Keogh, B. Hu, N. Begum, A. Bagnall, A. Mueen, and G. Batista, “The ucr time series classification archive,” July 2015.
- [20] V. Y. Pan and Z. Q. Chen, “The complexity of the matrix eigenproblem,” in *Proceedings of the thirty-first annual ACM symposium on Theory of computing*. ACM, 1999, pp. 507–516.

ABClass : Une approche d'apprentissage multi-instances pour les séquences

Manel Zoghlami^{1,2}
Mondher Maddouri⁴

Sabeur Aridhi³
Engelbert Mephu Nguifo¹

¹ Université Clermont Auvergne, CNRS, LIMOS, BP 10125, 63173 Clermont Ferrand, France

² Université Tunis El Manar, Faculté des sciences de Tunis, LIPAH, 1060 Tunis, Tunisie

³ Université de Lorraine, CNRS, Inria, LORIA, F-54000 Nancy, France

⁴ Université de Jeddah, Faculté d'administration des affaires, BP 80327, 21589 Jeddah, KSA

manel.zoghlami@etu.uca.fr

Résumé

Dans le cas du problème de l'apprentissage multi-instances (MI) pour les séquences, les données d'apprentissage consistent en un ensemble de sacs où chaque sac contient un ensemble d'instances/séquences. Dans certaines applications du monde réel, comme la bioinformatique, comparer un couple aléatoire de séquences n'a aucun sens. En fait, chaque instance de chaque sac peut avoir une relation structurelle et/ou fonctionnelle avec d'autres instances dans d'autres sacs. Ainsi, la tâche de classification doit prendre en compte la relation entre les instances sémantiquement liées à travers les sacs. Dans cet article, nous présentons ABClass, une nouvelle approche de classification MI des séquences. Chaque séquence est représentée par un vecteur d'attributs extraits à partir de l'ensemble des instances qui lui sont liées. Pour chaque séquence du sac à prédire, un classifieur discriminant est appliqué afin de calculer un résultat de classification partiel. Ensuite, une méthode d'agrégation est appliquée afin de générer le résultat final. Nous avons appliqué ABClass pour résoudre le problème de la prédiction de la résistance aux rayonnements ionisants (RRI) chez les bactéries. Les résultats expérimentaux sont satisfaisants.

Mots clés

apprentissage multi-instances, séquences protéiques, prédiction de la résistance aux rayonnements ionisants chez les bactéries

Abstract

In Multiple Instance Learning (MIL) problem for sequence data, the learning data consist of a set of bags where each bag contains a set of instances/sequences. In some real world applications such as bioinformatics comparing a random couple of sequences makes no sense. In fact, each instance of each bag may have structural and/or functional relationship with other instances in other bags. Thus, the classification task should take into account the relation between semantically related instances across bags. In this

paper, we present ABClass, a novel MIL approach for sequence data classification. Each sequence is represented by one vector of attributes extracted from the set of related instances. For each sequence of the unknown bag, a discriminative classifier is applied in order to compute a partial classification result. Then, an aggregation method is applied in order to generate the final result. We applied ABClass to solve the problem of bacterial Ionizing Radiation Resistance (IRR) prediction. The experimental results were satisfactory.

Keywords

multiple instance learning, protein sequences, prediction of bacterial ionizing radiation resistance

1 Introduction

L'apprentissage multi-instances (MI) est une variante des méthodes d'apprentissage classiques qui peut être utilisée pour résoudre des problèmes dans lesquels les étiquettes sont affectées à des sacs, c'est-à-dire un ensemble d'instances, plutôt que des instances individuelles. Une hypothèse majeure de la plupart des méthodes d'apprentissage MI existantes est que chaque sac contient un ensemble d'instances qui sont distribuées indépendamment. De nombreuses applications du monde réel telles que la bioinformatique, la fouille de sites Web et la fouille de textes doivent traiter des données séquentielles (données sous forme de séquences). Lorsque le problème abordé peut être formulé comme un problème MI, chaque instance de chaque sac peut avoir une relation structurelle et/ou fonctionnelle avec d'autres instances dans d'autres sacs. Le problème que nous voulons résoudre dans ce travail est le problème d'apprentissage MI pour les séquences qui présentent des relations / dépendances à travers les sacs.

Ce travail a été initialement proposé pour résoudre le problème de la prédiction de la résistance aux rayonnements ionisants (RRI) chez les bactéries [21] [4]. L'objectif est d'apprendre un classifieur qui classe une bactérie soit comme bactérie résistante aux radiations ionisantes

(BRR) soit comme bactérie sensible aux radiations ionisantes (BSRI). Les BRR sont importantes en biotechnologie. Elles pourraient être utilisées pour la bioremédiation de déchets radioactifs et dans l'industrie thérapeutique [5] [10]. Cependant, un nombre limité de travaux d'apprentissage automatique a été proposé pour résoudre le problème de la prédiction de la BRR chez les bactéries. Ce problème pourrait être formalisé en tant que problème d'apprentissage MI : chaque bactérie est représentée par un ensemble de séquences protéiques. Les bactéries représentent les sacs et les séquences protéiques représentent les instances. En particulier, le contenu de chaque séquence protéique peut être différent d'une bactérie à une autre, par exemple, chaque sac contient la protéine appelée *endonuclease III*, mais elle est exprimée différemment d'un sac à un autre : il s'agit de protéines orthologues [7]. Afin d'apprendre l'étiquette d'une bactérie inconnue, comparer un couple aléatoire de séquences n'a pas de sens, il est plutôt préférable de comparer les séquences protéiques qui ont une relation/dépendance fonctionnelle : les protéines orthologues. Ce travail traite donc le problème d'apprentissage MI ayant les trois critères suivants : (1) les instances à l'intérieur des sacs sont des séquences, nous devons donc traiter le format de la représentation des données, (2) les instances peuvent avoir des dépendances entre les sacs et (3) toutes les instances à l'intérieur d'un sac contribuent à définir l'étiquette du sac.

L'hypothèse standard de l'apprentissage MI indique qu'un sac est positif si au moins une de ses instances est positive alors que dans chaque sac négatif toutes les instances sont négatives [6]. Ceci n'est pas garanti dans certains domaines donc des hypothèses alternatives ont été proposées [9]. Particulièrement, l'hypothèse standard n'est pas adaptée au problème relatif à la BRR car une instance positive n'est pas suffisante pour classer un sac comme positif. Nous optons plutôt à l'hypothèse collective [2] [9] : toutes les instances contribuent à la définition de l'étiquette du sac.

Nous proposons dans ce travail une formalisation du problème de l'apprentissage MI pour les séquences. Nous présentons aussi une approche naïve et une nouvelle approche appelée ABCClass. ABCClass effectue d'abord une étape de prétraitement des séquences en entrée qui consiste à extraire des motifs à partir de chaque ensemble de séquences liées. Ces motifs seront utilisés comme attributs pour construire une matrice binaire pour chaque ensemble où chaque ligne correspond à une séquence. Ensuite, un classifieur discriminant est appliqué aux séquences d'un sac inconnu afin de prédire son étiquette. Nous décrivons l'algorithme de notre approche et nous présentons une étude expérimentale en l'appliquant au problème de la prédiction de la BRR chez les bactéries.

Le reste de cet article est organisé comme suit. La section 2 définit le problème de l'apprentissage MI pour les séquences. Dans la section 3, nous présentons un aperçu de quelques travaux relatifs aux problèmes MI. Dans la section 4, nous décrivons l'approche que nous proposons pour

la classification MI des séquences ayant des dépendances à travers les sacs. Dans la section 5, nous décrivons notre environnement expérimental et nous discutons les résultats obtenus.

2 Contexte

Dans cette section, nous présentons les notions de base relatives à l'apprentissage MI pour les séquences. Nous décrivons d'abord la terminologie et la formulation de notre problème. Ensuite, nous présentons un cas d'utilisation simple qui sert d'exemple illustratif tout au long de ce document.

2.1 Formulation du problème

Une séquence est une liste ordonnée d'événements. Un événement peut être représenté comme une valeur symbolique, une valeur numérique, un vecteur de valeurs ou un type de données complexe [19]. Il existe de nombreux types de séquences comme les séquences symboliques, les séries temporelles simples et les séries temporelles multivariées [19]. Dans notre travail, nous nous intéressons aux séquences symboliques puisque les séquences protéiques sont décrites à l'aide de symboles (acides aminés). On note Σ l'*alphabet* défini par un ensemble fini de caractères ou de symboles. Une séquence symbolique simple est donc définie comme une liste ordonnée de symboles de Σ .

Soit BD une base de données d'apprentissage qui contient un ensemble de n sacs étiquetés : $BD = \{(B_i, Y_i), i = 1, 2, \dots, n\}$ où $Y_i = \{-1, 1\}$ est l'étiquette du sac B_i . Les instances de B_i sont des séquences et sont notées B_{ij} . Formellement $B_i = \{B_{ij}, j = 1, 2, \dots, m_{B_i}\}$, où m_{B_i} est le nombre total d'instances dans le sac B_i . Nous notons que les sacs ne contiennent pas nécessairement le même nombre d'instances. Le problème étudié dans ce travail consiste à apprendre un classifieur MI à partir de BD . Étant donné un sac inconnu $Q = \{Q_k, k = 1, 2, \dots, q\}$, où q est le nombre total d'instances dans Q , le classifieur doit utiliser les séquences dans ce sac et celles dans chaque sac de BD afin de prédire l'étiquette de Q .

Nous notons qu'il existe une relation notée $\mathcal{R}$ qui relie les instances à travers les différents sacs. Elle est définie en fonction du domaine d'application. Pour représenter cette relation, nous optons pour une représentation à base d'indice. Nous notons que cette notation ne signifie pas que les instances sont ordonnées. En fait, une étape de prétraitement attribue un indice aux instances de chaque sac selon la façon suivante : chaque instance B_{ij} d'un sac B_i est reliée par $\mathcal{R}$ à l'instance B_{hj} d'un autre sac B_h dans BD . Une instance peut ne pas avoir d'instances correspondantes dans certains sacs, c'est-à-dire qu'une séquence est liée à zéro ou une séquence par sac. La relation $\mathcal{R}$ pourrait être généralisée pour traiter les problèmes où chaque instance a plus d'une instance cible par sac. La notation d'indice telle que décrite précédemment ne sera pas appropriée dans ce cas.

2.2 Exemple illustratif

Afin d'illustrer notre approche, nous nous appuyons sur l'exemple suivant. Soit $\Sigma = \{A, B, \dots, Z\}$ un alphabet. Soit $BD = \{(B_1, +1), (B_2, +1), (B_3, -1), (B_4, -1), (B_5, -1)\}$ une base d'apprentissage contenant 5 sacs (B_1 et B_2 sont des sacs positifs, B_3, B_4 et B_5 sont des sacs négatifs). Initialement, les sacs contiennent les séquences suivantes :

$B_1 = \{\text{ABMSCD}, \text{EFNOGH}, \text{RUV R}\}$

$B_2 = \{\text{CCGHDDEF}, \text{EABZQCD}\}$

$B_3 = \{\text{GHW MY}, \text{ACDXYZ}\}$

$B_4 = \{\text{ABIJYZ}, \text{KLSSO}, \text{EFYRTAB}\}$

$B_5 = \{\text{EFFVGH}, \text{KLSNAB}\}$

Nous utilisons d'abord la relation inter-sacs $\mathfrak{R}$ pour représenter les instances reliées en utilisant la notation d'indice décrite précédemment.

$$B_1 = \begin{cases} B_{11} = \text{ABMSCD} \\ B_{12} = \text{EFNOGH} \\ B_{13} = \text{RUV R} \end{cases} \quad B_2 = \begin{cases} B_{21} = \text{EABZQCD} \\ B_{22} = \text{CCGHDDEF} \end{cases}$$

$$B_3 = \begin{cases} B_{31} = \text{ACDXYZ} \\ B_{32} = \text{GHW MY} \end{cases} \quad B_4 = \begin{cases} B_{41} = \text{ABIJYZ} \\ B_{42} = \text{EFYRTAB} \\ B_{43} = \text{KLSSO} \end{cases}$$

$$B_5 = \begin{cases} B_{52} = \text{EFFVGH} \\ B_{53} = \text{KLSNAB} \end{cases}$$

Le but ici est de prédire l'étiquette d'un sac inconnu $Q = \{Q_1, Q_2, Q_3\}$ où :

$$Q = \begin{cases} Q_1 = \text{ABWXCD} \\ Q_2 = \text{EFXYGHN} \\ Q_3 = \text{KLOF} \end{cases}$$

3 Travaux existants

Plusieurs algorithmes d'apprentissage MI ont été proposés incluant Diverse Density (DD) [15], MI-SVM [3], Citation-kNN [18] et MILKDE [8]. Un état de l'art des approches d'apprentissage MI avec une étude comparative pourrait être trouvé dans [1] et [2].

L'idée principale de l'approche DD [15] est de trouver les points qui sont proches d'au moins une instance de chaque sac positif et qui sont loin des instances des sacs négatifs. On cherche ensuite le point optimal selon une mesure définie par les auteurs et qui porte le même nom que l'algorithme. Cette mesure prend en considération le nombre de sacs positifs ayant des instances proches du point en question et la distance entre ce point et les instances négatives. L'approche MI-SVM [3] est une adaptation des machines à vecteurs de support au problème d'apprentissage MI. Elle reformule la maximisation de la marge des sacs en prenant

en considération les contraintes du MI. Pour un sac positif, seule l'instance ayant la plus grande marge a un impact sur l'apprentissage. Les autres instances sont ignorées. MISMO utilise l'algorithme SMO [16] basé sur l'apprentissage par machines à vecteurs de support en conjonction avec un noyau MI [11]. L'algorithme Citation-kNN [18] est basé sur la règle des plus proches voisins et sur le concept de citation et de référence. Un sac est étiqueté non seulement selon ses voisins (les références), mais aussi selon les sacs qui le reconnaissent comme leur voisin (les citeurs). Dans [8], les auteurs présentent l'algorithme MILKDE qui se base sur l'identification des instances les plus représentatives dans chaque sac positif en utilisant un calcul de vraisemblance. Dans [20], les auteurs traitent l'apprentissage MI sur des données structurées. Ils décrivent trois scénarios des relations existantes entre les données : I-MILSD où les relations sont disponibles au niveau des instances, B-MILSD où les relations sont au niveau des sacs et BI-MILSD où les relations sont disponibles au niveau des instances et des sacs.

L'application des algorithmes MI présentés ci-dessus sur des sacs de séquences entraîne deux problèmes. Le premier problème réside dans la représentation des données à traiter. Elles sont représentées dans un format attribut-valeur. Dans le cas des séquences, la technique la plus utilisée pour transformer les données en un format attribut-valeur consiste à extraire des motifs qui servent comme attributs. Nous notons que trouver une description uniforme de toutes les instances en utilisant un ensemble de motifs n'est pas toujours une tâche facile. En apprentissage MI, cela pourrait conduire à une matrice énorme et creuse. Par exemple, dans [3], l'évaluation empirique est effectuée sur la base de données TREC9 pour la catégorisation des documents. Les termes sont utilisés pour présenter le texte. On obtient alors une matrice creuse et de grande dimension. Dans [17] et [20], un ensemble de séquences protéiques a été utilisé dans l'évaluation empirique. L'objectif est d'identifier les protéines dites *Trx-fold*. Chaque séquence est considérée comme un sac et certaines de ses sous-séquences sont considérées comme des instances. Ces sous-séquences sont alignées et représentées en utilisant 8 attributs numériques [13] [17]. Le deuxième problème avec les algorithmes présentés est qu'ils ne traitent pas les relations inter-sacs qui peuvent exister entre les instances, à l'exception des algorithmes dans [20]. Dans [20], le score d'alignement est utilisé pour identifier les relations entre les protéines : si le score entre une paire de protéines dépasse 25, alors les auteurs considèrent qu'il existe un lien entre elles. Seul l'algorithme B-MILSD qui traite l'information au niveau du sac a été utilisé dans l'étude expérimentale. Dans un travail antérieur, nous avons proposé l'algorithme MIL-ALIGN [4] qui traite le problème de la prédiction de la RRI chez les bactéries. Il utilise une technique d'alignement pour discriminer les séquences puis applique une méthode d'agrégation pour générer le résultat de prédiction final. Dans ABCClass, nous représentons les sé-

quences en utilisant un format attribut-valeur qui prend en compte les dépendances entre les instances liées à travers les sacs.

4 Approches d'apprentissage MI pour les séquences

Dans cette section, nous présentons d'abord l'approche naïve pour traiter le problème de l'apprentissage MI pour les séquences ayant des relations à travers les sacs. Ensuite, nous présentons notre approche nommée ABClass.

4.1 Approche MI naïve pour les séquences

La méthode la plus simple qu'on peut utiliser pour résoudre le problème de l'apprentissage MI pour les séquences est d'utiliser des classifieurs MI standards. Cependant, les algorithmes MI couramment utilisés nécessitent une description en format attribut-valeur uniforme pour toutes les instances des différents sacs. L'approche naïve contient deux étapes. La première est une étape de prétraitement qui transforme l'ensemble des séquences en une matrice attribut-valeur où chaque ligne correspond à une séquence et chaque colonne correspond à un attribut. La deuxième étape consiste à appliquer un classifieur MI existant. La Figure 1 illustre l'approche naïve pour l'apprentissage MI pour les séquences. La technique la plus utilisée pour transformer les séquences en un format attribut-valeur consiste à extraire des motifs qui seront utilisés comme attributs.

Nous notons que trouver une description uniforme de toutes les instances utilisant un ensemble de motifs n'est pas toujours une tâche facile. Puisque notre approche naïve prend en compte les relations entre les instances à travers les sacs, l'étape de prétraitement extrait les motifs à partir de chaque ensemble d'instances reliées. L'union de ces motifs est ensuite utilisée comme attributs pour construire une matrice attribut-valeur où chaque ligne correspond à une séquence. La présence ou l'absence d'un attribut dans une séquence est notée respectivement par 1 ou 0. Il est utile de mentionner que seul un sous-ensemble des attributs utilisés est représentatif pour chaque séquence. Par conséquent, nous pouvons avoir une grande matrice creuse.

Nous appliquons l'approche naïve sur notre exemple illustratif. Nous supposons que les attributs sont des sous-séquences (longueur minimale = 2) qui se trouvent au moins dans deux instances. Soit $listeMotifs_1 = \{AB, CD, YZ\}$ la liste des motifs extraits à partir des instances $\{B_{i1}, i = 1, \dots, 4\}$. $listeMotifs_2 = \{EF, GH\}$ est la liste des motifs extraits à partir de $\{B_{i2}, i = 1, \dots, 5\}$ et $listeMotifs_3 = \{KL\}$ est la liste des motifs extraits à partir de $\{B_{i3}, i \in \{1, 4, 5\}\}$. L'union de $listeMotifs_1$, $listeMotifs_2$ et $listeMotifs_3$ produit la liste $listeMotifs = \{AB, CD, YZ, EF, GH, KL\}$. Afin de coder les séquences de la base d'apprentissage, nous générons la matrice attribut-valeur suivante notée M :

$$M = \begin{pmatrix} \begin{array}{cccccc|cccccc|cccccc} & \text{Instance1} & & & & & & \text{Instance2} & & & & & & \text{Instance3} & & & & & \\ 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & \\ 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 1 & 0 & - & - & - & - & - & - \\ 0 & 1 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 & 0 & - & - & - & - & - & - \\ 1 & 0 & 1 & 0 & 0 & 0 & 1 & 0 & 0 & 1 & 0 & 0 & 0 & 0 & 0 & 0 & 0 & 1 \\ - & - & - & - & - & - & 0 & 0 & 0 & 1 & 1 & 0 & 1 & 0 & 0 & 0 & 0 & 1 \end{array} \begin{array}{l} B_1 \\ B_2 \\ B_3 \\ B_4 \\ B_5 \end{array} \end{pmatrix}$$

Le taux d'éparsité de M est 77.2%.

4.2 Approche proposée : ABClass

Afin d'éviter l'utilisation d'un grand vecteur d'attributs pour décrire les séquences, nous présentons ABClass (pour Across Bag Sequences Classification), une nouvelle approche qui prend en compte les relations entre les instances à travers les sacs. La Figure 2 illustre le principe de ABClass. Durant la phase d'apprentissage, chaque ensemble d'instances reliées sera représenté par son propre vecteur de motifs. Cela réduit le nombre d'attributs qui ne sont pas représentatifs de la séquence traitée. En appliquant un classifieur classique (mono-instance) sur chaque vecteur d'attributs, un modèle de classification est construit. Durant la phase de prédiction, des résultats de prédiction partiels sont produits pour chaque instance du sac inconnu. Ces résultats sont ensuite agrégés pour avoir le résultat final. Lors de l'exécution de l'algorithme, nous allons utiliser les va-

riables suivantes :

- Une matrice M pour stocker les données codées de la base d'apprentissage.
- Un vecteur QV pour stocker les données codées du sac à prédire.
- Un vecteur PV pour stocker les résultats de prédiction partiels.

ABClass est décrit dans l'Algorithme 1. La fonction *SéqLiéesEntreSacs* regroupe les instances liées à travers les sacs dans des listes. Informellement, les principales étapes de l'algorithme ABClass sont les suivantes :

1. Pour chaque séquence Q_k du sac inconnu Q , les instances reliées à travers les sacs sont regroupées dans une liste (lignes 1 et 2).
2. L'algorithme extrait des motifs de la liste des instances regroupées. Ces motifs sont utilisés pour coder les instances afin de créer un modèle discriminant (lignes 3 à 5).

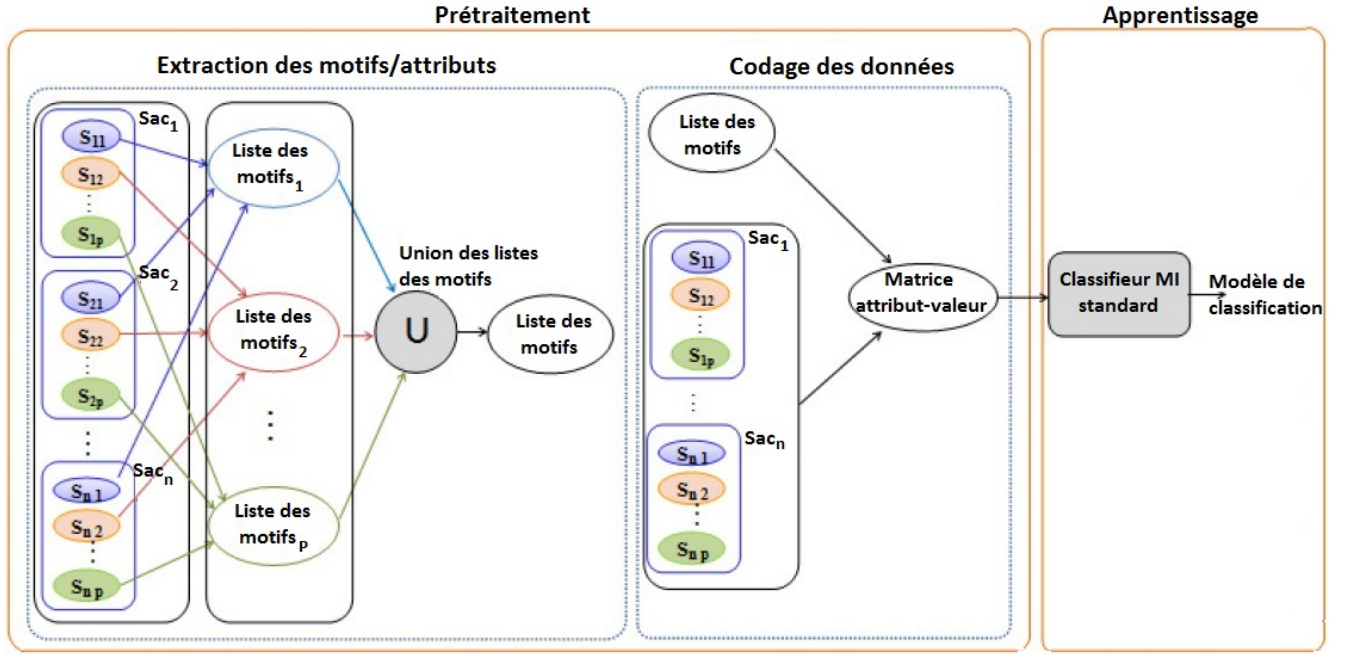


FIGURE 1 – Vue d'ensemble de l'approche MI naïve.

3. *ABClass* utilise les motifs extraits pour représenter l'instance Q_k du sac inconnu dans un vecteur QV_k , puis le compare avec le modèle correspondant. Le résultat de la comparaison est stocké dans le k^{ime} élément du vecteur PV (lignes 6 et 7).
4. Une méthode d'agrégation est appliquée à PV pour calculer le résultat final P (ligne 9) qui consiste en une étiquette positive ou négative.

Algorithm 1 L'algorithme *ABClass*

Entrée: Base d'apprentissage $BD = \{(B_i, Y_i) | i = 1, 2, \dots, n\}$, Sac inconnu $Q = \{Q_k | k = 1, 2, \dots, q\}$

Sortie: Résultat de prédiction P

- 1: **Pour** chaque $Q_k \in Q$ **Faire**
 - 2: $listeSéqLiées_k \leftarrow SéqLiéesEntreSacs(k, DB)$
 - 3: $listeMotifs_k \leftarrow extraireMotifs(listeSéqLiées_k)$
 - 4: $M_k \leftarrow coderDonnées(listeMotifs_k, listeSéqLiées_k)$
 - 5: $Modèle_k \leftarrow générerModèle(M_k)$
 - 6: $QV_k \leftarrow coderDonnées(listeMotifs_k, Q_k)$
 - 7: $PV_k \leftarrow appliquerModèle(QV_k, Modèle_k)$
 - 8: **Fin**
 - 9: $P \leftarrow Agrégation(PV)$
 - 10: **Retourner** P
-

Nous appliquons l'approche *ABClass* sur notre exemple illustratif. Puisque le sac à prédire contient 3 instances Q_1 , Q_2 et Q_3 , nous avons besoin de 3 itérations suivies d'une étape d'agrégation.

Itération 1 : L'algorithme regroupe l'ensemble des instances reliées et extrait les motifs correspondants.

$$listeSéqLiées_1 = \{B_{11}, B_{21}, B_{31}, B_{41}\}$$

$$listeMotifs_1 = \{AB, CD, YZ\}$$

Ensuite, il génère la matrice attribut-valeur M_1 décrivant les séquences liées à Q_1 .

$$M_1 = \begin{pmatrix} AB & CD & YZ \\ 1 & 1 & 0 \\ 1 & 1 & 0 \\ 0 & 1 & 1 \\ 1 & 0 & 1 \end{pmatrix} \begin{matrix} B_{11} \\ B_{21} \\ B_{31} \\ B_{41} \end{matrix}$$

Le pourcentage d'éparsité de la matrice M_1 est réduit à 33% car il n'est pas nécessaire d'utiliser les motifs extraits des instances $\{B_{i2}, i = 1, \dots, 5\}$ et $\{B_{i3}, i \in \{1, 4, 5\}\}$ pour décrire les instances $\{B_{i1}, i = 1, \dots, 4\}$. Un modèle est ensuite créé en utilisant les données codées et un vecteur QV_1 est généré pour décrire Q_1 .

$$QV_1 = \begin{pmatrix} 1 \\ 1 \\ 0 \end{pmatrix}$$

En appliquant le modèle au vecteur QV_1 , on obtient le premier résultat de prédiction partiel et on le stocke dans le vecteur PV .

$$PV_1 \leftarrow appliquerModèle(QV_1, Modèle_1)$$

Itération 2 : La deuxième itération concerne la deuxième instance Q_2 du sac à prédire. Nous appliquons les mêmes instructions que celles décrites dans la première itération.

$$listeSéqLiées_2 = \{B_{21}, B_{22}, B_{32}, B_{42}, B_{52}\}$$

$$listeMotifs_2 = \{EF, GH\}$$

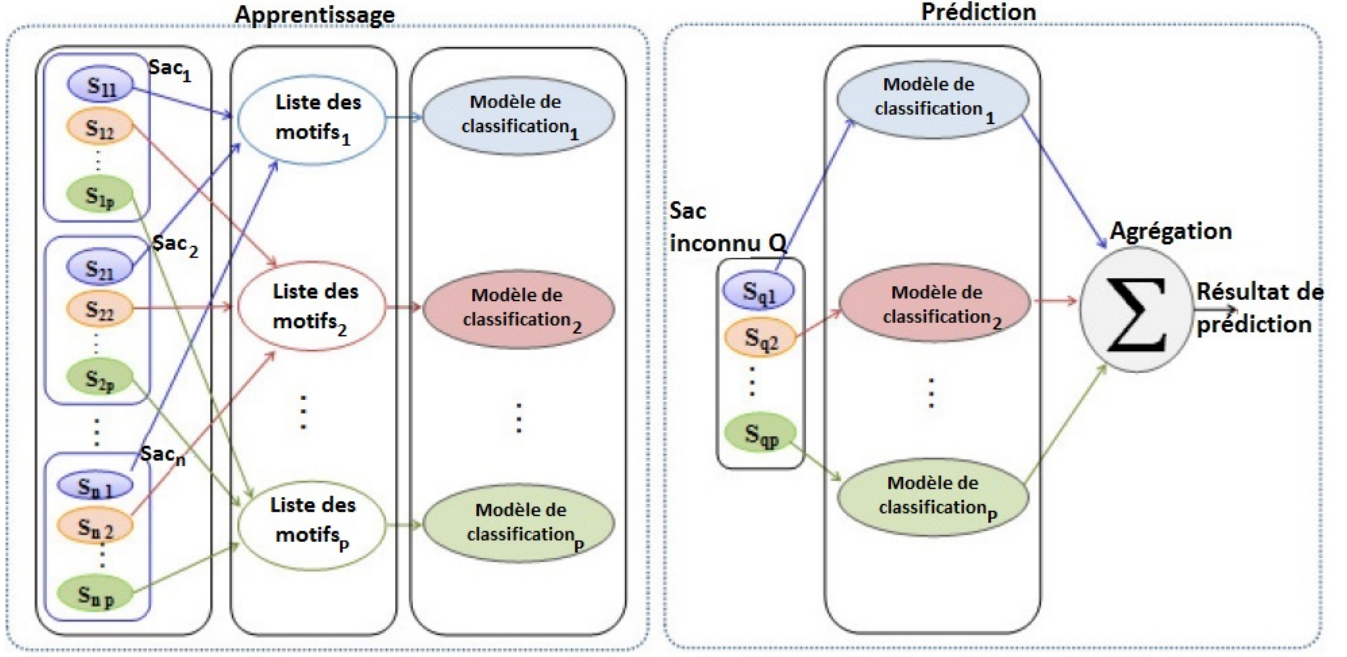


FIGURE 2 – Vue d'ensemble de l'approche ABClass.

$$M_2 = \begin{pmatrix} EF & GH \\ 1 & 1 \\ 1 & 1 \\ 0 & 1 \\ 1 & 0 \\ 1 & 1 \end{pmatrix} \begin{matrix} B_{12} \\ B_{22} \\ B_{32} \\ B_{42} \\ B_{52} \end{matrix}$$

$$QV_2 = \begin{pmatrix} 1 \\ 1 \end{pmatrix}$$

$$PV_2 \leftarrow \text{appliquerModèle}(QV_2, \text{Modèle}_2)$$

Itération 3 : Seuls les sacs B_1 , B_4 et B_5 ont des instances reliées à Q_3 .

$$\begin{aligned} \text{listeSéqLiées}_3 &= \{B_{13}, B_{43}, B_{53}\} \\ \text{listeMotifs}_3 &= \{KL\} \end{aligned}$$

$$M_3 = \begin{pmatrix} KL \\ 0 \\ 1 \\ 1 \end{pmatrix} \begin{matrix} B_{13} \\ B_{43} \\ B_{53} \end{matrix}$$

$$QV_3 = (1)$$

$$PV_3 \leftarrow \text{appliquerModèle}(QV_3, \text{Modèle}_3)$$

L'étape d'agrégation est finalement utilisée pour générer la décision finale de prédiction en utilisant les résultats partiels. Nous optons pour le vote majoritaire.

5 Etude expérimentale

Nous appliquons ABClass et l'approche naïve au problème de la prédiction de la RRI chez les bactéries qui peut être formulé comme un problème MI pour séquences. Les bactéries représentent les sacs et la structure primaire des protéines de réparation de l'ADN représentent les séquences. Une bactérie inconnue est classée comme BRRI ou BSRI. Pour nos tests, nous avons utilisé la base de données décrite dans [4]. Cet ensemble de données comprend 28 sacs (14 BRRI et 14 BSRI). Chaque bactérie/sac contient 25 à 31 instances qui correspondent aux protéines. Nous avons utilisé des classificateurs implémentés dans l'outil de fouille de données WEKA [12] afin de tester les approches proposées.

5.1 Protocole expérimental

Nous utilisons la technique d'évaluation Leave-One-Out (LOO) dans nos expérimentations. Afin d'évaluer l'approche naïve et l'approche ABClass, nous codons d'abord les séquences protéiques de chaque sac en utilisant un ensemble de motifs générés par une méthode d'extraction de motifs existante. Nous utilisons la méthode DMS [14] pour l'extraction des motifs. DMS permet de construire des motifs pouvant discriminer une famille de protéines d'une autre. Elle identifie d'abord les motifs dans les séquences protéiques. Ensuite, les motifs extraits sont filtrés afin de ne conserver que les motifs discriminants et minimaux. Une sous-chaîne est considérée comme discriminante entre la famille F et les autres familles si elle apparaît dans F significativement plus que dans les autres familles. DMS extrait des motifs selon deux seuils α et β où α est le taux minimum d'occurrences des motifs dans les séquences d'une

TABLE 1 – Éparsité de la matrice attribut-valeur utilisée dans l’approche naïve.

Paramètre d’extraction des motifs	Nombre total des motifs	Éparsité (%)
S1	671	73.9
S2	1490	73.8
S3	4562	85.7
S4	8077	91.1

famille F et β est le taux maximum d’occurrences des motifs dans toutes les séquences sauf celles de la famille F . Dans ce qui suit, nous présentons les paramètres d’extraction des motifs utilisés en fonction des valeurs de α et β :

- **Le paramètre S1** : ($\alpha = 1$ et $\beta = 0.5$) : pour extraire des motifs fréquents avec une discrimination moyenne.
- **Le paramètre S2** : ($\alpha = 1$ et $\beta = 1$) : pour extraire des motifs fréquents et non discriminants.
- **Le paramètre S3** : ($\alpha = 0.5$ et $\beta = 1$) : pour extraire des motifs non discriminants avec des fréquences moyennes.
- **Le paramètre S4** : ($\alpha = 0$ et $\beta = 1$) : pour extraire des motifs non fréquents et non discriminants.
- **Le paramètre S5** : ($\alpha = 1$ et $\beta = 0$) : pour extraire des motifs fréquents et strictement discriminants.

5.2 Résultats

La Table 1 représente pour chaque paramètre d’extraction le nombre de motifs (longueur minimale = 3) extraits à partir de chaque ensemble de séquences de protéines orthologues. Pour le paramètre S5 ($\alpha = 1$ et $\beta = 0$), aucun motif fréquent et strictement discriminant n’a été trouvé pour la plupart des protéines. C’est pourquoi nous n’allons pas utiliser ces valeurs de α et β pour le reste des expérimentations. Nous notons que le nombre de motifs extraits augmente pour les valeurs élevées de β et les valeurs faibles de α . Comme présenté dans la Table 1, le nombre de motifs non fréquents et non discriminants est très élevé. Afin de coder les données dans l’approche naïve, on utilise l’union des motifs comme attributs. Par conséquent, la matrice attribut-valeur est grande et creuse. Nous montrons dans la Table 1 le taux d’éparsité de la matrice qui mesure le pourcentage des éléments nuls par rapport au nombre total des éléments. L’éparsité est généralement proportionnelle au nombre de motifs utilisés. Par exemple, Elle va de 73.9% avec 671 motifs à 91.1% avec 8077 motifs.

La Figure 3 montre les taux de bonne classification obtenus en appliquant l’approche naïve et l’approche ABCClass. Elle montre l’impact de l’ensemble des motifs utilisés dans l’étape de prétraitement sur les résultats de prédiction. Par exemple, en utilisant le classifieur MISVM, la précision varie de 53.5% à 78.5%. Bien que les motifs extraits en utilisant le paramètre S1 soient discriminants, l’approche naïve ne fournit pas globalement de bons résultats pour ce

paramètre. La raison pourrait être que le nombre de motifs discriminants pour certaines protéines est très faible (limité à 10 motifs au maximum). En utilisant l’approche naïve, le meilleur résultat est fourni par le classifieur MISMO. Les résultats des autres classifieurs MI dépendent des motifs utilisés. La plupart d’entre eux fournissent un bon résultat en utilisant le paramètre S3 (motifs non discriminants avec une fréquence moyenne). Le nombre de motifs extraits par protéine en utilisant ce paramètre est compris entre 228 et 1505. Ce nombre de motifs est acceptable pour coder une séquence protéique.

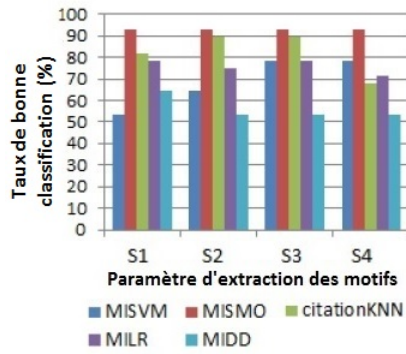
L’approche ABCClass fournit globalement de bons résultats puisque le taux de bonne classification le plus faible est 89.2%. Cela montre que notre approche est efficace. Le meilleur résultat est atteint en utilisant le classifieur J48 et les paramètres d’extraction de motifs S3 et S4. En utilisant ces deux paramètres, un grand nombre de motifs non discriminants est extrait. La Table 2 représente le taux des modèles de classification qui contribuent à prédire la bonne classe de chaque bactérie en utilisant l’approche ABCClass. Nous présentons ce taux pour les deux paramètres d’extraction de motifs qui fournissent les meilleurs résultats, à savoir S3 et S4. Le taux des modèles réussis pour B1, B11 et B15 est marqué en gras parce que ces trois bactéries génèrent toujours des taux faibles comparés au taux des autres bactéries. Nous notons que les résultats sont similaires à ceux trouvés dans [4]. Ces résultats peuvent aider à comprendre certaines caractéristiques des bactéries étudiées. En particulier *M. radiotolerans* (B11) et *B. abortus* (B15) présentent les taux les plus bas. Une explication biologique possible est fournie dans [4] et [21].

6 Conclusion

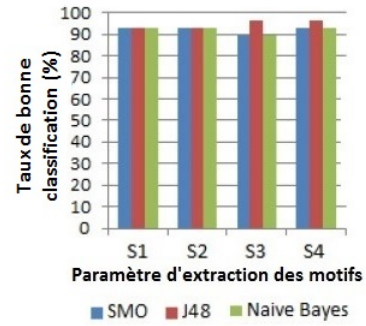
Dans cet article, nous avons abordé le problème d’apprentissage MI dans le cas où les instances sont des séquences. Nous nous sommes concentrés sur les données qui présentent des dépendances entre les instances des différents sacs. Nous avons décrit notre nouvelle approche nommée ABCClass et nous l’avons appliquée au problème de la prédiction de la RRI chez les bactéries. Dans l’étude expérimentale, nous avons montré que l’approche proposée est efficace. Dans le futur travail, nous étudierons comment utiliser la connaissance du domaine afin d’améliorer l’efficacité de notre algorithme. Nous voulons spécifiquement définir des poids pour les séquences dans la phase d’apprentissage en utilisant la connaissance du domaine.

Remerciements

Ce travail a été partiellement soutenu par le projet franco-tunisien : Direction Générale de la Recherche Scientifique en Tunisie (DGRST) / Centre National de la Recherche Scientifique en France (CNRS) [IRRB11 / R-14-09], par la Région française d’Auvergne et par la Fédération de Recherche en Environnement [UBP / CNRSFR-3467].



(a) L'approche naïve



(b) L'approche ABCClass

FIGURE 3 – Résultats de la classification en utilisant l'approche naïve et l'approche ABCClass.

TABLE 2 – Taux des modèles de classification réussis en utilisant l'approche ABCClass et la méthode d'évaluation LOO

ID de la bactérie	paramètre d'extraction des motifs S3			paramètre d'extraction des motifs S4		
	SMO	J48	Naive Bayes	SMO	J48	Naive Bayes
B1	44	60	36	60	64	68
B2	100	100	100	100	100	100
B3	100	90.3	100	100	90.3	100
B4	100	96.6	100	100	93.3	100
B5	100	90	100	100	90	100
B6	100	83.3	100	100	83.3	100
B7	100	93.5	100	100	93.5	100
B8	96.5	96.5	96.5	100	93.1	100
B9	100	84	100	100	84	100
B10	100	82.1	92.8	100	82.1	100
B11	17.8	50	17.8	21.4	50	35.7
B12	100	92.8	96.4	100	92.8	100
B13	88.8	66.6	70.3	88.8	66.6	77.7
B14	90	73.3	100	93.3	70	96.6
B15	3.5	32.1	35.7	0	32.1	14.2
B16	100	96.6	96.6	100	96.6	100
B17	96.2	96.2	96.2	96.2	96.2	96.2
B18	100	100	100	100	100	100
B19	100	100	100	100	100	100
B20	89.6	62	96.5	82.7	62	86.2
B21	96.5	82.7	96.5	93.1	82.7	93.1
B22	100	100	96.6	100	96.6	100
B23	100	96.7	93.5	100	96.7	100
B24	100	100	93.5	100	100	100
B25	100	96.7	93.5	100	96.7	100
B26	100	100	93.5	100	100	100
B27	100	100	100	100	100	100
B28	96.6	100	96.6	96.6	93.3	96.6

Références

- [1] E. Alpaydin, V. Cheplygina, M. Loog, and D. M. Tax, Single-vs. multiple-instance classification, *Pattern Recognition*, Vol. 48, pp. 2831–2838, 2015.
- [2] J. Amores, Multiple instance classification : Review, taxonomy and comparative study, *Artificial Intelligence*, Vol. 201, pp. 81–105, 2013.
- [3] S. Andrews, I. Tsochantaridis, and T. Hofmann, Support vector machines for multiple-instance learning. *Advances in Neural Information Processing Systems*, pp. 561–568, MIT Press, Cambridge, MA, 2003.

- [4] S. Aridhi, H. Sghaier, M. Zoghliami, M. Maddouri, and E. Mephu Nguifo, Prediction of ionizing radiation resistance in bacteria using a multiple instance learning model, *Journal of Computational Biology*, Vol. 23, pp. 10–20, 2016.
- [5] H. Brim, A. Venkateswaran, H. M. Kostandarithes, J. K. Fredrickson, and M. J. Daly, Engineering *Deinococcus Geothermalis* for bioremediation of high-temperature radioactive waste environments, *Applied and environmental microbiology*, Vol. 69, pp. 4575–4582, 2003.
- [6] T. G. Dietterich, R. H. Lathrop, and T. Lozano-Pérez, Solving the multiple instance problem with axis-parallel rectangles, *Artificial Intelligence*, Vol. 89, pp. 31–71, 1997.
- [7] G. Fang, N. Bhardwaj, R. Robilotto, and M. B. Gerstein, Getting started in gene orthology and functional analysis, *PLoS computational biology*, Vol. 6, e1000703, 2010.
- [8] A. W. Faria, F. G. F. Coelho, A. Silva, H. Rocha, G. Almeida, A. P. Lemos, and A. P. Braga, MILKDE : A new approach for multiple instance learning based on positive instance selection and kernel density estimation, *Engineering Applications of Artificial Intelligence*, Vol. 59, pp. 196–204, 2017.
- [9] J. Foulds and E. Frank, A review of multi-instance learning assumptions, *The Knowledge Engineering Review*, Vol. 25, pp. 1–25, 2010.
- [10] P. Gabani and O. V. Singh, Radiation-resistant extremophiles and their potential in biotechnology and therapeutics, *Applied microbiology and biotechnology*, Vol. 97, pp. 993–1004, 2013.
- [11] T. Gärtner, P. A. Flach, A. Kowalczyk, and A. J. Smola, Multi-instance kernels, *In Proceedings of the 19th International Conference on Machine Learning*, 2002.
- [12] M. Hall, E. Frank, G. Holmes, B. Pfahringer, P. Reutemann, and I. H. Witten, The weka data mining software : an update, *ACM SIGKDD explorations newsletter*, Vol. 11, pp. 10–18, 2009.
- [13] J. Kim, E. N. Moriyama, C. G. Warr, P. J. Clyne, and J. R. Carlson, Identification of novel multitransmembrane proteins from genomic databases using quasi-periodic structural properties, *Bioinformatics*, Vol. 16, pp. 767–775, 2000.
- [14] M. Maddouri and M. Elloumi, Encoding of primary structures of biological macromolecules within a data mining perspective, *Journal of Computer Science and Technology*, Vol. 19, pp. 78–88, 2004.
- [15] O. Maron and T. L. Pérez, A framework for multiple-instance learning, *Advances in Neural Information Processing Systems*, Vol. 10, pp. 570–576, 1998.
- [16] J. C. Platt, Fast training of support vector machines using sequential minimal optimization, *Advances in kernel methods*, pp. 185–208, 1999.
- [17] Q. Tao, S. Scott, N. Vinodchandran, and T. T. Osugi, SVM-based generalized multiple-instance learning via approximate box counting, *In Proceedings of the 21st international conference on Machine learning*, Vol. 10, pp. 799–806, 2004.
- [18] J. Wang and J. D. Zucker, Solving multipleinstance problem : A lazy learning approach, *In Proceedings of the 17th International Conference on Machine Learning*, pp. 1119–1125, 2000.
- [19] Z. Xing, J. Pei, and E. Keogh, A brief survey on sequence classification, *ACM SIGKDD Explorations Newsletter*, Vol. 12, pp. 40–48, 2010.
- [20] D. Zhang, Y. Liu, L. Si, J. Zhang, and R. D. Lawrence, Multiple instance learning on structured data, *Advances in Neural Information Processing Systems*, pp. 145–153, 2011.
- [21] M. Zoghliami, S. Aridhi, M. Maddouri, and E. Mephu Nguifo, An overview of in silico methods for the prediction of ionizing radiation resistance in bacteria, *Ionizing Radiation : Advances in Research and Applications, Physics Research and Technology Series*, pp. 241–256, Nova Science Publishers Inc., 2018.

Méthode non paramétrique pour l'analyse et la classification des données fonctionnelles

Papa MBAYE^{1,2}

Anne-Françoise YAO¹

Chafik SAMIR²

¹ Laboratoire de Mathématiques Blaise Pascal CNRS UMR 6620

² Laboratoire d'Informatique, de Modélisation et d'Optimisation des Systèmes CNRS UMR 6158

papa_alioune_meissa.mbaye@uca.fr, anne.yao@uca.fr, chafik.samir@uca.fr

Résumé

L'analyse de données fonctionnelles joue un rôle important dans beaucoup de domaines de la santé publique et des applications biomédicales. En particulier, de telles méthodes statistiques fournissent des outils permettant de recaler, de comparer et de modéliser des données constituées de mesures corrélées. Dans ce travail, nous présenterons une nouvelle approche d'analyse de régression pour la classification de données fonctionnelles. D'abord, nous commencerons par analyser les observations fonctionnelles en faisant un recalage temporel. Ensuite, nous allons étudier différentes représentations standards de la littérature et estimer le modèle de régression appropriée comme une fonction de densité. Enfin, un exemple d'application, constitué de personnes ayant l'Arthrite Rhumatoïde (AR) et de personnes bien portantes comme groupe de référence, est présenté.

Mots Clef

Analyse de données fonctionnelles, Régression non paramétrique, Recalage.

Abstract

Functional data analysis plays an increasingly important role in many public health and biomedical applications. In particular, such statistical methods provide tools for warping, comparing, averaging, and modeling data involving correlated measurements. In this paper, we present a new approach of regression analysis for classification of functional data. First, we analyze functional observations to capture their key spatio-temporal patterns by searching optimal warping and then estimate the regression function. Next, we investigate different standard representations from literature and estimate the appropriate regression model as a density function. Finally, an example of application involving patients with Rheumatoid Arthritis and healthy subjects as a reference group, is presented.

Keywords

Functional Data Analysis, Nonparametric Regression, Registration, Time Warping.

1 Introduction

Analyser des données constituées de fonctions (courbes, surfaces ou d'autres fonctions), au lieu de vecteurs de scalaires, devient de plus en plus populaire [1, 2]. De tels problèmes nécessitent de considérer les courbes comme des fonctions continues et d'utiliser des représentations et analyses appropriées. Les méthodes de régression fonctionnelle ont été largement utilisées pour résoudre ce genre de problèmes [1, 3]. Récemment, différentes méthodes ont été proposées pour les régressions linéaires fonctionnelles. Cependant une étape clé pour analyser les données fonctionnelles temporelles est la capacité de capturer la variabilité temporelle, qui peut être considérée comme une transformation aléatoire du temps. En effet les variations obtenues au niveau des données collectées sont dues à plusieurs facteurs, incluant les outils de mesure et le comportement humain ; ce qui fait que les mêmes personnes observées peuvent donner lieu à différentes observations. La procédure de recalage pourrait ainsi être utilisée pour traiter cette variabilité temporelle qui est considérée comme une nuisance. Plusieurs alternatives ont été introduites pour représenter les courbes ou pour les comparer d'une manière invariante [2, 4].

L'arthrite est une maladie polymorphe qui est souvent caractérisée par un gonflement d'un ou de plusieurs articulations. D'après [5], l'arthrite est l'une des principales causes de l'incapacité physique qui affecte les jeunes et les personnes âgées, où les femmes sont plus touchées que les hommes. Malheureusement, il n'y a actuellement aucun remède pour l'arthrite et les traitements coûteux sont disponibles selon le type d'arthrite. Il y a plusieurs formes d'arthrites, dans lesquelles l'Arthrite Rhumatoïde, que l'on notera par la suite par AR, est la forme la plus commune

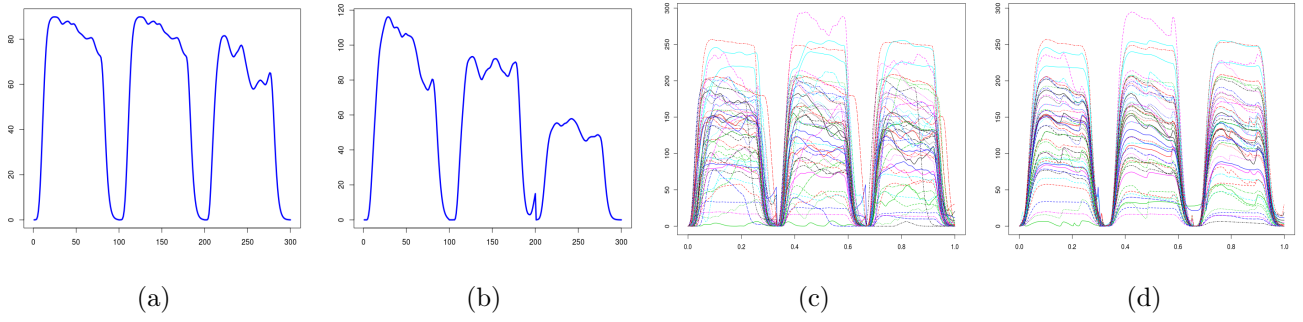


Figure 1: Exemples de données fonctionnelles qui mesurent l'intensité de la force musculaire: (a) Fonction de la force de la main d'une personne bien portante, (b) Fonction de la force de la main d'une personne malade (AR modéré), (c) L'ensemble des fonctions musculaires avant recalage, et (d) fonctions musculaires après recalage utilisant notre méthode.

d'inflammation chronique [6].

Dans les diagnostics quotidiens, l'examen clinique est utilisé pour reconnaître les modèles spécifiques et les symptômes, et si nécessaire, il est confirmé par d'autres tests, i.e. imagerie IRM et les tests du sang. Malheureusement, de tels tests sont très chers pour les patients et longs pour les médecins. Les types de diagnostics mentionnés précédemment peuvent être utilisés pour automatiser la classification de la maladie, mais au stade précoce de l'AR, ces critères ne sont pas habituellement satisfaisants. Dans les années récentes, la recherche médicale a entraîné une nouvelle compréhension de l'AR ; en particulier, il est indiqué que les mesures de force de la main sont une technique bonne et peu coûteuse pour une évaluation préopératoire de personnes malades [7]. Bien que quelques des caractéristiques discrètes citées précédemment puissent être utiles pour cet objectif, la fonction de force de la main contient plus d'informations de diagnostic et s'avère être un indicateur significatif sur la présence et le stade de la maladie. Dans cet article, nous nous concentrons sur cette nouvelle procédure de diagnostic. La fonction de force de la main d'une personne bien portante est donnée au niveau de la Figure 1(a) et celle d'une personne atteinte de l'AR au niveau de la Figure 1(b). Cette dernière montre un modèle clair de personnes malades où toutes les amplitudes de la force de la main ne sont pas très fortes et décroissent avec le deuxième et le troisième test. Cependant, en regardant les Figures 1(a) et 1(b), nous remarquons que le problème de classification entre les personnes bien portantes et les personnes malades est très difficile. Par ailleurs, pour illustrer l'importance du recalage, nous affichons les courbes originales avant recalage en 1(c) et après recalage en 1(d). Les fonctions considérées ici appartiennent à l'ensemble $\mathbb{L}^2([0, 1], \mathbb{R}^+)$ car ces intensités sont enregistrées de manière continue durant un intervalle de temps $T = [0, 1]$ et sont à valeurs dans $\mathbb{R}^+$. Récemment, l'analyse de données fonctionnelles a été proposée pour une étude plus générale. Bien qu'il

existe une large littérature sur l'analyse statistique de fonctions, voir par exemple [2, 4, 8], quand on se limite sur l'analyse de fonctions qui nécessite le recalage temporel, la littérature est toujours relativement limitée [1, 9–13].

Dans ce travail, nous proposons un modèle de régression fonctionnelle non paramétrique pour diagnostiquer l'AR. Autrement dit, on cherchera d'abord à apprendre une fonction de régression et à partir de cette fonction utiliser un seuil pour faire la classification, c'est à dire pour prédire la présence ou l'absence de la maladie. A notre connaissance, l'analyse de la régression sur des données fonctionnelles complètes sous forme de signaux de force de la main, pour diagnostiquer l'AR, n'a pas été précédemment étudiée. Un modèle statistique approprié est nécessaire dans cette application pour modéliser ces données fonctionnelles. En particulier, nous nous intéressons à l'étude de la variabilité au sein des groupes de personnes malades et de personnes bien portantes en utilisant la méthode de régression fonctionnelle complète. Une difficulté au niveau de la main est le fait que les signaux bruts des forces de la main ne sont pas alignés dans le temps. Autrement dit, différents patients exerceront leur force à des temps différents, et alors, il devient important de découpler la quantité de force exercée (amplitude de fonction) et combien de temps la force a été exercée (phase de fonction). Ainsi, nous avons besoin d'un modèle statistique global pour l'analyse des données fonctionnelles de force de la main qui permet la séparation des variabilités d'amplitude et de phase. Le modèle récent dans [14] fournit une approche mathématique et statistique efficace pour la séparation amplitude-phase de données fonctionnelles, et par la suite l'analyse statistique de ces deux composantes. Nous adaptons cette méthode pour étudier les signaux de la force de la main et pour définir un nouveau modèle (représentation fonctionnelle couplée à un modèle de régression) basé sur des données fonctionnelles complètes dans le but de caractériser la

maladie par des méthodes d'apprentissage statistique. Donc à partir de variables (signaux), qui sont à valeurs dans un espace de fonctions, le modèle prédira si la personne est bien portante ou malade.

2 Modélisation et Analyse de données fonctionnelles

Nous proposons une nouvelle représentation de la force de la main qui exploite le stade de la maladie comme une distance appropriée des observations de référence. Ce travail est inspiré en premier par les diagnostics classiques basés sur le maximum des mesures de la force (et éventuellement de la vitesse d'atteinte) qui entraînaient une énorme perte d'informations pertinentes pour la classification de la maladie d'AR. Par conséquent, les précisions de la classification décroissent significativement quand la variabilité entre les personnes bien portantes croît. Ainsi pour améliorer l'analyse statistique complète, nous prenons une approche d'analyse de données fonctionnelles pour analyser les fonctions de force de la main qui représentent l'effort continu, répétitif fait par les personnes. Cette nouvelle représentation utilisée pour la classification des personnes atteintes de l'AR apporte des informations complémentaires et indispensables sur l'état de la maladie. L'intensité de la force de la main est représentée par une fonction absolument continue x définie sur un intervalle $I = [0, 1]$, pour simplifier. Comme montré dans la Figure 1, $x(t) = 0$ là où il n'y a pas d'effort : au début du test ($t = 0$), au temps de repos et à la fin du test ($t = 1$). On peut remarquer dès à présent que la variance des intensités des personnes bien portantes est faible, alors que celle des personnes malades est forte, dû aux changements progressifs causés par la maladie en évolution. Pour arriver à une telle conclusion, on a à définir un modèle approprié, qui fournit une distance appropriée et des outils d'analyses statistiques en vue d'obtenir des classifications précises (i.e. séparation du groupe des personnes bien portantes et du groupe des personnes malades). Une qualité importante d'un tel modèle est d'être capable de résumer efficacement et de capturer la variabilité dans les deux classes. De plus, on espère que la distance définie pourra fournir une mesure naturelle entre les signaux de la force de la main, permettant ainsi aux rhumatologues de quantifier le stade de gravité de la maladie d'AR, en se basant sur une personne bien portante (référence). Par la suite, nous décrirons les éléments nécessaires qui seront utilisés pour recalibrer les données fonctionnelles.

Supposons un échantillon de variables aléatoires $\{x_i, i = 1, \dots, n\}$, où x_i est une fonction assez lisse définie dans un domaine unité de $\mathbb{R}$, et $\{y_i, i = 1, \dots, n\}$ une suite de variables binaires. $y_i = 0$ si la personne est bien portante et $y_i = 1$ si la personne est malade. x_i et y_i ne sont pas généralement directement

observables, au lieu de cela nous observons leurs discrétisations, avec du bruit aléatoire supplémentaire. Ainsi les données observées sont des vecteurs finis $(x_1; y_1); \dots; (x_n; y_n)$. Nous supposons que ces erreurs sont gaussiennes de moyenne nulle et qu'elles sont indépendantes.

Supposons un ensemble de fonctions de force $\{x_i, i = 1, \dots, n\}$, notre but est de trouver un ensemble de fonctions de reparamétrisation $\{\gamma_i^*, i = 1, \dots, n\}$ (variabilité de phase) tel que les fonctions $\{x_i \circ \gamma_i^*, i = 1, \dots, n\}$ soient alignées de manière optimale et alors ne varient qu'au niveau des amplitudes. γ est une fonction qui est définie par $\{\gamma : [0, 1] \rightarrow [0, 1]; \dot{\gamma} \succ 0\}$. Dans plusieurs publications précédentes, c'est la norme $\mathbb{L}^2$ pénalisée (norme $\mathbb{L}^2$ qui mesure l'écart entre deux fonctions plus une pénalisation sur la fonction de reparamétrisation) qui a été utilisée pour le recalage. Ces approches sont connues de ne pas bien fonctionner pour les fonctions de *pinching* (similaire au surapprentissage) et pour l'*asymétrie* de solutions [3]. Ce qui crée un effet sévère sur les analyses qui y découlent et cet effet vient du fait que la norme $\mathbb{L}^2$ n'est pas une métrique sur l'espace de fonctions modulo le groupe des reparamétrisations Γ . Ainsi dans ce papier, chaque fonction sera représentée par sa fonction q définie par $q(t) = \text{sign}(\dot{x}(t))\sqrt{|\dot{x}(t)|}$, où $\dot{x} = dx/dt$. Nous restreignons x d'être absolument continue parce que l'espace des résultats des fonctions q est $\mathbb{L}^2([0, 1], \mathbb{R})$, qui est l'ensemble des fonctions définies sur $[0, 1]$ et de carré intégrable. Si une fonction x est reparamétrisée par une fonction γ en $x \circ \gamma$, alors sa fonction q change et devient $(q \circ \gamma)\sqrt{\dot{\gamma}}$ et on la notera par $(q * \gamma)$. La propriété la plus importante de cette transformation est que $\|q\| = \|q * \gamma\|$ pour tout $\gamma \in \Gamma$, où $\|\cdot\|$ est la norme $\mathbb{L}^2$ de la fonction. Cette propriété permet de résoudre le problème de recalage optimal entre deux fonctions de force de la main x_1 et x_2 comme suit. Soit q_1 et q_2 leurs fonctions q . Alors la fonction de reparamétrisation optimale de x_2 à x_1 est donnée par $\gamma^* = \arg \inf_{\gamma \in \Gamma} \|q_1 - q_2 * \gamma\|$. La quantité à droite forme une distance appropriée dans l'espace quotient $\mathbb{L}^2/\Gamma$. Cette distance peut être utilisée pour définir des statistiques comprenant la moyenne de la fonction de force de la main, qui agira comme un modèle pour plusieurs recalages.

Le problème de phase et de séparation d'amplitude est lié aux fonctions de recalage non linéaire. Supposons $x : [0, 1] \rightarrow \mathbb{R}$ une fonction absolument continue et Γ l'ensemble de toutes les frontières préservant le difféomorphisme de $[0, 1]$ à lui-même. Alors pour tout $\gamma \in \Gamma$, la composition $x \circ \gamma$ représente le temps recalé de la fonction originale x . La phase est plus qu'un concept relatif. Si une fonction de reparamétrisation γ est utilisée pour recalibrer la fonction x_2 à x_1 , alors ce γ est nommé la *phase relative* de x_1 à x_2 . Notons que l'inverse de ce γ est la phase relative de x_2 à x_1 . En cas de plusieurs fonctions, comme dans le cas de notre

application, les composantes de phase sont définies en cherchant une moyenne de fonction et alors en évaluant la phase relative de chaque fonction donnée par rapport à la moyenne. Voir Algorithme 1 pour plus de détails.

Data : fonctions x_i .

Result : Moyenne de Fréchet μ_f , fonction de reparamétrisation γ_i^* , fonctions recalées x_i^* .

1. **Initialisation:** calculer les q_i correspondant à chaque $\{x_i\}$ et $\mu_q = \frac{1}{n} \sum_{i=1}^n q_i$.
2. **Recalage:** Pour $i = 1, 2, \dots, n$ calculer $\gamma_i^* = \arg \inf_{\gamma \in \Gamma} \|\mu_q - q_i * \gamma\|^2$.
3. **Actualisation:** Actualiser μ_q en utilisant $\mu_q \leftarrow \frac{1}{n} \sum_{i=1}^n (q_i * \gamma_i^*)$. Tant qu'il n'y a pas de convergence, on retourne à l'étape 2.
4. **Centrer:** Calculer la moyenne de la fonction de recalage $\bar{\gamma}$ et actualiser μ_q en utilisant $\mu_q \leftarrow \mu_q * \bar{\gamma}^{-1}$.
5. **Recalage final:** Répéter l'étape 2. Calculer μ_x et $x_i^* = x_i \circ \gamma_i^*$.

Algorithme 1 : Algorithme de séparation Phase-Amplitude

Ainsi, nous pourrions attribuer une amplitude et une composante de phase à chaque fonction d'un ensemble donné, et utiliser ces composantes pour définir les caractéristiques de l'AR nécessaires à la classification des personnes.

Supposons que nos fonctions x_i sont de classe C^k , $k \in \{0, 1, 2\}$. Pour le reste du papier, au lieu de x_i , nous allons utiliser une variable globale notée z_i , globale dans la mesure où elle sera utilisée pour différentes représentations comme l'intensité de la force $z_i = x_i$, sa vitesse $z_i = \dot{x}_i$, son accélération $z_i = \ddot{x}_i$ et la fonction de courbure correspondante $z_i = c_i$, pour la régression. Nous montrons dans la Figure 2 l'allure de ces différentes représentations fonctionnelles z_i . Il est important de noter que pour un signal parfait, on s'attend à ce que l'intensité de la force soit nulle au début et à la fin de chaque test. Ainsi, nous comptons sur les dérivées et la courbure pour capturer la distance entre une observation donnée et une observation ayant un comportement normal. Etant donné que les observations réelles, même les groupes des patients, ne sont pas parfaits, nous ferons un test répétitif (presque périodique) pour améliorer cette partialité. Nous rappelons que notre objectif est d'utiliser les variables fonctionnelles d'intensité, ou une des

représentations, pour prédire l'état d'une personne. Pour obtenir ceci, la méthode d'estimation de la régression fonctionnelle à noyau est utilisée. Notre analyse se fera sur des données déjà recalées, avec toutes les représentations citées précédemment.

3 Régression fonctionnelle à noyau avec réponse binaire

Différents estimateurs non paramétriques de régression ont été proposés dans la littérature quand la variable aléatoire explicative z_i prend ces valeurs dans un espace de dimension finie. Il y a beaucoup de travaux dans la littérature qui traitent les limites de ces estimateurs et d'autres questions qui y sont liées, comme la sélection de la fenêtre optimale dans les cas dépendants et indépendants. Pour plus de détails, on peut se référer aux [15, 16] et aux références citées dedans. Les résultats asymptotiques des données fonctionnelles ont récemment eu un intérêt croissant, on peut se référer aux [17, 18] et à la récente monographie faite par Ferraty et Vieu [19] et les références citées dedans.

Pour formuler le problème de l'estimateur de la régression fonctionnelle, supposons $(z_i, y_i)_{i \in \mathbb{N}}$ une séquence de couple de variables aléatoires (Z, Y) où z_i prend ces valeurs dans un espace métrique $(E, d(\cdot, \cdot))$ et y_i est binaire. Nous considérons le modèle

$$Y = r(Z) + \epsilon \quad (1)$$

D'après (1), $r(z_i) = \mathbb{E}[Y|Z = z_i]$. Considérons d'abord E comme étant un espace d'Hilbert $\mathcal{H}$ muni de sa métrique associée d . z_i étant de dimension infinie, nous allons la décomposer dans la base de fonction $\phi = (\phi_1(t), \dots, \phi_p(t)) : z_i(t) = \sum_{j=1}^p \alpha_{ij} \phi_j(t) = \alpha_i^T \phi$ avec $\alpha_i = (\alpha_{i1}, \dots, \alpha_{ip})$.

Pour des raisons pratiques, au lieu de travailler avec les z_i , nous allons travailler avec les coefficients α_i , qui sont de dimension finie, issus des décompositions des z_i dans la base de fonction ϕ .

L'estimateur de type Nadaraya-Watson a été introduit par Ferraty et Vieu [20]. Dans notre cas, il est défini par :

$$\hat{r}_n(z_i) = \frac{\sum_j y_j K_h(d(\alpha_i, \alpha_j))}{\sum_j K_h(d(\alpha_i, \alpha_j))}$$

où le dénominateur est différent de zéro et $K_h(d(\alpha_i, \alpha_j)) = K\left(\frac{d(\alpha_i, \alpha_j)}{h}\right)$. Ici K est une fonction noyau à valeurs réelles, h est le paramètre de la fenêtre (qui tend vers zéro quand n tend vers l'infini) et d est la métrique associée à $\mathcal{H}$.

Puisque Y est binaire, on cherchera plutôt à modéliser

$$g(Y) = r(Z) + \epsilon \quad (2)$$

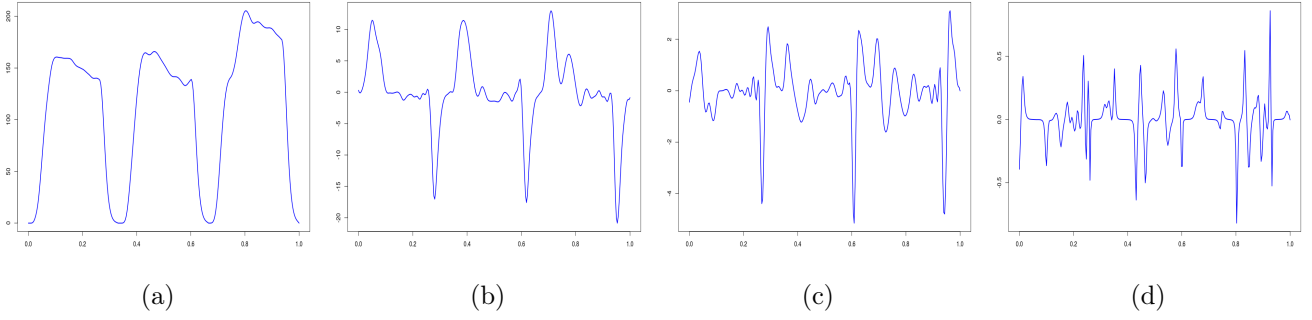


Figure 2: Exemples de différentes fonctions représentant l'intensité de la force de la main: (a) une courbe originale, (b) la vitesse, (c) l'accélération, et (d) la courbure.

où g est la fonction logit. La fonction réciproque de cette logit, appliquée à $\hat{r}_n(z_i)$, renvoie des valeurs de probabilités auxquelles nous allons appliquer un seuil pour faire la classification.

Les taux de convergence presque sûre, sur un ensemble compact de l'estimateur $\hat{r}_n$, sont établis dans [21] pour les processus asymptotiquement indépendants, alors que Masry [22] obtient la convergence de la moyenne quadratique. De plus la normalité asymptotique a été obtenue par Ferraty et al. [23]

Toujours dans l'optique de trouver le meilleur modèle, nous changeons d'espace et on choisit E comme étant la sphère de Hilbert. $\alpha_i \in \mathbb{R}^p$, nous nous restreignons à la sphère $\mathbb{S}^{p-1}$. Ainsi nous avons utilisé la distance géodésique s définie sur cette sphère par :

$$s(\alpha_i, \alpha_j) = \arccos\left(\frac{\alpha_i^T \alpha_j}{\|\alpha_i\| \|\alpha_j\|}\right).$$

La question qui se pose maintenant c'est quelles sont les valeurs optimales de h et de seuil qu'il faut prendre pour classer les malades et les personnes bien portantes. Dans la Figure 3, nous affichons des exemples de distribution des distances géodésiques sur la sphère et les h optimales retenues. Pour calibrer la performance de notre modèle d'estimation (régression fonctionnelle à noyau), nous considérons les critères : MSE (Mean Squared Error) et MCC (Matthews Coefficient Correlation).

- **MSE:** C'est l'erreur quadratique moyenne. Elle est définie par :

$$\frac{1}{n} \sum_i^n (y_i - \hat{y}_i)^2$$

avec n le nombre d'observation prédite et $\hat{y}_i$ la valeur prédite de la i -ème observation.

- **MCC:** Basé sur les Vrais et Faux Positifs (VP, FP), et sur les Vrais et Faux Négatifs

(VN, FN), il est généralement considéré comme une mesure équilibrée [24]. *MCC*:

$$\frac{TP \times TN - FP \times FN}{\sqrt{(TP + FP)(TP + FN)(TN + FP)(TN + FN)}}$$

MCC ou coefficient de corrélation entre les valeurs observées et les valeurs prédites, renvoie des valeurs comprises entre -1 et 1. Plus la valeur est proche de +1, plus la classification est bonne. Plus elle est proche de -1, plus la classification est mauvaise.

4 Résultats expérimentaux

Dans cette section, nous décrivons notre approche qui vise à classer les observations dans le groupe des personnes malades ou dans celui des personnes bien portantes, en utilisant les signaux de force de la main. Les résultats présentés dans cette section sont les résultats moyens obtenus après 100 itérations. Chaque itération consiste à générer aléatoirement 1500 observations composées de personnes malades et de personnes bien portantes, dont 900 constituent la base d'apprentissage et de validation et les 600 restantes la base test. Nous utilisons un modèle de régression fonctionnelle avec différents critères et différentes représentations. Chaque observation est représentée par une seule fonction de force de la main, combinant les 3 tests consécutifs. De ces fonctions, dérivent différentes représentations utiles pour la classification. Notre modèle de régression fonctionnelle à noyau utilise le noyau gaussien. Ces paramètres sont choisis grâce à la base d'apprentissage et les optimaux sont retenus grâce à la base de validation, avec le critère *MCC*. Cela assure et améliore la précision de la classification. Avant de présenter les principaux résultats de ce travaux, nous montrons une comparaison d'une méthode proposée et une simple approche d'analyse de données fonctionnelles, qui utilise la métrique $\mathbb{L}^2$ entre les fonctions et qui ne tient pas en compte des variabilités de phase. Nous calculons la matrice de distance

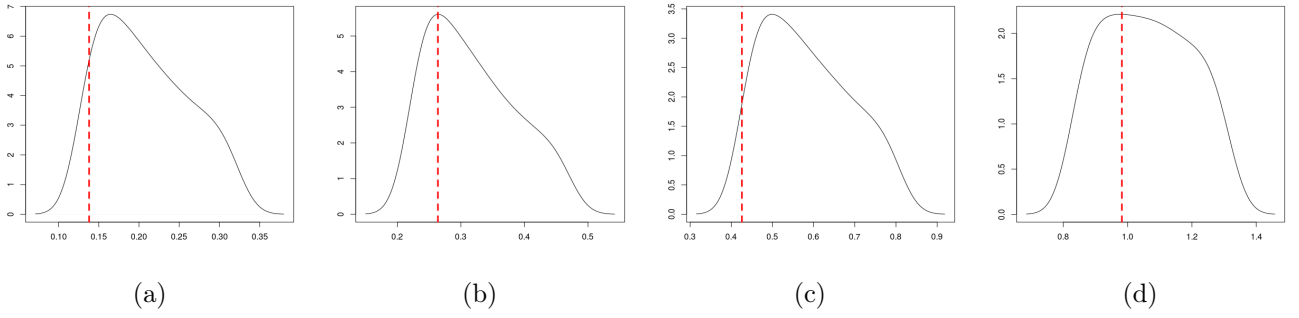


Figure 3: Exemples de densités des distances géodésiques sur la sphère pour chaque représentation, en rouge la valeur de la h optimale retenue pour le modèle : (a) Initial, (b) Vitesse, (c) Accélération, et (d) Courbure.

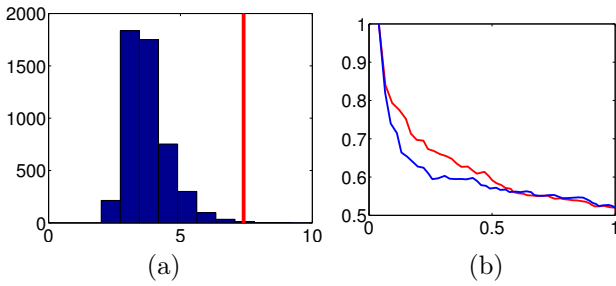


Figure 4: (a) Distribution empirique du test statistique de distance d'amplitude, avec la distance entre les 2 groupes marquée en rouge. (b) Précision (ordonnée) vs. Rappel (abscisse) la courbe de la méthode proposée (rouge) et la courbe obtenue avec la méthode d'analyse de données fonctionnelles non élastique basée sur la métrique $\mathbb{L}^2$ (blue).

pour chaque méthode et nous affichons la courbe Rappel/Précision dans la Figure 4(b). De cette figure, on peut dire que le fait de prendre en compte la variabilité de phase de signaux de force de la main est important et a le potentiel d'améliorer considérablement la performance de classification.

Nous évaluons maintenant la performance de notre modèle en calculant, après avoir prédit les variables réponses de la base test, les valeurs du critère utilisé (MCC).

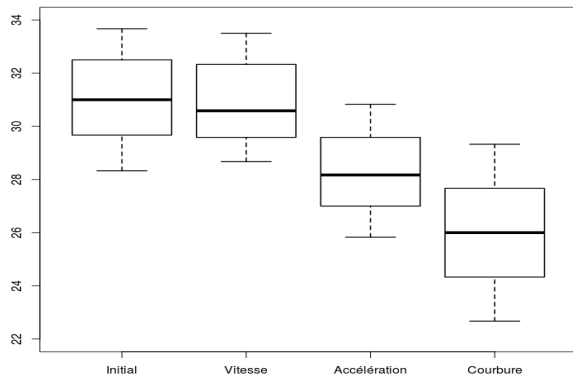
Si on utilise la métrique $L2$ dans notre modèle, c'est avec la représentation courbure qu'on obtient un plus petit taux d'erreur, comme on peut le visualiser au niveau de la Figure 5(a). Par ailleurs, si on utilise la distance géodésique sur la sphère, notée ici par s , c'est la vitesse qui nous donne une meilleure classification des deux groupes, voir Figure 5(b). Nous pouvons aussi remarquer qu'avec la métrique s , c'est la représentation vitesse qui nous donne la meilleure valeur de spécificité (plus petite erreur de première es-

pèce) et la courbure nous donne une meilleure valeur de sensibilité (plus grande valeur de la puissance du test), voir Figure 6.

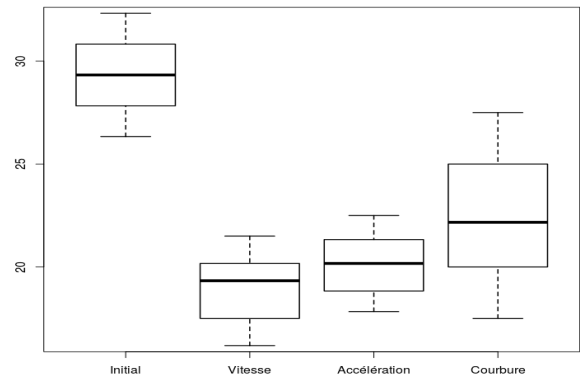
Comme nous l'avons énoncé précédemment, les personnes ayant un AR avancé montrent une décroissance significative de leurs forces de main durant les tests, comparés aux personnes bien portantes. Et cet aspect était le plus utilisé par les rhumatologues dans leurs diagnostics. Cependant, une telle procédure n'est pas applicable pour toutes les personnes malades à cause de différents facteurs comme l'âge, le genre, et plus important encore, le niveau de sévérité de la maladie. Les patients ayant un niveau d'AR moyen étaient difficiles à détecter avec le diagnostic classique. Ainsi il est important de rappeler que le fait d'utiliser les mesures continues de force de la main est une méthode bénéfique, rapide et facile, et plus encore, il est très efficace pour diagnostiquer le degré de la maladie. De plus, les informations extraites de la force de la main ont une interprétation clinique naturelle et donc plus intéressantes pour les médecins.

5 Conclusion

Ce travail présente une nouvelle approche permettant de caractériser les données fonctionnelles pour la classification de l'Arthrite Rhumatoïde (AR). Cette méthode a l'avantage d'utiliser les courbes recalées et de capturer ainsi plus d'informations des signaux, contrairement aux diagnostics classiques utilisés précédemment. Une fois que les courbes sont recalées, différentes représentations fonctionnelles ont été utilisées et la fonction de densité conditionnelle a été utilisée pour estimer la régression. Que ça soit la métrique d ou s , le fait d'utiliser la représentation standard (courbes initiales) ne nous permet pas d'avoir une meilleure classification. Ceci est dû au fait que la représentation standard ne capte pas bien la variabilité des différences de forces émises par les personnes. D'où l'importance d'utiliser d'autres représentations fonctionnelles, comme la vitesse, l'accélération ou la cour-

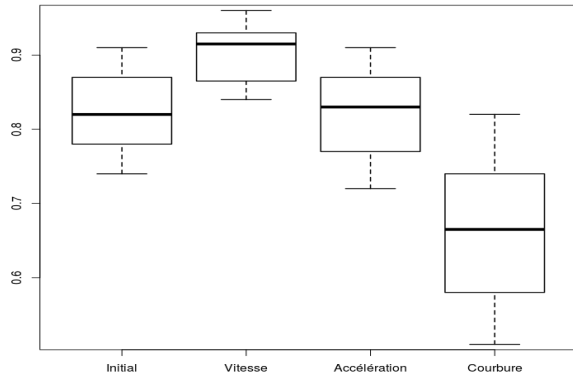


(a)

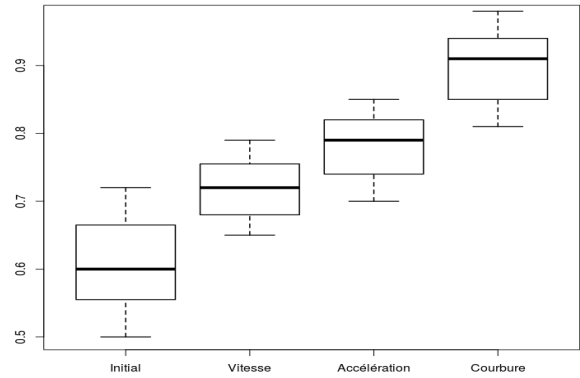


(b)

Figure 5: Dans chaque figure, on a représenté les taux d’erreurs obtenus en fonction des représentations fonctionnelles utilisées : (a) en utilisant la métrique d , i.e. la métrique $\mathbb{L}^2$ et (b) en utilisant la métrique s .



(a)



(b)

Figure 6: Ces résultats sont obtenus avec la métrique s (distance géodésique sur la sphère) : (a) Spécificité en fonction des représentations et (b) Sensibilité en fonction des représentations.

bure. On voit par exemple qu’en utilisant la métrique d ($\mathbb{L}^2$), c’est la courbure qui nous donne la meilleure classification. Et si on utilise la distance géodésique sur la sphère (s), c’est la vitesse qui nous donne les meilleurs résultats de classification. Ces résultats expérimentaux nous montrent que les diagnostics utilisés précédemment sont insuffisants et que notre modèle est très prometteur pour ce sujet.

References

- [1] G. James, “Curve alignment by moments,” *Annals of Applied Statistics*, pp. 480–501, 2007.
- [2] J. O. Ramsay and B. W. Silverman, *Functional Data Analysis, Second Edition*. Springer Series in Statistics, 2005.
- [3] J. D. Tucker, “Functional component analysis and regression using elastic methods,” *Electronic Theses, Treatises and Dissertations, Florida State University*, 2014.

- [4] R. Tang and H. G. Müller, "Pairwise curve synchronization for functional data," *Biometrika*, vol. 95, no. 4, pp. 875–889, 2008.
- [5] D. L. Scott, F. Wolfe, and T. W. Huizinga, "Rheumatoid arthritis," *The Lancet*, vol. 376, no. 9746, pp. 1094–1108, 2010.
- [6] S. J. Bigos, J. Holland, C. H. and J. S. Webster, M. Battie, and J. A. Malmgren, "High-quality controlled trials on preventing episodes of back problems: systematic literature review in working-age adults," *Spine Journal*, vol. 9, no. 2, pp. 147–68, 2009.
- [7] G. Michael and W. Richard, "A systematic exploration of distal arm muscle activity and perceived exertion while applying external forces and moments," *Ergonomics*, vol. 51, no. 8, pp. 1238–1257, 2008.
- [8] A. Kneip and T. Gasser, "Statistical tools to analyze data representing a sample of curves," *The Annals of Statistics*, vol. 20, pp. 1266–1305, 1992.
- [9] C. Samir, S. Kurtek, A. Srivastava, and N. Borges, "An elastic functional data analysis framework for preoperative evaluation of patients with rheumatoid arthritis," in *Applications of Computer Vision (WACV), 2016 IEEE Winter Conference on*. IEEE, 2016, pp. 1–8.
- [10] J. O. Ramsay and X. Li, "Curve registration," *Journal of the Royal Statistical Society, Series B*, vol. 60, p. 351–363, 1998.
- [11] D. Gervini and T. Gasser, "Self-modeling warping functions," *Journal of the Royal Statistical Society, Series B*, vol. 66, pp. 959–971, 2004.
- [12] X. Liu and H. G. Müller, "Functional convex averaging and synchronization for time-warped random curves," *Journal of the American Statistical Association*, vol. 99, pp. 687–699, 2004.
- [13] A. Kneip and J. O. Ramsay, "Combining registration and fitting for functional models," *Journal of the American Statistical Association*, vol. 103, no. 483, pp. 1155–1165.
- [14] A. Srivastava, W. Wu, S. Kurtek, E. Klassen, and J. S. Marron, "Registration of functional data using fisher-rao metric," *arXiv: 1103.3817v2*, 2011.
- [15] L. T. Tran, "Density estimation for time series by histograms," *Journal of statistical planning and inference*, vol. 40, no. 1, pp. 61–79, 1994.
- [16] N. Lai *et al.*, "Kernel estimates of the mean and the volatility functions in a nonlinear autoregressive model with arch errors," *Journal of statistical planning and inference*, vol. 134, no. 1, pp. 116–139, 2005.
- [17] D. N. Politis and J. P. Romano, "Limit theorems for weakly dependent hilbert space valued random variables with application to the stationary bootstrap," *Statistica Sinica*, pp. 461–476, 1994.
- [18] J. O. Ramsay and B. W. Silverman, *Applied functional data analysis: methods and case studies*. Springer New York, 2002, vol. 77.
- [19] F. Ferraty and P. Vieu, *Nonparametric functional data analysis: theory and practice*. Springer Science & Business Media, 2006.
- [20] —, "Dimension fractale et estimation de la régression dans des espaces vectoriels semi-normés," *Comptes Rendus de l'Académie des Sciences-Series I-Mathematics*, vol. 330, no. 2, pp. 139–142, 2000.
- [21] —, "Nonparametric models for functional data, with application in regression, time series prediction and curve discrimination," *Nonparametric Statistics*, vol. 16, no. 1-2, pp. 111–125, 2004.
- [22] E. Masry, "Nonparametric regression estimation for dependent functional data: asymptotic normality," *Stochastic Processes and their Applications*, vol. 115, no. 1, pp. 155–177, 2005.
- [23] F. Ferraty, A. Mas, and P. Vieu, "Nonparametric regression on functional data: inference and practical aspects," *Australian & New Zealand Journal of Statistics*, vol. 49, no. 3, pp. 267–286, 2007.
- [24] P. Baldi, S. Brunak, Y. Chauvin, C. A. Andersen, and H. Nielsen, "Assessing the accuracy of prediction algorithms for classification: an overview," *Bioinformatics*, vol. 16, no. 5, pp. 412–424, 2000.

Expression verbale des tendances à l'action en langue française

Alya Yacoubi^{1,2}

Nicolas Sabouret¹

¹ LIMSI, CNRS, Univ. Paris-Sud, Université Paris-Saclay

² DAVI les humaniseurs, 92800 Puteaux

Bat. 508, Campus universitaire, 91405 Orsay, **email** : prenom.nom@limsi.fr

Résumé

Cet article présente une étude expérimentale sur l'expression verbale des tendances à l'action. Notre travail se place dans le contexte des interactions entre un homme et un agent conversationnel expressif capable de comportements émotionnels. L'une des dimensions de ces comportements est le mécanisme de la tendance à l'action mis en évidence par plusieurs chercheurs en psychologie cognitive. Nous avons extrait 30 tendances à l'action prises de la littérature en psychologie et nous avons défini des actes locutoires en langue française indépendants du contexte pour les exprimer. Les tendances exprimées ont été évaluées auprès de 80 participants dans deux contextes différents. Cela nous a permis de valider le choix des actes locutoires pour 21 tendances à l'action. Nous avons également mis en évidence des similarités fortes entre certaines tendances dans le cadre d'un dialogue.

Mots Clef

Tendances à l'action, expression verbale, dialogue, interaction homme-agent.

Abstract

This article presents an experimental study on the verbal expression of action tendencies. Our work concerns dialogical interactions between a human and an conversational agent capable of expressing emotional behavior. One of the dimensions of these behaviors is the mechanism of the action tendency highlighted by several researchers in the field of cognitive psychology. We extracted 30 action tendencies taken from the literature in psychology and we defined context-free speech acts in French. The tendencies expressed were evaluated with 80 participants in two different contexts. This allowed us to validate the choice of locutionary acts for 21 tendencies to action. We also found strong similarities between some action tendencies in the dialogical context. In this article, we detail our experimental approach and the results obtained.

Keywords

Actions tendencies, verbal expression, dialogue, human-agent interaction

1 Introduction

Les émotions sont un mécanisme déterminant dans les interactions entre les individus [11, 6, 14], mais aussi dans le contexte des interactions entre un utilisateur et un agent conversationnel, comme l'ont montré plusieurs travaux en informatique affective [10, 17]. Ainsi, certains d'entre eux ont tenté de modéliser le mécanisme émotionnel chez des agents conversationnels par des comportements non-verbaux [15] ou par le choix du contenu verbal de leurs énoncés [3].

Pour ce qui est du comportement verbal, il a été montré dans les travaux de [13] portant sur le corpus de l'agent conversationnel Laura (l'agent du site web d'EDF), que les utilisateurs expriment leurs émotions à l'écrit de la même façon face à un agent virtuel que face à un humain, avec des choix similaires de termes et de typologie de phrase. Malheureusement, la plupart des agents conversationnels utilisés dans l'industrie ne sont pas conçus pour exprimer en retour des émotions. Ainsi, lorsque l'utilisateur exprime sa satisfaction, il est assez perturbant de n'observer aucune réaction de la part de l'agent qui continue sa tâche comme si de rien n'était.

En effet, la plupart de ces agents conversationnels, développés par des entreprises spécialisées pour des clients (comme par exemple l'agent libraire [23] ou encore la conseillère virtuelle Anna de IKEA) sont orientés vers la réalisation d'une tâche plus que vers l'expression de compétences sociales. L'effort est mis sur la constitution d'une base de connaissances et de règles de dialogues bien adaptées au domaine considéré. Cela conduit à des agents très efficaces pour trouver une information adaptée mais assez pauvres en terme de comportements sociaux.

Nos travaux visent à permettre de doter de tels agents conversationnels orienté tâche d'un mécanisme émotionnel pour adapter leur comportement dialogique au contexte social de l'interaction. Pour cela, nous avons construit un modèle informatique qui définit le processus émotionnel et son impact sur le comportement de l'agent à deux niveaux : la poursuite de la tâche (interruption, changement de sujet, etc) et l'expression verbale de composantes émotionnelles qui peuvent s'insérer dans le dialogue orienté tâche. Ainsi, notre modèle peut être branché sur un agent orienté tâche

sans modification de sa base de connaissance ni de la stratégie commerciale qui a été implémentée à travers la base de règles.

Pour modéliser le processus émotionnel, nous nous sommes basés sur la théorie des *tendances à l'action* introduite par Frijda [7] qui définit les émotions comme étant des mécanismes heuristiques incitant ainsi l'individu à faire une action dirigée vers une source d'émotion. La colère par exemple est associée à la tendance à être agressif envers l'autre [21, 8]. Il est important de noter qu'une tendance à l'action ne s'exprime pas nécessairement, qu'elle peut être inhibée ou redirigée. Ainsi, dans le contexte d'un dialogue, la tendance à être agressif pourrait se traduire par le fait d'insulter l'interlocuteur [2], de parler plus vulgairement [19] (tendance redirigée) ou simplement être masquée derrière une politesse plus froide (tendance inhibée). Dans cet article, nous nous intéressons à l'expression verbale de tendances à l'action par l'agent conversationnel. Nous présentons une étude qui a été faite en forme de questionnaire en ligne dans lequel nous présentons différents dialogues illustrant chacun une situation d'interaction et donc une expression d'une tendance à l'action particulière. Notre objectif est de pouvoir brancher notre modèle affectif sur n'importe quel agent orienté tâche, indépendamment du domaine. Nous avons donc fait le choix de phrases génériques qui expriment une tendance à l'action indépendamment du contexte du dialogue. Nous avons sélectionné 30 tendances à l'action présentes dans la littérature en psychologie et nous leur avons attribué une expression verbale. Nous avons ensuite mené cette étude expérimentale pour valider l'expression verbale de chaque tendance à l'action et évaluer la perception des utilisateurs des tendances à l'action au cours d'un dialogue.

Cet article est structuré comme suit : la deuxième section est dédiée aux travaux théoriques sur l'expression verbale des émotions aussi bien dans le dialogue humain-humain que humain-agent. Ensuite nous présentons la méthodologie de construction de corpus et sélection de tendances à l'action. La quatrième section décrit le protocole expérimental où nous présentons le questionnaire qui a été construit et ses caractéristiques. À la fin nous présentons les résultats obtenus, leurs interprétations et nous proposons un sous-ensemble d'actes locutoires validés expérimentalement pour exprimer des tendances à l'action dans différents contextes.

2 Travaux connexes

Comme l'ont montré Ochs et Schieffelin [16], les émotions se manifestent dans le langage par différents moyens : le vocabulaire, la construction de la phrase, etc.. De nombreux chercheurs en informatique et en psychologie du langage ont ainsi tenté de caractériser les actes de langage émotionnels. Dans le contexte de la détection automatique des émotions, les auteurs de [5] ont identifié, au moyen d'un système d'apprentissage sur des dialogues annotés, le champ lexical de 5 émotions (la colère, la peur,

la satisfaction, l'excuse, neutre). Dans un cadre plus général, Argaman [1] a étudié le lien entre lexique et intensité émotionnelle. Il a proposé de mesurer les marqueurs linguistiques utilisés par 500 sujets pour exprimer leurs opinions à propos de films mettant en jeu des émotions positives (joie) et négatives (tristesse) plus ou moins intenses. Il a identifié dix caractéristiques des actes de dialogues émotionnels. Cependant, les travaux de [18] ont montré qu'il reste difficile, pour la majorité des personnes, de contrôler complètement le choix des mots employés (répétitions, expression à la première personne). Dans notre étude, nous nous sommes basés sur quelques-unes des caractéristiques d'Argaman pour choisir les actes locutoires : l'emploi des « intensificateurs » (« très », « le plus », « extrêmement ») et « d'atténuateurs » lexicaux (« assez », « plutôt »), l'utilisation de la première personne du singulier (« je », « moi »), les expressions exclamatives (par exemple « wow », « zut »).

Dans le domaine des agents virtuels, de nombreux chercheurs se sont intéressés à la modélisation de l'affect dans une interaction humain-agent. Cependant, seulement quelques uns ont modélisé l'expression verbale des émotions. Nous pouvons citer comme exemple ALMA [9] où l'expression des émotions chez l'agent se traduit par la sélection d'un libellé et d'une phrase en langage naturel. Dans ce travail, les énoncés des personnages virtuels sont scriptés en fonction de l'émotion à exprimer et de la situation de la tâche. Les énoncés des personnages dont l'humeur est positive sont ainsi plus longs et plus positifs. Au contraire, lorsqu'ils sont d'humeur hostile, les énoncés sont plus courts, plus précis et plus durs. Ces principes doivent être respectés pour la construction d'actes locutoires affectifs. Cependant, nous souhaitons que les énoncés puissent être indépendants de la tâche. C'est pourquoi nous n'avons pas pu reprendre directement les actes de langages proposés dans ALMA.

À notre connaissance, notre étude est la première à proposer une validation systématique d'actes locutoires expressifs indépendants du contexte applicatif et en relation avec l'expression de tendances à l'action.

3 Construction du corpus

Avant de présenter le protocole expérimental lui-même, nous détaillons la méthodologie que nous avons suivie en amont pour choisir les tendances à l'action de la littérature et pour construire les actes locutoires associés.

3.1 Choix des tendances à l'action

Une première étape de construction de notre modèle affectif était de définir un nombre fini de tendances à l'action possibles. La notion de tendance à l'action est présente dans les travaux de plusieurs chercheurs en psychologie des émotions [6, 11]. Plusieurs dizaines de tendances à l'action ont été étudiées dans la littérature avec des différences subtiles en terme de comportement et de conditions d'activation.

Nous avons choisi un sous-ensemble de tendances à l'action adaptées au modèle cognitif que nous avons conçu [25]. Ce modèle se base sur un processus d'évaluation cognitive qui consiste à attribuer des valeurs à des variables d'évaluation cognitive comme la désirabilité, la contrôlabilité, etc. Roseman [21] a présenté un travail qui lie le processus de l'évaluation cognitive aux tendances à l'action générées. Ce lien s'effectue à travers des stratégies émotionnelles, notion qui a été introduite par Frijda et qui permet de regrouper les tendances à l'action en tenant compte de leur direction. Une stratégie émotionnelle est définie par Roseman comme étant le but ultime d'une tendance à l'action. Par exemple, *attack other* est une tendance à l'action qui répond à la stratégie émotionnelle *move against* et qui est orientée vers *other* (l'interlocuteur). Elle apparaît lorsque la situation est indésirable mais contrôlable (ce qui caractérise la stratégie émotionnelle *move against*) et que la cause de l'émotion est l'interlocuteur.

Notre choix de tendances à l'action couvre l'ensemble des stratégies émotionnelles définies par Roseman, à l'exception de la surprise. Pour sélectionner ces tendances à l'action parmi celles proposées dans la littérature, nous utilisons les critères suivants :

- La tendance à l'action doit être exprimable dans un contexte dialogique (la tendance de violence physique par exemple est difficilement exprimée par la modalité verbale) ;
- La tendance à l'action doit appartenir à une stratégie émotionnelle. Ce lien "tendance à l'action-stratégie émotionnelle" doit être mentionné dans la littérature en psychologie ;
- Les tendances doivent être distinctes.

3.2 Choix des termes désignant les tendances à l'action

Notre objectif est d'interroger des sujets sur les tendances à l'action qu'ils reconnaissent dans un dialogue. Pour cela, il faut bien choisir les termes qui désignent ces tendances à l'action dans le questionnaire (et ce, indépendamment des actes locutoires que nous souhaitons valider).

Les travaux en psychologie sur les stratégies émotionnelles et les tendances ont principalement été faits en anglais. Les termes désignant ces notions sont donc des termes anglais. Nous avons fait un travail de traduction et d'interprétation en langue française en essayant d'être le plus proche possible du sens. Pour cela, nous nous sommes basés sur la description des tendances dans les travaux de Frijda, Roseman et des autres auteurs du domaine [24, 12]. Les traductions en français ont été faites par des chercheurs francophones, en collaboration avec des collègues anglophones. Nous avons aussi dû tenir compte de la différence entre les termes scientifiques utilisés par les chercheurs pour désigner les tendances à l'action et notre besoin applicatif. Les termes scientifiques, pensés dans les années 80, ne sont pas adaptés pour un questionnaire orienté grand public des années 2010 et pourraient entraîner des incompréhensions ou

des ambiguïtés d'interprétation chez les participants. C'est pour cette raison que nous avons tenté de simplifier certains termes pour que cela soit accessible à tous les participants. Par exemple, le terme "expier" a été remplacé par "subir les conséquences de son erreur".

3.3 Choix des actes locutoires expressifs

Le choix du contenu locutoire pour exprimer verbalement une tendance à l'action dans un dialogue a été fait avec des natifs pour chaque langue (le questionnaire étant disponible en français et anglais) dans notre laboratoire. Les chercheurs se sont projetés dans des situations correspondant aux différentes stratégies émotionnelles considérées et ont proposé des phrases qu'ils auraient pu dire spontanément dans cette situation. Nous avons alors extrait les phrases les plus générales possibles (i.e. indépendantes du contexte) en respectant les principes linguistiques énoncés dans la section 2.

Le tableau 1 résume les tendances à l'action choisies avec leur stratégies émotionnelles et les phrases choisies. Le nombre de tendances à l'action est variable d'une stratégies à l'autre : cela provient de la variété de tendances que nous avons pu trouver dans la littérature pour couvrir les six stratégies de Roseman et les différentes directions (ou causes) possibles. La première colonne correspond à la stratégie émotionnelle ; la deuxième correspond à la cause vers laquelle serait orientée la tendance à l'action. La troisième colonne représente les codes que nous avons choisi pour avoir une notation uniforme tout au long de l'article. Les deux autres colonnes sont dédiées respectivement aux tendances à l'action et à leur acte locutoire choisi.

3.4 Choix des situations

Afin d'évaluer les phrases correspondant aux différentes tendances à l'action, nous avons créé des scénarios où nous présentons des situations qui déclenchent des réponses émotionnelles chez un agent nommé Alex. Pour chaque situation, nous avons défini un objectif de l'agent (par exemple, orienter un touriste). Un énoncé d'un autre agent, Bob, impacte directement ou indirectement cet objectif (par exemple, Bob dit qu'il est content des renseignements fournis par Alex).

Les situations sont choisies de manière à couvrir les différentes valeurs possibles pour les variables d'appraisal proposées par Roseman [22] qui sont : la désirabilité, la contrôlabilité, et la certitude de la situation, la nature de l'objectif touché et la nature du problème. En particulier, nous distinguons les cas où l'impact sur l'objectif est direct ou indirect i.e. l'énoncé impacte un sous-objectif (par exemple, donner une réponse utile est un sous-objectif de « orienter un touriste »). Chaque situation se traduit ainsi par une attribution de valeurs aux variables d'appraisal et c'est la combinaison de ces valeurs qui détermine la stratégie émotionnelle.

Nous avons construit un scénario pour chaque stratégie émotionnelle et pour chaque cause possible parmi : l'agent (*self*), son interlocuteur (*other*) ou des circonstances exté-

rieures aux deux agents (*circ*). Nous avons ensuite choisi un acte locutoire indépendant du contexte pour chaque tendance à l'action faisant partie d'une stratégie émotionnelle. Reprenons notre exemple de conseiller touristique Alex qui a pour objectif que le touriste soit bien renseigné. Si Bob qui affirme qu'il s'est perdu à cause des instructions erronées d'Alex, cela est contraire au but. Cela déclencherait donc la stratégie émotionnelle *Move Against* avec la cause *self*. Une des tendance à l'action déclenchées par cette stratégie émotionnelle est de « subir les conséquences de son erreur ». Cela peut s'exprimer, par exemple, par la phrase « Que puis-je faire pour me faire pardonner ? ».

4 Méthodologie

A travers cette étude expérimentale, nous avons pour objectif de valider notre choix d'actes locutoires exprimant des tendances à l'action indépendamment du contexte. Pour cela nous avons établi un protocole expérimental pour vérifier trois hypothèses que nous présentons dans cette section.

4.1 Protocole expérimental

Afin de valider notre choix d'actes locutoires, nous avons construit un questionnaire dans lequel nous présentons différents dialogues illustrant chacun une situation d'interaction et donc une tendance à l'action particulière chez l'agent Alex (suivant la désirabilité, la contrôlabilité, et la certitude de la situation, en plus de la nature de l'objectif touché et la nature du problème).

Pour chaque situation, nous proposons alors 5 tendances à l'action dont une est celle que nous pensons la plus adéquate. Les autres tendances proposées ont été choisies comme suit :

- les tendances appartenant à la même stratégie émotionnelle et orientées vers une même cause
- les tendances à l'action orientée vers la même cause que la tendance à l'action voulue mais appartenant à une stratégie émotionnelle différente.
- les tendances à l'action appartenant à la même stratégie émotionnelle que la tendance à l'action voulue mais orientée vers une cause différente.

Exemple : Dans la figure 1, nous présentons une question extraite du questionnaire. Le fait que le touriste ne s'est pas perdu dans la ville est un événement désirable pour le guide touristique Alex. La cause de cet événement est "self" vu que le touriste ne s'est pas perdu grâce à ses conseils. La tendance que nous voulions exprimer dans ce dialogue est : S'afficher. Les seules autres tendances qui répondent à la même stratégie émotionnelle et orientée vers la même cause "self" sont "rechercher la reconnaissance" et "faire savoir autour de soi". "Continuer comme ça" répond à la même stratégie mais orientée vers la cause "circ" et "s'approcher de l'autre" fait partie de la même stratégie mais orientée vers "other".

Variantes A et B : Pour chacune des 30 tendances à l'action, deux situations différentes ont été construites, illus-

Stratégie émotionnelle	Cause	Code	Tendance à l'action	Acte locutoire
Moving Against	Circ	TA01	Tout casser	Rah ! Y'en a marre !
		TA02	Retirer l'obstacle	Qu'est-ce qu'on peut faire ?
	Self	TA03	Se punir	C'est vrai je suis nul.
		TA04	Attaquer l'autre	Tout ça est à cause de vous !
	Other	TA05	Insulter l'autre	Vous êtes vraiment stupide !
		TA06	Blessar l'autre	Vous n'êtes bon à rien !
Move Away from	Circ	TA07	Éviter le problème	Et si on parlait d'autre chose ?
		TA08	S'enfuir de la situation	Je ne veux pas parler de cela. Arrêtons la discussion !
		TA09	Se protéger	Je ne suis pas ici pour entendre ce genre de remarques. Arrêtons là notre discussion !
		TA10	Pleurer	Snif !
	Self	TA11	Demander de l'aide	Ce n'est pas possible. Il faut que vous m'aidiez !
		TA12	Corriger l'erreur	Pardon. Je vais tenter de corriger.
		TA13	Minimiser	Ce n'est pas très grave.
		TA14	Subir les conséquences de son erreur	Qu'est-ce que je peux faire pour me faire pardonner ?
	Other	TA15	S'éloigner de l'autre	Je ne veux plus vous parler.
		TA16	Sauter de joie	Youpi !
Moving Toward	Circ	TA17	Partager sa bonne humeur	ça me fait très plaisir.
		TA18	Continuer dans le même plan	Cool !
	Self	TA19	Faire savoir autour de soi	Tu as vu comme je suis bon ! Il faut que vous le disiez à tout le monde !
		TA20	S'afficher	je sais je suis parfait !
	Other	TA21	Chercher la reconnaissance	J'aime recevoir des compliments.
		TA22	Se rapprocher de l'autre	J'aimerais continuer à vous parler.
	Other	TA23	Chercher le contact de l'autre	J'aimerais vous parler plus souvent.
Moving it away	Circ	TA24	Rejeter la situation	Je n'aime pas ça !
	Self	TA25	Se soumettre	Je suis prêt à faire ce que vous voulez.
		TA26	Se désister	Je préfère laisser la place à mon collègue.
		TA27	Se cacher	J'aimerais disparaître
		TA28	Repousser l'autre	Agh ! Je ne veux plus discuter avec vous !
	Other	TA29	Renoncer	Je ne peux rien faire de plus.
Stop Moving away from	Circ	TA30	Se détendre	Ouf ! c'est mieux comme ça !

TABLE 1 – Les tendances à l'action regroupées par stratégies émotionnelles et causes et leurs actes locutoires correspondants

trées dans deux dialogues distincts et présentées à des participants différents. Nous avons donc deux versions du questionnaire nommées A et B, ce qui permettra d'évaluer la sensibilité des actes locutoires au contexte applicatif. Pour chaque participant qui commence le test, une des deux versions du questionnaire est sélectionnée d'une manière aléa-

Lisez attentivement le dialogue. La phrase en gras et soulignée exprime une tendance à effectuer une action. Indiquez votre degré d'accord pour chaque tendance à l'action proposée.

- Alex : Bonjour, que puis-je faire pour vous ?
 - Bob : Grâce à vos conseils, je ne me suis pas perdu !
 - Alex : **Je sais je suis parfait !** Dites, qu'est-ce que je peux faire pour vous ?

À votre avis, Alex avait tendance à :

	Pas du tout d'accord	Pas d'accord	Neutre	D'accord	Tout à fait d'accord
S'afficher	☐	☐	☐	☐	☐
Continuer comme ça	☐	☐	☐	☐	☐
Faire savoir autour de soi	☐	☐	☐	☐	☐
S'approcher de l'autre	☐	☐	☐	☐	☐
Chercher de la reconnaissance	☐	☐	☐	☐	☐

FIGURE 1 – Une question extraite du questionnaire d'évaluation de choix d'actes locutoires

toire et automatique.

4.2 Hypothèses

Nous voulons vérifier que des actes locutoires choisis sont bien perçus comme associés à la tendance à l'action qu'ils sont censés représenter. Par exemple, nous faisons l'hypothèse que l'acte locutoire "Y'en a marre !" sera bien perçu comme exprimant la tendance à l'action "Détruire" qui appartient à la stratégie émotionnelle "Move Against" face à une cause extérieure (circ).

Cependant, l'étude de la littérature a révélé qu'il n'y a pas de consensus quant au nombre exact des tendances à l'action. L'émotion "colère" par exemple peut conduire à une tendance à attaquer l'autre, critiquer sa conduite ou l'insulter. Nous formulons donc l'hypothèse que certains actes locutoires seront perçus comme relevant de deux tendances à l'action distinctes, que nous pourrions alors envisager de regrouper.

Enfin, nous voulons nous assurer que la perception des tendances à l'action est indépendante du contexte du dialogue. Ces différents objectifs nous ont amenés à formuler les deux groupes d'hypothèses suivantes :

H1a : Nous faisons l'hypothèse initiale que, pour chaque acte locutoire al_i , la tendance à l'action associée ta_i est significativement mieux reconnue que les 4 autres tendances à l'action proposées.

Reprenons l'exemple sur la figure 1. Le score donné par les participants à la tendance "s'afficher" devrait être significativement plus élevé que celui attribué à "chercher la reconnaissance", "faire savoir autour de soi", "continuer comme ça" ou "s'approcher de l'autre".

Pour vérifier cette hypothèse, nous calculons la p-value sui-

vant un test de Wilcoxon¹ entre ta_i et chacune des autres tendances à l'action proposées. Cette mesure se fait sans prendre en compte les éventuelles différences de contexte. L'hypothèse H1a est considérée comme valide pour l'acte locutoire si et seulement si le score de p-value est inférieur à 0,05 pour chacune des 4 comparaisons.

H1b : Lorsque H1a n'est pas vérifiée, c'est-à-dire lorsque la tendance à l'action proposée n'a pas été significativement mieux reconnue, nous faisons l'hypothèse qu'une autre tendance à l'action "proche" est peut-être associée à cet acte locutoire.

Nous regardons s'il existe une tendance à l'action ta_j parmi les 4 autres tendances à l'action proposées qui est significativement mieux reconnue par les participants. Cette mesure se fait sans prendre en compte les éventuelles différences de contexte et l'hypothèse H1b est valide si le score de p-value entre ta_j et chacune des 4 autres tendances à l'action proposées est inférieur à 0,05.

H1c : Lorsque ni H1a, ni H1b ne sont vérifiées, nous faisons l'hypothèse que plusieurs tendances à l'action différentes peuvent être représentées par ce même acte locutoire.

Nous cherchons s'il existe une tendance à l'action ta_j parmi celles proposées telle que, lorsque nous fusionnons les scores obtenus pour ta_j et ta_i , nous obtenons un score

1. Nous avons effectué un test non-paramétrique vu que la distribution des données ne suit aucune loi. Nous avons également eu besoin d'un test sur des échantillons appariés vu que nous allons comparer les réponses d'un même groupe de participants. Le test de Wilcoxon signé (Wilcoxon signed rank test)[20] permet de prendre en compte le niveau de différence à l'intérieur des paires. Nous avons opté pour ce test appliqué pour chaque scénario, entre la réponse qui a eu le score plus élevé et toutes les autres tendances proposées dans le même scénario.

significativement plus élevé que pour les trois autres tendances à l'action proposées (au sens de la p-value).

H2a : Lorsque H1a ou H1b est vérifiée, c'est-à-dire lorsqu'il existe bien une tendance à l'action qui est mieux reconnue que les autres, nous formulons l'hypothèse que cette propriété est valide indépendamment du contexte.

L'hypothèse H2a est considérée comme valide si et seulement si H1a (resp. H1b) peut être vérifiée séparément dans les deux sous-groupes associés aux contextes A et au contexte B.

H2b : Lorsque H1a ou H1b est vérifiée, nous formulons aussi l'hypothèse que le score attribué à la tendance à l'action est indépendant du contexte. Pour cela, nous calculons la valeur de corrélation inter-groupe entre le contexte A et le contexte B en utilisant un test de κ de Cohen [4]. L'hypothèse est considérée comme valide si et seulement si l'accord inter-annotateur obtenu par le score de κ est supérieur à 0,4.

4.3 Procédure expérimentale

Nous avons réalisé une étude inter-sujets avec 80 participants francophones (moyenne d'âge = 42 ans, 52 femmes, 28 hommes). Chaque participant a répondu à 30 questions correspondant aux 30 tendances à l'action dans un ordre aléatoire. Nous avons eu 29 participants pour la version A du contexte et 51 pour la version B (le choix était aléatoire et nous n'avons gardé que les participants ayant répondu à l'ensemble du questionnaire). Les participants ont eu accès à l'expérimentation à travers un lien qui a été partagé et relayé sur les réseaux sociaux et les listes de participants RISC ainsi que le personnel du laboratoire et de l'entreprise. La seule condition requise pour les participants est la maîtrise de la langue française. En effet, comme les actes de langage expriment une réaction impulsive et non réfléchie, il est difficile de comprendre les subtilités de la langue pour des personnes non-francophones.

Le questionnaire est anonyme et les participants n'ont pas été rémunérés. Le questionnaire commence par deux questions exemples qui n'ont pas été considérées (il s'agit juste de familiariser l'utilisateur avec l'interface Limesurvey et avec les situations de dialogue présentées par la suite). Les 30 situations présentées à chaque participant sont choisies de manière aléatoire et changent d'un participant à l'autre, ainsi que l'ordre des tendances à l'action proposées, pour réduire les biais éventuels.

5 Analyse des résultats

Les résultats obtenus pour les différentes hypothèses sont résumés dans le tableau 2. Ce tableau se lit de la manière suivante : la première colonne indique le nom de la tendance à l'action que nous souhaitons exprimer par l'acte locutoire présenté aux participants (cf. table 1). La deuxième colonne indique si l'hypothèse H1a a été vérifiée, c'est-à-dire si la tendance à l'action a été reconnue de manière significative.

Numéro	H1a	H1b	H1c	H2a	H2b
TA01	Non	Non	-	Non	Non
TA02	Oui	-	-	Oui	Non
TA03	Oui	-	-	Oui	Oui
TA04	Non	Non	TA06	Non	Oui
TA05	Oui	-	-	Oui	Oui
TA06	Non	Non	TA05	Non	Oui
TA07	Non	Non	TA08	Oui	Oui
TA08	Non	Non	TA07	Oui	Non
TA09	Non	TA08	-	Non	Oui
TA10	Non	Non	TA07	Non	Oui
TA11	Oui	-	-	Oui	Oui
TA12	Oui	-	-	Oui	Oui
TA13	Oui	-	-	Oui	Oui
TA14	Non	TA13	-	Non	Oui
TA15	Non	Non	-	Oui	Non
TA16	Non	TA17	-	Oui	Oui
TA17	Oui	-	-	Oui	Oui
TA18	Non	TA17	-	Non	Oui
TA19	Non	TA21	-	Non	Oui
TA20	Oui	-	-	Oui	Non
TA21	Oui	-	-	Oui	Oui
TA22	Non	TA23	-	Oui	Oui
TA23	Non	Non	TA22	Non	Oui
TA24	Oui	-	-	Oui	Oui
TA25	Oui	-	-	Oui	Oui
TA26	Non	Non	-	Non	Oui
TA27	Oui	-	-	Oui	Oui
TA28	Non	Non	-	Non	Oui
TA29	Non	Non	-	Non	Oui
TA30	Oui	-	-	Oui	Oui

TABLE 2 – La vérification des hypothèses pour toutes les tendances à l'action

Si la réponse est "Non", la colonne suivante indique si l'hypothèse H1b a été vérifiée, c'est-à-dire si une autre tendance à l'action proposée aux participants a été mieux reconnue, de manière significative. Dans les deux cas, les colonnes H2a et H2b indiquent l'impact du contexte sur les résultats. La colonne H2a indique si la tendance a été significativement reconnue dans les deux groupes de participants. La colonne H2b indique si les réponses des participants pour les cinq tendances proposées sont similaires (Kappa de Cohen).

Enfin, la colonne H1c indique, lorsque ni H1a ni H1b ne sont vérifiées, s'il est possible de confondre la tendance à l'action souhaitée avec une autre tendance à l'action, dans l'optique de les fusionner. Dans ce cas, la colonne H2a indique si le score obtenu en regroupant ces deux tendances est significativement reconnu dans chacun des deux groupes de participants (test de Wicoxon) et la colonne H2b si les réponses obtenues sont similaires (Kappa).

Étude des hypothèses

Les actes locutoires associés aux tendances à l'action TA03, TA05, TA11, TA12, TA13, TA17, TA21, TA24, TA25, TA27 et TA30 vérifient l'hypothèse H1a et l'hypothèse H2. Ils sont significativement bien associés à la tendance à l'action souhaitée, et ce indépendamment du contexte (11 actes sur 30). Il faut y ajouter les actes locutoires TA02 et TA20 qui sont bien reconnus mais pour lesquels l'accord inter-annotateur suggère une influence du contexte.

Pour les tendances à l'action TA16 et TA22, l'acte locutoire

proposé est associé à une autre tendance à l'action que celle souhaitée, et ce de manière significative et indépendante du contexte. Cela suggère que les actes locutoires peuvent être utilisés comme formulation alternative pour exprimer ces tendances à l'action (respectivement TA17 et TA23).

Il est intéressant de noter que, de manière symétrique, l'acte locutoire choisi pour TA23 vérifie l'hypothèse H1c avec la proposition TA22. Or ces deux tendances à l'action correspondent à la même stratégie émotionnelle et à la même direction. Ce résultat suggère qu'il est possible de fusionner ces deux tendances à l'action et d'utiliser indifféremment l'un des deux actes locutoires proposés.

L'acte locutoire associé à TA18 est lui aussi associé à TA17, tout comme TA16, même si le résultat n'est pas valable dans tous les contextes. Cela suggère, puisqu'il s'agit de la même stratégie émotionnelle et de la même direction pour ces trois tendances à l'actions, qu'il est préférable de se limiter à la tendance à l'action TA17.

De même, l'acte locutoire TA19 est associé à TA21 qui est dans le même groupe, et TA09 est associé à TA08, ce qui milite en faveur de leur regroupement dans la formulation la mieux reconnue. Cependant, dans le cas de TA08, le tableau montre une confusion avec TA07 qui suggère un regroupement avec deux formulations possibles. De plus, le locutoire TA10 semble se confondre avec TA07 et pourrait être accepté comme formulation alternative.

L'acte locutoire associé à la tendances à l'action TA04 semble se confondre avec TA06, laquelle est aussi confondue avec TA05. Cela suggère un regroupement autour de TA05, peut-être en utilisant TA06 comme formulation alternative, voire de TA04. Une étude statistique plus poussée permettrait de déterminer si le regroupe de ces trois tendances à l'action conduit à un score significatif.

L'acte locutoire associé à la tendance à l'action TA14 semble exprimer la tendance TA13. Comme ces deux tendances font partie de la même stratégie émotionnelle et dirigée vers la même source, il est possible de les regrouper autour de la tendance TA13.

Les actes locutoires associés aux tendances TA01, TA15, TA26, TA28 et TA29 sont plus problématiques : ils n'ont pas été reconnus par les participants et il est impossible de les regrouper avec d'autres tendances choisies par les participants. Nous pensons donc que le choix d'acte locutoire n'est pas satisfaisant pour ces tendances à l'action et qu'il faut proposer d'autres formulations.

C'est particulièrement important pour la tendance TA15 qui est la seule à exprimer la stratégie émotionnelle "Move Away From" dirigée vers autrui, pour la tendance TA28 (stratégie émotionnelle "Moving it Away From" dirigée vers l'interlocuteur) et pour la tendance TA29 (stratégie "Stop Moving Toward").

Les tendances à l'action retenues et les actes locutoires que nous considérons comme validés sont listés dans le tableau 3.

Tendance à l'action	Formulations possibles
TA02 (retirer l'obstacle)	Qu'est-ce qu'on peut faire ?
TA03 (se punir)	C'est vrai je suis nul !
TA05 (Insulter l'autre)	"Vous êtes vraiment stupide !" ou "Vous n'êtes bon à rien !" ou "Tout ça est à cause de vous !"
TA07 (Éviter le problème)	"Et si on parlait d'autre chose ?" ou "Je ne veux pas parler de cela. Arrêtons là la discussion" ou "Snif !"
TA11 (Demander de l'aide)	Ce n'est pas possible. Il faut que vous m'aidiez !
TA12 (Corriger l'erreur)	"Pardon je vais tenter de corriger."
TA13 (Minimiser)	"Ce n'est pas très grave." ou "Qu'est-ce que je peux faire pour me faire pardonner ?"
TA17 (Partager sa bonne humeur)	"ça me fait très plaisir !" ou "Youpi" ou "Cool !"
TA20 (S'afficher)	Je sais je suis parfait !
TA21 (Chercher la reconnaissance)	"J'aime recevoir des compliments" ou "Tu as vu comme je suis bon ! Il faut que vous le disiez à tout le monde !"
TA23 (Chercher le contact de l'autre)	"J'aimerais vous parler plus souvent." ou "J'aimerais continuer à vous parler."
TA24 (Rejeter la situation)	Je n'aime pas ça !
TA25 (Se soumettre)	Je suis prêt à faire ce que vous voulez.
TA27 (Se cacher)	J'aimerais disparaître.
TA30 (Se détendre)	Ouf ! C'est mieux comme ça !

TABLE 3 – Actes locutoires retenus après expérimentation.

6 Conclusion et perspectives

Nous avons mené une étude expérimentale sur un ensemble de situations dialogiques où des tendances à l'action ont été déclenchées et exprimées verbalement et textuellement (sans prosodie). Nous nous sommes basés sur des travaux antérieurs en linguistique et psychologie du langage pour choisir les actes locutoires correspondant aux 30 tendances à l'action sélectionnées dans la littérature. Notre objectif à travers cette étude était de valider notre choix d'actes locutoires pour chaque tendance à l'action et de bien s'assurer que ces derniers sont indépendants de la situation. Nous avons également fait l'hypothèse que, dans un contexte purement dialogique, certaines tendances à l'action vont être perçues d'une façon similaire.

Les résultats obtenus sont globalement encourageants. Nous avons retenus 16 tendances à l'action parmi les 30. Cette étude nous a permis de conclure que certaines tendances à l'action peuvent être regroupées. Ce qui fait que certaines peuvent s'exprimer par plusieurs actes locutoires. Nous avons 6 tendances à l'action qui n'ont pas été perçues par les participants et dont les résultats ne nous permettent pas de conclure quant à nos hypothèses. Nous envisageons de faire une deuxième étude expérimentale en modifiant notre choix d'actes locutoires pour étudier ces 6 cas.

Nous sommes conscients que la limitation d'un seul acte locutoire par tendance à l'action entraîne un risque de répétition dans les dialogues. Nous envisageons d'élargir cette

expérimentation avec d'autres actes locutoires pour diversifier le choix de l'acte à mettre dans un dialogue. Aussi, nous estimons que l'indépendance au contexte pourrait être mieux validée par la présentation d'autres variantes avec différentes situations possibles.

Ce travail a été fait dans le cadre d'une évaluation d'un modèle affectif chez un agent conversationnel. Il pourrait être utilisé dans des situations similaires où des tendances à l'action sont exprimées verbalement.

7 Remerciements

Nos remerciements à Rachel Bawden pour son aide précieuse lors de la traduction des termes de tendances à l'action de l'anglais au français.

Références

- [1] Osnat Argaman. Linguistic markers and emotional intensity. *Journal of psycholinguistic research*, 39(2) :89–99, 2010.
- [2] Evelien Bossuyt. *Experimental studies on the influence of appraisal on emotional action tendencies and associated feelings*. PhD thesis, Ghent University, 2012.
- [3] Sabrina Campano and Nicolas Sabouret. A socio-emotional model of impoliteness for non-player characters. In *Affective Computing and Intelligent Interaction and Workshops, 2009. ACII 2009. 3rd International Conference on*, pages 1–7. IEEE, 2009.
- [4] Jacob Cohen. Weighted kappa : Nominal scale agreement provision for scaled disagreement or partial credit. *Psychological bulletin*, 70(4) :213, 1968.
- [5] Laurence Devillers, Sophie Rosset, Hélène Bonneau-Maynard, and Lori Lamel. Annotations for dynamic diagnosis of the dialog state. In *In Proc. LREC*, 2002.
- [6] Nico H Frijda. The emotions : Studies in emotion and social interaction. *Paris : Maison de Sciences de l'Homme*, 1986.
- [7] Nico H Frijda. The laws of emotion. *American psychologist*, 43(5) :349, 1988.
- [8] Nico H Frijda, Peter Kuipers, and Elisabeth Ter Schure. Relations among emotion, appraisal, and emotional action readiness. *Journal of personality and social psychology*, 57(2) :212, 1989.
- [9] Patrick Gebhard. Alma : a layered model of affect. In *Proceedings of the fourth international joint conference on Autonomous agents and multiagent systems*, pages 29–36. ACM, 2005.
- [10] Jonathan Gratch and Stacy Marsella. A domain-independent framework for modeling emotion. *Cognitive Systems Research*, 5(4) :269–306, 2004.
- [11] Richard S Lazarus. Cognition and motivation in emotion. *American psychologist*, 46(4) :352, 1991.
- [12] Robert W Levenson. Basic emotion questions. *Emotion review*, 3(4) :379–386, 2011.
- [13] Irina Maslowski. Quelles sont les caractéristiques des interactions problématiques entre des utilisateurs et un conseiller virtuel ? In *JEP-TALN-RECITAL 2016*, volume 3, pages 94–107, 2016.
- [14] Albert Mehrabian. Pleasure-arousal-dominance : A general framework for describing and measuring individual differences in temperament. *Current Psychology*, 14(4) :261–292, 1996.
- [15] Radoslaw Niewiadomski, Elisabetta Bevacqua, Maurizio Mancini, and Catherine Pelachaud. Greta : an interactive expressive eca system. In *Proceedings of The 8th International Conference on Autonomous Agents and Multiagent Systems-Volume 2*, pages 1399–1400. International Foundation for Autonomous Agents and Multiagent Systems, 2009.
- [16] Elinor Ochs and Bambi Schieffelin. Language has a heart. *Text-Interdisciplinary Journal for the Study of Discourse*, 9(1) :7–26, 1989.
- [17] Magalie Ochs, Nicolas Sabouret, and Vincent Corruble. Simulation of the dynamics of nonplayer characters' emotions and social relations in games. *IEEE Transactions on Computational Intelligence and AI in Games*, 1(4) :281, 2009.
- [18] Sally Planalp. *Communicating emotion : Social, moral, and cultural processes*. Cambridge University Press, 1999.
- [19] Kaska Porayska-Pomsta and Helen Pain. Providing cognitive and affective scaffolding through teaching strategies : applying linguistic politeness to the educational context. In *International Conference on Intelligent Tutoring Systems*, pages 77–86. Springer, 2004.
- [20] Ronald H Randles. Wilcoxon signed rank test. *Encyclopedia of statistical sciences*, 1988.
- [21] Ira J Roseman. Emotional behaviors, emotivational goals, emotion strategies : Multiple levels of organization integrate variable and consistent responses. *Emotion Review*, 3(4) :434–443, 2011.
- [22] Ira J Roseman. Appraisal in the emotion system : Coherence in strategies for coping. *Emotion Review*, 5(2) :141–149, 2013.
- [23] Giovanni Semeraro, Hans HK Andersen, Verner Andersen, Pasquale Lops, and Fabio Abbattista. Evaluation and validation of a conversational agent embodied in a bookstore. In *ERCIM Workshop on User Interfaces for All*, pages 360–371. Springer, 2002.
- [24] Robert C Solomon. The logic of emotion. *Nous*, pages 41–49, 1977.
- [25] Alya Yacoubi and Nicolas Sabouret. TEATIME : a Formal Model of Action Tendencies in Conversational Agents. In *In Proc. 10th International Conference on Agents and ARTificial intelligence (ICAART)*, 2018.

Classification d'images en apprenant sur des échantillons positifs et non labélisés avec un réseau antagoniste génératif

F. Chiaroni¹²

M-C. Rahal¹

F. Dufaux²

N. Hueber³

¹ Institut VEDECOM, équipe Perception du véhicule à conduite déléguée

² L2S, CNRS, CentraleSupélec, Univ Paris-Sud, Univ Paris-Saclay

³ Institut Saint-Louis Franco-Allemand (ISL), équipe ELSI

florent.chiaroni@l2s.centralesupelec.fr

Abstract

In this article, we suggest a novel approach for image classification task from a positive unlabeled dataset. Its proper functioning is based on generative adversarial networks (GANs) abilities. These allow us to generate fake images whose distribution is close to the distribution of the negative samples included in the unlabelled dataset available, while remaining different from the distribution of positive samples that are not labeled. Then we train a CNN classifier with the positive samples and the fake samples generated, as it would have been done with a classical Positive Negative dataset. Tests performed on three different image classification datasets show that the system is stable in its behavior with a non negligible fraction of positive samples present in the unlabeled dataset. Although very different, this method outperforms the state of the art in PU learning on the RGB CIFAR-10 dataset.

Résumé

Dans ce document, nous proposons une nouvelle approche répondant à la tâche de classification d'images à partir d'un apprentissage sur données positives et non-labélisées. Son bon fonctionnement repose sur certaines particularités des réseaux antagonistes génératifs (GANs). Ces derniers nous permettent de générer des fausses images dont la distribution se rapproche de la distribution des échantillons négatifs inclus dans le jeu de données non labélisé disponible, tout en restant différente de la distribution des échantillons positifs non labélisés. Ensuite, nous entraînons un classifieur convolutif avec les échantillons positifs et les faux échantillons générés, tel que cela aurait été fait avec un jeu de données classique de type Positif Négatif. Les tests réalisés sur trois jeux de données différents de classification d'images montrent que le système est stable dans son comportement jusqu'à une fraction conséquente d'échantillons positifs présents dans le jeu de données non labélisé. Bien que très différente, cette méthode surpasse l'état de l'art PU learning sur le jeu de données RVB CIFAR-10.

Mots Clef

Apprentissage Positif Non Labélisé (PU learning), Classification d'Images, Apprentissage Profond, Apprentissage de Représentations, Modèles Génératifs.

1 Introduction

Les méthodes d'apprentissage utilisant des filtres à noyaux de convolution ont démontré de bonnes performances de prédiction dans le domaine du traitement d'image, et plus particulièrement pour la tâche de classification d'images. Pour réaliser de telles performances, de grands jeux de données entièrement labélisés sont requis. De nos jours, plusieurs jeux de données distincts peuvent être amenés à être fusionnés pour cette raison afin d'augmenter la capacité de généralisation d'un modèle d'apprentissage tel que cela est proposé dans YOLO9000 [18]. Par ailleurs, pour atténuer ce besoin de grands jeux de données labélisés, des méthodes d'apprentissage semi-supervisé existent [16]. Mais, si un objet n'appartenant à aucune classe labélisée du jeu de données d'entraînement doit être traité, il reste difficile de prédire le comportement du modèle entraîné à son égard.

Néanmoins, une idée pouvant répondre à ce problème consiste à se focaliser principalement sur les données qui nous intéressent. Cela est le cas pour les méthodes de One-Class Classification (OCC) [8], détection de nouveauté [14] où il est utilisé uniquement des échantillons de la classe d'intérêt ; la classe positive. Cependant, à notre connaissance, les méthodes OCC ont une performance limitée lorsqu'elles sont appliquées à des tenseurs de données de grande dimensionnalité tels que des images. De plus, il est souvent facile d'acquérir des échantillons non labélisés susceptibles de contenir des informations pertinentes à propos des contre-exemples de la classe d'intérêt. De cette manière, nous abordons le problème d'apprentissage Positif Non labélisé (apprentissage PU). Il se trouve que les méthodes d'apprentissage Positif Non labélisé ont été appliquées récemment à des données de type images tel que la méthode Rank Pruning (RP) [13]. Cette méthode est la plus performante de l'état de l'art dans un contexte où l'on n'a pas de connais-

sances à priori sur les fractions d'échantillons bruités. Elle est cependant coûteuse en calculs car elle consiste à réaliser plusieurs entraînements consécutifs du même classifieur de manière à éliminer les échantillons les moins pertinents pendant la phase d'entraînement. De plus, selon [12], ces méthodes deviennent compétitives lorsque le nombre d'échantillons non labélisés dans le jeu de données d'entraînement augmente considérablement. Cela est un avantage lorsque l'on peut obtenir facilement des données non labélisées. Par ailleurs, les réseaux génératifs antagonistes (GANs) ont attiré notre attention en raison de leur capacité à générer de faux échantillons x_F qui ont une distribution $p_G(x_F)$ qui tend vers la distribution $p_{data}(x_R)$ des échantillons réels x_R utilisés pendant son entraînement. Le GAN originel [5] contient un modèle génératif G et un modèle discriminatif D . Ces deux modèles possèdent une structure de type perceptron multi-couche. Un vecteur de bruit z , composé de variables aléatoires continues, est placé en entrée de G . D est entraîné à distinguer les échantillons réels des faux échantillons générés par G , pendant que ce dernier est entraîné à produire des faux échantillons qui doivent sembler réels au possible. Cet entraînement adversaire consiste à utiliser la fonction d'évaluation minimax $V(G, D)$:

$$\min_G \max_D V(G, D) = \mathbb{E}_{x_R \sim p_{data}(x_R)} \log D(x_R) + \mathbb{E}_{z \sim p_z(z)} \log[1 - D(G(z))].$$

Lorsque D ne peut plus distinguer les vrais échantillons des faux, nous obtenons la propriété suivante, avec y_D le scalaire de sortie prédit :

$$p_G(x_F) \xrightarrow[y_D \rightarrow \frac{1}{2}]{} p_{data}(x_R).$$

D'autres variantes du GAN sont apparues telles que le DC-GAN [15], qui adapte sa structure au traitement d'images en intégrant des couches convolutives. Le Wasserstein GAN (WGAN) [1] utilise d'une part la distance *Earth - Mover* (*EM*) dans sa fonction de coût, et d'autre part limite les valeurs des poids de son modèle à un certain intervalle, afin de rectifier le problème d'instabilité des précédentes versions du GAN.

En raison de leur capacité à apprendre des représentations pertinentes d'un point de vue sémantique et de leur efficacité déjà démontrée en apprentissage semi-supervisé [20], nous avons décidé d'exploiter d'une certaine manière leurs avantages pour une application d'apprentissage PU. Parallèlement à notre étude, l'approche [7] est apparue pour répondre à la même problématique en utilisant un modèle d'apprentissage de type GAN. [7] requière deux modèles génératifs et trois discriminateurs pour l'étape générative contre un générateur et un discriminateur pour notre approche. Cela devrait nous conférer un temps de calcul moindre et un apprentissage mieux maîtrisé. Mais leur étude s'est arrêtée à la description fonctionnelle de leur modèle¹. Ici, l'approche proposée que l'on nomme Positive-GAN ("PGAN" par la

suite) a été testée sur trois jeux de données différents et dont les résultats sont très prometteurs en termes robustesse et de performance de prédiction pour le traitement d'images complexes. Il surpasse l'état de l'art sur le jeu de données le plus difficile que nous avons testé.

Le document est organisé tel que ci-dessous. Dans la section suivante nous présentons la méthode. Les expérimentations et les résultats sont présentés dans la troisième section. Pour finir, une conclusion est faite sur notre approche et de futures directions de recherche sont suggérées.

2 Méthode d'apprentissage proposée : Le Positive-GAN

Dans cette section, nous décrivons notre système d'apprentissage PU de manière générique et focalisons la description sur la méthode d'entraînement. Notre méthode d'apprentissage Positive-GAN (PGAN) consiste à substituer l'absence d'échantillons négatifs labélisés x_N avec les faux échantillons x_F générés par notre GAN, dont la distribution est proche au possible de celle des x_N , tout en étant différente de celle des échantillons positifs x_P . La figure 1 illustre le fonctionnement du système.

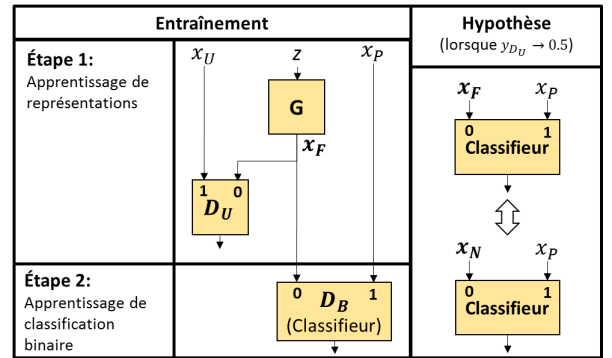


FIGURE 1 – Système d'apprentissage PU proposé : Positive-GAN.

Lors de l'étape 1, le GAN est entraîné avec les échantillons non labélisés x_U à partir du jeu de données d'entraînement

RJCIA, une nouvelle version de la publication [7] est apparue, en version pre-print, où l'étude de leur méthode a été complétée : Leur nouvelle version datant du 4 avril 2018 inclue une partie théorique et quelques tests comparatifs sur les jeux de données USPS et MNIST auxquels il aurait été pertinent de se comparer, s'il n'y avait pas eu de conflits de dates. Leur nouvelle version nécessite la connaissance à priori de la fraction d'échantillons positifs inclus dans le jeu de données non labélisé pour fonctionner. Les tests ont été réalisés avec une faible proportion d'échantillons positifs labélisés, ce qui met en avant l'intérêt des méthodes génératives pour l'augmentation de dataset. Cependant, des détails sur l'initialisation des hyper-paramètres λ_P , λ_N et λ_U auraient été appréciés, ainsi que la réalisation de tests sur des bases de données plus complexes telles que CIFAR-10 avec des structures de type réseaux convolutifs. En effet, bien que leur méthode semble fonctionnelle avec une structure de type perceptron multi-couche, les étapes de convolution rendent plus difficile l'effondrement d'un générateur et donc la divergence de leur générateur G_N pour l'apprentissage de la distribution des échantillons négatifs à partir d'échantillons non labélisés et d'échantillons positifs labélisés.

1. À noter qu'après la date de soumission de notre article à CNIA-

PU qui contient une fraction $\pi \in (0, 1)$ d'échantillons positifs et une fraction $1 - \pi$ d'échantillons négatifs x_N . Le système Positif-Non labélisé inclue trois modèles convolutifs avec différents rôles respectifs :

- Le modèle discriminateur D_U est entraîné à distinguer les vrais échantillons x_U des faux non labélisés générés x_F , avec $y_{D_U} \in (0, 1)$ sa valeur de sortie prédite.
- Le modèle génératif G prend en entrée un vecteur de bruit z constitué de variables aléatoires continues, et fournit en sortie, dans le même format que x_U , les faux échantillons $x_F = G(z)$. G est entraîné de manière antagoniste à D_U afin de générer des faux échantillons tels que leur distribution $p(x_F)$ tend vers $p(x_U)$.
- Lors de la deuxième étape, une fois que l'entraînement du GAN est considéré comme terminé, le classifieur binaire convolutif D_B est entraîné à distinguer les échantillons réels positifs x_P des faux échantillons x_F .

Les explications présentées ci-dessous ont pour objectif de développer l'intuition derrière le système proposé.

Nous rappelons que le jeu de données non labélisé est composé d'une fraction π d'échantillons positifs x_P et d'une fraction $1 - \pi$ d'échantillons négatifs x_N . Ainsi, si le GAN est correctement entraîné sur les échantillons non labélisés x_U , on peut en déduire que :

$$\begin{aligned} p(x_F) &\xrightarrow{y_{D_U} \rightarrow \frac{1}{2}} p(x_U) \\ \Leftrightarrow p(x_F) &\xrightarrow{y_{D_U} \rightarrow \frac{1}{2}} \pi p(x_P) + (1 - \pi) p(x_N), \end{aligned}$$

et l'on admet alors comme forte hypothèse pour la suite que $p(x_F) = \pi p(x_{FP}) + (1 - \pi) p(x_{FN})$, avec :

$$\begin{cases} p(x_{FP}) \xrightarrow{y_{D_U} \rightarrow \frac{1}{2}} p(x_P) \\ p(x_{FN}) \xrightarrow{y_{D_U} \rightarrow \frac{1}{2}} p(x_N). \end{cases}$$

Lorsque $y_{D_U} \rightarrow \frac{1}{2}$, nous démarrons la deuxième étape du PGAN. Par ailleurs, un GAN n'est pas parfait dans son fonctionnement lorsqu'il se voit être appliqué à des tenseurs de grandes dimensions, ainsi :

$$p(x_{FP}) \neq p(x_P), \text{ et } p(x_{FN}) \neq p(x_N). \quad (1)$$

Il est donc alors possible d'estimer une distance d non nulle dans la fonction de coût du classifieur D_B , tel que :

$$d(p(x_P), p(x_F)) \Leftrightarrow \begin{cases} d(p(x_P), p(x_{FP})) \\ d(p(x_P), p(x_{FN})). \end{cases}$$

Mais, bien que calculée, la distance $d(p(x_P), p(x_{FP}))$ n'est pas exploitée dans l'application finale où nous traitons uniquement des échantillons réels avec le classifieur D_B . Ainsi, lorsque $p(x_{FN}) \xrightarrow{y_{D_U} \rightarrow \frac{1}{2}} p(x_N)$ et que D_B a été également correctement entraîné, nous obtenons l'équivalence :

$$d(p(x_P), p(x_{FN})) \Leftrightarrow d(p(x_P), p(x_N)). \quad (2)$$

Nous sommes donc capable de calculer la distance qui nous intéresse. En transférant ce raisonnement dans notre méthode PU, cela revient à affirmer les équivalences suivantes à la sortie de la fonction de coût L_{D_B} du classifieur D_B lorsque $y_{D_U} \rightarrow \frac{1}{2}$:

$$\begin{aligned} L_{D_B} &= \mathbb{E}_{x_P \sim p(x_P)} \log D_B(x_P) \\ &\quad + \mathbb{E}_{z \sim p_z(z)} \log [1 - D_B(G(z))] \\ \Leftrightarrow L_{D_B} &= \mathbb{E}_{x_P \sim p(x_P)} \log D_B(x_P) \\ &\quad + \mathbb{E}_{x_N \sim p(x_N)} \log [1 - D_B(x_N)]. \end{aligned}$$

Ainsi, grâce à l'hypothèse de fonctionnement proposée ci-dessus on peut affirmer que la méthode PGAN devient similaire à un entraînement sur des échantillons positifs et négatifs respectivement labélisés, tout en s'éloignant d'un entraînement de type apprentissage PU, malgré le fait que le jeu de données que nous utilisons contienne uniquement des échantillons positifs labélisés et des échantillons non labélisés. De plus, intuitivement, calculer $d(p(x_P), p(x_{FP}))$ peut favoriser dans une certaine mesure, l'apprentissage des frontières de $p(x_P)$. Cependant, deux risques peuvent survenir avec cette méthode :

- Si les échantillons x_U contiennent majoritairement des échantillons x_P , alors il est possible que G ne soit plus apte à générer suffisamment de faux échantillons similaires aux échantillons x_N .
- Si G génère des faux échantillons ayant une distribution égale à celle des vrais échantillons, contrairement à l'inégalité 1, alors le PGAN deviendrait équivalent en termes de performances à un entraînement classique PU. Mais lorsque la dimensionalité des images à traiter devient large, alors ce risque disparaît. De plus, il peut être atténué en utilisant une structure pour le classifieur D_B dont les performances en prédiction sont meilleures que celles de la structure utilisée pour le discriminateur D_U .

Trouver une solution permettant d'éviter ces deux risques de se produire reste une question ouverte très intéressante pour améliorer la fiabilité de la méthode.

3 Expérimentations

3.1 Réglages des tests

Les expériences ont été réalisées sur les trois jeux de données MNIST [10], Fashion-MNIST [21] et CIFAR-10 [9]. Nous avons comparé notre approche à RP [13], qui est à notre connaissance la meilleure méthode d'apprentissage bruité (noisy learning) et PU ne nécessitant pas de connaissances a priori de la fraction π . De plus, l'implémentation de l'auteur est disponible². Nous indiquons aussi la performance du classifieur entraîné sur le jeu de données d'entraînement initial contenant des échantillons entièrement labélisés positifs et négatifs, et nous appelons évidemment cette méthode PN, que nous considérons comme la référence du cadre idéal. Nous comparons aussi le PGAN à un

2. <https://github.com/cgnorthcutt/rankpruning>

TABLE 1 – Résultats comparatifs en fonction des F1-Scores mesurés sur MNIST, Fashion-MNIST et CIFAR-10 avec le classifieur entraîné sur 20 époques.

	ref	$\rho = 0.5, \pi = 0.1$			$\rho = 0.5, \pi = 0.3$			$\rho = 0.5, \pi = 0.5$			$\rho = 0.5, \pi = 0.7$		
Jeux de données	PN	PU	PGAN	RP	PU	PGAN	RP	PU	PGAN	RP	PU	PGAN	RP
0	0.997	0.633	0.974	0.992	0.445	0.973	0.955	0.320	0.973	0.991	0.689	0.902	0.880
1	0.998	0.774	0.971	0.995	0.642	0.979	0.996	0.851	0.958	0.994	0.884	0.863	0.993
2	0.990	0.395	0.972	0.975	0.658	0.959	0.923	0.795	0.947	0.936	0.692	0.914	0.987
3	0.996	0.716	0.963	0.991	0.620	0.953	0.991	0.766	0.934	0.882	0.729	0.885	0.829
4	0.997	0.512	0.964	0.972	0.802	0.952	0.995	0.717	0.945	0.933	0.551	0.914	0.977
5	0.993	0.701	0.974	0.985	0.725	0.950	0.943	0.799	0.949	0.910	0.626	0.873	0.973
6	0.992	0.708	0.962	0.928	0.758	0.959	0.992	0.699	0.971	0.993	0.613	0.944	0.990
7	0.995	0.603	0.962	0.947	0.433	0.960	0.991	0.620	0.926	0.988	0.783	0.737	0.979
8	0.995	0.741	0.949	0.929	0.506	0.941	0.982	0.339	0.922	0.941	0.651	0.849	0.818
9	0.981	0.785	0.959	0.954	0.442	0.956	0.979	0.561	0.939	0.941	0.750	0.865	0.904
AVG_{MNIST}	0.993	0.657	0.965	0.967	0.603	0.958	0.975	0.647	0.946	0.951	0.697	0.875	0.933
T-shirt/top	0.908	0.724	0.926	0.899	0.206	0.91	0.937	0.821	0.873	0.947	0.695	0.802	0.91
Trouser	0.993	0.815	0.983	0.993	0.247	0.969	0.989	0.938	0.953	0.99	0.681	0.911	0.984
Pullover	0.932	0.635	0.9	0.887	0.29	0.885	0.925	0.695	0.865	0.917	0.657	0.842	0.888
Dress	0.952	0.601	0.941	0.948	0.312	0.925	0.955	0.852	0.893	0.914	0.631	0.853	0.882
Coat	0.882	0.614	0.909	0.847	0.252	0.889	0.942	0.788	0.845	0.92	0.686	0.83	0.918
Sandal	0.995	0.793	0.945	0.977	0.444	0.964	0.98	0.819	0.923	0.985	0.67	0.919	0.981
Shirt	0.818	0.446	0.852	0.758	0.398	0.846	0.847	0.797	0.819	0.873	0.554	0.792	0.853
Sneaker	0.983	0.73	0.973	0.973	0.271	0.952	0.979	0.865	0.943	0.967	0.64	0.922	0.977
Bag	0.989	0.772	0.978	0.976	0.536	0.947	0.99	0.837	0.96	0.965	0.685	0.757	0.977
Ankle boot	0.985	0.704	0.964	0.979	0.354	0.973	0.986	0.824	0.963	0.976	0.609	0.942	0.976
$AVG_{F-MNIST}$	0.944	0.683	0.937	0.924	0.331	0.926	0.953	0.824	0.904	0.945	0.651	0.857	0.935
Plane	0.727	0.341	0.818	0.669	0.557	0.784	0.795	0.295	0.758	0.743	0.621	0.731	0.718
Auto	0.78	0.506	0.801	0.695	0.492	0.737	0.829	0.414	0.789	0.798	0.521	0.734	0.783
Bird	0.447	0.175	0.688	0.56	0.439	0.744	0.68	0.184	0.694	0.644	0.359	0.688	0.542
Cat	0.5	0.125	0.658	0.384	0.272	0.722	0.651	0.249	0.718	0.67	0.446	0.69	0.698
Deer	0.698	0.272	0.68	0.605	0.232	0.708	0.708	0.3	0.708	0.64	0.43	0.633	0.602
Dog	0.567	0.2	0.632	0.539	0.37	0.756	0.648	0.258	0.746	0.733	0.514	0.678	0.712
Frog	0.691	0.35	0.837	0.666	0.418	0.793	0.794	0.256	0.788	0.769	0.693	0.75	0.749
Horse	0.786	0.373	0.693	0.653	0.515	0.757	0.723	0.26	0.751	0.759	0.611	0.675	0.711
Ship	0.832	0.313	0.821	0.764	0.565	0.809	0.831	0.324	0.775	0.785	0.623	0.716	0.755
Truck	0.771	0.462	0.822	0.685	0.367	0.786	0.637	0.272	0.754	0.617	0.539	0.724	0.564
$AVG_{CIFAR-10}$	0.680	0.312	0.745	0.622	0.423	0.760	0.730	0.281	0.748	0.716	0.536	0.702	0.684

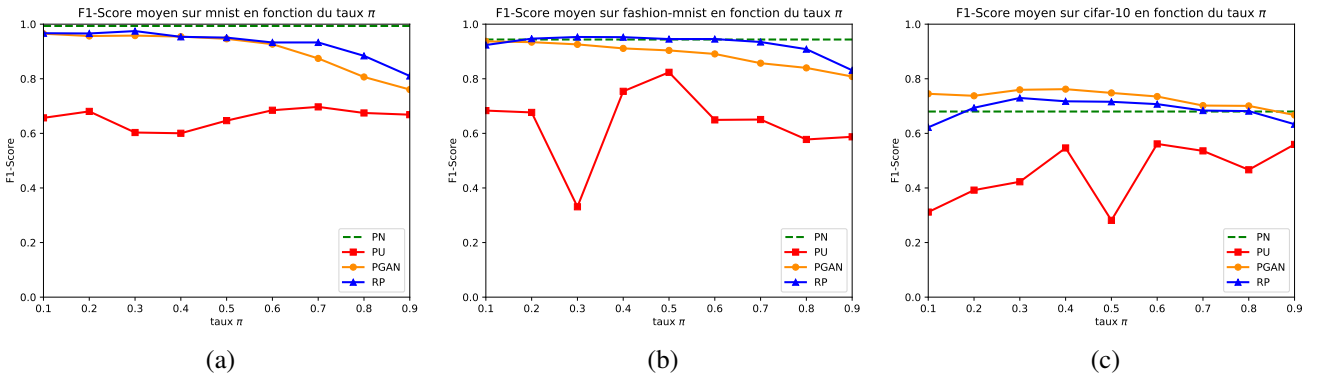


FIGURE 2 – F1-Scores moyens après 20 époques d’entraînement pour le classifieur en fonction du taux π qui varie entre 0.1 et 0.9 avec un pas de 0.1, pour PN (vert), PU (rouge), RP (bleu) et PGAN (orange) sur MNIST (a), Fashion-MNIST (b) et CIFAR-10 (c).

entraînement que l'on nomme PU, qui est équivalent à PN, mais avec une substitution des échantillons négatifs par des échantillons non labélisés.

Pour ces expérimentations, les méthodes PN, PU, RP et PGAN sont testées avec exactement le même classifieur convolutif afin d'être impartial. Nous avons utilisé le modèle convolutif de classification d'images proposé par tensorflow³ pour rester générique. Il contient deux couches convolutives successives, respectivement suivies d'une étape de max-pooling, et se finit par deux couches entièrement connectées consécutives. La fonction d'activation en sortie de chaque couche est ReLU, mis à part pour la dernière où softmax est appliquée. Nous avons uniquement modifié la dimension de sortie de la dernière couche que l'on fait passer de 10 neurones à 2, afin d'être adapté à notre tâche de classification binaire. Le classifieur est entraîné sur 20 époques. Pour les images 32x32x3 de CIFAR-10, les largeurs et longueurs des tenseurs d'entrée et de sortie des deux couches convolutives sont adaptées, et la profondeur des filtres à noyaux de la première couche convolutive est établie à 3 afin de correspondre aux trois canaux de ces images RVB. Mais le nombre de filtres et leurs largeur et longueur restent inchangés.

Pour l'étape générative du PGAN, nous avons associé la méthode d'entraînement du WGAN [1] à l'architecture du DCGAN [15] en raison de leurs performances. À noter qu'avec la distance EM , $p(x_F)$ tend vers $p(x_U)$ lorsque y_{D_U} tend vers 0. Bien que cela ne soit pas une nécessité, dans le cadre de ces expériences, le vecteur de bruit d'entrée z est constitué de variables aléatoires continues de distributions uniformes. La durée d'entraînement de notre modèle génératif dépend de la complexité du jeu de données à traiter : 10 époques pour MNIST, 20 pour Fashion-MNIST, et 100 pour CIFAR-10. Pour ce dernier, nous faisons les mêmes modifications dans la structure de D_U et G , tel que cela a été expliqué précédemment pour le classifieur.

Au sujet de la création de notre jeu de données d'entraînement, ρ correspond à la fraction d'échantillons positifs du jeu de données initial qui contient n_P échantillons positifs. Ces $\rho \times n_P$ échantillons collectés sont ensuite introduits dans notre jeu de données non-labellisé U_{train} , qui contient initialement uniquement des échantillons négatifs N dont le nombre total est n_N . π est la fraction d'échantillons positifs P que l'on impose dans le jeu de données non labellisé d'entraînement U_{train} . Pour se faire, nous retirons de U_{train} un certain nombre d'échantillons négatifs que nous n'exploitons pas, de manière à respecter π . U_{train} contient alors à la fois des N et des P selon les paramètres ρ et π . Nous établissons qu'avec $\pi \in [\frac{1}{\rho \frac{n_N}{n_P} + 1}, 1)$ et $\rho \in (0, 1)$, nous pouvons alors obtenir consécutivement, avec P_{train} l'ensemble d'échantillons positifs d'entraînement, les deux jeux de données d'entraînement suivants :

$$P_{train} = \{(1 - \rho) n_P P ; 0 N\},$$

3. https://github.com/tensorflow/tensorflow/blob/master/tensorflow/examples/tutorials/mnist/mnist_softmax.py

$$U_{train} = \{\rho n_P P ; \frac{1 - \pi}{\pi} \rho n_P N\},$$

où les notations $a P$ et $b N$ désignent respectivement a éléments positifs et b éléments négatifs.

Pour trouver les équations définissant U_{train} selon les paramètres π et ρ , et l'intervalle des valeurs possibles pour π , nous utilisons n_U qui représente le nombre total d'échantillons non labélisés contenus dans l'ensemble U_{train} , tel que :

$$\begin{aligned} U_{train} &= \{\pi n_U P ; (1 - \pi) n_U N\} \\ &= \{\rho n_P P ; (1 - \pi) \frac{\rho n_P}{\pi} N\}, \end{aligned}$$

car nous imposons $\rho n_P = \pi n_U$.

Or, pour que cela soit réalisable, il faut que $(1 - \pi) \frac{\rho n_P}{\pi}$ soit inférieur ou égal à n_N . Cela revient donc à dire que $\pi \in [\frac{1}{\frac{n_N}{\rho n_P} + 1}, 1)$.

Les résultats présentés ci-dessous sont tous réalisés avec $\rho = 0.5$ et pour plusieurs valeurs de π .

3.2 Résultats

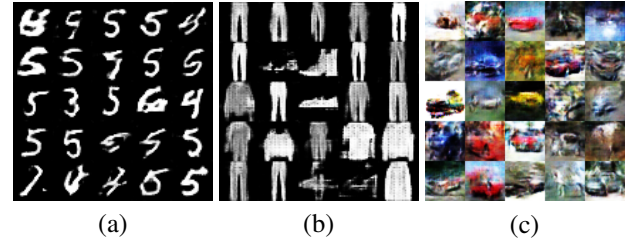


FIGURE 3 – Images générées par G avec $\rho = 0.5$ et $\pi = 0.5$ après 10 époques sur MNIST (a), 20 sur Fashion-MNIST (b), et 100 sur CIFAR-10 (c). Les classes positives respectives sont ici "5", "trouser" and "automobile".

Dans la figure 3, nous présentons quelques fausses images générées par G , respectivement pour MNIST, Fashion-MNIST et CIFAR-10. Nous pouvons remarquer que les images générées par G semblent visuellement acceptables, ce qui indique d'un point de vue qualitatif le bon fonctionnement du modèle génératif. Afin d'obtenir un tel résultat, plus les images sont grandes et complexes, et plus il faut entraîner le GAN sur un grand nombre d'époques.

Pour calculer le F1-Score, la fonction ArgMax est appliquée aux deux neurones de sortie du classifieur. Si l'indice du premier neurone est renvoyé par ArgMax, alors l'échantillon traité est classifié comme négatif. Sinon, il est considéré comme positif. De plus, étant donné que les jeux de données de test contiennent 9 fois plus d'échantillons négatifs que d'échantillons positifs, une fois toutes les prédictions de test réalisées, nous adaptons les proportions des échantillons négatifs à celle des positifs de manière à obtenir un F1-Score pertinent. Le tableau 1 montre une partie des F1-Scores comparatifs mesurés pour chaque classe pour chacun des trois jeux de données exploités, et respectivement pour chaque méthode testée. Sur Fig. 2, il peut être

observé que la méthode PN est une bonne référence sur MNIST et Fashion-MNIST. Nous trouvons que l’efficacité de la méthode d’apprentissage PGAN est équivalente à celle de la méthode RP jusqu’à $\pi = 0.5$ sur MNIST et $\pi = 0.3$ sur Fashion-MNIST. Son efficacité décline ensuite un peu plus vite que pour RP, mais tout en conservant un score acceptable. Sur CIFAR-10 le F1-Score moyen est systématiquement meilleur pour notre méthode PGAN. Aussi, notre méthode présente de meilleurs résultats que la référence PN jusqu’à $\pi = 0.8$, ce qui est très intéressant. Cela est probablement dû au fait que les échantillons générés représentent une plus grande variété de distributions d’échantillons négatifs que celles incluses dans le jeu de données initial. De plus, le F1-Score du PGAN est significativement et systématiquement meilleur que la méthode PU sur l’ensemble des trois jeux de données, même avec seulement 10% d’échantillons positifs parmi les échantillons non labélisés, autrement dit avec $\pi = 0.1$.

La figure 4 présente l’étude de la robustesse de l’approche PGAN. Les figures 4.a et 4.b montrent que la méthode PGAN a comparativement à RP une meilleure stabilité dans son fonctionnement de manière à permettre de prédire plus facilement l’évolution de son F1-Score en fonction de π pour chaque classe du jeu de données. La figure 4.c nous montre que le classifieur se stabilise et converge après 10 époques d’entraînement. Pour réaliser l’histogramme de la figure. 4.d, nous avons récupéré la valeur du deuxième neurone de sortie du classifieur qui correspond à la probabilité prédite pour une image d’appartenir à la classe positive. On peut observer que les distributions respectives des échantillons de test positifs et négatifs estimées par le PGAN sont de forme gaussienne, ce qui est une caractéristique intéressante pour des applications réelles.

TABLE 2 – Stabilité moyenne des performances (F1-Scores) des méthodes PGAN et RP en fonction de π sur les jeux de données MNIST, Fashion-MNIST et CIFAR-10

jeux de données	PGAN	RP	$\frac{E_{RP}}{E_{PGAN}}$
MNIST	0.00039	0.00172	4.410
Fashion-MNIST	0.00016	0.00025	1.563
CIFAR-10	0.00044	0.00182	4.136

En complément aux figures 4.a et 4.b, le tableau 2 présente la quantification des robustesses réalisée pour RP et PGAN en fonction de π en ce qui concerne leurs performances de prédiction, pour chacun des trois jeux de données de test. Pour ce faire, nous avons lissé respectivement la courbe $s(\pi)$ de chaque classe représentant l’évolution du F1-Score en fonction de π . Les courbes lissées $\tilde{s}(\pi)$ ont été obtenues en appliquant un filtre moyen avec un noyau de taille 3. Ensuite, l’erreur quadratique moyenne MSE pour chaque classe est calculée entre $\tilde{s}(\pi)$ et $s(\pi)$ tel que ci-dessous,

avec k le nombre d’échantillons de $\tilde{s}(\pi)$:

$$MSE = \frac{1}{k} \sum_{i=1}^k (\tilde{s}^{(i)} - s^{(i+1)})^2. \quad (3)$$

Puis, nous calculons les erreurs moyennes E_{PGAN} et E_{RP} pour chaque jeu de données, ainsi que le ratio $E_{RP} : E_{PGAN}$. On peut de cette manière constater que notre méthode a systématiquement un comportement plus stable, qui est d’un facteur 4 sur MNIST et CIFAR-10.

4 Conclusion

Ainsi, nous avons démontré que l’approche d’apprentissage PU proposée surpasse l’état de l’art sur les images RVB complexes du jeu de données CIFAR-10, et a un comportement plus stable sur l’ensemble des jeux de données testés jusqu’à une fraction acceptable π d’échantillons positifs dans le jeu de données non labélisé d’entraînement. Ces résultats sont en cohérence avec le raisonnement formulé et permettent ainsi d’envisager des applications PU sur des données de plus grandes dimensions. Le PGAN ne nécessite pas de connaissances a priori sur la fraction d’échantillons positifs non labélisés. Cependant, il reste à étudier plus en profondeur les risques de fonctionnement indiqués dans la section 1, afin de garantir un fonctionnement idéal pour cette approche.

L’optimisation du système peut se prolonger en testant d’autres récentes variantes du GAN tels que le BEGAN [3], le WGAN-GP [6], ou bien d’autres modèles génératifs de type auto-encodeurs variationnels (VAEs) par exemple, afin de généraliser l’approche proposée aux réseaux génératifs. Une autre idée peut être d’exploiter le vecteur latent z du GAN pour réaliser des opérations arithmétiques linéaires, tel que dans [4], afin de générer des faux échantillons dont on pourrait peut-être ainsi mieux gérer la distribution. Dans cette même idée, trouver un moyen d’exploiter les données positive labélisées pour la phase d’entraînement du générateur est envisagé.

Étant données les performances prometteuses obtenues, une future orientation certaine est d’étendre cette méthode à l’analyse de plus grandes images et donc permettre la réalisation de tâches plus complexes telles que la détection d’objets [19], [11], [17] ou la segmentation sémantique [2].

Références

- [1] M. Arjovsky, S. Chintala, and L. Bottou. Wasserstein generative adversarial networks. In *International Conference on Machine Learning*, pages 214–223, 2017.
- [2] V. Badrinarayanan, A. Kendall, and R. Cipolla. Segnet : A deep convolutional encoder-decoder architecture for image segmentation. *IEEE transactions on pattern analysis and machine intelligence*, 39(12) :2481–2495, 2017.

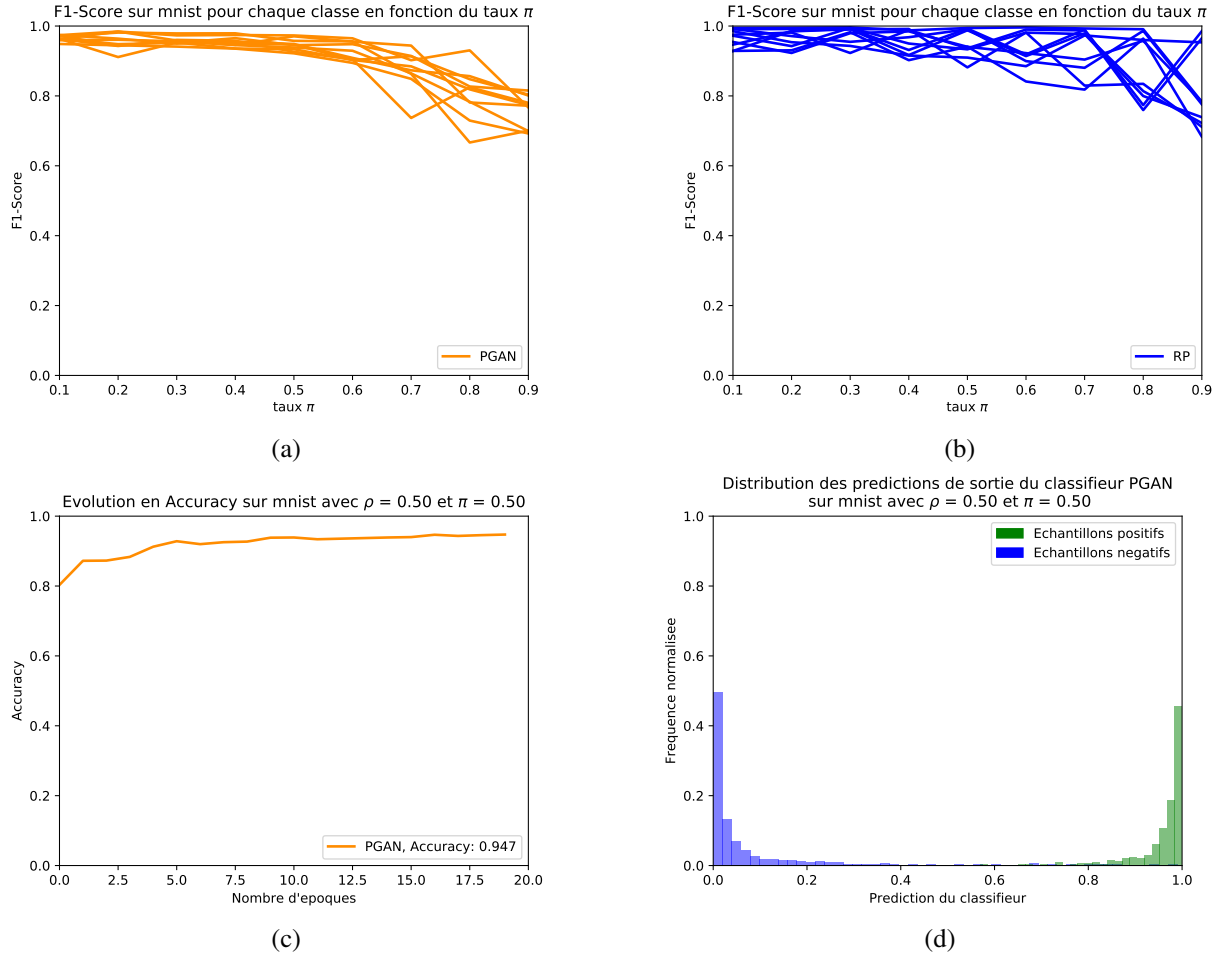


FIGURE 4 – Analyse de la robustesse sur MNIST. évolution du F1-Score pour chaque classe en fonction de π pour PGAN (a), et pour RP [13] (b). (c) montre l'évolution de l'Accuracy pendant l'entraînement du PGAN avec la classe positive "5" et $\pi = 0.5$. (d) est l'histogramme des distributions des valeurs de sortie du deuxième neurone du classifieur à son 20ème époque d'entraînement de (c) pour les échantillons de test positifs (vert) et négatifs (bleu).

- [3] D. Berthelot, T. Schumm, and L. Metz. Began : Boundary equilibrium generative adversarial networks. *arXiv preprint arXiv :1703.10717*, 2017.
- [4] P. Bojanowski, A. Joulin, D. Lopez-Paz, and A. Szlam. Optimizing the latent space of generative networks. *arXiv preprint arXiv :1707.05776*, 2017.
- [5] I. Goodfellow, J. Pouget-Abadie, M. Mirza, B. Xu, D. Warde-Farley, S. Ozair, A. Courville, and Y. Bengio. Generative adversarial nets. In *Advances in neural information processing systems*, pages 2672–2680, 2014.
- [6] I. Gulrajani, F. Ahmed, M. Arjovsky, V. Dumoulin, and A. C. Courville. Improved training of wasserstein gans. In *Advances in Neural Information Processing Systems*, pages 5769–5779, 2017.
- [7] M. Hou, Q. Zhao, C. Li, and B. Chaib-draa. A generative adversarial framework for positive-unlabeled classification. *arXiv preprint arXiv :1711.08054*, 2017.
- [8] S. S. Khan and M. G. Madden. One-class classification : taxonomy of study and review of techniques. *The Knowledge Engineering Review*, 29(3) :345–374, 2014.
- [9] A. Krizhevsky and G. Hinton. Learning multiple layers of features from tiny images. 2009.
- [10] Y. LeCun, L. Bottou, Y. Bengio, and P. Haffner. Gradient-based learning applied to document recognition. *Proceedings of the IEEE*, 86(11) :2278–2324, 1998.
- [11] W. Liu, D. Anguelov, D. Erhan, C. Szegedy, S. Reed, C.-Y. Fu, and A. C. Berg. SSD : Single shot multibox detector. In *European Conference on Computer Vision*, pages 21–37. Springer, 2016.
- [12] G. Niu, M. C. du Plessis, T. Sakai, Y. Ma, and M. Sugiyama. Theoretical comparisons of positive-unlabeled learning against positive-negative learning. In *Advances in Neural Information Processing Systems*, pages 1199–1207, 2016.

- [13] C. G. Northcutt, T. Wu, and I. L. Chuang. Learning with confident examples : Rank pruning for robust classification with noisy labels. In *Proceedings of the Thirty-Third Conference on Uncertainty in Artificial Intelligence*, UAI'17. AUAI Press, 2017.
- [14] M. A. Pimentel, D. A. Clifton, L. Clifton, and L. Tarassenko. A review of novelty detection. *Signal Processing*, 99 :215–249, 2014.
- [15] A. Radford, L. Metz, and S. Chintala. Unsupervised representation learning with deep convolutional generative adversarial networks. *arXiv preprint arXiv :1511.06434*, 2015.
- [16] A. Rasmus, M. Berglund, M. Honkala, H. Valpola, and T. Raiko. Semi-supervised learning with ladder networks. In *Advances in Neural Information Processing Systems*, pages 3546–3554, 2015.
- [17] J. Redmon, S. Divvala, R. Girshick, and A. Farhadi. You only look once : Unified, real-time object detection. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 779–788, 2016.
- [18] J. Redmon and A. Farhadi. YOLO9000 : Better, Faster, Stronger. *arXiv preprint arXiv :1612.08242*, 2016.
- [19] S. Ren, K. He, R. Girshick, and J. Sun. Faster R-CNN : Towards real-time object detection with region proposal networks. In *Advances in neural information processing systems*, pages 91–99, 2015.
- [20] T. Salimans, I. Goodfellow, W. Zaremba, V. Cheung, A. Radford, and X. Chen. Improved techniques for training gans. In *Advances in Neural Information Processing Systems*, pages 2234–2242, 2016.
- [21] H. Xiao, K. Rasul, and R. Vollgraf. Fashion-mnist : a novel image dataset for benchmarking machine learning algorithms, 2017.

Inférence bayésienne pour l'estimation de déformations larges par champs gaussien: application au recalage d'images multi-modales

Thomas Deregnaucourt¹

Chafik Samir¹

Anne-Françoise Yao²

¹LIMOS, CNRS UMR 6620, Université Clermont Auvergne, France

²LMBP, CNRS UMR 6158, Université Clermont Auvergne, France

thomas.deregnaucourt@uca.fr - chafik.samir@uca.fr - anne.yao@uca.fr

Résumé

Le problème de recalage d'images consiste à estimer la déformation globale entre une image source I_1 et une image cible I_2 . Dans ce contexte, nous nous intéressons à l'estimation d'un champ de déformation U sur le domaine $\Omega = [0, 1]^2$ de I_1 sachant U sur un ensemble fini de courbes $\beta \in \Omega$. Pour ce faire, nous proposons une nouvelle méthode basée sur des modèles gaussiens pour recalage des images multimodales. La méthode proposée commence par résoudre le problème de correspondance entre les courbes β s puis estime le déplacement sur tout Ω . La solution optimale est calculée à l'aide du maximum de vraisemblance et de l'inférence bayésienne. D'après les résultats obtenus sur des données réelles et simulées, la déformation résultante a l'avantage d'être exacte sur les observations et d'être lisse sur Ω .

Mots Clef

Recalage d'images, Statistique spatiale, Processus gaussiens, Inférence bayésienne.

Abstract

Image registration aims to estimate the global deformation between a target image I_1 and a reference image I_2 . In this context, we will focus on estimating a random field U on the I_1 domain $\Omega = [0, 1]^2$ based on observations of U on a finite set of curves $\beta \in \Omega$. Indeed, we present a new multimodal image registration method based on Gaussian random fields. The proposed method first find the optimal correspondences between curves β s then estimate the deformation vector field on Ω . The optimal solution is computed using Maximum Likelihood and Bayesian inference. Based on results using both real and simulated data, the resulting deformation has the advantage of being exact on the observations as being sufficiently smooth over the whole Ω .

Keywords

Image registration, Spatial statistics, Gaussian process, Bayesian inference.

1 Introduction

Le recalage d'images est une méthode qui vise à estimer la transformation, soumise à certaines contraintes, d'une image source I_1 vers une image cible I_2 , afin de fusionner leurs informations complémentaires. Cette méthode est utilisée dans de nombreux domaines d'applications [11, 15, 2]. En imagerie médicale on utilise le recalage d'images pour détecter des maladies, valider un traitement, comparer les données du patient avec des atlas anatomiques, etc. [12]. L'estimation de cette déformation est basée soit sur les intensités, soit sur des caractéristiques géométriques, soit sur les deux [12, 10]. Dans le premier cas, on cherche une transformation conservant la correspondance entre les niveaux de gris, et dans le second la correspondance entre des points, des courbes, etc.

Pour cet article, nous nous sommes intéressés au problème de l'endométrie. Cette maladie est provoquée par l'apparition de muqueuse utérine, aussi appelé endomètre, en dehors de la cavité utérine. Elle touche approximativement 10% des femmes en âge de procréer, et peut provoquer divers symptômes tels que des douleurs pelviennes chroniques, une dysenterie sévère, une infertilité, etc. [3]. Il est alors nécessaire de pouvoir détecter si une patiente à l'endométrie afin de la traiter efficacement, que ce soit par des antalgiques, des traitements hormonaux, ou par chirurgie dans les cas les plus sévères. L'endomètre pouvant pénétrer d'autres tissus et organes, les méthodes de détection de l'endométrie se basent sur plusieurs modalités d'images, donnant des informations complémentaires. Plus précisément, l'échographie permet d'avoir une estimation de l'infiltration de l'endomètre dans d'autres tissus, et l'imagerie par résonance magnétique (IRM) une position précise des kystes [3]. La fusion des données IRM/échographie permet alors d'avoir un diagnostic précis, mais nécessite un recalage entre ces deux modalités.

Dans notre cas, les deux images I_1 et I_2 représentent respectivement l'échographie et l'IRM d'un même organe. Cependant, comme le montre la Figure 1, les modalités d'images ont des distributions d'intensités différentes, ce qui rend inefficace les méthodes de recalage basées sur les intensités. D'autre part, un spécialiste peut extraire le con-

tour des organes présents dans les deux images. Nous allons alors utiliser ces dernières pour effectuer le recalage. Pour ce faire, nous estimerons le champ de déformation entre les deux images, à l'aide des champs gaussiens.

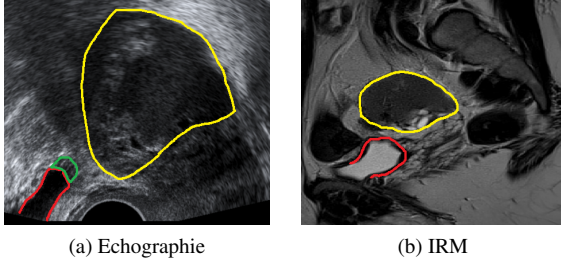


Figure 1: Exemple de courbes extraites manuellement issues d'images multimodales représentant les mêmes organes (en rouge et en jaune respectivement) : échographie à gauche et IRM à droite.

Le reste du papier est organisé de la manière suivante. Dans la section 2 nous formaliserons notre problème, et expliquerons en détail notre méthode de résolution. Puis nous présenterons nos résultats, sur des données synthétiques et réelles, dans la section 3. Enfin, la section 4 conclura ce papier.

2 Méthodologie

2.1 Formulation du problème

Soient Ω un domaine fermé de $\mathbb{R}^2$, $\beta_1 \subset I_1$ les courbes de l'échographie et $\beta_2 \subset I_2$ celles de l'IRM. On cherche à estimer un champ de déformation Ψ sur Ω transformant I_1 en I_2 . Ce champ doit déformer β_1 en β_2 , tout en étant lisse. Pour résoudre ce problème, il nous faut :

1. Trouver une correspondance optimale entre β_1 et β_2 .
2. Estimer une déformation Ψ induite par un champ de déplacement U , c'est-à-dire telle que :

$$\begin{aligned} \Psi : \Omega &\rightarrow \Omega \\ X &\mapsto \Psi(X) = X + U(X) \end{aligned}$$

et vérifiant la contrainte $\Psi(\beta_1) = \beta_2$

Nous allons tout d'abord nous intéresser au premier problème.

2.2 Correspondance optimale entre courbes

Afin de trouver une correspondance optimale entre les courbes, nous adaptons les travaux de Srivastava et al. [13]. Dans ce papier, les auteurs s'intéressaient à l'analyse des formes, et cherchaient une invariance aux transformations préservant la forme, c'est-à-dire à la translation, la rotation, la mise à l'échelle et la re-paramétrisation. Dans notre cas, la translation et la rotation sont déjà fixées pour

toute l'image, et la mise à l'échelle n'est pas une nuisance. Ainsi, nous cherchons seulement l'invariance à la re-paramétrisation. Par brièveté, nous ne décrivons le processus que pour les courbes ouvertes, mais celui-ci peut être étendu simplement à des courbes fermés [7].

Soit $\beta : [0, 1] \rightarrow \mathbb{R}^2$ une courbe ouverte paramétrisée. On utilise par la suite la représentation square-root velocity function (SRVF) q de β , défini par :

$$\begin{aligned} q : [0, 1] &\rightarrow \mathbb{R}^2 \\ t &\mapsto \frac{\beta'(t)}{\sqrt{\|\dot{\beta}(t)\|_2}} \end{aligned}$$

Le mapping $\beta \iff (\beta(0), q)$ étant une bijection, il est possible de revenir aux courbes originales en stockant le premier point de ces dernières. On note $\mathcal{C}$ l'espace de SRVFs :

$$\mathcal{C} = \left\{ q \in \mathbb{L}^2([0, 1], \mathbb{R}^2) \mid \int_0^1 \|q(t)\|_2^2 dt = 1 \right\}$$

Comme on recherche une représentation des courbes in-variantes aux re-paramétrisations, nous allons utiliser des classes d'équivalence. Nous définissons d'abord le groupe des re-paramétrisations Γ :

$$\Gamma = \{ \gamma : [0, 1] \rightarrow [0, 1] \mid \gamma(0) = 0, \gamma(1) = 1, 0 < \dot{\gamma} < \infty \}$$

La re-paramétrisation d'une courbe β par $\gamma \in \Gamma$ est donnée par $\beta \circ \gamma$, et la SRVF de cette courbe re-paramétrisée est alors $(q \circ \gamma)\sqrt{\dot{\gamma}}$. Ainsi, pour unifier tous les éléments de $\mathcal{C}$ représentant la même courbe, on définit nos classes d'équivalence par $[q] = \{(q \circ \gamma)\sqrt{\dot{\gamma}} \mid \gamma \in \Gamma\}$. On note $\mathcal{S} = \mathcal{C}/\Gamma = \{[q], q \in \mathcal{C}\}$ l'ensemble des classes d'équivalence. Afin de comparer deux courbes, on impose la métrique $\mathbb{L}^2$ à $\mathcal{S}$. Sous la représentation SRVF, la métrique $\mathbb{L}^2$ correspond à une métrique élastique sur l'espace original des courbes [8], ce qui permet de déformer les courbes pour avoir la correspondance optimale. Cette déformation optimale entre deux points sur $\mathcal{S}$ est obtenue par le chemin géodésique, et la distance entre elles est définie par la longueur du chemin.

Afin de voir comment ceci peut résoudre notre problème, on note q_1 et q_2 la représentation SRVF respective des deux courbes β_1 et β_2 . Afin de calculer la géodésique entre leurs classes d'équivalence $[q_1]$ et $[q_2]$, on fixe q_1 , et on cherche la re-paramétrisation optimale de q_2 en résolvant $\hat{\gamma} = \arg \inf_{\gamma \in \Gamma} \|q_1 - (q_2 \circ \gamma)\sqrt{\dot{\gamma}}\|_2^2$. La re-paramétrisation $\hat{\gamma}$ donnera alors la correspondance optimale entre les courbes.

Par la suite, on note $\{X_i, i = 1, \dots, N\}$ l'ensemble de N points représentant la discrétisation de β_1 et $\{U_i = U(X_i), i = 1, \dots, N\}$ les déplacements correspondants donnés par $\beta_2 \circ \hat{\gamma}$. Afin de résoudre le second problème, nous supposons que U est un champ gaussien.

2.3 Champ de déformation gaussien

On rappelle qu'un champ gaussien U est défini par :

$$U(X) \sim \mathcal{N}(\mu(X), C(X))$$

où $\mu(X)$ et $C(X)$ sont respectivement la moyenne et la variance de $U(X)$. Une loi gaussienne étant entièrement décrite par sa moyenne et sa variance, il nous suffit de trouver ces derniers pour définir notre champ. Pour ce faire, nous supposons tout d'abord que U est un champ stationnaire, c'est-à-dire que $\mu(X) = \mu, \forall X \in \Omega$. On a alors:

$$\mathcal{N}(\mu(X), C(X)) = \mu + \mathcal{N}(0, C(X))$$

ce qui implique que μ est une translation sur l'image I_1 . Nous pouvons alors supposer que $\mu = 0$. Nous estimons ensuite $C(X)$ par une méthode paramétrique. Il existe un grand choix de fonctions de covariance candidates dans la littérature [14, 1]. Dans ce travail, nous avons choisi C comme fonction de covariance de Matérn :

$$C(h) = \tau \frac{2^{1-\nu}}{\Gamma(\nu)} (h\alpha)^\nu K_\nu(h\alpha) \quad (1)$$

où h est la corrélation spatiale et K la fonction de Bessel modifiée de seconde espèce. Généralement, $\tau > 0$ est appelé le paramètre de variance (marginale), $\alpha > 0$ le paramètre d'échelle et $\nu > 0$ le paramètre de lissage. Si $\nu = \frac{1}{2} + k, k \in \mathbb{N}$, l'équation 1 se réduit au produit d'une exponentielle et d'un polynôme[5]:

$$C(h) = \tau e^{-h\alpha} \sum_{l=0}^k \frac{(k+l)!}{(2k)!} \binom{k}{l} (2h\alpha)^{k-l}$$

Lorsque $\nu = \frac{1}{2}$, C est la fonction de covariance exponentielle, et elle devient gaussienne pour $\nu = +\infty$. De manière plus générale, pour $\nu = \frac{1}{2} + k$, le champ U sera de classe C^k . Par conséquent, définir notre champ revient à estimer l'hyperparamètre $\theta = (\tau, \alpha, \nu)$ de la fonction de covariance. On note C_θ la fonction de covariance d'hyperparamètre θ , et Σ_θ la matrice de covariance associée aux points $\{X_i, i = 1, \dots, N\}$, donnée par:

$$\Sigma_\theta = \begin{pmatrix} C_\theta(\|X_1 - X_1\|) & \dots & C_\theta(\|X_1 - X_N\|) \\ \vdots & \ddots & \vdots \\ C_\theta(\|X_N - X_1\|) & \dots & C_\theta(\|X_N - X_N\|) \end{pmatrix}$$

2.4 Estimation des paramètres de la fonction de covariance

Estimation par maximum de vraisemblance. Le premier estimateur considéré est obtenu par maximum de vraisemblance. Dans notre modèle, la fonction de vraisemblance est définie par:

$$L(\theta | U_X) = f(U_X | \theta) = \frac{1}{(2\pi)^{N/2} |\Sigma_\theta|^{1/2}} e^{-\frac{U_X^T \Sigma_\theta^{-1} U_X}{2}}$$

où $U_X = (U_1 \dots U_N)$. Pour ce faire, nous devons minimiser la fonction de log-vraisemblance négative, donnée par:

$$\begin{aligned} -\log(L(\theta | U_X)) &= \frac{N}{2} \ln(2\pi) + \frac{N}{2} \ln(\tau) + \frac{\ln |V_{\alpha, \nu}|}{2} \\ &+ \frac{U_X^T V_{\alpha, \nu}^{-1} U_X}{2\tau} \end{aligned}$$

où $\tau V_{\alpha, \nu} = \Sigma_\theta$. Il n'existe cependant pas de solution analytique à cette équation, et devons alors utiliser des méthodes numériques pour déterminer $\hat{\theta}$. Pour ce faire, nous avons choisi de comparer les méthodes de Nelder-Mead [9], la descente du gradient et de Newton. L'optimisation sur ν étant difficile dans notre cas, nous avons estimé ce paramètre par validation croisée.

Estimation par inférence bayésienne. Malgré l'utilisation de méthodes itératives, notre estimation de l'estimateur du maximum de vraisemblance peut converger vers des maximums locaux. Afin d'éviter cela, nous avons choisi d'utiliser d'autres estimateurs basés sur l'inférence bayésienne. On va alors trouver des estimateurs à partir de la loi a posteriori de nos paramètres $f(\theta | U_X)$. Cette dernière est construite, avec une loi a priori $\pi(\theta)$ sur nos paramètres, à l'aide de la règle de Bayes:

$$f(\theta | U_X) = \frac{f(U_X | \theta)\pi(\theta)}{\pi(U_X)} \propto L(\theta | U_X)\pi(\theta)$$

Cependant, la loi de la densité a posteriori de nos paramètres n'étant pas calculable, nous allons alors échantillonner cette distribution. Pour ce faire nous utilisons une méthode de Monte-Carlo par chaîne de Markov (MCMC): l'algorithme de Metropolis-Hastings [6]. Ce dernier construit, à partir d'une loi de proposition $q(\cdot | \theta)$, une chaîne de Markov de la manière suivante.

- Choisir $\theta^1 \sim \pi(\theta)$
- Créer θ^{t+1} à l'aide de θ^t
 1. Générer $\theta^* \sim q(\cdot | \theta^t)$
 2. Calculer la probabilité d'acceptation p :
$$p = \min \left(1, \frac{\pi(\theta^{t+1})L(\theta^{t+1} | U_X)q(\theta^t | \theta^{t+1})}{\pi(\theta^t)L(\theta^t | U_X)q(\theta^{t+1} | \theta^t)} \right)$$
 3. Choisir $\theta^{t+1} = \theta^*$ avec probabilité p , sinon choisir $\theta^{t+1} = \theta^t$

Afin de construire notre chaîne de Markov, nous devons définir les lois de $\pi(\theta)$ et $(q(\theta | \cdot))$. Pour ce faire, nous supposons tout d'abord que ces lois sont séparables, c'est-à-dire que $q(\theta | \cdot) = q(\tau | \cdot)q(\alpha | \cdot)q(\nu | \cdot)$ et $\pi(\theta) = \pi(\tau)\pi(\alpha)\pi(\nu)$. N'ayant que peu d'informations a priori sur nos paramètres, nous choisissons de mettre des lois peu informatives, dont un résumé est présenté en Tableau 1.

Lois a priori $\pi(\cdot)$	Lois de proposition $q(\cdot \tilde{\theta})$
$\alpha \sim \mathcal{U}[0, 1]$	$\alpha \sim \mathcal{U}[\tilde{\alpha} - 0.05, \tilde{\alpha} + 0.05]$
$\tau \sim \mathcal{U}[0, 500]$	$\tau \sim \mathcal{U}[\tilde{\tau} - 50, \tilde{\tau} + 50]$
$\nu \sim \mathcal{U}[\frac{1}{2}, \dots, \frac{11}{2}]$	$\nu \sim \mathcal{U}[\tilde{\nu} - 1, \dots, \tilde{\nu} + 1]$

Tableau 1: Distribution a priori et de proposition sur le paramètre θ .

Une fois la chaîne MCMC construite, nous avons choisi d'estimer $\hat{\theta}$ à partir de la loi a posteriori $f(\alpha | U_X)$, plutôt que de $f(\theta | U_X)$. Pour ce faire, on estime tout d'abord $\hat{\alpha}$, puis on sélectionne l'hyperparamètre $\hat{\theta}$ correspondant. Il existe trois estimateurs bayésiens, que nous utiliserons par la suite: le maximum a posteriori (MAP), la moyenne et la médiane.

2.5 Interpolation sur une nouvelle position

Une fois l'hyperparamètre $\hat{\theta}$ estimé, la fonction de covariance $\hat{C}$ est connue, et ainsi la loi de U également. Il nous reste alors à interpoler le déplacement $U(X^*)$ sur une nouvelle position X^* . Pour ce faire, nous allons calculer l'espérance conditionnelle. Comme U est un processus gaussien, ceci revient à effectuer un krigeage simple[4], défini par:

$$U(X^*) = (U_1 \quad \dots \quad U_N) \hat{\Sigma}^{-1} \begin{pmatrix} \hat{C}(\|X^* - X_1\|) \\ \dots \\ \hat{C}(\|X^* - X_N\|) \end{pmatrix}$$

3 Applications

Pour chaque exemple, on discrétise les courbes en 300 points. Avant de présenter nos résultats, nous présentons les critères de qualité du recalage.

3.1 Critères sur la qualité du recalage

Un bon recalage doit avoir une faible erreur d'interpolation, et doit être lisse. Pour évaluer la qualité d'interpolation de notre méthode, on utilise 200 points pour effectuer le recalage, et les 100 restants pour l'évaluation. La qualité est alors estimée entre les points d'évaluation à l'aide de la racine carrée des erreurs en moyenne quadratique (RMSE). Comme on utilise des courbes pour effectuer notre recalage, on peut aussi utiliser ces dernières pour estimer la qualité d'interpolation. Nous avons choisi d'utiliser la distance de Fréchet (FD) définie par:

$$d_F(F_1, F_2) = \inf_{\gamma_1, \gamma_2 \in \Gamma} \max \{d(F_1 \circ \gamma_1, F_2 \circ \gamma_2)\}$$

où γ_1 et γ_2 sont des reparamétrisations. Concernant la régularité du champ de déformation, nous utilisons la carte de la norme du laplacien. Un maximum de la norme du laplacien (MaxLap) faible signifie alors que le champ de déformation est lisse.

3.2 Application à des données synthétiques

Afin d'évaluer la performance de notre approche, nous nous intéressons à des données synthétiques présentant différents degrés de déformation. Un exemple de résultat obtenu est présenté en Figure 2. En ce basant sur la courbe de la norme du laplacien, on remarque que la déformation est lisse et locale. De plus, d'après le Tableau 2, qui résume l'évaluation de chaque méthode sur cet exemple, le champ de déformation admet de faibles erreurs d'interpolation.

Méthode	RMSE	DF	MaxLap
Nelder-Mead	0.0713	0.5813	0.2902
Gradient	0.0699	0.5541	0.2889
Newton	0.0698	0.5477	0.2878
MAP	0.0718	0.5909	0.2982
Moyenne	0.0721	0.5947	0.3003
Mediane	0.0718	0.5918	0.2990

Tableau 2: Evaluation quantitative des différentes méthodes d'estimation des paramètres sur un exemple de données synthétique présenté en Figure 2.

3.3 Application sur des données réelles

Contrairement aux données synthétiques, plusieurs courbes peuvent être en correspondance pour le même recalage d'images IRM/échographie. La Figure 3 montre un exemple de ce type de données, dont les résultats obtenus sont résumés dans le Tableau 3. Sur cet exemple, les méthodes d'optimisation de Nelder-Mead et de Newton ne donnent pas de bons résultats. Pour la descente du gradient et l'inférence bayésienne, le recalage est lisse et donne de faibles erreurs d'interpolation.

Afin d'étudier la stabilité de notre méthode, nous l'avons appliqué sur 7 données réelles. Les résultats obtenus sont résumés sur la Figure 4. En moyenne, l'erreur d'interpolation est faible pour tous les estimateurs, avec un champ de déformation relativement lisse. Les estimateurs bayésiens sont plus performants car ils fournissent un champ de déformation plus lisse, tout en gardant une petite erreur d'interpolation.

Méthode	RMSE	DF	MaxLap
Nelder Mead	0.1272	0.1796	0.7550
Gradient	0.1353	0.1132	0.4091
Newton	0.1268	0.1755	0.7058
MAP	0.1383	0.1200	0.3667
Moyenne	0.1380	0.1190	0.3680
Mediane	0.1381	0.1201	0.3673

Tableau 3: Évaluation quantitative des différentes méthodes d'estimation des paramètres sur un exemple de données réelle présenté en Figure 3.

4 Conclusion

Nous avons construit un outil de recalage d'images, basé sur les champs gaussiens. Cette méthode est efficace car le champ estimé est lisse et admet de faibles erreurs d'interpolation.

Afin d'améliorer l'estimation des paramètres, et donc la qualité du recalage, plusieurs approches sont envisagées

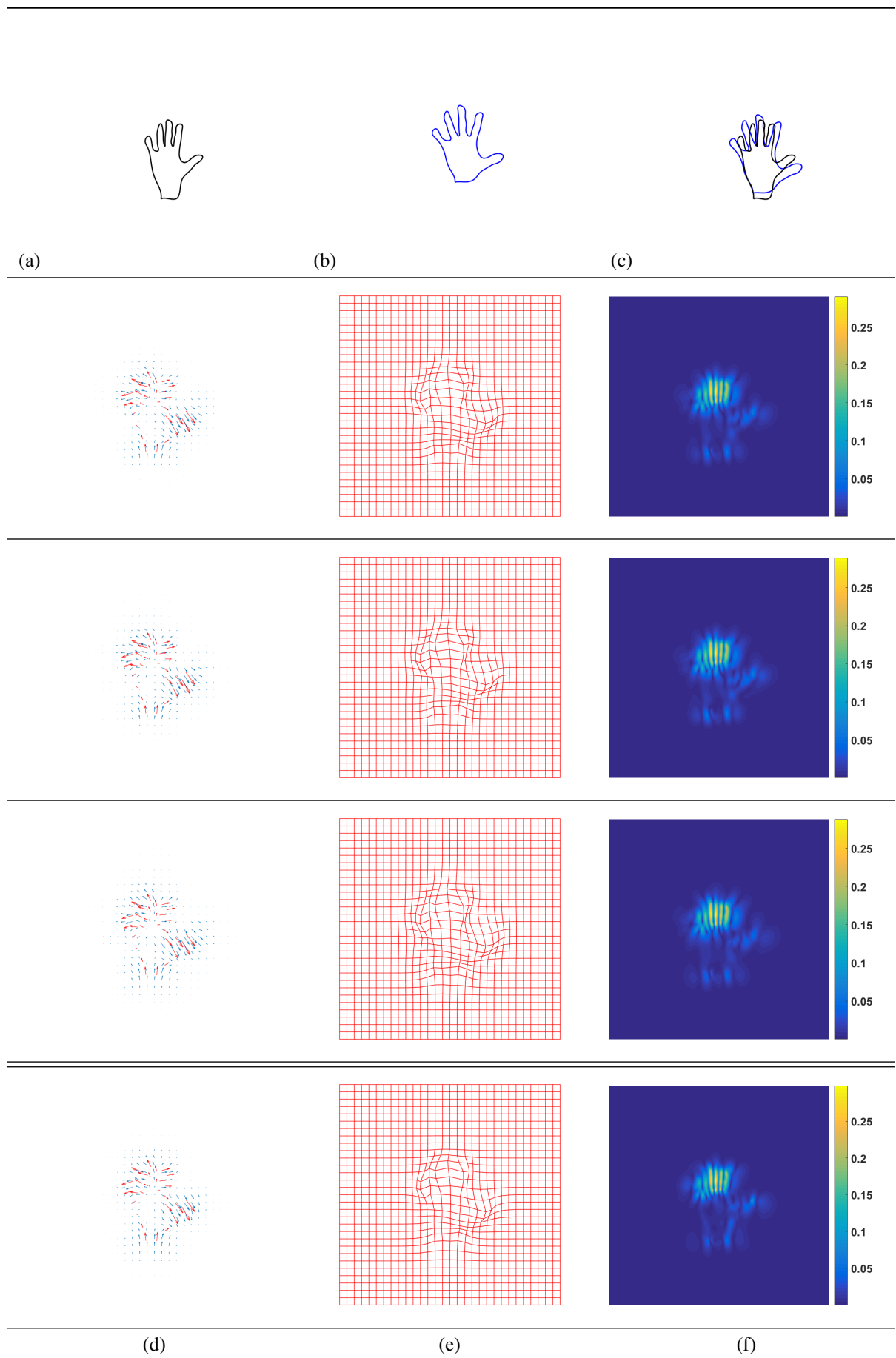


Figure 2: Exemple de champ de déformation estimé sur des données synthétiques: (a) courbes sources, (b) courbes cible, (c) courbes sources (en noir) et cible (en bleu) superposées, (d) le champ de déformation connu (rouge) et estimé (bleu), (e) la grille uniforme déformé par le champ estimé, et (f) la carte de la norme du laplacien. De haut en bas, les résultats sont obtenus après utilisation de la méthode d'optimisation suivante: Nelder-Mead, descente du gradient, Newton, Map.

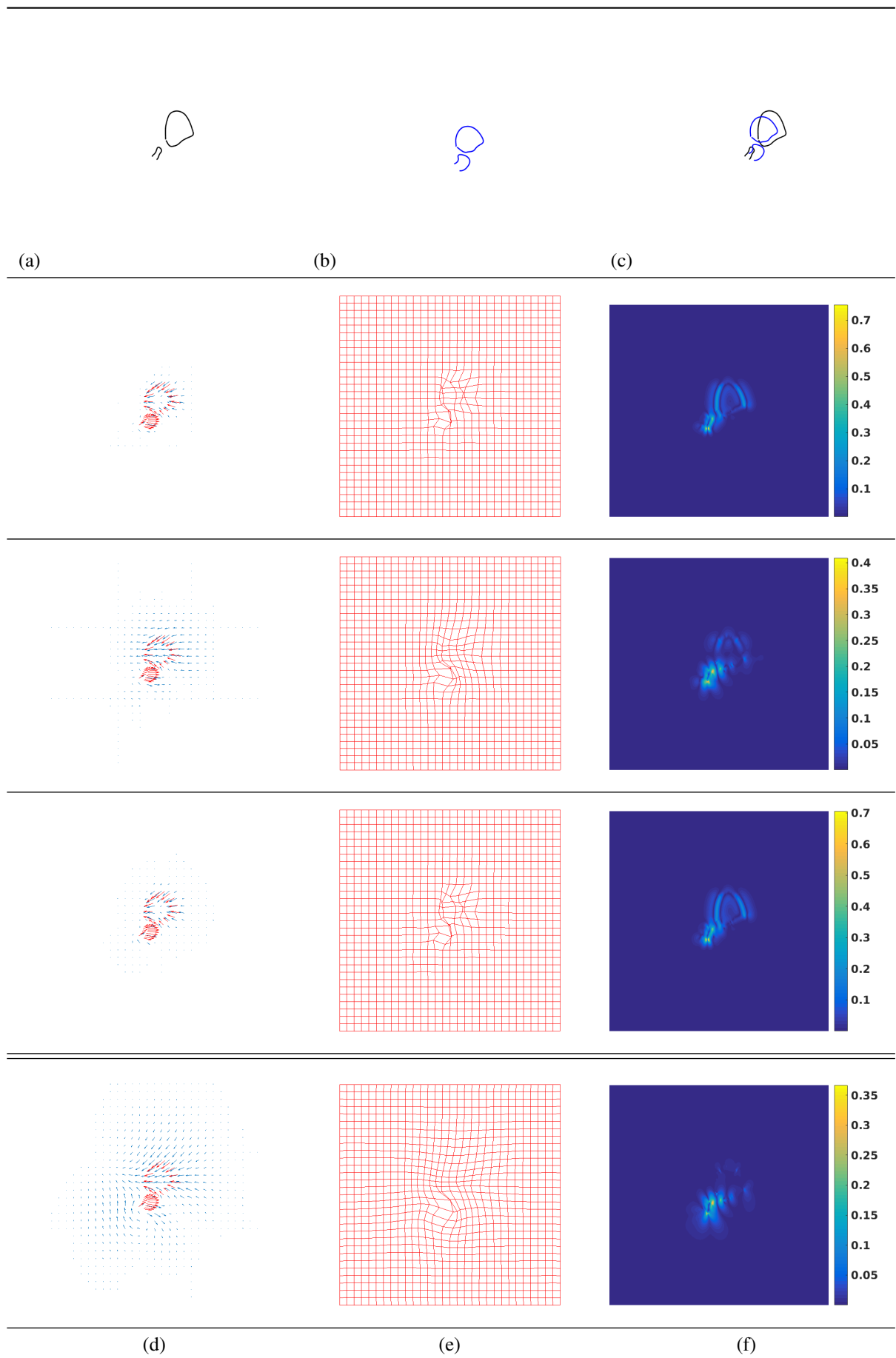


Figure 3: Exemple de champ de déformation estimé sur un recalage IRM/échographie: (a) courbes sources, (b) courbes cible, (c) courbes sources (en noir) et cible (en bleu) superposées, (d) le champ de déformation connu (rouge) et estimé (bleu), (e) la grille uniforme déformé par le champ estimé, et (f) la carte de la norme du laplacien. De haut en bas, les résultats sont obtenus après utilisation de la méthode d'optimisation suivante: Nelder-Mead, descente du gradient, Newton, Map.

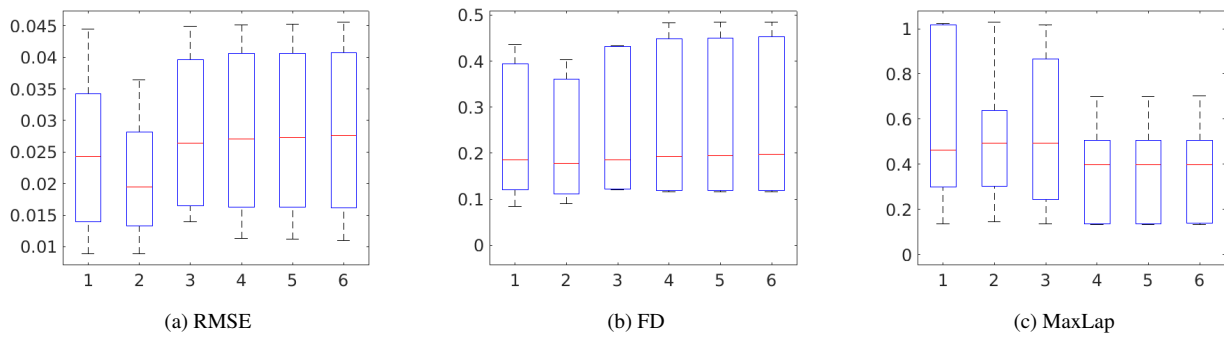


Figure 4: Évaluation quantitative des différentes méthodes d'estimation des paramètres sur 7 données réelles: (1) Nelder-Mead, (2) descente du gradient, (3) Newton, (4) MAP, (5) moyenne, et (6) médiane. Pour chaque méthode, nous présentons un boxplot de (a) la RMSE, (b) la distance de Fréchet, et (c) le maximum de la norme du laplacien.

et feront l'objet de futurs travaux. Tout d'abord, pour les méthodes de la descente du gradient et de Newton, une approximation des dérivées partielles semble nécessaire pour éviter les erreurs numériques. De plus, un pas adaptatif pourrait améliorer la convergence de ces algorithmes. Dans un second temps, on pourrait utiliser des lois a priori plus informatives sur nos paramètres, ainsi que des lois de propositions plus restreintes pour améliorer la convergence des chaînes MCMC.

Références

- [1] R. J. Adler and J. E. Taylor. *Random Fields and Geometry*. Springer Monographs in Mathematics, 2007.
- [2] O. Arandjelović, D.-S. Pham, and S. Venkatesh. Efficient and accurate set-based registration of time-separated aerial images. *Pattern Recognition*, 48(11):3466–3476, 2015.
- [3] L. P. Chamié, R. Blasbalg, R. M. A. Pereira, G. Warmbrand, and P. C. Serafini. Findings of pelvic endometriosis at transvaginal us, mr imaging, and laparoscopy. *Radiographics*, 31(4):E77–E100, 2011.
- [4] N. Cressie. *Statistics for Spatial Data, Revised Edition*. Wiley, 1993.
- [5] T. Gneiting, W. Kleiber, and M. Schlather. Matérn cross-covariance functions for multivariate random fields. *Journal of the American Statistical Association*, 105:1167–1177, 2010.
- [6] W. K. Hastings. Monte carlo sampling methods using markov chains and their applications. *Biometrika*, 57(1):97–109, 1970.
- [7] S. Kurtsek, A. Srivastava, E. Klassen, and Z. Ding. Statistical modeling of curves using shapes and related features. *Journal of the American Statistical Association*, 107(499):1152–1165, 2012.
- [8] W. Mio, A. Srivastava, and S. Joshi. On shape of plane elastic curves. *International Journal of Computer Vision*, 73(3):307–324, 2007.
- [9] J. Nelder and R. Mead. A simplex method for function minimization. *The Computer Journal*, 7, issue. 4:308–313, 1965.
- [10] K. Rohr. *Landmark-based image analysis: Using geometric and intensity models*. Kluwer Academic Publishing, 2001.
- [11] J. E. Roos, D. Weishaupt, S. Wildermuth, J. K. Willmann, B. Marincek, and P. R. Hilfiker. Experience of 4 years with open mr defecography: pictorial review of anorectal anatomy and disease. *Radiographics*, 22(4):817–832, 2002.
- [12] A. Sotiras, C. Davatzikos, and N. Paragios. Deformable medical image registration : A survey. *INRIA Report*, september 2012.
- [13] A. Srivastava, E. Klassen, S. Joshi, and I. Jermyn. Shape analysis of elastic curves in Euclidean spaces. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 33:1415–1428, 2011.
- [14] M. L. Stein. *Interpolation of Spatial Data*. Springer Series in Statistics, 1999.
- [15] Z. Tu, W. Xie, J. Cao, C. Van Gemeren, R. Poppe, and R. C. Velkamp. Variational method for joint optical flow estimation and edge-aware image restoration. *Pattern Recognition*, 65:11–25, 2017.

Mécanisme de négociation multilatérale pour la prise de décision collective

Ndeye Arame DIAGO ¹

Samir AKNINE ¹

Onn SHEHORY ²

Mbaye SENE ³

¹ Université de Lyon ² Bar Ilan University ³ Université de Dakar

aramesdiago@yahoo.fr

Résumé

Dans cet article, nous proposons un modèle de négociation multilatérale dans lequel l'interaction des agents est totalement décentralisée. Le protocole proposé s'inspire du principe "Diviser pour régner". Les agents sont répartis en plusieurs groupes appelés anneaux dans lesquels ils négocient. Nous proposons des politiques d'interaction permettant la communication entre des agents appartenant à des anneaux différents. Nous cherchons à travers cette approche distributive à limiter les champs d'interaction des agents pour réduire la complexité de raisonnement des agents et faciliter la recherche d'accords collectifs. Nous avons évalué les performances de notre protocole en le comparant à une approche centralisée, c'est-à-dire que tous les agents négocient dans un seul groupe.

Mots Clef

Négociation multilatérale, Diviser pour régner, Décision collective

Abstract

In this paper, we propose a multilateral negotiation model in which the agents interact in a fully decentralized way. The proposed protocol draws on "divide and rule" approach. The agents are divided into different groups called "rings" in which negotiations take place simultaneously. We propose interaction policies allowing the communication between agents from different rings. We seek through this approach to limit the agents' scope of interaction and hence to reduce agents reasoning complexity to facilitate agreement research. We have evaluated our protocol by comparing it with a model where all of the agents negotiate in the same group.

Keywords

Multilateral negotiation, Divide and rule, Collective decision.

1 Introduction

La négociation multilatérale est un processus d'échange qui permet aux participants ayant des

intérêts différents pour la réalisation d'un but commun, de construire ensemble des solutions collectivement satisfaisantes. Ces participants ont le plus souvent un comportement égoïste, et ne partagent pas les informations concernant leurs fonctions d'utilité et leurs préférences. Chacun tente de conduire la négociation vers un accord qui satisfait ses objectifs. Dans un tel contexte, la négociation devient très complexe lorsque le nombre de participants est important. Dans cet article, nous nous intéressons aux négociations multilatérales fondées sur des approches heuristiques[4]. Les agents construisent la solution à leur problème à travers leurs interactions, à la différence des modèles basés sur la théorie des jeux dont l'espace des solutions est supposé connu par tous les agents [1].

La plupart des travaux effectués sur la négociation multilatérale ne fournissent pas de protocoles spécifiques ou s'appuient sur des hypothèses difficilement applicables dans un contexte réel. Certains étendent les mécanismes proposés pour le cas d'une négociation bilatérale [1]. Par exemple, [5] propose une généralisation du protocole de concession monotone [12] pour la négociation multilatérale, il ne fournit pas un protocole spécifique. [1] propose une généralisation du mécanisme de négociation bilatérale basé sur les offres alternées pour le cas de la négociation multilatérale. D'autres mécanismes de négociation multilatérale désignent un médiateur qui se charge, par exemple, de superviser et de détecter les conflits entre les parties négociatrices, et de suggérer des offres [10, 3]. Le problème de ces protocoles de négociation avec médiateur réside dans leur structure centralisée. Ils ne garantissent pas toujours une solution satisfaisante. Notre objectif est de proposer des protocoles de négociation applicables dans un contexte réel et qui permettent aux agents de construire des solutions collectives de façon totalement décentralisée.

[2] présente deux extensions du protocole des offres alternées : *Stacked Alternating Offers Protocol* (SAOP) et *Alternating Multiple Offers Protocol* (AMOP). Dans ces protocoles, chaque agent ne participe à la négociation que quand son tour arrive. Dans

le protocole SAOP, celui-ci démarre toujours la négociation en soumettant une offre, et lorsque son tour arrive, il n'évalue que la dernière offre soumise. Par conséquent, toutes les offres ne sont pas évaluées par tous les agents, et certains d'entre eux peuvent passer à côté des offres qui les auraient intéressées. Dans le protocole AMOP, tous les agents soumettent séquentiellement leurs offres. Ensuite, chacun vote pour toutes les offres. Dans ce protocole, chaque agent a une meilleure visibilité sur l'espace des solutions, et il peut évaluer toutes les offres des autres à la différence du protocole SAOP. Cependant, ce protocole est très coûteux en nombre de messages comparé au protocole SAOP. Le processus de raisonnement des agents devient plus complexe car les agents doivent évaluer toutes les offres en même temps, et prennent des décisions pour chacune de ces offres.

Notre modèle s'appuie sur les hypothèses suivantes : (1) les agents négocient sur un seul problème (attribut) et doivent trouver une solution collective et acceptable, c'est-à-dire, acceptée par la majorité des agents ; (2) les agents sont égoïstes et chacun cherche à défendre ses propres intérêts ; (3) la négociation s'effectue sans médiateur, les agents interagissent entre eux de façon totalement décentralisée ; (4) les agents ne partagent pas leur fonction d'utilité ; (5) la négociation s'effectue sur une durée limitée. De telles négociations nécessitent des protocoles spécifiques pour faciliter la prise de décision et la définition de concepts de solution de négociation équitable et collectivement satisfaisante.

Le mécanisme de négociation que nous proposons, dans cet article, est fondé sur une approche incrémentale et itérative permettant de distribuer la négociation mais aussi, d'appréhender la complexité du raisonnement des agents et ainsi, de faciliter la convergence de la négociation tout en limitant le temps. Pour aborder ces problèmes, nous nous intéressons à l'organisation du système multi-agents. La forme organisationnelle du système peut significativement affecter la complexité de la recherche d'accords, la flexibilité, la réactivité et peut induire des coûts de calcul et de communication [9, 13]. Le modèle de négociation proposé s'inspire de l'approche "Diviser pour régner". Il s'agit, d'abord, de diviser l'ensemble des agents en groupes dans lesquels les négociations vont prendre place. Nous cherchons à travers cette répartition à limiter les champs d'interaction des agents, à réduire la complexité de leur raisonnement et à faciliter la recherche d'accords. Ensuite, nous proposons des politiques d'interaction qui s'adaptent aux modèles d'organisation des agents proposés. Nous proposons différentes stratégies de négociation qui permettent aux agents de prendre des décisions qui répondent

au mieux à leurs objectifs.

2 Le modèle de négociation

Pour illustrer le mécanisme de négociation proposé, nous considérons un scénario de négociation basé sur un projet de mise en place d'une nouvelle ligne de tramway pour une agglomération. Ce projet implique par exemple, les maires des communes concernées, le préfet de la région, les conseillers régionaux et les compagnies de transports. Leur but est de déterminer le tracé de la nouvelle ligne de tramway, c'est-à-dire, les coordonnées de ses arrêts. Ces participants ayant des intérêts différents vont négocier pour s'accorder sur une solution. Par exemple, le but des maires est d'optimiser le coût de financement du projet et la desserte des établissements publics en préférant que la ligne desserve plusieurs communes. Le préfet s'intéresse à la sécurité des personnes en préférant que certains arrêts soient proches des postes de police. Les conseillers régionaux s'intéressent à la desserte des établissements publics tels que les lycées. Quant aux compagnies de transport, elles souhaitent que les arrêts soient proches des stations de métro et bus pour faciliter les correspondances.

Le mécanisme de négociation proposé consiste à structurer les agents en petits groupes appelés anneaux dans lesquels les agents s'échangent des propositions. La négociation s'effectue simultanément au sein de chaque anneau. Dans chaque anneau, les agents soumettent leurs propositions. Ils reçoivent celles des autres et prennent des décisions. Un agent peut accepter, refuser, renforcer ou attaquer une proposition soumise par un autre agent. La négociation s'effectue de manière décentralisée sans médiation, autrement dit les agents décident eux-mêmes de leurs propositions et réagissent avec les autres selon leurs propres croyances tout en respectant les règles du protocole. Chaque proposition soumise par un agent est une suggestion de solution de négociation. Elle devient une *solution acceptable* si elle est acceptée par la majorité des agents. Le choix de la règle de la majorité est justifié par le fait que nous considérons les négociations impliquant un grand nombre d'agents dans lesquelles l'application de la règle de l'unanimité rend difficile la recherche d'accords collectifs. Pour rendre flexible le mécanisme de négociation proposé, le protocole permet aux agents qui le souhaitent de prendre connaissance via l'écoute flottante des interactions qui s'effectuent dans les autres anneaux. Cela permet aux agents d'appréhender les négociations menées au-delà de leur anneau. Ainsi, un agent peut réagir sur les propositions écoutées tout en restant dans son anneau. Comme il peut décider de rejoindre un autre anneau dans lequel ses intérêts seront mieux défendus. Cependant, ils utilisent différentes stratégies pour choisir le meilleur

anneau qu'ils vont rejoindre s'ils en écoutent plus d'un. Lorsqu'un agent se déplace dans un nouvel anneau, il peut soit soumettre de nouvelles propositions, soit défendre ses propositions ayant déjà obtenu un certain taux d'acceptation en vue d'atteindre la majorité. Les interactions inter-anneaux et les déplacements des agents d'un anneau à un autre sont régis par les règles du protocole et la politique d'interaction proposée. Une politique d'interaction inter-anneaux spécifie *comment* et *quand* les agents appartenant à des anneaux différents interagissent entre eux.

La négociation s'effectue sur une durée limitée et en plusieurs tours séparés par des points de contrôle. Ces derniers permettent d'évaluer à chaque tour de négociation les échanges effectués et de vérifier si un accord a été trouvé. À l'issue d'un point de contrôle, une solution peut être trouvée, sinon une autre phase de négociation est effectuée si le temps imparti à la négociation n'a pas expiré. Nous considérons qu'il existe une solution de négociation dès lors qu'une proposition est acceptée par la majorité des agents. Les points de contrôle s'effectuent à des périodes régulières prédéfinies. En résumé, notre modèle de négociation se déroule en trois phases :

Phase 1 : c'est la répartition initiale des agents dans les anneaux. Elle consiste à définir le nombre d'anneaux à créer et à affecter les agents dans les anneaux.

Phase 2 : c'est la négociation proprement dite. Les agents formulent et soumettent leurs propositions. Ils reçoivent celles des autres et prennent leurs décisions.

Phase 3 : c'est le point de contrôle. Il consiste à vérifier s'il existe, des propositions dont les taux d'acceptations ont atteint le seuil de la majorité fixé par le concepteur du système.

3 Formalisation

Soient $\mathcal{A} = \{a_1, \dots, a_n\}$ l'ensemble des agents du système et $\mathcal{G} = \{g_1, \dots, g_m\}$ l'ensemble des anneaux dans lesquels les agents sont affectés. Nous désignons par $\{r_1, \dots, r_Q\}$ l'ensemble des tours de négociation, par $\{C_{pt_1}, \dots, C_{pt_Q}\}$ l'ensemble des points de contrôle. Chaque tour r_q de négociation est suivi d'une phase de contrôle C_{pt_q} . Soient ϕ_{maj} le seuil de majorité et d_l le temps imparti à la négociation.

Chaque agent a_i possède un ensemble de propositions $\mathcal{P}_{a_i}$ qu'il peut soumettre et une fonction d'utilité $u_{a_i} \in [0, 1]$ lui permettant d'évaluer les propositions qu'il reçoit.

Fonction d'utilité et Aspirations : Soit $\mathcal{X} = \{x_1, \dots, x_r\}$ l'ensemble des critères pour évaluer une proposition. Chaque agent a_i possède un sous-ensemble de critères $\mathcal{X}_i \subseteq \mathcal{X}$ avec lesquels il définit sa fonction d'utilité. L'objectif de chaque agent est de maximiser son utilité mais à un certain moment

de la négociation, il peut être appelé à faire des concessions pour établir un compromis avec les autres agents. Ainsi, chaque agent fixe ses aspirations représentées ici par l'utilité minimale $U_{min_{a_i}}$ qu'il peut espérer, autrement dit son utilité de réserve en dessous de laquelle aucune proposition n'est acceptable, et l'utilité maximale $U_{max_{a_i}}$ en dessus de laquelle toute proposition est acceptable. Lorsque l'utilité d'une proposition se trouve entre $U_{min_{a_i}}$ et $U_{max_{a_i}}$, l'acceptation de cette proposition dépend des stratégies de concession de l'agent. Nous rappelons que dans le protocole proposé, les agents négocient sur la base d'informations incomplètes. Les fonctions d'utilité sont des informations privées. Chaque agent génère ses propositions, et fixe ses aspirations en se basant sur ses propres connaissances qui sont souvent limitées. Cependant, au cours de ses échanges avec les autres agents, il peut recevoir des propositions qui peuvent être au-delà de ses aspirations.

Actes de communication : Soit $\mathcal{O} = \{Propose, Accepte, Refuse, Renforce, Attaque\}$ l'ensemble des actes du langage que les agents utilisent pour interagir.

Propose($a_i, Ag_k, p_\alpha^i, argp_\alpha^{i+}$) : l'agent $a_i \in Ag_k$ soumet une proposition p_α^i avec un ensemble d'arguments positifs $argp_\alpha^{i+} \subseteq Argp_\alpha^{i+}$ à tout agent $a_j \in Ag_k$.

Accepte(a_i, Ag_k, p_α^r) : l'agent $a_i \in \mathcal{A}$ accepte la proposition p_α^r soumise par $a_r \in Ag_k$ et envoie un message à tout agent $a_j \in Ag_k$.

Refuse(a_i, Ag_k, p_α^r) : l'agent $a_i \in \mathcal{A}$ refuse la proposition p_α^r soumise par $a_r \in Ag_k$ et envoie un message à tout $a_j \in Ag_k$.

Attaque($a_i, Ag_k, p_\alpha^r, argp_\alpha^{r-}$) : l'agent $a_i \in \mathcal{A}$ attaque la proposition p_α^r soumise par $a_r \in Ag_k$ et envoie un message à tout $a_j \in Ag_k$ avec un ensemble d'arguments négatifs $argp_\alpha^{r-} \subseteq Argp_\alpha^{r-}$.

Renforce($a_i, Ag_k, p_\alpha^r, argp_\alpha^{r+}$) : l'agent $a_i \in \mathcal{A}$ renforce la proposition p_α^r soumise par $a_r \in Ag_k$ et envoie un message à tout $a_j \in Ag_k$, avec un ensemble d'arguments positifs $argp_\alpha^{r+} \subseteq Argp_\alpha^{r+}$.

p_α^i signifie que la proposition p_α est soumise par a_i . Chaque proposition p_α soumise par un agent est associée à un tuple $(T_{p_\alpha}^{ac}, T_{p_\alpha}^{re}, T_{p_\alpha}^{rf}, T_{p_\alpha}^{at}, v_{p_\alpha}, ws(p_\alpha), \epsilon_{p_\alpha})$ dont les éléments représentent, respectivement, le taux d'acceptation, le taux de renforcement, le taux de refus, le taux d'attaque, sa valeur de support avec $v_{p_\alpha} = \frac{T_{p_\alpha}^{ac} + T_{p_\alpha}^{re}}{T_{p_\alpha}^{rf} + T_{p_\alpha}^{at} + 1}$, sa valeur de satisfaction sociale

détaillée plus loin et l'estampille. Les taux ci-dessus sont calculés par rapport au nombre d'agents dans le système. Par exemple, $T_{p_\alpha}^{ac}$ est le rapport entre le nombre d'acceptations et le nombre d'agents n . Les autres taux tels que le taux de renforcement, d'attaque, et de refus sont calculés de la même façon. Les agents formulent des arguments positifs ou négatifs pour respectivement renforcer ou attaquer

les propositions. Les renforcements et les attaques sont pris en compte dans l'évaluation de la valeur de support d'une proposition.

4 Mécanisme de négociation

4.1 Répartition des agents en anneaux

La répartition de $\mathcal{A} = \{a_1, \dots, a_n\}$ en plusieurs groupes $\mathcal{G} = \{g_1, \dots, g_m\}$ s'effectue en deux étapes : (1) le choix du nombre de groupes ; (2) l'affectation des agents dans les groupes. Chaque agent doit appartenir à un et un seul groupe, chaque groupe doit avoir au moins deux agents. Ainsi, pour un nombre n d'agents on peut former au plus $m = \lfloor n/2 \rfloor$ anneaux ou groupes. L'affectation des agents dans les anneaux s'effectue telle que : $\bigcap_{k \geq 1}^m \mathcal{A}g_k = \emptyset, \bigcup_{k \geq 1}^m \mathcal{A}g_k = \mathcal{A}, |\mathcal{A}g_k| \geq 2$.

4.2 Négociation

Les agents négocient sur la base des objectifs qu'ils souhaitent atteindre à la fin de la négociation, désignés ici par le terme *aspiration*. L'aspiration d'un agent désigne la valeur d'utilité qu'il estime obtenir de la solution finale de négociation. Dans ce protocole, pour faciliter l'obtention d'une solution de négociation, nous incitons les agents à avoir une certaine flexibilité sur leurs aspirations. Chaque agent a_i définit sa zone d'accord possible, désignée ici par l'intervalle $[U_{min_{a_i}}, U_{max_{a_i}}]$, c'est-à-dire l'intervalle de valeurs d'utilité acceptables. Au début de la négociation, chaque agent a_i cherche à satisfaire au maximum son objectif en soumettant ou en n'acceptant que des propositions dont les utilités sont supérieures ou égales à $U_{max_{a_i}}$. Au cours de la négociation si l'agent n'atteint pas son premier objectif ($U_{max_{a_i}}$) au bout d'un certain temps, il doit réviser à la baisse son utilité maximum $U_{max_{a_i}}$. Ainsi, l'agent diminue de manière incrémentale et stratégique ses aspirations selon l'état d'avancement de ses négociations. La question qui se pose est : *à quel moment, et sous quelles conditions, l'agent va diminuer sa valeur d'utilité maximale $U_{max_{a_i}}$ espérée et de combien ?*.

Tactiques de mise-à-jour des aspirations. Nous considérons ici, que la révision à la baisse des aspirations d'un agent dépend des facteurs tels que le temps et le résultat qu'il obtient à cet instant, c'est-à-dire le nombre d'acceptations obtenues par la proposition qu'il supporte. Une proposition supportée par un agent, peut être celle qu'il a soumise ou celle qu'il a acceptée. Nous définissons des tactiques permettant à chaque agent a_i de calculer à chaque tour r_q , une nouvelle valeur d'utilité $U_{max_{a_i}}^q$ qu'il espère avoir. Cette valeur est telle que $U_{min_{a_i}} \leq U_{max_{a_i}}^q \leq U_{max_{a_i}}$. Cela permet à l'agent de faire des concessions, soit en acceptant, soit en soumettant des propositions qu'il

n'aurait pas acceptées au tour précédent. Un agent a_i définit une nouvelle zone d'accords possibles, lorsqu'il estime que sa proposition courante notée $p_c(i)$ (acceptée ou soumise), n'a aucune chance d'être acceptée par au moins la majorité des agents, lorsque le temps imparti à la négociation sera atteint. La tactique d'un agent a_i dépend de la façon dont il estime l'évolution du taux d'acceptation d'une proposition qui peut avoir la chance d'atteindre au moins le seuil de la majorité ϕ_{maj} à la date limite d_i de la négociation.

Nous désignons par $\alpha^p(t)$ une fonction d'estimation de l'évolution du taux d'acceptation d'une proposition en fonction du temps t tel que : $0 \leq \alpha^p(t) \leq 1, 0 \leq t \leq d_i$. La fonction α^p est continue et croissante. Elle peut être définie de différentes façons, représentant chacune une tactique de négociation [7, 6]. α^p permet de calculer à chaque tour r_q de négociation commençant à la date t_q une valeur $\alpha^p(t_q)$ que l'agent va comparer avec le taux d'acceptation réel $\mathcal{T}_{p_c(i)}^{ac}$ de sa proposition courante à cette même date t_q . Ainsi, l'agent décidera de réviser à la baisse ou non ses aspirations.

La fonction α^p guide l'agent dans son processus de concession. Elle permet à l'agent de décider s'il doit continuer à défendre sa proposition courante ou s'il doit diminuer ses aspirations, et donc d'accepter ou de soumettre de nouvelles propositions. Dans ce mécanisme de négociation, nous considérons un exemple de définition de α^p telle que : $\alpha^p(t) = \frac{\phi_{maj}}{d_i} \times t$. α^p est une fonction linéaire, croissante et continue. Ici nous considérons que les temps des phases de contrôles sont négligeables. Lorsqu'un agent a_i décide de réviser à la baisse ses aspirations à la date t_q à laquelle commence le tour r_q de négociation, il doit savoir de combien il va les réduire. Le calcul de la nouvelle valeur d'utilité $U_{max_{a_i}}^q$ au tour r_q dépend de l'écart entre le taux d'acceptation de sa proposition courante $\mathcal{T}_{p_c(i)}^{ac}$, et du taux $\alpha^p(t_q)$ qu'il estimait avoir pour espérer que sa proposition atteigne au moins à la fin de la négociation le seuil de majorité ϕ_{maj} . L'équation 4.2 ci-dessous montre comment $U_{max_{a_i}}^q$ est évaluée.

- si $\alpha^p(t_q) > \mathcal{T}_{p_c(i)}^{ac}$ alors l'agent met à jour ses aspirations et $U_{max_{a_i}}^q = U_{min_{a_i}} + (U_{max_{a_i}}^{q-1} - U_{min_{a_i}}) \times (1 - (\alpha^p(t_q) - \mathcal{T}_{p_c(i)}^{ac}))$.
- si $\alpha^p(t_q) < \mathcal{T}_{p_c(i)}^{ac}$ alors l'agent maintient ses aspirations et $U_{max_{a_i}}^q = U_{max_{a_i}}^{q-1}$.

$U_{max_{a_i}}^{q-1}$ est la précédente valeur d'utilité que l'agent souhaitait avoir pendant le tour r_{q-1} . Au premier tour de la négociation r_1 , $U_{max_{a_i}}^1 = U_{max_{a_i}}$. Au tour r_q , $U_{max_{a_i}}^q$ est recalculée que si $\alpha^p(t_q) > \mathcal{T}_{p_c(i)}^{ac}$. Si

non l'agent maintient la précédente valeur. Par souci de simplification, nous considérons dans cet article que les agents utilisent la même fonction d'estimation (fonction linéaire) mais chaque agent a_i définit ses seuils d'utilité $[U_{min_{a_i}}, U_{max_{a_i}}]$ qu'il ne partage pas avec les autres agents. Il s'agit d'une information privée. Dans d'autres variantes de ce protocole, chaque agent peut chacun définir sa propre fonction d'estimation en se basant par exemple, sur ses propres croyances.

Dans cet article, nous nous sommes limités à un exemple de définition simple de α^p où nous n'avons considéré que trois paramètres : le seuil de la majorité ϕ_{maj} fixé, la date limite d_l de la négociation et le temps t . Ces paramètres sont des informations publiques et connues par tous les agents. D'autres définitions de α^p peuvent être mises en œuvre dans d'autres contextes prenant en compte d'autres paramètres relevant de croyances propres à chaque agent.

La prise de décision. Un agent prend des décisions en fonction de ses aspirations $[U_{min_{a_i}}, U_{max_{a_i}}]$ qu'il met à jour à chaque tour r_q de la négociation, et aussi en fonction de certaines règles d'interaction détaillées plus loin.

◇ *Accepter une proposition* : un agent a_j accepte une proposition p_i d'un autre agent a_i au tour r_q de la négociation si $u_{a_j}(p_i) > u_{a_j}(p_c(j))$. Ainsi, sa nouvelle proposition courante devient $p_c(j) = p_i$. Nous distinguons deux situations avec des niveaux d'acceptation différents :

- Si $u_{a_j}(p_i) > U_{max_{a_j}}$, il s'agit d'une acceptation accompagnée d'arguments positifs renforçant cette proposition car sa valeur d'utilité est au-delà des attentes de a_j .
- Si $U_{max_{a_i}}^q \leq u_{a_j}(p_i) \leq U_{max_{a_j}}$ il s'agit d'une acceptation sans argument de renforcement car cette proposition est admise par concession.

Règle 1 : *Un agent peut accepter plusieurs propositions au cours de la négociation.*

◇ *Refuser une proposition* : un agent a_j refuse une proposition p_i d'un autre agent a_i au tour r_q de la négociation si $u_{a_j}(p_i) < U_{max_{a_j}}^q$. Il refuse et attaque avec des arguments négatifs toute proposition dont l'utilité est inférieure à $U_{min_{a_j}}$ à n'importe quel tour de négociation.

Règle 2 : *Un agent ne peut pas refuser une proposition qu'il a déjà acceptée.*

◇ *Soumettre une nouvelle proposition* : un agent a_i soumet une nouvelle proposition p_i avec une valeur d'utilité $U_{max_{a_i}}^q \leq u_{a_i}(p_i) \leq U_{max_{a_i}}^{q-1}$ au tour r_q de la

négociation à l'instant δ_t , si le taux d'acceptation $T_{p_c(i)}^{ac}$ de sa proposition courante est plus petit que $\alpha^p(\delta_t)$ et si les taux d'acceptation de toutes les propositions précédentes qu'il a soumises sont inférieurs à $\alpha^p(\delta_t)$.

Négociation au sein d'un groupe. La négociation s'effectue simultanément dans chaque anneau, chaque agent soumet ses propositions aux membres de son anneau. Chaque agent a_i évalue avec sa fonction d'utilité u_{a_i} chaque proposition qu'il reçoit. Il peut ensuite soit l'accepter, soit la refuser, soit la renforcer, soit l'attaquer. Les agents interagissent selon les règles du protocole définies comme suit :

R_{i1} : chaque agent peut soumettre un nombre limité de propositions δ_p et peut faire un nombre limité de refus δ_r au cours de la négociation.

R_{i2} : chaque agent ne peut soumettre ses propositions qu'aux membres de son anneau à tout moment de la négociation.

R_{i3} : chaque agent peut accepter une proposition faite par un autre agent à n'importe quel moment de la négociation.

R_{i4} : une proposition ne peut être acceptée, refusée, attaquée ou renforcée plusieurs fois par le même agent.

Ces règles permettent de gérer les interactions entre les agents durant la négociation en les rendant plus cohérentes. Elles permettent aussi de limiter les taux de messages et de traitements. Le fait de limiter le nombre de refus facilite la recherche d'accords en évitant certaines stratégies des agents. Les agents peuvent aussi faire des concessions en acceptant, par exemple, une proposition qu'ils auraient refusée.

Les politiques d'interaction inter-anneaux. Nous proposons la politique, FIC (*Free Inter-rings Communication*), qui permet la communication libre et réglementée entre les agents appartenant à des anneaux différents. FIC autorise dans un premier temps à chaque agent qui le souhaite d'écouter à tout moment de la négociation les propositions soumises en dehors de son anneau. Il peut réagir sur ces propositions soit en les renforçant, soit en les attaquant tout en restant dans son anneau. Cependant, un agent ne peut soumettre ses propositions que dans son propre anneau. Dans un deuxième temps, l'agent peut se déplacer dans un autre anneau pour soumettre à nouveau ses propositions afin qu'elles atteignent la majorité. Lorsqu'il rejoint un anneau il est autorisé à accepter, refuser, attaquer, ou renforcer toutes propositions soumises dans cet anneau. Les migrations des agents entre les anneaux sont régies par des règles de mobilité limitant le nombre de déplacements et le nombre de fois qu'un agent peut visiter un anneau.

R_{m1} : un agent a_i ne peut visiter plus de δ_v fois le même anneau pendant toute la négociation.

R_{m2} : un agent a_i ne peut effectuer plus de δ_d déplacements pendant toute la négociation.

Chaque agent qui décide de changer d'anneau, doit identifier l'anneau dans lequel ses intérêts seront mieux défendus. La politique d'interaction FIC augmente la flexibilité de communication entre agents dans différents anneaux. Elle permet aux agents de négocier progressivement en se déplaçant entre les anneaux. Nous précisons que l'interaction des agents appartenant à des anneaux différents est limitée par les intérêts des agents. Nous supposons que les surcoûts de communication qu'induit cette politique d'interaction, sont négligeables par rapport à la complexité du raisonnement des agents, et au coût de traitement d'informations dans le cas où ils sont tous réunis dans un seul groupe.

Migrations des agents. Compte tenu des règles de mobilité présentées ci-dessus, un agent a besoin de déterminer le moment où il doit changer d'anneau et quel anneau il doit rejoindre. Cela consiste à identifier d'abord l'anneau dans lequel ses propositions seront mieux défendues, et ensuite de rejoindre cet anneau. Les désaccords, dans un anneau, offrent la possibilité à un agent en dehors de cet anneau de présenter une proposition qui peut satisfaire au mieux les objectifs des agents de cet anneau. Nous considérons ici, qu'un anneau est faible à l'instant δ_t de la négociation lorsque toutes les propositions soumises dans cet anneau ont des taux d'acceptation qui sont inférieurs à un certain seuil $\varphi_{weak} = \alpha^p(\delta_t)$. Chaque agent évalue les anneaux qu'il écoute selon les critères suivants : (1) le nombre d'agents dans l'anneau, car l'agent vise à convaincre autant d'agents que possible ; (2) l'écart entre les taux d'acceptation des propositions soumises dans l'anneau et ceux de l'agent souhaitant se déplacer dans cet anneau. Ces critères permettent à l'agent d'évaluer son degré d'influence dans l'anneau qu'il souhaite rejoindre et les chances que ses propositions soient acceptées. Le choix de l'anneau à rejoindre est non-trivial lorsque l'agent écoute plusieurs anneaux en même temps. Chaque agent associe un vecteur d'utilité à chaque anneau écouté selon les critères ci-dessus. Ainsi, la *dominance de Lorenz* [8] est utilisée pour comparer les vecteurs d'utilité des anneaux écoutés par un agent. Cela permet à chaque agent d'établir un ordre de préférence totale sur l'ensemble des anneaux écoutés.

Les points de contrôle. Soit $C_{pt} = \{C_{pt_1}, \dots, C_{pt_q}\}$ l'ensemble des points de contrôle prévus. Chaque point de contrôle C_{pt_r} est représenté par un tuple (t_r, S_r) , t_r est la date, S_r est l'ensemble des propositions acceptables. $C_{pt} = \{C_{pt_1} = (t_1, S_1), \dots, C_{pt_q} = (t_q, S_q)\}$ avec $(t_1 < \dots, t_q < d_l)$. Lorsqu'un point de contrôle C_{pt_r} s'effectue, chaque agent $a_i \in \mathcal{A}$ communique ses propositions qui ont atteint le seuil de la majorité.

Si $|S_r| = \emptyset$ alors il n'y a aucune solution trouvée. Si $r \neq q$ alors la négociation continue au tour suivant et le résultat sera évalué au prochain point de contrôle $C_{pt_{r+1}}$. Si $r = q$ alors la négociation se termine en ÉCHEC. Si $|S_r| = 1$ alors il existe une seule proposition majoritaire qui sera retenue comme la solution finale de la négociation. La négociation se termine alors en SUCCÈS. Si $|S_r| > 1$ alors il y a plusieurs propositions acceptables comme solution. Le processus de décision social présenté dans section 4.3 est utilisé pour sélectionner la proposition qui sera retenue comme solution finale.

4.3 Concept de solution équitable et juste

Le degré de satisfaction des agents. La satisfaction d'un agent a_i pour une proposition p est désignée ici par une valeur de score qu'il attribue à cette proposition. Cette valeur de score dépend de l'écart entre ses aspirations, c'est-à-dire, les valeurs d'utilité attendues et la valeur d'utilité de la proposition p . Le score exprime le degré de satisfaction d'un agent a_i que nous évaluons à trois niveaux : -1 lorsque p est au-delà de l'objectif fixé, $u_{a_i}(p) \geq U_{max_{a_i}}$; -2 lorsque p est dans la zone de concessions, $u_{a_i}(p) \in [U_{min_{a_i}}, U_{max_{a_i}}]$; -3 lorsque la proposition p sera refusée, c'est-à-dire, $u_{a_i}(p) < U_{min_{a_i}}$ ou $u_{a_i}(p) \in [U_{min_{a_i}}, U_{max_{a_i}}^q]$. Nous désignons alors trois valeurs de score notées $\sigma_2, \sigma_1, \sigma_0$ représentant, respectivement, les valeurs de score des trois niveaux mentionnés ci-dessus. Ces valeurs sont telles que $\sigma_2 > \sigma_1 > \sigma_0 = 0$ et sont des informations publiques, communes à tous les agents.

Notons que chaque agent a_i définit sa propre fonction d'utilité u_{a_i} (information privée) et ses seuils d'utilité $U_{min_{a_i}}, U_{max_{a_i}}$ (informations privées). Les valeurs d'utilité des agents ne sont pas comparables puisque chaque agent définit sa propre fonction d'utilité. Par conséquent, deux agents ayant la même valeur d'utilité pour une proposition peuvent avoir des degrés de satisfaction différents. Pour évaluer le résultat de la négociation, nous définissons une nouvelle méthode d'évaluation du bien-être social différente de celles utilisant directement les valeurs d'utilité des agents. Par exemple, les fonctions de bien-être sociale de Nash ou utilitariste effectuent respectivement le produit et la somme des utilités individuelles des agents.

La satisfaction sociale d'une proposition. Nous désignons par $w_s : \mathcal{P} \mapsto \mathbb{R}_+$ la fonction d'agrégation des valeurs de score que nous proposons. Elle attribue une valeur réelle positive à chaque proposition soumise désignant sa valeur de satisfaction sociale. La définition de w_s s'inspire du principe de Rawls (*maximin*) [11] indiquant que le bien-être social d'une allocation dépend du bien-être de l'individu qui a le niveau de satisfaction le plus bas, c'est-à-dire la personne avec l'utilité minimale. Dans notre

mécanisme, la fonction w_s est définie comme suit : $w_s(p_\alpha) = \frac{\sum \sigma_{a_i}}{n - n_{ap_\alpha}}$. n est le nombre d'agents du système, n_{ap_α} désigne le nombre d'agents qui ont accepté la proposition p_α sans prendre en compte l'agent ayant soumis la proposition. Nous avons donc $n_{ap_\alpha} \leq n - 1$. σ_{a_i} est la valeur de score de chaque agent a_i . w_s garantit le choix de la proposition respectant à la fois la règle de la majorité et celle de l'équité selon Rawls. Ainsi, maximiser w_s consiste à minimiser le nombre d'agents ayant une satisfaction nulle, c'est-à-dire, les agents ayant refusé la proposition.

Proposition 3 *La proposition ayant la plus grande valeur de satisfaction sociale a aussi le taux d'acceptation le plus élevé lorsque σ_1 et σ_2 prennent certaines valeurs (voir preuve).*

Preuve Soit $f, g, f(k) = \frac{(k+1)\sigma_1}{n-k}, g(k) = \frac{(k+1)\sigma_2}{n-k}$ deux fonctions, $k \in [1, n-1]$ and $\mathcal{P}_k$ est l'ensemble des propositions avec un nombre k acceptations.

$\forall p_\alpha \in \mathcal{P}_k, \forall \sigma_1 \in]0, \sigma_2[, f(k) \leq w_s(p_\alpha) \leq g(k)$. $\forall p_\alpha \in \mathcal{P}_k, p_\beta \in \mathcal{P}_{k+1}, \exists \sigma_1 \in]0, \sigma_2[\quad / \quad f(k) \leq w_s(p_\alpha) \leq g(k) < f(k+1) \leq w_s(p_\beta) \leq g(k+1)$. $f(k+1) > g(k) \Leftrightarrow \frac{(k+2)}{n-k-1}\sigma_1 > \frac{(k+1)}{n-k}\sigma_2$.

Soit $z(k) = \frac{(k+1)}{n-k}\sigma_2 \times \frac{n-k-1}{k+2}$ alors $\forall \sigma_2 > 0, \sigma_1 \in]M, \sigma_2[, f(k+1) > g(k)$.

M le maximum de la fonction $z(k)$ avec $k \in [1, n-1]$. En résumé, la proposition 3 est vraie $\forall \sigma_2 > 0$ et $\sigma_1 \in]M, \sigma_2[$.

Nous avons montré que pour toute proposition p_α qui a k acceptations, $f(k) \leq w_s(p_\alpha) \leq g(k)$. S'il existe, à l'instant $t+1$, une proposition p_β qui a $k+1$ acceptations alors quelles que soient les valeurs de satisfaction des agents pour p_α , $w_s(p_\beta) > w_s(p_\alpha)$, nous avons donc $f(k) \leq w_s(p_\alpha) \leq g(k) < f(k+1) \leq w_s(p_\beta) \leq g(k+1)$.

La fonction de satisfaction sociale w_s permet de classer toutes les propositions dont les taux d'acceptation sont supérieurs ou égaux à ϕ_{maj} . La proposition qui a la plus grande valeur de satisfaction sociale devient la solution. Nous avons montré que si cette proposition existe, elle a aussi le plus grand taux d'acceptation. Notons que pour obtenir ce résultat, les valeurs de σ_1 et σ_2 doivent respecter une certaine contrainte telle que : $\sigma_1 \in]M, \sigma_2[$. Les résultats de ce travail nous ont permis aussi d'identifier certaines conditions dans lesquelles un système de vote mixte combinant le vote par approbation et le vote par valeurs peut garantir un choix collectif respectant à la fois la règle de la majorité et la somme des valeurs de score. Il peut arriver que plusieurs propositions obtiennent la même valeur de satisfaction. Ainsi pour garantir l'unicité de la solution de négociation, nous utilisons les valeurs de support des propositions dont les valeurs de satisfaction sont égales pour désigner

la proposition gagnante. S'il existe encore des ex aequo leurs estampilles sont utilisées. Il existe toujours un écart de temps ϵ entre les dates d'obtention de la majorité pour deux propositions.

5 Résultats empiriques

Nous avons implémenté en Java le protocole proposé. Les simulations sont réalisées en faisant varier le nombre d'agents et le nombre d'anneaux. Le scénario de négociation se base sur l'exemple du projet de tramway. Une liste de propositions et un ensemble de critères d'évaluation sont prédéfinis. Chaque agent sélectionne aléatoirement un sous-ensemble de propositions et un sous-ensemble de critères. Nous avons effectué plusieurs expériences, dans le cas où les agents forment un seul anneau et dans le cas où les agents forment plusieurs anneaux. Pour chaque expérience, nous avons mesuré le taux de convergence. Ce taux de convergence est le rapport entre le nombre de fois qu'une solution est trouvée et le nombre d'itérations de négociation effectuées. Nous avons fait varier le nombre d'agents n de 5 à 50 et le nombre d'anneaux m de 1 à 25. Pour un nombre n d'agents, on peut créer au plus $\lfloor n/2 \rfloor$. Chaque anneau doit avoir au moins deux agents. $m = 1$, correspond au modèle de négociation dans lequel tous les agents négocient dans un seul anneau. La figure 1 montre que le nombre d'anneaux

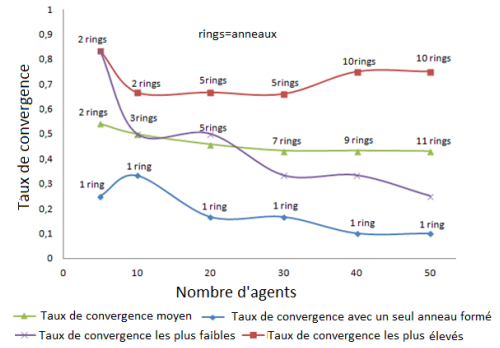


FIGURE 1 – Convergence de la négociation selon le nombre d'anneaux formés

formés par les agents a un impact sur le résultat de la négociation. Par exemple, lorsque le nombre d'agents est supérieur à 20, nous observons que la formation de $\lfloor n/2 \rfloor$ anneaux conduit à des taux de convergence plus faibles mais qui sont toujours meilleurs qu'un seul anneau formé. La courbe coloriée en violet représente le taux de convergence de la négociation lorsque $\lfloor n/2 \rfloor$ anneaux sont formés. Les taux sont plus faibles comparés au cas où le nombre d'anneaux est plus petit que $\lfloor n/2 \rfloor$. Les courbes de la figure 2 montrent l'évolution du nombre d'acceptations des propositions pour $n = 5$ agents, $n = 10$ agents et $n = 30$ agents. Les courbes en pointillés représentent

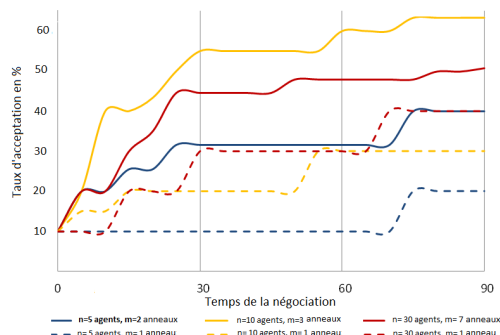


FIGURE 2 – Acceptation moyenne des propositions soumises en fonction du temps.

les résultats obtenus lorsque tous les agents forment un seul groupe et celles en traits pleins représentent les résultats lorsque les agents sont répartis en plusieurs groupes. Ces courbes montrent que les propositions sont acceptées plus facilement dans le protocole proposé. Cela prouve que notre protocole favorise la convergence de la négociation.

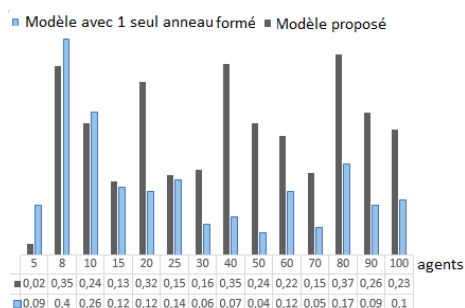


FIGURE 3 – Qualité de la solution de négociation

La figure 3 montre la comparaison entre notre modèle de négociation distribuée et le modèle centralisé (un seul anneau est formé) selon la qualité moyenne de la solution. La qualité de la solution pour chaque agent est l'écart entre l'utilité de la solution et son aspiration maximale. La qualité moyenne est le rapport entre la somme des qualités individuelles et le nombre total d'agents. Nous observons que les solutions obtenues lorsque le nombre d'agents augmente sont meilleures dans notre modèle que dans le modèle centralisé.

6 Conclusion

Cet article a présenté un modèle de négociation multilatérale dans lequel les agents interagissent de façon incrémentale et itérative. Les agents sont répartis en plusieurs groupes (anneaux) dans lesquels ils mènent leurs négociations. L'objectif de ce travail est de faciliter la recherche d'accords. Nous avons évalué les performances de notre mécanisme de négociation en uti-

lisant un scénario de négociation. Les résultats empiriques montrent que ce modèle de négociation fondé sur une approche distributive permet aux agents de négocier plus facilement comparé aux modèles dans lesquels tous les agents sont dans un seul groupe. Dans nos futurs travaux, nous testerons ce mécanisme avec d'autres exemples de négociation pour montrer d'autres propriétés.

Références

- [1] B. An, N. Gatti, and V. Lesser. Alternating-offers bargaining in one-to-many and many-to-many settings. *AMAI*, 77(1-2) :67–103, 2016.
- [2] R. Aydoğan, D. Festen, K. V. Hindriks, and C. M. Jonker. Alternating offers protocols for multilateral negotiation. In *Modern Approaches to ACAN*, pages 153–167. Springer, 2017.
- [3] R. Aydoğan, K. V. Hindriks, and C. M. Jonker. Multilateral mediated negotiation protocols with feedback. In *Novel Insights in ACAN*, pages 43–59. Springer, 2014.
- [4] D. De Jonge and C. Sierra. Nb3 : a multilateral negotiation algorithm for large, non-linear agreement spaces with limited time. *AAMAS*, 29(5) :896–942, 2015.
- [5] U. Endriss. Monotonic concession protocols for multilateral negotiation. In *AAMAS*, pages 392–399. ACM, 2006.
- [6] P. Faratin, C. Sierra, and N. R. Jennings. Negotiation decision functions for autonomous agents. *RAS*, 24(3-4) :159–182, 1998.
- [7] P. Faratin, C. Sierra, and N. R. Jennings. Using similarity criteria to make negotiation trade-offs. In *ICMAS*, pages 119–126. IEEE, 2000.
- [8] B. Golden and P. Perny. Infinite order lorenz dominance for fair multiagent optimization. In *AAMAS*, pages 383–390, 2010.
- [9] B. Horling and V. Lesser. A survey of multi-agent organizational paradigms. *The Knowledge engineering review*, 19(4) :281–316, 2004.
- [10] M. Klein, P. Faratin, H. Sayama, and Y. Bar-Yam. Protocols for negotiating complex contracts. *IEEE Intelligent Systems*, 18(6) :32–38, 2003.
- [11] J. Rawls. *Théorie de la justice*. 1997.
- [12] J. S. Rosenschein and G. Zlotkin. *Rules of encounter : designing conventions for automated negotiation among computers*. MIT press, 1994.
- [13] A. Schuldt, J. O. Berndt, and O. Herzog. The interaction effort in autonomous logistics processes : potential and limitations for cooperation. In *ACCL*, pages 77–90. 2011.

Choisir un encodage CNF de contraintes de cardinalité performant pour SAT

T. Delacroix

IMT Atlantique - Dépt. LUSSI, Brest, France

thomas.delacroix@imt-atlantique.fr

Résumé

Cet article répond à une double problématique : (1) comment choisir un encodage CNF pour des contraintes de cardinalité de type $\#k(x_1, \dots, x_n)$ où $\#$ peut être l'un des symboles $\leq, =, \geq$; (2) déterminer un encodage CNF performant pour les contraintes de cardinalité plus générales $\in \mathcal{K}(x_1, \dots, x_n)$ où $\mathcal{K} \subset \llbracket 0, n \rrbracket$. Pour ce faire, on introduit d'abord un nouvel encodage séquentiel bidirectionnel. On décrit alors un processus pour choisir l'encodage le plus performant pour une contrainte de cardinalité donnée s'appuyant sur une comparaison de différents encodages pour tous les cas possibles de valeurs n et k . Enfin, on montre que l'encodage séquentiel bidirectionnel permet de répondre à la problématique (2).

Mots Clef

CNF, SAT, encodage, contraintes de cardinalité.

Abstract

This article has a double aim : (1) define a process for choosing the most efficient CNF encoding for cardinality constraints of type $\#k(x_1, \dots, x_n)$ where $\#$ is one of the following symbols $\leq, =, \geq$; (2) determine an efficient CNF encoding for the more general cardinality constraints of type $\in \mathcal{K}(x_1, \dots, x_n)$ where $\mathcal{K} \subset \llbracket 0, n \rrbracket$. In order to do this, we introduce a new sequential bidirectional encoding. We then describe a process for choosing the most efficient encoding for a given cardinality constraint based on a comparison of different encodings for all possible values of n and k . Finally, we show that the sequential bidirectional encoding can be used to reach our second aim.

Keywords

CNF, SAT, encoding, cardinality constraints.

1 Introduction

Parmi les solveurs modernes en programmation sous contrainte les plus performants, on trouve aujourd'hui un certain nombre de solveurs CNF-SAT. Les palmarès récents du MiniZinc Challenge en témoignent [13, 11].

Un solveur CNF-SAT permet d'obtenir une valuation pour laquelle une expression logique sous forme conjonctive normale (CNF) donnée est satisfaite lorsqu'il en existe une. Pour une contrainte particulière, la performance du solveur dépend donc à la fois de l'algorithme de résolution du sol-

veur et à la fois de la façon dont la contrainte est exprimée sous forme CNF.

De nombreux travaux se sont donc penchés sur la question de savoir comment exprimer des contraintes classiques des CSP sous forme CNF de manière performante pour les solveurs CNF-SAT [1, 2, 3, 4, 5, 7, 10, 14, 15]. C'est le cas notamment pour les contraintes de cardinalité¹ de type $\#k(x_1, \dots, x_n)$ où $\#$ peut être l'un des symboles $\leq, =, \geq$. En effet, un encodage naïf de ces contraintes contient, dès que n augmente, beaucoup trop de clauses pour pouvoir être utilisé de manière raisonnable en pratique. De nombreux encodages CNF ont donc été proposés, fonctionnant tous sur le même principe général : des variables supplémentaires sont introduites de manière à réduire drastiquement le nombre de clauses.

Parmi les travaux existants, on trouve des comparaisons des différents encodages proposés [2, 9, 12]. Toutefois, ces comparaisons ne sont pas exhaustives. En effet, l'accent y est généralement mis sur le comportement des encodages lorsque n tend vers l'infini. Or, en pratique, on peut également être amené à considérer de très nombreuses contraintes de cardinalité de faible dimension. Une comparaison exhaustive des encodages existants s'impose donc afin d'essayer d'optimiser au maximum l'étape du choix de l'encodage dans la résolution SAT et cela constitue un des éléments central de cet article. On s'aperçoit alors qu'il n'y a pas un encodage plus performant que tous les autres mais de nombreux encodages performants selon les paramètres du problème. On montre également qu'il est possible de combiner des encodages pour obtenir de meilleures performances.

Par ailleurs, on introduit dans cet article un nouvel encodage : l'encodage séquentiel bidirectionnel. Cet encodage, dont la définition est assez naturelle, est particulièrement adapté pour considérer des contraintes de cardinalité plus complexes, notamment les contraintes de cardinalité correspondant à un intervalle. On montre par ailleurs qu'il permet de considérer des contraintes de cardinalité de type $\in \mathcal{K}(x_1, \dots, x_n)$ où $\mathcal{K} \subset \llbracket 0, n \rrbracket$ de manière performante ce qui représente une nouveauté.

1. Le terme contrainte de cardinalité est utilisé ici selon la nomenclature standard dans le contexte SAT et ne doit pas être confondu avec d'autres notions telles que celle de *global cardinality constraint* utilisée dans le contexte de la programmation sous contrainte.

2 Encodage séquentiel bidirectionnel

Dans la suite, on considère des entiers n et k tels que $n \geq 2$ et $k \in \llbracket 1, n-1 \rrbracket$, les autres cas étant évidemment triviaux. La démarche qui mène à définir l'encodage proposé dans cette section s'apparente à la démarche qui mène à l'encodage séquentiel proposé par Carsten Sinz dans [12]. En effet, dans l'article précité, l'auteur définit les sommes partielles $s_i = \sum_{m=1}^i x_m$ et considère le j -ième bit $s_{i,j}$ de la représentation unaire de s_i . Il transpose alors ces bits en variable booléenne dans un encodage CNF pour aboutir à l'encodage ci-dessous pour la contrainte de cardinalité $\leq k(x_1, \dots, x_n)$. Par la suite, on le désigne par le nom d'encodage séquentiel unidirectionnel et on le note $SeqU_{\leq k}^n$.

$$\left. \begin{array}{l} (\neg x_1 \vee s_{1,1}) \\ (\neg s_{1,j}) \quad \forall j \in \llbracket 1, k \rrbracket \\ (\neg x_i \vee s_{i,1}) \\ (\neg s_{i-1,1} \vee s_{i,1}) \\ (\neg x_i \vee \neg s_{i-1,j-1} \vee s_{i,j}) \\ (\neg s_{i-1,j} \vee s_{i,j}) \\ (\neg x_i \vee \neg s_{i-1,k}) \\ (\neg x_n \vee \neg s_{n-1,k}) \end{array} \right\} \forall j \in \llbracket 1, k \rrbracket \quad \forall i \in \llbracket 1, n \rrbracket$$

$(SeqU_{\leq k}^n)$

Cependant, cette transposition contient une réduction qui aboutit à une perte d'information entre la variable $s_{i,j}$ tel qu'elle est encodée dans $SeqU_{\leq k}^n$ par rapport au bit $s_{i,j}$ décrit précédemment. En effet, le bit $s_{i,j}$ est équivalent à $s_i \geq j$. Or l'encodage $SeqU_{\leq k}^n$ donne $(s_i \geq j) \implies s_{i,j}$ mais pas l'implication réciproque. Cette perte d'information est volontaire car elle entraîne un encodage plus restreint de la contrainte $\leq k(x_1, \dots, x_n)$. Toutefois, ce choix n'est pas forcément judicieux lorsque l'on considère une contrainte $= k(x_1, \dots, x_n)$ ou deux contraintes $\geq k_1(x_1, \dots, x_n)$ et $\leq k_2(x_1, \dots, x_n)$ définissant un intervalle.

L'encodage CNF qui suit permet d'encoder exactement l'ensemble des équivalences $(s_i \geq j) \iff s_{i,j}$ pour tout $i \in \llbracket 1, n \rrbracket$ et pour tout $j \in \llbracket 1, k+1 \rrbracket$. On appelle encodage séquentiel bidirectionnel cet encodage et on le note $SeqB_{\#k}^n$.

$$\left. \begin{array}{l} (x_1 \vee \neg s_{1,1}) \\ (\neg x_i \vee s_{i,1}) \quad \forall i \in \llbracket 1, n \rrbracket \\ (\neg s_{j-1,j}) \quad \forall j \in \llbracket 1, k+1 \rrbracket \\ (\neg s_{i-1,j} \vee s_{i,j}) \\ (x_i \vee \neg s_{i-1,j} \vee \neg s_{i,j}) \\ (s_{i-1,j-1} \vee \neg s_{i,j}) \\ (\neg x_i \vee \neg s_{i-1,j-1} \vee s_{i,j}) \end{array} \right\} \forall j \in \llbracket 1, k+1 \rrbracket \quad \left. \right\} \forall i \in \llbracket 1, n \rrbracket$$

$(SeqB_{\#k}^n)$

À partir de cet encodage, il est très facile d'obtenir la contrainte de cardinalité $\leq k(x_1, \dots, x_n)$. En effet, il suffit de rajouter la clause $\neg s_{n,k+1}$. De même, la contrainte $\geq k(x_1, \dots, x_n)$ s'obtient simplement par le rajout de la clause $s_{n,k}$. Enfin, la contrainte $= k(x_1, \dots, x_n)$ s'obtient par le

rajout de ces deux clauses. On note $SeqB_{\leq k}^n$, $SeqB_{\geq k}^n$ et $SeqB_{=k}^n$ les encodages respectifs correspondants.

3 Choisir son encodage

3.1 Comparaisons des encodages

En plus des encodages décrits précédemment, on va également considérer l'encodage naïf $(N_{\leq k}^n)$ défini par :

$$\bigwedge_{i \in \mathcal{C}_n^{k+1}} \bigvee_{j=1}^{k+1} \neg x_{i_j} \quad (N_{\leq k}^n)$$

où $\mathbf{i} = (i_1, \dots, i_{k+1})$ est une combinaison appartenant à l'ensemble $\mathcal{C}_n^{k+1}$ des combinaisons de $k+1$ éléments de $\llbracket 1, n \rrbracket$. On rajoute l'encodage proposé par Bailleux & Boufkhad dans [3] ($BB_{\geq k_1, \leq k_2}^n$). On note que les auteurs précités ne donnent pas d'expression explicite de leur encodage mais décrivent plutôt un algorithme permettant de le construire.

D'autres encodages de la littérature [15] ne sont pas considérés car ils ne satisfont pas la condition de performance relative à la propagation unitaire décrite initialement dans [3] (i.e. il ne permettent pas de vérifier la contrainte sur une valuation partielle des variables x_i).

Par ailleurs, par manque de temps, nous n'avons pas intégré ici d'encodage à base de réseaux [1, 2, 7]. En effet, nous souhaitons d'abord vérifier que la génération de tels encodages est bien linéaire en leur nombre total de clauses (ou de littéraux), de manière à ce que la comparaison soit valable. Ce travail reste donc à compléter sur ce point. En effet, de tels encodages peuvent comporter un nombre total de clauses inférieur à ceux des encodages considérés ici pour un certain nombre de valeurs de n et k . Ceci est notamment le cas de l'encodage de la contrainte $\leq k(x_1, \dots, x_n)$ proposé par Asín et al. dans [2] dont le nombre total de clauses est égal à :

$$-3m + 6mK + \frac{3}{4}mK \log_2(K) + \frac{3}{4}mK \log_2(K) - 3K - \frac{3}{2}K \log_2(K)$$

avec $K = 2^{\lceil \log_2(k) \rceil}$ et $m = \lceil \frac{n}{K} \rceil$.

Le tableau 1 donne le nombres de clauses (avec le détail en fonction de la taille en nombre de littéraux de ces clauses) ainsi que le nombre de variables auxiliaires pour chacun de ces encodages. Les valeurs figurant dans ce tableau ont été recalculées à partir des descriptions de ces encodages dans les articles précités [12, 3] ainsi que le présent article.

Dans la suite de cet article, on utilise les informations de ce tableau pour permettre de choisir l'encodage le mieux adapté aux différents cas étudiés. On fera également usage de la règle suivante pour obtenir un encodage d'une contrainte $\geq k(x_1, \dots, x_n)$ en considérant l'encodage pour la contrainte $\leq k(x_1, \dots, x_n)$ via l'utilisation de la règle suivante :

$$\geq k(x_1, \dots, x_n) \iff \leq (n-k)(\neg x_1, \dots, \neg x_n) \quad (1)$$

De même, on pourra obtenir un encodage de la contrainte $= k(x_1, \dots, x_n)$ en combinant différents encodages via l'utilisation de la constation suivante :

$$= k(x_1, \dots, x_n) \iff (\leq k(x_1, \dots, x_n) \wedge \geq k(x_1, \dots, x_n)) \quad (2)$$

Encodage	Nombre de clauses composées de m littéraux		Nombre de variables auxiliaires
$SeqU_{\leq k}^n$	$m = 1$	$k - 1$	$nk - k$
	$m = 2$	$nk + 2n - 2k - 2$	
	$m = 3$	$nk - n - 2k + 2$	
	Total	$2nk + n - 3k - 1$	
$SeqB_{\#k}^n$	$m = 1$	k	$nk + n$
	$m = 2$	$2nk + 2n - 2k$	
	$m = 3$	$2nk + n - 2k - 1$	
	Total	$4nk + 3n - 3k - 1$	
$N_{\leq k}^n$	$m = k + 1$	$\binom{n}{k+1}$	0
$BB_{\geq k_1, \leq k_2}^n$ ²	$m = 1$	$n - k_2 + k_1$	$n \log_2(n)$
	$m = 3$	$n^2 + 2n \log_2(n) + n - 2$	
	Total	$n^2 + 2n \log_2(n) + 2n - k_2 + k_1 - 2$	

TABLE 1 – Nombre et tailles de clauses pour chaque encodage considéré

3.2 Contrainte $\leq k(x_1, \dots, x_n)$

Dans cette section, on compare différents encodages de $\leq k(x_1, \dots, x_n)$.

Conditions sur n et k par encodage			
$SeqU_{\leq k}^n$	$BB_{\geq 0, \leq k}^n$	$SeqB_{\geq n-k}^n$	$N_{\leq k}^n$
$n \leq 5$ et $k \in$			
$\emptyset$	$\emptyset$	$\emptyset$	$\llbracket 1, n \rrbracket$
$6 \leq n \leq 8$ et $k \in$			
$\llbracket 1, k_1 \rrbracket$	$\emptyset$	$\emptyset$	$\llbracket k_1, n \rrbracket$
où $k_1 = n - 4$.			
$9 \leq n \leq 13$ et $k \in$			
$\llbracket 1, k_1 \rrbracket$	$\emptyset$	$\emptyset$	$\llbracket k_1, n \rrbracket$
où $k_1 = n - 3$.			
$14 \leq n \leq 30$ et $k \in$			
$\llbracket 1, k_2 \rrbracket$	$\emptyset$	$\llbracket k_2, k_1 \rrbracket$	$\llbracket k_1, n \rrbracket$
où $k_1 = n - 3$ et $k_2 = \lceil \frac{2}{3}(n+1) \rceil$.			
$31 \leq n \leq 36$ et $k \in$			
$\llbracket 1, k_2 \rrbracket$	$\emptyset$	$\llbracket k_2, k_1 \rrbracket$	$\llbracket k_1, n \rrbracket$
où $k_1 = n - 2$ et $k_2 = \lceil \frac{2}{3}(n+1) \rceil$.			
$37 \leq n$ et $k \in$			
$\llbracket 1, k_3 \rrbracket$	$\llbracket k_3, k_2 \rrbracket$	$\llbracket k_2, k_1 \rrbracket$	$\llbracket k_1, n \rrbracket$
où $k_1 = n - 2$, $k_2 = \lceil \frac{3n^2 - 2n \log_2(n) - 2n + 2}{4(n-1)} \rceil$ et $k_3 = \lceil \frac{n^2 + 2n \log_2(n) + n - 1}{2(n-1)} \rceil$.			

TABLE 2 – Encodage de $\leq k(x_1, \dots, x_n)$ donnant le nombre total minimal de clauses en fonction de n et k

2. Les valeurs données ici sont exactes si n est une puissance de 2.

On remarque d'abord que pour une telle contrainte, l'encodage $SeqU_{\leq k}^n$ est clairement toujours plus performant que l'encodage $SeqB_{\leq k}^n$ donc on peut exclure ce dernier de notre comparaison. Par contre, en utilisant (1), on peut considérer l'encodage $SeqB_{\geq n-k}^n$ appliqué à $(\neg x_1, \dots, \neg x_n)$. On note $SeqB_{\geq n-k}^n$ cet encodage.

On cherche donc à comparer les encodages $N_{\leq k}^n$, $SeqU_{\leq k}^n$, $SeqB_{\geq n-k}^n$, $BB_{\geq 0, \leq k}^n$. Une analyse complète des différents nombres totaux de clauses pour chacun de ces encodages permet de déterminer l'encodage offrant le plus petit nombre de clauses en fonction de n et k . On a réalisé cette analyse ici et regroupé les résultats dans le tableau 2. Ce tableau décrit, en fin de compte, une partition de l'ensemble des valeurs potentielles de n et k en 4 parties, chacune correspondant aux valeurs de n et k pour lesquelles l'encodage en colonne est optimal (pour le critère du nombre de clauses considéré ici).

Conditions sur n et k par encodage			
$SeqU_{\leq k}^n$	$BB_{\geq 0, \leq k}^n$	$SeqB_{\geq n-k}^n$	$N_{\leq k}^n$
$n \leq 5$ et $k \in$			
$\emptyset$	$\emptyset$	$\emptyset$	$\llbracket 1, n \rrbracket$
$6 \leq n \leq 7$ et $k \in$			
$\llbracket 1, k_1 \rrbracket$	$\emptyset$	$\emptyset$	$\llbracket k_1, n \rrbracket$
où $k_1 = n - 3$.			
$8 \leq n \leq 10$ et $k \in$			
$\llbracket 1, k_1 \rrbracket$	$\emptyset$	$\emptyset$	$\llbracket k_1, n \rrbracket$
où $k_1 = n - 2$.			
$11 \leq n \leq 27$ et $k \in$			
$\llbracket 1, k_2 \rrbracket$	$\emptyset$	$\llbracket k_2, k_1 \rrbracket$	$\llbracket k_1, n \rrbracket$
où $k_1 = n - 2$ et $k_2 = \lceil \frac{10n^2 - 3n - 3}{3(5n-6)} \rceil$.			
$28 \leq n \leq 148$ et $k \in$			
$\llbracket 1, k_2 \rrbracket$	$\emptyset$	$\llbracket k_2, k_1 \rrbracket$	$\llbracket k_1, n \rrbracket$
où $k_1 = n - 1$ et $k_2 = \lceil \frac{10n^2 - 3n - 3}{3(5n-6)} \rceil$.			
$149 \leq n$ et $k \in$			
$\llbracket 1, k_3 \rrbracket$	$\llbracket k_3, k_2 \rrbracket$	$\llbracket k_2, k_1 \rrbracket$	$\llbracket k_1, n \rrbracket$
où $k_1 = n - 1$, $k_2 = \lceil \frac{7n^2 - 6n \log_2(n) - 6n + 4}{10(n-1)} \rceil$ et $k_3 = \lceil \frac{3n^2 + 6n \log_2(n) + 3n - 7}{5n-8} \rceil$.			

TABLE 3 – Encodage de $\leq k(x_1, \dots, x_n)$ donnant le nombre total minimal de littéraux en fonction de n et k

Le tableau 2 permet de choisir un encodage en fonction du nombre de clauses de cet encodage mais il ne prend pas en compte la taille de ces clauses (i.e. le nombre de littéraux par clause). Or la taille des clauses dans une contrainte CNF peut avoir une influence importante sur la rapidité d'un solveur SAT sur cette contrainte. Comme l'évaluation d'une valuation d'une contrainte CNF est linéaire par rapport à son nombre total de littéraux, on pourrait également considérer le nombre total de littéraux de chacun de ces encodages plutôt que leur nombre de clauses comme critère pour choisir un encodage. Même si le nombre total de clauses est généralement utilisé comme critère de comparaison des encodages dans l'état de l'art précité, on penche plutôt pour l'utilisation du nombre total de littéraux. En

tout état de cause, on présentera systématiquement par la suite les valeurs obtenues pour chacun de ces deux critères. Le tableau 3 permet de déterminer l'encodage avec le plus petit nombre total de littéraux.

3.3 Contrainte $\geq k(x_1, \dots, x_n)$

Conditions sur n et k par encodage			
$N_{\leq n-k}^{\neg}$	$SeqB_{\geq k}^n$	$BB_{\geq k, \leq n}^n$	$SeqU_{\leq n-k}^{\neg}$
$n \leq 5$ et $k \in$			
$\llbracket 1, n \llbracket$	$\emptyset$	$\emptyset$	$\emptyset$
$6 \leq n \leq 8$ et $k \in$			
$\llbracket 1, 5 \llbracket$	$\emptyset$	$\emptyset$	$\llbracket 5, n \llbracket$
$9 \leq n \leq 13$ et $k \in$			
$\llbracket 1, 4 \llbracket$	$\emptyset$	$\emptyset$	$\llbracket 4, n \llbracket$
$14 \leq n \leq 30$ et $k \in$			
$\llbracket 1, 4 \llbracket$	$\llbracket 4, k_1 \llbracket$	$\emptyset$	$\llbracket k_1, n \llbracket$
où $k_1 = \lceil \frac{n-1}{3} \rceil$.			
$31 \leq n \leq 36$ et $k \in$			
$\llbracket 1, 3 \llbracket$	$\llbracket 3, k_1 \llbracket$	$\emptyset$	$\llbracket k_1, n \llbracket$
où $k_1 = \lceil \frac{n-1}{3} \rceil$.			
$37 \leq n$ et $k \in$			
$\llbracket 1, 3 \llbracket$	$\llbracket 3, k_2 \llbracket$	$\llbracket k_2, k_1 \llbracket$	$\llbracket k_1, n \llbracket$
où $k_1 = \lceil \frac{n^2 - 2n \log_2(n) - n + 1}{2(n-1)} \rceil$ et $k_2 = \lceil \frac{n^2 + 2n \log_2(n) - 2n - 2}{4(n-1)} \rceil$.			

TABLE 4 – Encodage de $\geq k(x_1, \dots, x_n)$ donnant le nombre total minimal de clauses en fonction de n et k

Conditions sur n et k par encodage			
$N_{\leq n-k}^{\neg}$	$SeqB_{\geq k}^n$	$BB_{\geq k, \leq n}^n$	$SeqU_{\leq n-k}^{\neg}$
$n \leq 5$ et $k \in$			
$\llbracket 1, n \llbracket$	$\emptyset$	$\emptyset$	$\emptyset$
$n = 6$ et $k \in$			
$\{1, 2, 3, 5\}$	$\emptyset$	$\emptyset$	$k = 4$
$n = 7$ et $k \in$			
$\llbracket 1, 4 \llbracket$	$\emptyset$	$\emptyset$	$\llbracket 4, n \llbracket$
$8 \leq n \leq 10$ et $k \in$			
$\llbracket 1, 3 \llbracket$	$\emptyset$	$\emptyset$	$\llbracket 3, n \llbracket$
$11 \leq n \leq 27$ et $k \in$			
$\llbracket 1, 3 \llbracket$	$\llbracket 3, k_1 \llbracket$	$\emptyset$	$\llbracket k_1, n \llbracket$
où $k_1 = \lceil \frac{5n^2 - 15n + 3}{3(5n-6)} \rceil$.			
$28 \leq n \leq 148$ et $k \in$			
$\llbracket 1, 2 \llbracket$	$\llbracket 2, k_1 \llbracket$	$\emptyset$	$\llbracket k_1, n \llbracket$
où $k_1 = \lceil \frac{5n^2 - 15n + 3}{3(5n-6)} \rceil$.			
$149 \leq n$ et $k \in$			
$\llbracket 1, 2 \llbracket$	$\llbracket 2, k_2 \llbracket$	$\llbracket k_2, k_1 \llbracket$	$\llbracket k_1, n \llbracket$
où $k_1 = \lceil \frac{2n^2 - 6n \log_2(n) - 11n + 7}{5n-8} \rceil$ et $k_2 = \lceil \frac{3n^2 + 6n \log_2(n) - 4n - 4}{10(n-1)} \rceil$.			

TABLE 5 – Encodage de $\geq k(x_1, \dots, x_n)$ donnant le nombre total minimal de littéraux en fonction de n et k

Dans cette section, on compare 5 encodages pour la contrainte $\geq k(x_1, \dots, x_n)$. Par l'équivalence (1), on ob-

tient bien la contrainte souhaitée en appliquant les encodages $N_{\leq n-k}^n$, $SeqU_{\leq n-k}^n$ et $SeqB_{\leq n-k}^n$ à $(\neg x_1, \dots, \neg x_n)$. On note respectivement $N_{\leq n-k}^{\neg}$, $SeqU_{\leq n-k}^{\neg}$ et $SeqB_{\leq n-k}^{\neg}$ ces trois encodages. On considère par ailleurs les deux encodages $SeqB_{\geq k}^n$ et $BB_{\geq k, \leq n}^n$.

Comme précédemment, on détermine les encodages donnant le nombre total minimal de clauses (tableau 4) ainsi que les encodages donnant le nombre total minimal de littéraux (5) en fonction de n et k . On remarque que, dans ce cas, c'est l'encodage $SeqB_{\leq n-k}^n$ qui est écarté systématiquement.

3.4 Contrainte $= k(x_1, \dots, x_n)$

Dans cette section, on compare 7 encodages différents de $= k(x_1, \dots, x_n)$ parmi lesquels 2 sont des encodages mixtes entre deux encodages différents. Il s'agit des encodages :

1. $BB_{=k}^n$ (i.e. $BB_{\geq k, \leq k}^n$);
2. $SeqB_{=k}^n$;
3. $SeqB_{=n-k}^{\neg}$;
4. $SeqU_{=k}^n$ (i.e. $SeqU_{\leq k}^n$ et $SeqU_{\leq n-k}^{\neg}$);
5. $N_{=k}^n$ (i.e. $N_{\leq k}^n$ et $N_{\leq n-k}^{\neg}$);
6. $NS_{=k}^n$ (i.e. $N_{\leq k}^n$ et $SeqU_{\leq n-k}^{\neg}$);
7. $SN_{=k}^n$ (i.e. $SeqU_{\leq k}^n$ et $N_{\leq n-k}^{\neg}$).

Comme dans les deux sections précédentes, on a déterminé les encodages donnant le nombre total de clauses minimal en fonction de n et k ainsi que celui donnant le nombre total de littéraux minimal en fonction des mêmes paramètres. Chacun de ces résultats étant difficilement synthétisable en un seul tableau, on renvoie aux annexes pour le détail pour tous les entiers $n \in \llbracket 6, 17 \rrbracket$.

Encodage	Condition sur k	Proportion en $n \rightarrow +\infty$
$SN_{=k}^n$	$1 \leq k < 3$	0%
$SeqB_{=k}^n$	$3 \leq k < k_2$	25%
$BB_{=k}^n$	$k_2 \leq k < k_1$	50%
$SeqB_{=n-k}^{\neg}$	$k_1 \leq k < n-2$	25%
$NS_{=k}^n$	$n-2 \leq k < n$	0%
où $k_1 = \lceil \frac{3n^2 - 2n \log_2(n) - 2n + 3}{4n-3} \rceil$ et $k_2 = \lceil \frac{n^2 + 2n \log_2(n) - n - 3}{4n-3} \rceil$.		

TABLE 6 – Encodage de $= k(x_1, \dots, x_n)$ donnant le nombre total minimal de clauses en fonction de k pour $n \geq 18$

Pour $n \leq 5$, c'est l'encodage naïf $N_{=k}^n$ qui donne les plus petits nombres de clauses et de littéraux quel que soit k . Les deux tableaux 6 et 7 donnent les résultats pour $n \geq 18$. On remarquera que les encodages $SeqU_{=k}^n$ et $N_{=k}^n$ en sont absents. Par ailleurs, on a indiqué, pour chaque encodage, la limite de la proportion à n donné de valeurs différentes de k pour laquelle cet encodage est optimal lorsque n tend vers l'infini.

Encodage	Condition sur k	Proportion en $n \rightarrow +\infty$
$SN_{=k}^n$	$k = 1$	0%
$SeqB_{=k}^n$	$2 \leq k < k_2$	30%
$BB_{=k}^n$	$k_2 \leq k < k_1$	40%
$SeqB_{=n-k}^n$	$k_1 \leq k < n-1$	30%
$NS_{=k}^n$	$k = n-1$	0%
où $k_1 = \left\lceil \frac{7n^2 - 6n \log_2(n) - 6n + 1}{10n - 9} \right\rceil$ et $k_2 = \left\lceil \frac{3n^2 + 6n \log_2(n) - 3n - 1}{10n - 9} \right\rceil$.		

TABLE 7 – Encodage de $= k(x_1, \dots, x_n)$ donnant le nombre total minimal de littéraux en fonction de k pour $n \geq 18$

3.5 Contrainte $\geq k_1, \leq k_2(x_1, \dots, x_n)$

On étudie ici le cas d'une contrainte donnant un intervalle. On compare 7 encodages de cette contrainte, correspondant aux cas généraux des encodages de $= k(x_1, \dots, x_n)$ de la section précédente. Il s'agit des encodages :

1. $BB_{\geq k_1, \leq k_2}^n$;
2. $SeqB_{\geq k_1, \leq k_2}^n$;
3. $SeqB_{\geq n-k_2, \leq n-k_1}^n$;
4. $SeqU_{\geq k_1, \leq k_2}^n$ (i.e. $SeqU_{\leq k_2}^n$ et $SeqU_{\leq n-k_1}^n$) ;
5. $N_{\geq k_1, \leq k_2}^n$ (i.e. $N_{\leq k_2}^n$ et $N_{\leq n-k_1}^n$) ;
6. $SN_{\geq k_1, \leq k_2}^n$ (i.e. $SeqU_{\leq k_2}^n$ et $N_{\leq n-k_1}^n$) ;
7. $NS_{\geq k_1, \leq k_2}^n$ (i.e. $N_{\leq k_2}^n$ et $SeqU_{\leq n-k_1}^n$).

On ne fait pas ici une comparaison exhaustive de toutes les valeurs possibles de k et n comme on a pu le faire pour les contraintes précédentes car cela nous semble trop fastidieux et difficilement exposable de manière concise. On donne simplement les expressions des nombres totaux de clauses et nombres totaux de littéraux pour chacun de ces encodages. Ces expressions peuvent être évaluées au cas par cas pour déterminer une valeur minimale et choisir l'encodage associé lorsque n est faible.

Encodage	Nombre total de clauses
$BB_{\geq k_1, \leq k_2}^n$	$n^2 + 2n \log_2(n) + 2n - k_2 + k_1 - 2$
$SeqB_{\geq k_1, \leq k_2}^n$	$4nk_2 + 3n - 3k_2 + 1$
$SeqB_{\geq n-k_2, \leq n-k_1}^n$	$4n^2 - 4nk_1 + 3k_1 + 1$
$SeqU_{\geq k_1, \leq k_2}^n$	$2n^2 + 2n(k_2 - k_1) - n - 3(k_2 - k_1) - 2$
$N_{\geq k_1, \leq k_2}^n$	$\binom{n}{k_2+1} + \binom{n}{k_1-1}$
$SN_{\geq k_1, \leq k_2}^n$	$\binom{n}{k_1-1} + 2nk_2 + n - 3k_2 - 1$
$NS_{\geq k_1, \leq k_2}^n$	$\binom{n}{k_2+1} + 2n^2 - 2nk_1 - 2n + 3k_1 - 1$

TABLE 8 – Nombre total de clauses pour les encodages de la contrainte $\geq k_1, \leq k_2(x_1, \dots, x_n)$

Sinon, en dehors de certains effets de bord quand $k_1 \leq 2$ ou bien $k_2 \geq n-2$, on peut voir que, pour n suffisamment grand, ce sont les trois premiers encodages qui sont les plus performants. On peut alors utiliser le protocole sui-

Encodage	Nombre total de littéraux
$BB_{\geq k_1, \leq k_2}^n$	$3n^2 + 6n \log_2(n) + 4n - k_2 + k_1 - 6$
$SeqB_{\geq k_1, \leq k_2}^n$	$10nk_2 + 7n - 9k_2 - 1$
$SeqB_{\geq n-k_2, \leq n-k_1}^n$	$10n^2 - 10nk_1 - 2n + 9k_1 - 1$
$SeqU_{\geq k_1, \leq k_2}^n$	$5n^2 + 5n(k_2 - k_1) - 7n - 9(k_2 - k_1) + 2$
$N_{\geq k_1, \leq k_2}^n$	$(n - k_2) \binom{n}{k_2} + k_1 \binom{n}{k_1}$
$SN_{\geq k_1, \leq k_2}^n$	$k_1 \binom{n}{k_1} + 5nk_2 + n - 9k_2 + 1$
$NS_{\geq k_1, \leq k_2}^n$	$(n - k_2) \binom{n}{k_2} + 5n^2 - 5nk_1 - 8n + 9k_1 + 1$

TABLE 9 – Nombre total de littéraux pour les encodages de la contrainte $\geq k_1, \leq k_2(x_1, \dots, x_n)$

vant pour choisir un encodage :

Si $\min(k_2, n - k_1) \geq \frac{3}{10}n$ **choisir** $BB_{\geq k_1, \leq k_2}^n$
Sinon si $k_2 \leq \min(\frac{3}{10}n, n - k_1)$ **choisir** $SeqB_{\geq k_1, \leq k_2}^n$
Sinon **choisir** $SeqB_{\geq n-k_2, \leq n-k_1}^n$ (3)

Le protocole ci-dessus utilise le critère du nombre total de littéraux minimal. Pour utiliser le critère du nombre total de clauses, il suffit de remplacer $\frac{3}{10}$ par $\frac{1}{4}$.

3.6 Contrainte $\in \mathcal{K}(x_1, \dots, x_n)$

Dans cette section, on considère un sous-ensemble $\mathcal{K} = \{k_1, \dots, k_m\} \subset \llbracket 0, n \rrbracket$ avec $m \geq 2$ et $k_1 < k_2 < \dots < k_m$. Le cas général d'une contrainte d'appartenance du cardinal à un sous-ensemble $\mathcal{K}$ se distingue entièrement des cas précédents. En effet, la proposition $k \in \mathcal{K}$ est équivalente à la disjonction $\bigvee_{i=1}^m [k = k_i]$. Ainsi, à part le cas particu-

lier d'un intervalle traité précédemment, on ne peut pas se ramener directement à une conjonction des encodages CNF décrits précédemment. Or obtenir une contrainte CNF à partir d'une disjonction de contraintes CNF nécessite une opération supplémentaire. Effectuer cette opération est prohibitif si l'on considère directement des disjonctions des encodages précédents. Toutefois, on peut également considérer ces disjonctions au sein d'un encodage séquentiel bidirectionnel unique ce qui permet de rendre cette opération tout à fait envisageable.

En effet, la contrainte suivante est bien un encodage de la contrainte $\in \mathcal{K}(x_1, \dots, x_n)$.

$$SeqB_{\#}^n \wedge \left(\bigvee_{i=1}^m (s_{n,k_i} \wedge \neg s_{n,k_i+1}) \right) \quad (4)$$

Cette encodage n'est pas une contrainte CNF mais on peut se ramener à une contrainte CNF équisatisfaisable facilement en introduisant m variables supplémentaires $y_0, \dots, y_{m-1}$ [8]. On définit ainsi l'encodage $SeqB_{\in \mathcal{K}}^n$ de la contrainte $\in \mathcal{K}(x_1, \dots, x_n)$ comme étant la contrainte CNF ci-dessous :

$$\begin{aligned}
SeqB_{\#}^n &\wedge \bigwedge_{i=1}^m \left(s_{n,k_i} \vee \neg y_{i-1} \vee \left(\bigvee_{j=i}^{m-1} y_j \right) \right) \\
&\wedge \bigwedge_{i=1}^m \left(\neg s_{n,k_i+1} \vee \neg y_{i-1} \vee \left(\bigvee_{j=i}^{m-1} y_j \right) \right) \\
&\wedge y_0
\end{aligned}
\quad (SeqB_{\in \mathcal{K}}^n)$$

Cette contrainte a un nombre total de clauses égal à :

$$4nk_m + 3n - 3k_m + 2m \quad (5)$$

et un nombre total de littéraux égal à :

$$10nk_m + 7n - 9k_m + m^2 + 3m - 2 \quad (6)$$

Enfin, on peut remarquer que l'encodage $SeqB_{\in \bar{\mathcal{K}}}^n$ (où $\bar{\mathcal{K}}$ est le complémentaire de $\mathcal{K}$) appliqué à $(\neg x_1, \dots, \neg x_n)$ donne également la contrainte $\in \mathcal{K}(x_1, \dots, x_n)$. On note $SeqB_{\in \bar{\mathcal{K}}}^n$ cet encodage. Il peut éventuellement être plus performant que l'encodage précédent selon $\mathcal{K}$. Son nombre total de clauses et son nombre total de littéraux s'obtiennent en remplaçant k_m par $\max(\bar{\mathcal{K}})$ et m par $n - m$ dans (5) et (6).

Le pire des cas est atteint pour une contrainte de type k pair. On a alors un nombre de clauses de l'ordre de $4n^2$ et un nombre de littéraux de l'ordre de $10.5n^2$ ce qui est tout à fait raisonnable. Par ailleurs, à notre connaissance, il s'agit du seul encodage CNF de la littérature de la contrainte $\in \mathcal{K}(x_1, \dots, x_n)$.

4 Conclusion

Lorsque l'on recherche un algorithme optimisé pour un problème donné, il est rare de découvrir un seul algorithme qui soit le plus optimisé pour chacune des occurrences de ce problème. Pour le problème d'une résolution CNF-SAT d'une contrainte de cardinalité, on peut remarquer que chacun des encodages d'une contrainte de cardinalité considérés dans cet article est préférable aux autres pour au moins quelques valeurs différentes de k et n . Par ailleurs, les comportements des différentes solutions à l'infini ne préjugent en rien quant à leurs comportements en faible dimension. Par exemple dans le cas des contraintes $\leq k(x_1, \dots, x_n)$ et $\geq k(x_1, \dots, x_n)$, l'encodage proposé par Bailleux et Boufkhad est préférable (selon le critère du nombre total de littéraux) dans environ 10% des cas lorsque n est très grand mais n'est préférable en aucun cas lorsque $n < 149$. Afin de pouvoir déterminer l'encodage réellement le mieux adapté à une contrainte et des paramètres donnés, il est nécessaire de passer par une étude exhaustive des cas comme cela a été réalisé dans cet article.

On pourrait objecter que, si n est petit alors la taille de l'encodage le sera également et qu'en conséquence, même si on reste loin de l'optimum possible, on pourra négliger

l'accroissement. Cet argument ne tient toutefois pas longtemps car un gain fixe sur une opération peut être non négligeable, d'autant plus si cette opération est répétée de nombreuses fois au cours d'un même processus.

Par ailleurs, le nouvel encodage séquentiel bidirectionnel qui est proposé dans cet article offre une performance nettement améliorée dans certains cas ainsi que des perspectives nouvelles.

En effet, dans le cas d'une contrainte de type $= k(x_1, \dots, x_n)$, les nombres totaux de clauses et de littéraux sont au mieux quadratiques en n pour les autres encodages et ceci quel que soit k . Toutefois, dans un certain nombre de problèmes, k est fixé indépendamment de n . Les nombres totaux de clauses et de littéraux de l'encodage séquentiel bidirectionnel sont alors linéaires en n . Cet encodage a ainsi pu être mis à profit dans le cas de la résolution SAT d'un problème d'emploi du temps en BTS [6] qui comporte un certain nombre de contraintes de type $= 3(x_1, \dots, x_n)$.

Cet encodage offre également de nouvelles perspectives via son application aux contraintes de type $\in \mathcal{K}(x_1, \dots, x_n)$ pour lesquelles il n'existait pas avant, à notre connaissance, d'encodage performant.

Enfin, on peut noter que sa présentation explicite (semblable à celle de l'encodage séquentiel dans [12] et qui a sûrement contribué à sa popularité) en permet une implémentation directe.

Le travail de recherche présenté dans cet article est d'ordre théorique. Si les critères du nombre total de clauses ou du nombre total de littéraux sont généralement pertinents, ils ne suffisent pas à déterminer l'encodage le mieux adapté qui pourra dépendre non seulement du problème considéré mais également du solveur utilisé. Il appelle donc d'autres travaux qui permettront de mettre en pratique (tels que [6]) et de valider ces résultats. Il sera également prolongé de manière à intégrer au processus de choix des encodages à base de réseaux dont l'encodage de $\leq k(x_1, \dots, x_n)$ défini dans [2].

5 Annexe

$k \backslash n$	6	7	8	9	10	11	12	13	14	15	16	17
1	6	6	6	6	6	6	6	6	6	6	6	6
2	5	6	6	6	6	6	6	6	6	6	6	6
3	5	5	6	6	6	6	6	6	6	6	6	6
4	5	5	5	4	2	2	2	2	2	2	2	2
5	7	7	7	4	1	1	1	2	2	2	2	2
6		7	7	7	3	1	1	1	1	1	1	2
7			7	7	7	3	1	1	1	1	1	1
8				7	7	7	3	3	1	1	1	1
9					7	7	7	3	3	1	1	1
10						7	7	7	3	3	1	1
11							7	7	7	3	3	3
12								7	7	7	3	3
13									7	7	7	3
14										7	7	7
15											7	7
16												7

TABLE 10 – Encodage de la contrainte $= k(x_1, \dots, x_n)$ donnant le nombre total minimal de clauses pour $n \in \llbracket 6, 17 \rrbracket$

Les tableaux 10 et 11 donnent, en fonction de k et n , les encodages de $= k(x_1, \dots, x_n)$ pour $n \in \llbracket 6, 17 \rrbracket$ ayant un

$k \backslash n$	6	7	8	9	10	11	12	13	14	15	16	17
1	6	6	6	6	6	6	6	6	6	6	6	6
2	6	6	6	6	6	6	6	6	6	6	6	6
3	5	6	4	2	2	2	2	2	2	2	2	2
4	7	7	4	4	4	2	2	2	2	2	2	2
5	5	7	4	4	4	4	2	2	2	2	2	2
6		7	7	3	4	4	4	2	2	2	2	2
7			7	7	3	3	3	4	4	4	2	2
8				7	7	3	3	3	3	4	4	4
9					7	7	3	3	3	3	3	4
10						7	7	3	3	3	3	3
11							7	7	3	3	3	3
12								7	7	3	3	3
13									7	7	3	3
14										7	3	3
15											7	3
16												7

TABLE 11 – Encodage de la contrainte $= k(x_1, \dots, x_n)$ donnant le nombre total minimal de littéraux pour $n \in \llbracket 6, 17 \rrbracket$

nombre total minimal de clauses et de littéraux respectivement. Les encodages sont indiqués dans le tableau par un numéro qui correspond à leur numérotation dans la section 3.4.

Références

- [1] R. Asín, R. Nieuwenhuis, A. Oliveras, and E. Rodríguez-Carbonell. Cardinality networks and their applications. In *International Conference on Theory and Applications of Satisfiability Testing*, pages 167–180. Springer, 2009.
- [2] R. Asín, R. Nieuwenhuis, A. Oliveras, and E. Rodríguez-Carbonell. Cardinality networks : a theoretical and empirical study. *Constraints*, 16(2) :195–221, 2011.
- [3] O. Bailleux and Y. Boufkhad. Efficient cnf encoding of boolean cardinality constraints. In *International conference on principles and practice of constraint programming*, pages 108–122. Springer, 2003.
- [4] O. Bailleux, Y. Boufkhad, and O. Roussel. A translation of pseudo-boolean constraints to sat. *Journal on Satisfiability, Boolean Modeling and Computation*, 2 :191–200, 2006.
- [5] O. Bailleux, Y. Boufkhad, and O. Roussel. New encodings of pseudo-boolean constraints into cnf. In *International Conference on Theory and Applications of Satisfiability Testing*, pages 181–194. Springer, 2009.
- [6] T. Delacroix. Planifier l’épreuve e5 à l’aide d’un solveur sat. In *APIA, Conférence Nationale sur les Applications Pratiques de l’Intelligence Artificielle*, 2018.
- [7] N. Eén and N. Sorensson. Translating pseudo-boolean constraints into sat. *Journal on Satisfiability, Boolean Modeling and Computation*, 2 :1–26, 2006.
- [8] ENS. Concours d’admission - composition d’informatique - a -, 2016.
- [9] A. M. Frisch and P. A. Giannaros. Sat encodings of the at-most-k constraint, some old, some new, some fast, some slow. In *Proc. of the Tenth Int. Workshop of Constraint Modelling and Reformulation*, 2010.
- [10] J. Marques-Silva and I. Lynce. Towards robust cnf encodings of cardinality constraints. *Principles and Practice of Constraint Programming–CP 2007*, pages 483–497, 2007.
- [11] MiniZinc. Minizinc challenge 2017 results, 2017.
- [12] C. Sinz. Towards an optimal cnf encoding of boolean cardinality constraints. *CP*, 3709 :827–831, 2005.
- [13] P. J. Stuckey, T. Feydy, A. Schutt, G. Tack, and J. Fischer. The minizinc challenge 2008–2013. *AI Magazine*, 35(2) :55–60, 2014.
- [14] N. Tamura, M. Banbara, and T. Soh. Compiling pseudo-boolean constraints to sat with order encoding. In *Tools with Artificial Intelligence (ICTAI), 2013 IEEE 25th International Conference on*, pages 1020–1027. IEEE, 2013.
- [15] J. P. Warners. A linear-time transformation of linear inequalities into conjunctive normal form. *Information Processing Letters*, 68(2) :63–69, 1998.

Apprentissage par Analogies grâce à des outils de la Théorie des Catégories

L. Cordesses¹ O. Bentahar¹ K. Poulet¹ A. Laurent¹ S. Amar¹ T. Ehrmann¹ J. Page²

¹ Renault Innovation Silicon Valley, 1215 Bordeaux Drive, Sunnyvale, 94089 CA, USA

² CNRS-Université Paris Diderot, Laboratoire SPHERE, 5 rue Thomas Mann, 75205 Paris cedex 13

{lionel.cordesses, kevin.poulet, thomas.ehrmann}@renault.com,
{aude.laurent, sarah.amar, omar.bentahar}@nissan-usa.com,
ju.page@hotmail.fr

Résumé

Les algorithmes d'apprentissage automatique utilisés pour contrôler des systèmes physiques doivent pouvoir apprendre rapidement, avec peu d'exemples, notamment lorsque le coût ou la durée des expérimentations sont trop importants. Un tel objectif peut être atteint en utilisant des concepts provenant des sciences cognitives et des sciences sociales, et formalisés grâce à des outils mathématiques de la théorie des catégories et de la théorie du contrôle.

Ce formalisme aboutit à un système d'apprentissage automatique qui accumule les connaissances, puis les transpose à de nouvelles configurations. Cette approche est illustrée sur un système cyber-physique (un circuit de voitures), et sur les jeux Atari 2600.

Mots Clef

Apprentissage Automatique, Théorie des Catégories, Atari 2600, Circuit de voitures miniatures.

Abstract

Machine learning algorithms for controlling devices need to learn quickly, with few trials, when limited by the duration and cost of experiments. Such a goal can be attained with concepts borrowed from Cognitive Science and Social Science, and formalized mathematically using tools from Category Theory and Control Theory.

This leads to a machine learning system that accumulates knowledge, then transposes it to new configurations. Illustrations of this approach are presented on a cyber-physical system – the slot car game – and also on Atari 2600 games.

Keywords

Machine Learning, Category Theory, Atari 2600, Slot Car.

1 Introduction

Les algorithmes qui contrôlent des systèmes cyber-physiques doivent apprendre comment opérer rapidement dans un environnement partiellement connu, le tout avec de

moins en moins de données. Les solutions à base d'apprentissage par renforcement, avec des réseaux de neurones par exemple, ont fait leurs preuves. Cependant, ces méthodes à la pointe de la recherche ont besoin de beaucoup de données d'entraînement, données que l'on risque de ne pas pouvoir obtenir si l'on respecte des contraintes budgétaires et temporelles.

Nous proposons ici une approche complémentaire à l'apprentissage par renforcement et aux réseaux de neurones pour permettre aux machines d'apprendre rapidement avec des puissances de calcul réduites. L'objectif est d'être aussi performant que les solutions existantes avec environ un pour cent des données et du temps d'entraînement usuels.

L'aspect novateur de notre travail repose sur l'utilisation d'outils élémentaires de la *théorie des catégories* — une branche des mathématiques datant des années 1940, assez peu utilisée en apprentissage automatique (en anglais : Machine Learning) — pour formaliser notamment la notion d'*analogie*. La théorie des catégories a été conçue en lien avec la topologie algébrique, pour permettre de transférer des théorèmes et des concepts d'une branche des mathématiques à une autre [28]. Nous utilisons ces outils pour avoir un cadre générique pour notre approche d'apprentissage automatique, ce qui permet, par exemple, de formaliser rigoureusement le transfert de connaissances.

Par ailleurs, nous pensons qu'il est fondamental de rechercher des structures dynamiques pour construire une intelligence artificielle (IA) adaptative. Les sciences sociales et les sciences cognitives montrent que les structures dynamiques cadrent la cognition (cf. [22]), les langues (cf. [8]) et les interactions sociales (cf. [26]). En effet, ces structures permettent aux humains d'agir et de s'adapter en fonction des expériences passées.

Dans cet article, nous commençons par étudier les méthodes d'apprentissage automatique pour les processus décisionnels markoviens (MDP) appliqués à des jeux. Ensuite, nous expliquons les motivations qui supportent notre approche inspirée des sciences sociales. Grâce au concept d'*équivalence de catégories*, nous décrivons des classes de problèmes non bijectifs, sur lesquels nous centrons notre

approche. Il en résulte un comportement innovant de l'algorithme de contrôle. Enfin, nous illustrons cette approche avec des résultats sur un circuit de voitures miniatures, et sur des jeux vidéo Atari 2600. Notre méthode utilise également le savoir accumulé tout au long des expériences passées, et peut donc servir à valider des concepts modernes tels que le « lifelong machine learning » [5].

2 Travaux antérieurs

Nous avons décidé de tester notre approche sur un circuit de voitures miniatures et sur des jeux vidéo Atari, car ces deux expériences impliquent de prendre des décisions, tâche plus complexe que la classification pure. Dans ce contexte, les approches d'apprentissage automatique appliquées à la prise de décision reposent souvent sur de nombreux essais, ce qui implique de grandes quantités de données d'entraînement et requiert une grande puissance de calcul.

D'autre part, les jeux sont devenus un banc d'essai classique pour les algorithmes d'IA. Les circuits de voitures miniatures, par exemple, sont utilisés pour évaluer la performance des systèmes de prise de décision. Le traitement d'images basé sur de l'apprentissage par renforcement présenté dans [12] estime la position de la voiture sur le circuit grâce à un perceptron multicouche à convolution. L'entraînement prend 12 h et l'apprentissage de la stratégie de contrôle requiert 30 min supplémentaires. La solution plus rapide proposée par [24] pour le même système de circuit de voitures autonome dépend d'accéléromètres et d'un microcontrôleur rajoutés sur la voiture pour d'abord créer une cartographie du circuit, puis ensuite contrôler la vitesse de la voiture.

Les jeux vidéo sont aussi devenus de plus en plus utiles pour fournir une représentation cyber-physique de notre environnement. Même des systèmes anciens, tels que la console Atari 2600, fournissent une grande diversité de situations, allant des labyrinthes (jeux de style Pac-Man) et jeux d'action (tel que Space Invaders) aux jeux de balle et raquette (Breakout, Pong). Malgré le fait que ces différents problèmes réclament des stratégies variées pour être résolus par un joueur humain standard, ils impliquent tous des prises de décision et ont donc été modélisés par des MDPs [18]. Ce cadre permet l'implémentation de nombreuses méthodes différentes. Certains travaux, tel que celui présenté dans [18], utilisent une image de l'aire de jeu remise à l'échelle en entrée d'un « deep Q-network » (DQN), avec pour but de sélectionner la meilleure action à jouer. Une autre possibilité est l'utilisation de méthodes classiques de recherche et de planification, telles que le « Iterated Width algorithm » [16] ou les algorithmes de parcours d'arbre tels que l'algorithme de recherche arborescente Monte-Carlo [21] pour calculer la meilleure action possible. L'apprentissage par renforcement peu profond [15] repose sur une représentation linéaire simplifiée, mais obtient cependant des résultats similaires aux méthodes non-linéaires. Enfin, l'apprentissage maître-élève [3], l'apprentissage par

renforcement inverse [13], et l'apprentissage par imitation [20] peuvent également être utilisés pour entraîner un agent au comportement plus humain, qui a une efficacité et une réussite proches de celles des méthodes susmentionnées. Les « Schema Networks » [11] font un pas supplémentaire en direction du transfert de connaissances et de la réduction des besoins en données et en puissance de calcul pour l'entraînement, en utilisant un simulateur physique et en travaillant à un niveau plus élevé avec des entités plutôt que des pixels. Cette méthode obtient le même score qu'un joueur humain sur Breakout (30 points) avec 100 000 images d'entraînement.

3 Sciences Sociales, Sciences Cognitives et Théorie du contrôle pour l'apprentissage automatique

Pour atteindre notre objectif de réduction des besoins en données d'entraînement et en puissance de calcul, nous avons basé notre IA sur des concepts de haut niveau déjà étudiés par différentes sciences.

Tout ce qui est produit par l'humain contient, en partie, une certaine structure humaine. Par exemple, lorsque des humains conçoivent un jeu, ils utilisent des règles implicites qui leur permettent de jouer et de s'amuser. Cela signifie que ces jeux, comme les autres créations humaines, contiennent un minimum de structure pour que le joueur humain puisse les comprendre, et ils contiennent également de la dissonance cognitive pour les rendre attractifs [6][25][27]. La plupart des problèmes réels possèdent cette ambivalence.

Parmi les structures clés, on retrouve les schémas cognitifs universels auxquels les humains s'attendent quant aux objets et à l'environnement : état intérieur (intention et affect), propriétés matérielles (gravité, conservation de la forme, continuité de la trajectoire) ou propriétés naturelles non-humaines (mouvement et croissance). Ces schémas cognitifs sont acquis durant l'enfance quand l'enfant découvre le monde au fil de ses mouvements et en utilisant ses sens [22].

Nous supposons qu'une IA ayant la connaissance de ces schémas cognitifs serait capable de prendre les bonnes décisions pour jouer, et de transférer les connaissances acquises à une grande variété de problèmes similaires.

3.1 Entités dans le temps et dans l'espace

Nous avons conçu notre IA pour qu'elle raisonne au niveau syntaxique et non au niveau de l'échantillon ou du pixel. D'une certaine manière, cela correspond à raisonner au niveau du morphème en linguistique, défini dans [8] comme le plus petit élément significatif d'un langage. Nous appelons entités les éléments de base de ce niveau syntaxique. Ces entités sont détectées par des outils classiques de traitement du signal et d'apprentissage automatique. Elles sont comme les objets de notre vie de tous les jours : des tronçons de circuit (lignes droites et virages), des voitures, des

balles, des raquettes, des murs. Elles sont organisées géométriquement dans un espace et peuvent être décrites, par exemple, par les coordonnées cartésiennes de leur boîte englobante.

Les données caractérisant ces entités sont collectées à chaque échantillon de temps pour construire des modèles dynamiques décrits par les équations d'état (1) à temps discret pour un système d'ordre N_d , où x est le vecteur d'état, y la mesure, u la commande, A la matrice d'état, B la matrice de commande, C la matrice d'observation, D la matrice d'action directe, et k l'indice de l'échantillon.

$$\begin{aligned} x(k+1) &= Ax(k) + Bu(k) \\ y(k) &= Cx(k) + Du(k) \end{aligned} \quad (1)$$

3.2 Le Moi au sein du monde

Une étape clé pour se forger une identité propre lors du développement de l'enfant est appelée le *stade du miroir*. Nous avons conçu notre IA pour qu'elle commence à ce stade en identifiant le « Moi » parmi toutes les entités.

Cette identification du « Moi » est réalisée en utilisant des résultats de la théorie du contrôle, à savoir le critère de contrôlabilité d'un système dynamique, comme présenté, par exemple, dans [10]. Le système est complètement contrôlable si et seulement si (ssi) la matrice de contrôlabilité $C_{N_d} \triangleq [B|AB|\dots|A^{N_d-1}B]$ est de rang N_d . Nous définissons le « Moi » comme étant l'entité qui vérifie le critère de contrôlabilité.

Comme le « Moi » et son environnement peuvent tous les deux changer, les modèles établis peuvent ne plus être valides. En s'inspirant d'une idée provenant de l'approche scientifique décrite dans [23], selon laquelle il doit être possible d'invalidier un système scientifique empirique par l'expérience, nous avons conçu notre algorithme de telle sorte à ce qu'il puisse réfuter ses propres modèles pour les remplacer par de meilleurs. Le critère pour cette décision est basé sur l'erreur entre les prédictions (basées sur le modèle) et les mesures. Cette remise en cause ne s'effectue pas seulement pendant la phase d'apprentissage mais également pendant la phase de test.

Une fois que le « Moi » est identifié, notre IA classe les entités restantes entre amis et ennemis, d'une manière similaire à l'apprentissage par renforcement. Elle suit ensuite une simple stratégie de survie : essayer d'aller vers ses amis, à part si un ennemi est proche du « Moi », auquel cas la première chose à faire est de s'enfuir.

En résumé, notre IA prend en compte les récompenses de son environnement pour classer les entités, comme en apprentissage par renforcement. Comme dans des approches scientifiques empiriques, elle met à jour les classifications et les modèles quand les mesures ne sont pas compatibles avec les prédictions. Enfin, elle décrit au moins une entité comme étant le « Moi », avec toutes les conséquences que cela implique pour sa survie.

4 La Théorie des Catégories comme cadre pour des analogies non bijectives

L'un des outils les plus efficaces dont dispose l'humain pour évoluer dans une situation inconnue est sa capacité à faire des *analogies* entre cette nouvelle situation et ses expériences passées. Nous pressentons qu'une IA capable d'établir des analogies, et sachant, par exemple, comment jouer au jeu Breakout sera capable de transposer ses capacités au jeu Pong, même si les aires de jeu et les règles ne sont pas identiques, à l'instar d'un joueur de tennis qui saurait au moins partiellement comment manier une raquette de badminton même s'il n'a jamais pratiqué ce sport.

Les situations où deux problèmes ont exactement le même nombre d'états et des structures isomorphes sont rares. La notion d'analogie semble alors adaptée au rapprochement de structures conceptuellement proches mais mathématiquement différentes. Or, la plupart des tentatives de formalisation de cette notion sont fondées sur une perspective structuraliste qui s'articule autour de la *théorie des ensembles via* le concept d'isomorphisme qui préserve les structures des ensembles. Il existe des outils mathématiques pour identifier des structures non isomorphes tels que l'équivalence de catégories en théorie des catégories [17]. En effet, celle-ci apparaît comme un cadre de travail plus naturel pour aborder la notion structuraliste d'analogie qui considère que les relations entre éléments sont plus essentielles que les éléments eux-même pour définir des analogies [9].¹

Dans les problèmes qui reposent sur la théorie des ensembles, l'identification se réduit aux relations d'identité et aux cas bijectifs, tandis que la théorie des catégories apporte des descriptions plus riches des objets, ce qui permet de nouvelles sortes d'identifications. Dans le cadre de cet article, nous nous limitons à des concepts très élémentaires de la théorie des catégories, qui pourraient être décrits en termes ensemblistes, tout comme les situations que l'on expose². L'apport de ces travaux est le changement de perspective conceptuelle que nous croyons prometteur pour le domaine de l'apprentissage automatique³.

Théorie des catégories. Une catégorie $\mathcal{C}$ est une collection d'objets et de morphismes (ou flèches) entre certains de ces objets, munie d'une composition de morphismes, et peut ainsi être assimilée à un graphe orienté. Si A et B sont

1. Une autre approche pourrait être l'utilisation de la théorie des modèles, plus spécifiquement du concept d'équivalence élémentaire [4] dans la perspective de doter l'IA de raisonnements logiques. Pour un panorama philosophique des concepts d'analogie et de raisonnement analogique, voir [1].

2. Ce qui n'est pas étonnant puisque la théorie des ensembles et la théorie des catégories, en tant que propositions mathématiques fondationnelles de toutes les mathématiques, se fondent mutuellement l'une l'autre. Chaque concept catégorique doit pouvoir être décrit en termes ensemblistes et réciproquement.

3. Dans une perspective bayésienne, [7] utilise des catégories de probabilités conditionnelles : les objets sont des espaces mesurables et les flèches des noyaux de Markov.

des objets de $\mathcal{C}$, la flèche $a : A \rightarrow B$ est un isomorphisme si elle est inversible, i.e s'il existe une flèche $b : B \rightarrow A$, telle que $ba = Id_A$ et $ab = Id_B$. Dans ce cas, les objets A et B sont isomorphes. Les isomorphismes définissent une relation d'équivalence sur la classe des objets de $\mathcal{C}$, pour laquelle on note $\mathcal{C}/\simeq$ son quotient. Aussi, si $F : \mathcal{C} \rightarrow \mathcal{C}'$ est ce qu'on appelle une équivalence de catégories, celle-ci induit une bijection $F' : (\mathcal{C}/\simeq) \rightarrow (\mathcal{C}'/\simeq)$ entre les classes d'objets isomorphes même si F n'est pas bijective. Dès lors, on n'identifie plus les objets (ou états) un à un entre deux situations, mais les types (ou classes d'isomorphismes) de ces états.

Ce procédé peut être utilisé dans le cas de problèmes observables. En effet, considérons deux ensembles d'états non vides $\mathcal{C}$ et $\mathcal{C}'$ sans autre hypothèse sur leur cardinalité. Supposons aussi l'existence de deux fonctions d'observations $f : \mathcal{C} \rightarrow O$ et $f' : \mathcal{C}' \rightarrow O'$. Quitte à restreindre O et O' , supposons que ces deux fonctions sont surjectives. Pour tout $o \in O$, on dit que les états $x \in \mathcal{C}$ dont l'observation associée est o (i.e $f(x) = o$) sont de type T_o . Cela définit une relation d'équivalence R_f sur l'ensemble $\mathcal{C} : \forall x, y \in \mathcal{C}, x R_f y$ ssi $f(x) = f(y)$. Transposé dans le champ lexical des catégories, cela revient à placer une flèche inversible entre deux objets x et y de $\mathcal{C}$ ssi $x R_f y$. $\mathcal{C}$ devient alors une catégorie, où toutes les flèches sont inversibles et telle que $\mathcal{C}/\simeq$ est exactement le quotient $\mathcal{C}/R_f$, quotient qui définit aussi l'ensemble des types d'états de $\mathcal{C}$. Puisque ces types ont été définis via les observations, la surjection $f : \mathcal{C} \rightarrow O$ induit une bijection $\tilde{f} : (\mathcal{C}/\simeq) \rightarrow O$ entre l'ensemble de types et l'ensemble d'observations. $\tilde{f}$ est donc l'inverse de la fonction $T : O \rightarrow (\mathcal{C}/\simeq), o \mapsto T_o$ qui définit les types T_o . Le même raisonnement peut être reproduit à partir de $\mathcal{C}'$ et de f' .

Enfin, supposons qu'il existe une bijection $G : O \rightarrow O'$ entre les ensembles d'observation et que O et O' sont donc de même cardinalité. Ainsi, on peut définir une autre bijection $F' = \tilde{f}'^{-1} \circ G \circ \tilde{f} : (\mathcal{C}/\simeq) \rightarrow (\mathcal{C}'/\simeq)$ entre les ensemble de types d'objets, bijection en réalité induite par l'équivalence de catégories $F : \mathcal{C} \rightarrow \mathcal{C}'$ définie comme suit : pour tout $x \in \mathcal{C}$, soit $o = f(x)$ et soit $x' \in f'^{-1}(G(o))$, définissons F telle que $F(x) = x'$. Si $\mathcal{C}$ et $\mathcal{C}'$ ont des cardinaux différents, F ne peut pas être bijective, mais F' l'est. Par construction, F envoie chaque état x vers un état x' du même type (modulo G). Cela permet de relier les catégories $\mathcal{C}$ et $\mathcal{C}'$, de telle sorte que si l'on dispose d'une stratégie applicable dans $\mathcal{C}'$, elle peut être transposée dans $\mathcal{C}$ par la fonction F .

Si l'on remplace $\mathcal{C}$ et $\mathcal{C}'$ par des ensembles d'entités (option 1) ou de paires d'entités (option 2) dans deux (situations de) jeux différent(e)s, nous pouvons les comparer grâce à des fonctions d'observation du type f et f' comme décrit ci-dessus. Cela permet de transférer des connaissances relatives à des entités ou des paires d'entités. Nous avons ainsi expérimenté les cas particuliers où dans l'option 1, $f : \mathcal{C} \rightarrow \{0, 1\}$ et $f' : \mathcal{C}' \rightarrow \{0, 1\}$ valent 1 si l'entité est un « Moi » et 0 sinon ; et dans l'option 2, $f : \mathcal{C} \rightarrow \{0, 1\}$ et

$f' : \mathcal{C}' \rightarrow \{0, 1\}$ valent 1 si les trajectoires des deux entités d'une paire s'intersectent, et 0 sinon. Ce qui est observé du Moi d'une situation peut alors être supposé du moi d'une autre situation par transfert de connaissances. Des observations ultérieures permettront alors de confirmer ou infirmer ces suppositions.

Une explication plus détaillée et approfondie de la théorie des catégories peut être trouvée dans [17]. Son utilisation nous permet de formaliser une grande diversité de jeux et de situations.

Analogies entre circuits de voitures miniatures. Nous introduisons les notations suivantes : soient N et N' les nombres de segments par configuration du circuit, et $\mathcal{C} = \{s, s \in [1, N]\}$, $\mathcal{C}' = \{s', s' \in [1, N']\}$ les ensembles de positions possibles de la voiture sur le circuit, qui ne sont utiles que pour le raisonnement théorique et ne sont pas utilisés lors des expériences. Soient $(u, i)_s$ (resp. $(u', i')_{s'}$) la tension et le courant mesurés lorsque la voiture passe sur la section s (resp. s'). Soit $1 \leq s_0 \leq N$ (resp. $1 \leq s'_0 \leq N'$) la position initiale de la voiture dans la configuration $\mathcal{C}$ (resp. $\mathcal{C}'$). Nous notons k une ligne droite de $\mathcal{C}$ et l un virage. De même, soient k' une ligne droite et l' un virage de $\mathcal{C}'$.

Le joueur peut agir sur $(u', i')'_s$ en utilisant la manette de jeu, ce qui correspond à la stratégie π' définie par (2).

$$\pi'(s') = \begin{cases} (u', i')_{k'}, & \text{si } s' \text{ est une ligne droite} \\ (u', i')_{l'}, & \text{sinon} \end{cases} \quad (2)$$

Nous voulons identifier $\mathcal{C}$ et $\mathcal{C}'$, pour pouvoir transposer la stratégie π' de $\mathcal{C}'$ à $\mathcal{C}$. Les états de $\mathcal{C}$ sont les positions s de la voiture sur le circuit. De même, les états de $\mathcal{C}'$ sont les positions s' . Si $N = N'$ et $s_0 = s'_0$, nous pouvons définir une bijection entre $\mathcal{C}$ et $\mathcal{C}'$ et transposer π' de manière triviale. En revanche, si $N \neq N'$ ou $s_0 \neq s'_0$, il est impossible de définir une telle bijection.

Cependant, si nous transformons $\mathcal{C}$ et $\mathcal{C}'$ en catégories, en définissant des morphismes, nous serons capables de définir une équivalence de catégories $F : \mathcal{C} \rightarrow \mathcal{C}'$. Dès lors, nous définirons la stratégie π appliquée sur $\mathcal{C}$ par $\pi = \pi' \circ F$. Pour définir ces morphismes, nous utilisons la fonction observable f définie sur les états s de $\mathcal{C}$ comme suit : $f : \mathcal{C} \rightarrow \{1, 2\}$ avec $f = h \circ g$ où g et h sont telles que $g(s) = (u, i)_s$ et $h((u, i)_s) = 1$ si s est un virage, et 2 sinon. f' est définie de la même manière sur les s' de $\mathcal{C}'$, c'est-à-dire $f' : \mathcal{C}' \rightarrow \{1, 2\}$ avec $f' = h' \circ g'$ et g' et h' sont définies de la même manière que g et h , mais sur les états de $\mathcal{C}'$.

Nous définissons un isomorphisme entre deux états de $\mathcal{C}$ ssi ils ont la même image par f , et un isomorphisme entre deux états de $\mathcal{C}'$ ssi ils ont la même image par f' . Nous définissons ensuite $F : \mathcal{C} \rightarrow \mathcal{C}'$ par (3).

$$F(s) = \begin{cases} l', & \text{si } f(s) = 1 \\ k', & \text{si } f(s) = 2 \end{cases} \quad (3)$$

Il est facile de voir, si les définitions exactes sont connues, que (3) est une équivalence de catégories qui permet de transférer π de $\mathcal{C}'$ à $\mathcal{C}$. F induit une bijection F' entre les ensembles de classes (ou types de position – ici $\mathcal{C}, \mathcal{C}'$ sont les types de virage et $\mathcal{S}, \mathcal{S}'$ sont les types de lignes droites) : $F' : (\mathcal{C} / \simeq) \rightarrow (\mathcal{C}' / \simeq)$ où $(\mathcal{C} / \simeq)$ est égal à $\{\mathcal{C}, \mathcal{S}\}$ et $(\mathcal{C}' / \simeq)$ est égal à $\{\mathcal{C}', \mathcal{S}'\}$. On obtient finalement $F'(\mathcal{C}) = \mathcal{C}'$ et $F'(\mathcal{S}) = \mathcal{S}'$.

Cet exemple de catégorisation systématique et de généralisation prouve que nous ne travaillons pas à l'échelle des états, mais que nous considérons les types d'états.

Analogies entre jeux vidéo. Considérons un jeu $\mathcal{C}$ (par exemple le jeu Breakout pour Atari 2600) où une raquette horizontale est située en x_0 sur un segment $[0, N]$ d'états initiaux possibles ($0 \leq x_0 \leq N$) et peut se déplacer vers la gauche ou la droite afin d'atteindre une balle devant arriver à la position $x_1 \in [0, N]$. La stratégie π est la suivante : si $x_1 < x_0$, se déplacer vers la gauche, si $x_1 = x_0$, ne pas bouger, mais si $x_1 > x_0$, alors se déplacer vers la droite. Désormais, considérons un autre jeu $\mathcal{C}'$ (tel que Pong sur Atari 2600 par exemple) où cette fois-ci une raquette verticale se trouve en position y_0 sur le segment $[0, N']$ des positions atteignables ($0 \leq y_0 \leq N'$) et peut se déplacer de haut en bas pour toucher une balle qui atteindra la position $y_1 \in [0, N']$.

Nous souhaitons identifier $\mathcal{C}$ et $\mathcal{C}'$ afin de transposer la stratégie π de $\mathcal{C}$ à $\mathcal{C}'$. Les états (objets) de $\mathcal{C}$ sont les $N + 1$ positions x_1 de $[0, N]$. De la même façon, les états (objets) de $\mathcal{C}'$ sont les $N' + 1$ positions y_1 de $[0, N']$. Si $N = N'$ et $x_0 = y_0$, alors il est facile de définir une bijection de $\mathcal{C}$ à $\mathcal{C}'$ puis de transposer π , mais si $N \neq N'$ ou $x_0 \neq y_0$ alors il est impossible de définir une telle bijection. Pour résoudre ce problème, nous transformons $\mathcal{C}$ et $\mathcal{C}'$ en catégories à travers la définition suivante de nos flèches, et ce faisant, nous pourrions définir une équivalence de catégories $F : \mathcal{C} \rightarrow \mathcal{C}'$. Pour définir ces flèches, nous utilisons les fonctions d'observation suivantes définies sur les états de $\mathcal{C}$ et $\mathcal{C}'$: $f : \mathcal{C} \rightarrow \{0, 1, 2\}$ et $f' : \mathcal{C}' \rightarrow \{0, 1, 2\}$ telles que $f(x_1) = 0$ ssi $x_1 < x_0$, $f(x_1) = 1$ ssi $x_1 = x_0$ et $f(x_1) = 2$ ssi $x_1 > x_0$. De la même façon, $f'(y_1) = 0$ ssi $y_1 < y_0$, $f'(y_1) = 1$ ssi $y_1 = y_0$ et $f'(y_1) = 2$ ssi $y_1 > y_0$. Ensuite, nous définissons une flèche inversible entre deux états de $\mathcal{C}$ ssi ils ont la même image par f , ainsi qu'une flèche inversible entre deux états de $\mathcal{C}'$ ssi ils ont la même image par f' . Soit $F : \mathcal{C} \rightarrow \mathcal{C}'$, telle que si $x_1 < x_0$, alors $F(x_1) = 0$; si $x_0 = x_1$, alors $F(x_1) = y_0$; et si $x_1 > x_0$, alors $F(x_1) = N'$. Cette fonction F définit une équivalence de catégories et induit une bijection F' entre les ensembles de classes : $F' : (\mathcal{C} / \simeq) \rightarrow (\mathcal{C}' / \simeq)$. $(\mathcal{C} / \simeq)$ correspond (grâce à la stratégie π) aux actions disponibles $\{Gauche, Attendre, Droite\}$ dans $\mathcal{C}$ alors que $(\mathcal{C}' / \simeq)$ correspond aux actions disponibles $\{Haut, Attendre, Bas\}$ dans $\mathcal{C}'$ (en choisissant d'orienter l'axe des y du haut vers le bas). Dans ce cas-ci $F'(Gauche) = Haut$, $F'(Attendre) = Attendre$ et $F'(Droite) = Bas$.

5 Résultats Expérimentaux

5.1 Circuit de voitures miniatures

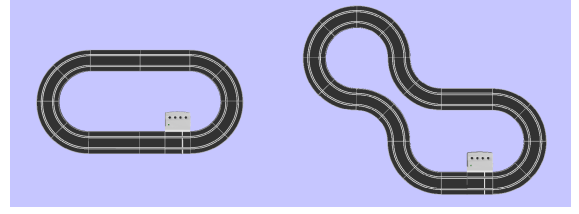


FIGURE 1 – Configuration des circuits : circuit 1 (à gauche), circuit 2 (à droite)

Matériel et Implémentation. Le montage expérimental est constitué d'un circuit MINI Challenge Set C1320 de la marque Scalextric dont le compteur de tours mécanique a été remplacé par un capteur à effet Hall omnipolaire digital. Le courant est mesuré par une résistance shunt placée en série avec les rails d'alimentation. Le courant et la tension sont d'abord filtrés par un filtre RC passif du second ordre puis échantillonnés à $f_s = 100$ Hz. Ce montage ne contient aucun capteur supplémentaire. Les algorithmes sont écrits en C et exécutés en temps réel sur un Arduino Mega 2560 qui dispose de 8192 octets de mémoire vive (RAM). Nous définissons la période d'échantillonnage par $t_s = 1/f_s$, et l'instant associé à l'échantillon k est kt_s où $k \in \mathbb{N}$.

Ce cas basé sur l'utilisation d'analogies, directement inspiré de notre utilisation de la théorie des catégories, repose sur deux modules : un module de récompense, ainsi qu'un module de prise de décision, tous deux décrits ci-dessous. Puisque qu'il n'y a ici qu'un seul système dynamique, c'est nécessairement le « Moi » et aucune identification n'est utile.

Comme en apprentissage par renforcement, notre approche s'appuie sur une récompense tirée de l'environnement. Celle-ci est le résultat de la combinaison de trois variables. La première variable est le temps du tour mesuré directement grâce au compteur de tours. La seconde variable binaire indique la présence de la voiture sur la piste, et la dernière variable binaire indique si la voiture est en mouvement. Ce module de récompense surveille constamment la voiture pour vérifier qu'elle n'est pas sortie de la piste à cause d'une vitesse trop élevée, mais aussi qu'elle ne s'est pas arrêtée à cause de frottements trop importants. Ces deux détecteurs reposent sur une classification par k plus proches voisins (k -NN) en utilisant les courants et tensions comme entrées.

Avec ce module de récompense, l'IA peut piloter la voiture sur des circuits qu'elle découvre sans rejouer ou manipuler des échantillons enregistrés lors d'une partie du joueur humain. La seule information conservée par l'algorithme est une vitesse de sécurité dont l'IA sait qu'elle ne provoque ni sortie de piste, ni arrêt de la voiture. Comme la configuration du circuit change, il n'y a pas bijection entre

TABLE 1 – Temps au tour en secondes (moyenne $\pm$ écart-type)(MLI = Modulation de Largeur d’Impulsion)

	HUMAIN (10 tours)	IA (jusqu’à 10 tours)	RÉGLAGES DE L’IA
CIRCUIT 1 (12 tronçons)			
Premier tour	2.99 ± 0.46	3.12 ± 0.09	MLI=39% de la vitesse maximale
Dernier tour	2.29 ± 0.14	2.52 ± 0.08	Analogies (adaptation de la vitesse)
CIRCUIT 2 (18 tronçons)			
Premier tour	4.30 ± 1.16	3.66 ± 0.03	MLI=39% de la vitesse maximale
Dernier tour	3.08 ± 0.54	3.13 ± 0.02	Analogies (adaptation de la vitesse)

les deux configurations et le cas bijectif ne peut pas s’appliquer. Il faut donc s’appuyer sur des analogies et transposer les connaissances acquises dans une première configuration en utilisant l’équation (2), conformément au formalisme des cas non bijectifs décrit dans la partie 4. La fonction $h((u, i)_s)$ est évaluée par un k-NN dont le coefficient a été choisi comme le meilleur compromis entre la performance possible sur la cible embarquée et l’obtention de résultats satisfaisants. Les deux classificateurs sont entraînés avec environ 1000 échantillons de (u, i) (ce qui correspond à seulement 10 s de jeu) et sont implantés en version condensée afin d’être exécutables sur l’Arduino.

En pratique, l’approche par analogies fonctionne comme suit : la voiture démarre sur le circuit inconnu avec cette vitesse de sécurité importée de la première configuration. La fonction $h((u, i)_s)$, évaluée par le k-NN à partir des mesures de courant et de tension, détermine si la voiture est dans une configuration que nous appelons courbe ou ligne droite. L’IA choisit ensuite la meilleure commande utilisée au cours des expériences précédentes, en vérifiant l’objectif de minimiser le temps au tour en maintenant la voiture sur la piste. L’algorithme permet donc de généraliser les connaissances acquises au cours d’expériences passées et de les appliquer dans une configuration radicalement différente. En effet, les deux circuits schématisés en figure 1 sont de forme et de taille différentes et une simple reproduction d’une commande préalablement enregistrée échouerait à fournir des résultats probants.

Résultats. Les expériences décrites dans cet article ont été menées sur deux configurations de complexité différente présentées figure 1, et leur résultats sont présentés tableau 1. L’IA démarre à vitesse constante (avec une Modulation de Largeur d’Impulsion de 39% de la vitesse maximale). Cette IA, qui repose sur l’établissement d’analogies entre configurations et ne rejoue pas d’enregistrement d’une quelconque partie précédente, est capable d’améliorer les temps au tour en moins de 10 tours sur une piste inconnue. L’algorithme atteint des temps similaires aux temps humains sur le circuit 2 : 3.08 s pour les derniers tours des joueurs humains contre 3.13 s pour l’algorithme. Les améliorations futures de l’algorithme sur un circuit inconnu incluront l’optimisation des vitesses transposées par l’algorithme. En effet, seule une vitesse de sécurité a été utilisée ici afin d’éviter les sorties de pistes des voitures

pilotées par l’algorithme, alors que certains tours humains ont occasionné des sorties de pistes et n’ont donc pas été comptabilisés.

Le cadre théorique détaillé en section 4 permet au système d’être au niveau des meilleurs joueurs sur ce jeu en moins d’une minute, sans ajout de capteurs embarqués. Ce cadre permet d’atteindre de tels résultats sur des circuits inconnus où la simple reproduction d’un tour humain conduirait immédiatement à une sortie de piste.

5.2 Jeux vidéo Atari 2600

Configuration et implémentation de la théorie. Tandis que le circuit de voitures miniatures nous a permis de valider l’approche sur des signaux analogiques réels, Arcade Learning Environment (ALE, cf [2]) nous a permis de la valider sur des configurations plus complexes, avec des signaux déjà échantillonnés provenant de l’émulateur. Les concepts d’entités avec le « Moi » introduits en 3.2 sont également utilisés pour jouer aux jeux Atari 2600. On détecte les entités grâce à des algorithmes de traitement d’images : le filtre de Sobel (image du centre en figure 2) et la détection des boîtes englobantes (image de droite en figure 2). Elle repose sur la librairie OpenCV [19]. Le « Moi » est trouvé en utilisant un algorithme d’identification de système. Des séquences d’actions pseudo-aléatoires [14] sont envoyées comme commande dans ALE pour d’abord identifier quelles entités sont affectées par ces signaux, puis pour construire les modèles dynamiques décrits par les équations (1) avec $N_d = 2$. Peu d’entités sont contrôlables : ce sont les « Moi ». Leurs formes peuvent changer en cours de jeu, comme la raquette dans Breakout, d’où la possibilité d’identifier plusieurs entités comme étant le « Moi ». Ces mesures mettent également à jour les fonctions $p(E, F)$ qui expriment la probabilité que le contact entre les entités E et F change le score, d’une manière similaire à une fonction de récompense en apprentissage par renforcement. À partir de ces fonctions p , nous déduisons les entités amies et ennemies, conduisant à la stratégie de survie basique décrite en partie 3.2.

Nous choisissons d’utiliser le DQN comme référence, car cette publication obtient de bons résultats par rapport à un joueur humain pour un grand nombre de jeux. Les tests ont donc été effectués avec les paramètres décrits dans [18] : l’IA joue pendant une durée maximale de 5 minutes. Même si notre objectif est de contrôler des systèmes cyber-

TABLE 2 – Scores après 10 000 images d’entraînement (moyenne $\pm$ écart type pour 8 parties)

JEU	ALÉATOIRE	JOUEUR HUMAIN	DQN	MLCT
Breakout	1.7	31.8	1.25 $\pm$ 1.02	414.37 $\pm$ 88.47
Pong	−20.7	9.3	−21.00 $\pm$ 0.00	0.25 $\pm$ 2.48

TABLE 3 – Effet du transfert de connaissances de Breakout vers Pong. Méthode (a) : Aucun transfert d’amis donc actions aléatoires. Méthode (b) : Stratégie de survie basique. Méthode (c) : Transfert de connaissances de Breakout à Pong

MÉTHODE	SCORE	NOMBRE D’IMAGES D’ENTRAÎNEMENT
(a)	−20	500 (Pong)
(b)	0	10 000 (Breakout) + 10 000 (Pong)
(c)	−2	10 000 (Breakout) + 500 (Pong)

physiques, nous souhaitons valider la versatilité de notre approche en la testant d’abord sur cet environnement de référence. Nous avons entièrement reproduit la configuration de cet article en utilisant le code mis à disposition par les auteurs, et avons obtenu des résultats similaires à ceux annoncés. Cela nous a permis de calculer un score moyen pour le DQN avec un faible nombre d’images d’entraînement, pour le comparer au score moyen obtenu avec notre approche.

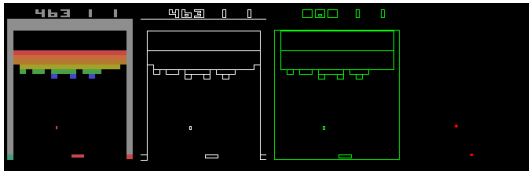


FIGURE 2 – Traitement d’images sur Breakout. De gauche à droite : original, contours, boîtes englobantes, prédictions des vitesses.

Résultats. Les résultats pour Breakout⁴ sont présentés dans la table 2 pour un temps d’entraînement de 10 000 images (moins de 3 minutes). La figure 3 compare cette approche avec le DQN et montre que des scores similaires sont atteignables en utilisant 20 000 fois moins d’images. Les étapes d’apprentissage sur Breakout sont les suivantes : durant les premières centaines d’images, l’IA utilise une séquence pseudo-aléatoire pour identifier le « Moi ». Une fois que cette identification a convergé vers le seul système que l’IA contrôle – la raquette – l’algorithme recherche des entités amies et ennemies. En quelques milliers d’images, il détecte que la balle est un ami car le score augmente quand elle casse une brique.

Les scores autour de 400 points correspondent à des parties où quasiment toutes les briques du premier niveau ont été cassées. L’écart-type élevé est dû à certaines parties atteignant des scores de plus de 600 points, obtenus grâce à la

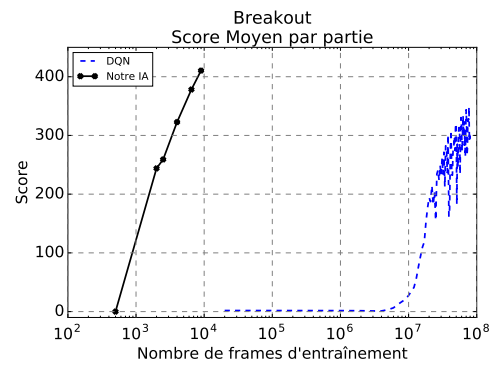


FIGURE 3 – Scores sur Breakout

capacité de l’IA à *réfuter* ses modèles.

Les résultats dans Pong ont été obtenus en utilisant trois approches : la première (a) identifie le « Moi » (mais pas la balle comme ami) en 500 images. La seconde (b) correspond à la procédure suivie pour Breakout : elle identifie le « Moi » puis les amis sur chacun des deux jeux en 10 000 images. La troisième (c) illustre le transfert de connaissances formalisé en utilisant les similarités décrites dans la partie 3.1 : l’IA commence par apprendre à jouer à Breakout en 10 000 images, puis elle identifie le « Moi » sur Pong en 500 images, et enfin elle transfère (sans images d’entraînement supplémentaires) la connaissance de l’ami vers Pong. L’expérience (c) permet d’obtenir un système capable de jouer à Breakout *et* Pong avec seulement 10 500 images d’entraînement comme mis en évidence en table 3.

6 Conclusion

À partir de concepts empruntés aux Sciences Sociales et Cognitives et d’enseignements tirés de caractéristiques propres aux méthodes scientifiques, ces travaux utilisent des outils élémentaires de la théorie des catégories pour faciliter l’établissement d’analogies et le transfert de connaissances entre situations non bijectives. Ces outils permettent aussi de formaliser une approche aboutissant à une IA capable de contrôler un agent dans un monde en

4. Les résultats pour le joueur humain et le mode aléatoire sont tirés de [18]. Le temps d’entraînement pour le joueur humain est de 2 heures soit 432 000 images.

perpétuelle évolution.

Les résultats des expériences conduites sur un système cyber-physique mais aussi sur des jeux vidéo Atari 2600 confirment nos attentes. En effet, ce système a été capable d'identifier, puis de contrôler le « Moi » dans différentes situations, que ce soient des circuits de configurations différentes, ou des jeux Atari 2600 différents, en s'appuyant sur un transfert de connaissances et une faible quantité de données.

References

- [1] P. Bartha, "Analogy and analogical reasoning," in *The Stanford Encyclopedia of Philosophy*, winter 2016 ed., E. N. Zalta, Ed. Metaphysics Research Lab, Stanford University, 2016.
- [2] M. G. Bellemare, Y. Naddaf, J. Veness, and M. Bowling, "The arcade learning environment: An evaluation platform for general agents," *Journal of Artificial Intelligence Research*, vol. 47, pp. 253–279, 2013.
- [3] M. Bogdanovic, D. Markovikj, M. Denil, and N. de Freitas, "Deep apprenticeship learning for playing video games," in *AAAI Workshop on Learning for General Competency in Video Games*, 2015.
- [4] C. Chang and H. J. Keisler, *Model Theory*, 3rd ed. North Holland, 1990.
- [5] Z. Chen and B. Liu, *Lifelong Machine Learning*. Morgan & Claypool Publishers, 2016.
- [6] M. Csikszentmihalyi, *Flow: The Psychology of Optimal Experience*, ser. Harper Perennial Modern Classics. Harper Collins, 2009.
- [7] J. Culbertson and K. Sturtz, "Bayesian machine learning via category theory," *arXiv:1312.1445v1 [math.CT]* 5 Dec 2013, 2013.
- [8] F. de Saussure, *Cours de linguistique générale*. Payot, 1916.
- [9] B. Falkenhainer, K. Forbus, and D. Gentner, "The structure-mapping engine: Algorithm and examples," *Artificial Intelligence*, vol. 41, pp. 2–63, 1989/90.
- [10] R. E. Kalman, "On the general theory of control systems," in *Proceedings of the First International Congress on Automatic Control*. Butterworth, London, 1960, pp. 481–492.
- [11] K. Kanksy, T. Silver, D. A. Mély, M. Eldawy, M. Lázaro-Gredilla, X. Lou, N. Dorfman, S. Sidor, D. S. Phoenix, and D. George, "Schema networks: Zero-shot transfer with a generative causal model of intuitive physics," in *ICML 2017*, 8 2017, pp. 1809–1818.
- [12] S. Lange, M. Riedmiller, and A. Voigtlander, "Autonomous reinforcement learning on raw visual input data in a real world application," in *The 2012 International Joint Conference on Neural Networks (IJCNN)*. IEEE, 2012, pp. 1–8.
- [13] G. Lee, M. Luo, F. Zambetta, and X. Li, "Learning a Super Mario controller from examples of human play," in *2014 IEEE Congress on Evolutionary Computation (CEC)*, July 2014, pp. 1–8.
- [14] W. S. Levine, *The Control Systems Handbook: Control System Advanced Methods*, 2011.
- [15] Y. Liang, M. C. Machado, E. Talvitie, and M. Bowling, "State of the art control of Atari games using shallow reinforcement learning," in *Proceedings of the 2016 ICAAMAS*, ser. AAMAS '16, 2016, pp. 485–493.
- [16] N. Lipovetzky, M. Ramirez, and H. Geffner, "Classical planning with simulators: Results on the Atari video games," in *IJCAI'15*, 2015, pp. 1610–1616.
- [17] S. Mac Lane, *Categories for the working mathematician*, 2nd ed. Springer-Verlag, 1998.
- [18] V. Mnih, K. Kavukcuoglu, D. Silver, A. A. Rusu, J. Veness, M. G. Bellemare, A. Graves, M. Riedmiller, A. K. Fidjeland, G. Ostrovski, S. Petersen, C. Beattie, A. Sadik, I. Antonoglou, H. King, D. Kumaran, D. Wierstra, S. Legg, and D. Hassabis, "Human-level control through deep reinforcement learning," *Nature*, vol. 518, no. 7540, pp. 529–533, February 2015.
- [19] OpenCV, "Open source computer vision library," <https://opencv.org/>, 2017.
- [20] J. Ortega, N. Shaker, J. Togelius, and G. N. Yannakakis, "Imitating human playing styles in Super Mario Bros," *Entertainment Computing*, vol. 4, no. 2, pp. 93 – 104, 2013.
- [21] T. Pepels, M. H. M. Winands, and M. Lanctot, "Real-Time Monte-Carlo Tree Search in Ms. Pac-Man," *IEEE Trans. on Computational Intelligence and AI in Games*, vol. 6, no. 3, pp. 245–257, Sept 2014.
- [22] J. Piaget, *The construction of reality in the child*. Basic Books, 1954.
- [23] K. Popper, *The Logic of Scientific Discovery*. Hutchinson & Co, London, 1959.
- [24] L. Pusman and K. Kosturik, "Control algorithm based on phase locked loop," in *Telecommunications Forum (TELFOR)*, 2013 21st, Nov 2013, pp. 605–607.
- [25] J. M. Schaeffer, *L'expérience esthétique*. Gallimard, 2015.
- [26] A. Schutz, *The phenomenology of the social world*, ser. Northwestern University studies in phenomenology & existential philosophy. Northwestern University Press, 1967.
- [27] D. Sperber, *La Contagion des idées*. Éditions Odile Jacob, 1996.
- [28] D. I. Spivak, *Category Theory for the Sciences*. The MIT Press, 2014.

Modélisation thématique à l'aide des plongements lexicaux issus de Word2Vec

Svitlana Galeshchuk¹

Bruno Chaves¹

¹ PSL Université Paris, Governance Analytics

{svitlana.galeshchuk,bruno.chavesferreira}@dauphine.fr

Résumé

Le papier étudie diverses approches pour la modélisation thématique et, en particulier, la méthode améliorée basée sur la paramétrisation des thèmes à partir d'une distribution continue sur l'espace des plongements lexicaux afin de tenir compte des interdépendances sémantiques. Ainsi, nous incorporons les représentations vectorielles des mots entraînés avec le réseau de neurones Word2Vec dans le processus génératif de la modélisation thématique. Nous proposons une approche alternative avec une approximation bêta de la distribution de l'information mutuelle et la comparons aux méthodes LDA standard et LDA Gaussien.

Mots Clefs

Modélisation Thématique, Latent Dirichlet Allocation, LDA Gaussien, Information Mutuelle, Plongements Lexicaux, Word2Vec.

Abstract

This paper discusses approaches for static topic modeling, in particular an improved method based on topic parametrization from a continuous distribution over the space of word embeddings. Word embeddings corpora proves to reflect semantic interdependences. Thus, we incorporate vectorized word representations trained with Word2Vec neural network in a generative process of topic modeling. The alternative approach with beta approximation of mutual information distribution over embeddings is proposed and compared with vanilla LDA and Gaussian LDA methods.

Keywords

Topic Modelling, Latent Dirichlet Allocation, Gaussian LDA, Mutual Information, Word Embedding, Word2Vec.

1 Introduction

La modélisation thématique est devenue une méthode de choix pour la fouille de données textuelles non structurées. De nombreux papiers (voir partie 2) se consacrent à l'implémentation de cette méthode, usuellement fondée sur la LDA (*Latent Dirichlet Allocation*), dans des domaines variés allant des textes juridiques aux papiers scientifiques. La méthode LDA définit les probabilités de mots sur une loi de Dirichlet qui appartient à la famille des distributions discrètes. Toutefois, l'usage de distributions discrètes em-

pêche la découverte de mots nouveaux émergeant des thématiques. Pour cela, il est recommandé d'employer des distributions continues permettant au modèle d'attribuer à un mot nouveau une probabilité élevée d'appartenance à une thématique simplement parce que ce dernier est similaire à un mot existant représentatif de la thématique en question. Dans ce contexte, nous étudions la méthode LDA statique et ses formes modifiées avec des distributions de mots continues suivant des lois normales et bêta.

Le papier a ainsi pour objectif d'améliorer l'algorithme du modèle LDA standard à l'aide de l'approche continue, notamment la loi bêta, sur les plongements lexicaux (*word embeddings*).

La suite du papier est structurée de la manière suivante : La section 2 présente une brève revue de la littérature de la modélisation thématique. La section 3 décrit l'approche LDA standard. La section 4 présente la méthodologie de l'approche faisant appel à une distribution continue avec les plongements lexicaux. La section 5 introduit les données utilisées. La section 6 décrit le dispositif expérimental. Enfin, la section 7 conclut en présentant les résultats et des pistes pour des recherches futures.

2 Revue de la littérature

La modélisation thématique, fondée sur la méthode LDA, est devenue l'une des principales méthodes de fouille textuelle de la dernière décennie. Elle fait partie de la famille des méthodes d'apprentissage non supervisées destinées à extraire les structures thématiques latentes des corpus textuels.

Cette approche a été appliquée avec succès à des domaines de recherche variés : le journalisme, pour analyser les structures et tendances thématiques des articles de presse [7], les corpus de brevets [5], ou encore, pour classifier les textes issus de la littérature scientifique [14]. Toutefois, la méthode LDA n'est pas exempte de défauts. En particulier, la représentation des thématiques définies comme des distributions discrètes de mots empêche la prise en considération de mots nouveaux. Cette limitation peut être contournée en mobilisant, à la place, une distribution continue de mots sur des plongements lexicaux. Ces derniers définissent la représentation vectorielle des mots basée sur le contexte de leur utilisation au sein du corpus.

Blei et al. [1] proposent d'utiliser une distribution gaus-

sienne dans le processus génératif de la modélisation thématique dynamique afin de suivre l'évolution des thématiques au cours du temps. Dans cette étude, nous nous focalisons sur l'usage de distributions continues pour l'amélioration de la modélisation thématique statique plutôt que dynamique. En poursuivant un objectif similaire au notre, certains travaux [6], [11] et [12] ont proposé l'utilisation de la distribution gaussienne. Toutefois, dans notre papier, nous nous fondons plutôt sur les contributions de Das et al. [3] et Xun et al. [13]. Ces auteurs font émerger les thématiques d'une distribution gaussienne sur des plongements lexicaux en utilisant le modèle Word2Vec. De manière similaire, nous utilisons le modèle Word2Vec pour présenter les mots sous la forme de vecteurs et déduire les thématiques de distributions continues. Toutefois, nous justifions l'usage de la distribution bêta plutôt que gaussienne sur la base des résultats de Levy et Goldberg [8]. En effet, ces auteurs montrent que le modèle Word2Vec estime de manière implicite les informations mutuelles des paires de mots.

3 Modélisation thématique LDA standard

Le modèle LDA (*Latent Dirichlet Allocation*) est un modèle Bayésien faisant partie de la famille des modèles non supervisés génératifs où les observations sont générées par des variables latentes. Dans le contexte de la modélisation thématique, on cherche à découvrir des thèmes latents, à partir d'une collection de documents (articles, ouvrages, etc.) considérés comme des « sacs de mots » (*bag-of-words*) dans le sens où l'on ne tient pas compte de l'ordre des mots. Chaque document est modélisé par un mélange de thèmes qui génère ensuite chaque mot du document. Blei et al. [2] décrivent le processus génératif de LDA de la manière suivante :

1. Pour $k = 1$ à K :
 - (a) Déduire la $\phi^{(k)} \sim \text{Dirichlet}(\beta)$
2. Pour chaque document d dans le corpus D :
 - (a) Déduire la distribution de thèmes $\theta_d \sim \text{Dirichlet}(\alpha)$
 - (b) Pour chaque index de mots n de 1 à N_d :
 - i. Déduire le thème $z_n \sim \text{Multinomiale}(\theta_d)$
 - ii. Déduire $w_{d,n} \sim \text{Multinomiale}(\phi^{z_n})$

Où $\phi^{(k)}$ est la distribution de mots dans le vocabulaire du k^{ieme} thème, θ_d est la distribution de thèmes dans le document d et z_n est le thème n associé au mot $w_{d,n}$.

4 Modélisation thématique à partir d'une distribution continue

Cette partie de l'article présente l'approche gaussienne de la modélisation thématique fondée sur le plongement lexical et le modèle Word2Vec. Nous commençons par décrire Word2Vec et son utilisation dans le cadre de la modélisation thématique. Ensuite, nous discutons de la possibilité

de représenter les thématiques à partir d'une distribution continue plutôt que discrète.

4.1 Le plongement lexical

Le modèle de Word2Vec est la représentation interne à partir d'un modèle de réseau de neurones de séquences de mots. Word2Vec utilise le perceptron monocouche pour apprendre le plongement lexical des mots ; c'est-à-dire que les mots sont appris à partir du contexte où ils sont mentionnés. Deux approches de Word2Vec sont proposées : CBOW et skip-gram. Nous implémentons le modèle skip-gram avec un échantillonnage négatif. Dans le processus d'apprentissage de Word2Vec, les mots avec des significations similaires convergent de manière graduelle vers les zones voisines de l'espace vectoriel [13]. Nous enrichissons les mots du corpus en les remplaçant par les mots correspondants de Word2Vec comme dans l'approche définie par Xun et al. [13].

4.2 La méthodologie de l'algorithme de LDA Gaussien et l'approche développée

La modélisation thématique de corpus textuels avec LDA est fondée sur les fréquences de types de mots. L'approche que nous utilisons est fondée sur l'idée selon laquelle les textes représentent des séquences de plongement lexicaux. Word2Vec transforme les mots en des vecteurs. Les mots, usuellement représentés par des valeurs discrètes, sont alors modifiés en des valeurs continues. Das et al. [3] font émerger les thématiques d'une distribution gaussienne sur ces plongements lexicaux et placent les *a priori conjugués* sur les valeurs suivantes : loi normale centrée à zéro pour la moyenne et la covariance.

Ils considèrent chaque document comme un mélange de thèmes de la loi de Dirichlet et décrivent le processus génératif de LDA Gaussien suivant :

1. Pour $k = 1$ à K :
 - (a) Déduire la covariance du thème $E_k \sim W^{-1}(\phi, v)$
 - (b) Déduire la moyenne du thème $\mu_k \sim N(\mu, \frac{1}{K} E_k)$
2. Pour chaque document d dans le corpus D :
 - (a) Déduire la distribution de thèmes $\theta_d \sim \text{Dirichlet}(\alpha)$
 - (b) Pour chaque index de mots n de 1 à N_d :
 - i. Déduire le thème $z_n \sim \text{Multinomiale}(\theta_d)$
 - ii. Déduire $v_{d,n} \sim N(\mu_{z_n}, E_{z_n})$

Ici $v_{d,n}$ est la représentation vectorielle du mot dans le document. W^{-1} est la loi de Wishart inverse pour la covariance.

Les auteurs justifient le choix de la paramétrisation gaussienne par les observations de Hermann et Blunsom [4] selon lesquelles les distances euclidiennes entre les plongements lexicaux sont corrélés avec la similarité sémantique. Pourtant, Levy et Goldberg [8] démontrent que le modèle

de Word2Vec factorise une matrice de contexte de mots (*co-occurrence matrix*) de manière implicite. Ses cellules sont les informations mutuelles des paires de mots et de contextes respectifs décalés d'une constante globale. Ainsi, les vecteurs de mots sont déduits de la distribution des informations mutuelles. Zaffalon et Hutter [15] montrent que la meilleure approximation de la loi de l'informations mutuelles conditionnelles est la loi bêta. Elle appartient à une famille de lois de probabilités continues. Dans notre approche nous suivrons les résultats de Zaffalon et Hutter [15]. Par suite nous proposons le processus génératif de LDA :

1. Pour $k = 1$ à K :
 - (a) Déduire la covariance du thème
 $E_k \sim W^{-1}(\phi, v)$
 - (b) Déduire la moyenne du thème
 $\mu_k \sim N(\mu, \frac{1}{K} E_k)$
2. Pour chaque document d dans le corpus D :
 - (a) Déduire la distribution de thèmes
 $\theta_d \sim \text{Dirichlet}(\alpha)$
 - (b) Pour chaque index de mots n de 1 à N_d :
 - i. Déduire le thème $z_n \sim \text{Multinomiale}(\theta_d)$
 - ii. Déduire $v_{d,n} \sim \text{bêta}(\alpha_n, \beta_{z_n})$

Où α et β sont les paramètres de forme de la distribution bêta.

5 Les données utilisées

Dans notre étude, nous utilisons un corpus composé des titres et résumés des articles présentés à la conférence SIOE (*Society for Institutional & Organizational Economics*) de 2008 à 2017. SIOE est une société savante internationale sur l'économie des institutions et organisations. Elle organise chaque année la principale conférence internationale consacrée à la recherche sur ces thématiques. Les données ont été récupérées à partir de la base de données MySQL du site web de la conférence (www.sioe.org).

6 Démarche expérimentale

Les résultats issus du modèle LDA standard, le modèle thématique reproduit à partir de Das et al., 2015 [3] et le modèle que nous avons développé sont présentés, respectivement, dans les tables 1, 2 et 3. Par ailleurs, la visualisation des résultats LDA avec librairie Python pyLDA est présentée dans la figure 1 ci-dessous. Les trois modèles sont présentés comme des clusters de mots sur 4 thématiques. Ces dernières sont assez proches dans les 3 modèles. On peut les représenter par les termes suivants : « Management », « Institutional framework », « Legal framework » et « Market environment ».

L'évaluation est l'un des principaux défis de la modélisation thématique. Des méthodes qualitatives et quantitatives peuvent être mobilisées comme dans [3] et [13]. Nous avons décidé d'utiliser une méthode quantitative et, plus

particulièrement, celle proposée par Röder et al. [10]. Les auteurs élaborent une méthodologie permettant de mesurer la cohérence thématique qui consiste à mesurer l'ajustement entre des paires les mots ou sous-ensemble de mots. L'algorithme commence par effectuer une segmentation par paires de mots. Ensuite, chaque paire de mots est évaluée à l'aide d'un score d'information mutuelle spécifique (*pointwise mutual information*) normalisée et les probabilités des mots sont calculées. La cohérence résulte de l'agrégation de la concordance des paires sur la base des probabilités calculées. Pour cela, nous avons utilisé la librairie Python Palmetto qui permet de calculer la cohérence thématique des ensembles de mots ci-dessous. Les résultats (arrondis) obtenus sont présentés dans la dernière ligne des tables 1, 2 et 3. Notre approche obtient le score agrégé, sur les 4 thèmes, le plus élevé (1.342). L'approche LDA standard arrive en second (1.315) et le LDA gaussien en dernier (1.246). Ces résultats n'en restent pas moins très proches. Par conséquent, nous envisageons d'utiliser d'autres méthodes qualitatives et quantitatives dans des recherches futures.

Management	Institutional framework	Legal framework	Market environment
firm market law industry innovation investment incentive cost organization model	institution country level development growth government effect state sector impact	policy state law model government court decision election party case	contract cost agent transaction model governance market property system party
0.442	0.311	0.316	0.250

TABLE 1: Modélisation thématique LDA standard

Management	Institutional framework	Legal framework	Market environment
firm market cost performance industry quality procurement incentive strategy agent	state corruption development industry market governance institution policy regime economy	innovation patent property patent regulation judge law crime rule firm	business analysis decision market right datum change innovation governance capital
0.390	0.433	0.323	0.098

TABLE 2: Modélisation thématique de Das et al. [3]

Management	Institutional framework	Legal framework	Market environment
firm contract cost price market governance procurement transaction agent strategy	government country level state market tax institution agent decision policy	enforcement law patent property right system rule crime judge firm	firm market country land investment capital innovation price level change
0.467	0.263	0.321	0.290

TABLE 3: Notre Modélisation (*distribution bêta*)

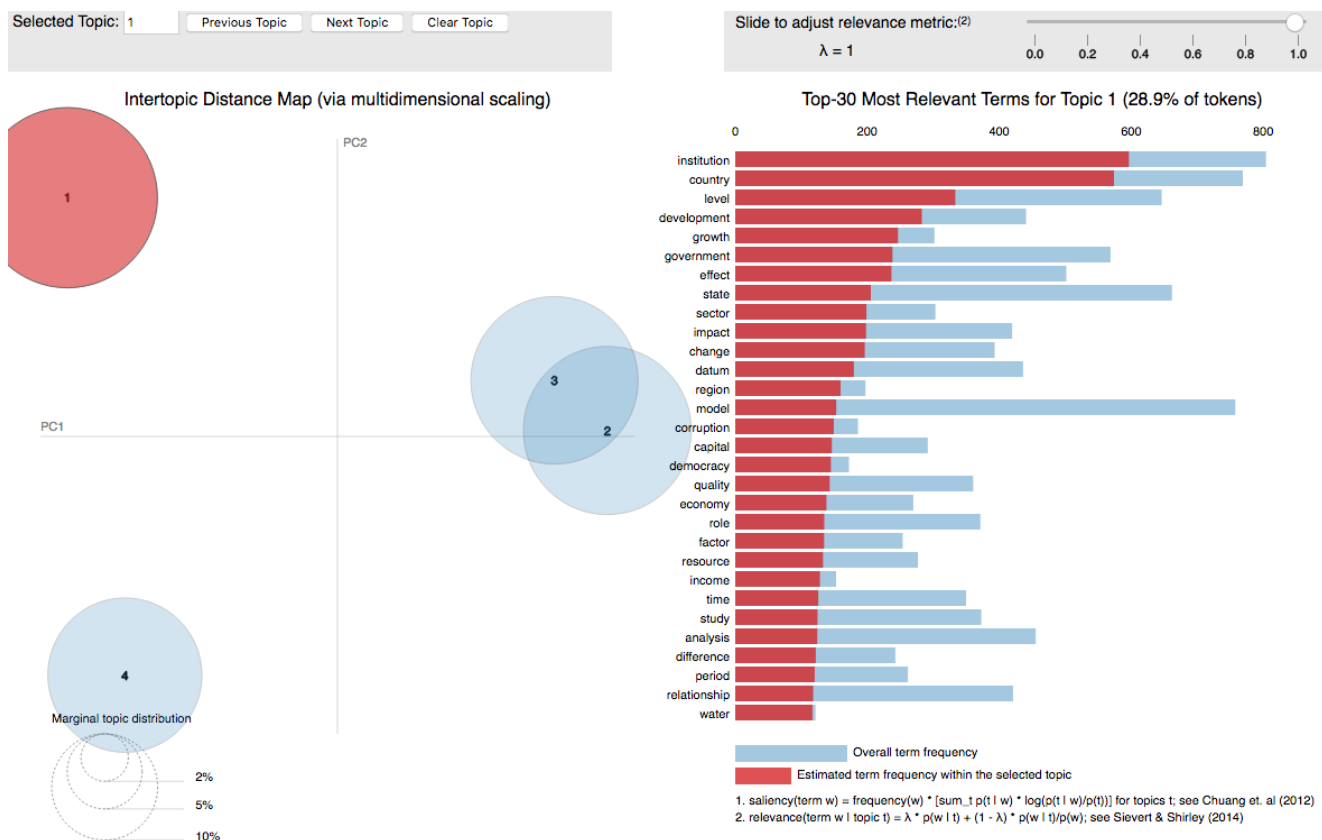


FIGURE 1 – Visualisation du thème « *Institutional framework* » avec le modèle LDA standard

Références

- [1] Blei, D. M., & Lafferty, J. D. (2006, June). Dynamic topic models. In *Proceedings of the 23rd international conference on Machine learning* (pp. 113-120). ACM.
- [2] Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent dirichlet allocation. *Journal of machine Learning research*, 3(Jan), 993-1022.
- [3] Das, R., Zaheer, M., & Dyer, C. (2015). Gaussian lda for topic models with word embeddings. In *Proceedings of the 53rd Annual Meeting of the Association for Computational Linguistics and the 7th International Joint Conference on Natural Language Processing* (Volume 1 : Long Papers) (Vol. 1, pp. 795-804).
- [4] Hermann, K. M., & Blunsom, P. (2014). Multilingual models for compositional distributed semantics. arXiv preprint arXiv :1404.4641.
- [5] Hu, Z., Fang, S., & Liang, T. (2014). Empirical study of constructing a knowledge organization system of patent documents using topic modeling. *Scientometrics*, 100(3), 787-799.
- [6] Hu, P., Liu, W., Jiang, W., & Yang, Z. (2012, September). Latent topic model based on Gaussian-LDA for audio retrieval. In *Chinese Conference on Pattern Recognition* (pp. 556-563). Springer, Berlin, Heidelberg.
- [7] Jacobi, C., van Atteveldt, W., & Welbers, K. (2016). Quantitative analysis of large amounts of journalistic texts using topic modelling. *Digital Journalism*, 4(1), 89-106.
- [8] Levy, O., & Goldberg, Y. (2014). Neural word embedding as implicit matrix factorization. In *Advances in neural information processing systems* (pp. 2177-2185).
- [9] Naili, M., Chaibi, A. H., & Ghezala, H. H. B. (2017). Comparative study of word embedding methods in topic segmentation. *Procedia Computer Science*, 112, 340-349.
- [10] Röder, M., Both, A., & Hinneburg, A. (2015, February). Exploring the space of topic coherence measures. In *Proceedings of the eighth ACM international conference on Web search and data mining* (pp. 399-408). ACM.
- [11] Wang, C., Blei, D., & Heckerman, D. (2012). Continuous time dynamic topic models. arXiv preprint arXiv :1206.3298.
- [12] Weinshall, D., Levi, G., & Hanukaev, D. (2013, February). LDA topic model with soft assignment of descriptors to words. In *International Conference on Machine Learning* (pp. 711-719).

- [13] Xun, G., Gopalakrishnan, V., Ma, F., Li, Y., Gao, J., & Zhang, A. (2016, December). Topic discovery for short texts using word embeddings. *In Data Mining (ICDM), 2016 IEEE 16th International Conference on* (pp. 1299-1304). IEEE.
- [14] Yau, C-K, Porter, A.L., Newman, N.C., and Suominen, A. (2014), Clustering scientific documents with topic modeling, *Scientometrics*, GTM special issue; 100 (3) 767-786.
- [15] Zaffalon, M. & Hutter M. (2002). Robust feature selection by mutual information distributions. *In Proceedings of the Eighteenth conference on Uncertainty in artificial intelligence (UAI'02)*, Adnan Darwiche and Nir Friedman (Eds.). Morgan Kaufmann Publishers Inc., San Francisco, CA, USA, 577-584.

Réseaux de Neurones Récurrents Multi-tâches pour l'Analyse Automatique d'Arguments

Jean-Christophe Mensonides¹
Jacky Montmain¹

Sébastien Harispe¹
Véronique Thireau²

¹ LIG2P, IMT Mines Ales, Univ Montpellier, Ales, France

² Université de Nîmes, CHROME, Rue du Dr Georges Salan, Nîmes, France

jean-christophe.mensonides@mines-ales.fr

Résumé

Dans cet article nous proposons une méthode d'extraction et d'analyse automatique d'arguments à partir de textes bruts, en nous affranchissant de l'utilisation de caractéristiques manuellement définies par des experts. Nous présentons un modèle multi-tâches faisant appel à des techniques d'apprentissage profond, composé de plusieurs couches de réseaux de neurones récurrents. Plus particulièrement, nous tirons parti de paramètres entraînés sur des tâches simples, comme l'étiquetage morpho-syntaxique ou le chunking, afin d'obtenir un modèle capable de traiter des tâches plus complexes nécessitant une compréhension fine du texte.

Mots Clef

Traitement Automatique du Langage Naturel, Extraction d'arguments, Réseaux de neurones récurrents, Apprentissage profond.

Abstract

In this article we propose a method performing automatic extraction and analysis of arguments from raw texts, without using handcrafted features. We introduce a multi-task deep learning model stacking several layers of recurrent neural networks. Specifically, we make use of weight parameters trained on simple tasks, such as Part-Of-Speech tagging and chunking, in order to obtain a model able to handle more complex tasks that require a detailed understanding of the text.

Keywords

Natural Language Processing, Argument mining, Recurrent neural networks, Deep learning.

1 Introduction

L'argumentation est un ensemble de techniques visant à faire adhérer un interlocuteur à un point de vue qui lui est présenté, en construisant un raisonnement à base d'arguments. Bien que l'étude de l'argumentation soit un champ étudié depuis longtemps dans des domaines tels que

la philosophie ou la linguistique, l'extraction et l'analyse automatique d'arguments au sein de corpus textuels (aussi appelé *argument mining*) forment des axes de recherche relativement nouveaux. Un système d'argument mining a pour objectif la génération automatique d'un graphe d'arguments à partir de textes non structurés, et peut généralement être divisé en une séquence d'étapes comportant notamment la détection d'arguments et la modélisation des liens unissant ces derniers [1]. Nous nous limitons à une étude de la micro-structure argumentative, consistant à analyser la manière dont différents composants argumentatifs interagissent entre eux au sein d'un même texte.

De manière plus spécifique, Stab et Gurevych [2] ont proposé le corpus Argument Annotated Essays (version 2), contenant 402 dissertations extraites de essayforum.com. La structure argumentative de chaque dissertation a été manuellement annotée suivant un modèle de graphe orienté acyclique connexe, dans lequel les noeuds représentent des composants argumentatifs et les arcs des liens entre ces derniers. Le schéma d'annotation utilisé permet de distinguer trois types de composants argumentatifs : (i) les conclusions majeures, reflétant le point de vue global de l'auteur sur le sujet disserté, (ii) les conclusions intermédiaires, représentant des affirmations qui ne pourraient être acceptées sans justifications complémentaires, et (iii) les prémisses, servant de justifications aux conclusions intermédiaires avancées. Les arcs du graphe sont porteurs d'une étiquette "support" ou "attaque" selon que le composant argumentatif source corrobore ou réfute la cible. Les arcs ne peuvent exister que a) d'une prémisse vers une autre prémisse, b) d'une prémisse vers une conclusion (majeure ou intermédiaire), et c) d'une conclusion intermédiaire vers une autre conclusion (majeure ou intermédiaire).

Afin d'obtenir automatiquement un graphe synthétisant la structure argumentative d'une dissertation, Stab et Gurevych [2] ont proposé une chaîne de traitement constituée de quatre étapes : (1) Délimitation des frontières des composants argumentatifs, (2) Détermination du type

de chaque composant argumentatif, (3) Détermination de l'existence d'un arc entre chaque paire ordonnée de composants argumentatifs, et (4) Etiquetage des arcs existants comme relation de support ou d'attaque.

Dans cet article, nous nous concentrons sur l'étude des tâches (1) et (2), en cherchant à nous affranchir de l'utilisation de caractéristiques définies manuellement par des experts. La section 2 présente un panorama des travaux antérieurs réalisés sur des tâches similaires à (1) et (2). La section 3 décrit le modèle que nous avons mis en place pour traiter les deux tâches évoquées ci-dessus. La section 4 présente les modalités d'entraînement du modèle. La section 5 est consacrée aux expérimentations que nous avons menées et aux résultats obtenus. La section 6 propose des directions et perspectives pour nos futures recherches.

2 Travaux antérieurs

La détection de composants argumentatifs consiste à déterminer les frontières séparant les unités textuelles porteuses d'arguments du reste du texte. Cette tâche est généralement considérée comme un problème de segmentation de texte supervisée au niveau du mot. Les modèles exploitant l'aspect séquentiel des mots, inhérent à la construction d'une argumentation convaincante, sont particulièrement adaptés et utilisés : Madnani et al (2012) utilisent un Conditional Random Field (CRF) afin d'identifier des segments non argumentatifs au sein de dissertations [3], Levy et al (2014) identifient les frontières d'unités textuelles représentant des conclusions supportant ou attaquant le sujet débattu dans des fils de discussions issus de Wikipedia [4], Ajjour et al (2017) utilisent des réseaux de neurones récurrents de type Long Short-Term Memory (LSTM) afin d'extraire des arguments issus de dissertations, d'éditoriaux et de commentaires générés par des internautes [5], Goudas et al (2014) identifient des phrases contenant des arguments avant de déterminer précisément leurs frontières au sein de médias sociaux à l'aide d'un CRF [6], Sardianos et al (2015) déterminent les limites de composants argumentatifs au sein d'articles de presse à l'aide d'un CRF [7], Stab et Gurevych (2017) utilisent un CRF afin d'isoler les composants argumentatifs au sein de dissertations [2], Eger et al (2017) ont recouru à des techniques d'apprentissage profond [8].

La tâche consistant à déterminer le type d'un composant argumentatif (prémisse, conclusion, etc.) a souvent été traité comme un problème de classification de texte supervisée. Eckle-Kohler et al (2015) distinguent des prémisses et des conclusions au sein d'articles de presse à l'aide de Naive Bayes, Random Forest et Support Vector Machine (SVM) [9], Park et Cardie (2014) utilisent un SVM pour déterminer à quel point des affirmations sont justifiées au sein de commentaires d'internautes relatifs à de nouveaux projets de législation [10], Stab et Gurevych (2017)

classifient des composants argumentatifs en prémisses, conclusions intermédiaires et conclusions majeures dans des dissertations en utilisant un SVM [2], Persing et Ng (2016) utilisent un classifieur d'entropie maximale afin de déterminer le type de composants argumentatifs [11], Potash et al (2016) utilisent des réseaux de neurones récurrents dits "séquence à séquence" dans l'objectif d'inférer le type de composants argumentatifs [12].

L'étude de modèles multi-tâches, capables de traiter plusieurs problèmes différents en partageant un sous-ensemble commun de paramètres, a fait l'objet d'un engouement récent au sein de la communauté du traitement automatique du langage. Ce type de modèles est bio-inspiré : un être humain est capable de réaliser une multitude de tâches différentes et peut exploiter, quand cela est nécessaire, son savoir-faire acquis concernant la résolution d'un type de problème pour apprendre plus vite à résoudre d'autres types de problèmes. Ruder (2017) énonce les raisons pour lesquelles ce type de modèle est efficace d'un point de vue apprentissage automatique [13] : l'utilisation de plusieurs corpus différents induit une augmentation implicite du nombre d'exemples disponibles pendant la phase d'entraînement. De plus, le modèle doit rechercher des caractéristiques utiles pour l'ensemble des tâches à traiter, ce qui limite la modélisation du bruit dans les données et permet une meilleur généralisation.

Søgaard et Goldberg (2016) montrent qu'induire de la connaissance a priori dans un modèle multi-tâches en hiérarchisant l'ordre des tâches à apprendre permet d'obtenir de meilleures performances [14]. Yang et al (2016) ont montré qu'entraîner un modèle multi-tâches et multi-langues permettait d'améliorer les performances sur des problèmes où les données ne sont que partiellement annotées [15], Hashimoto et al (2017) obtiennent des résultats compétitifs sur la majorité des tâches d'un même modèle [16]. Le bénéfice d'un modèle multi-tâches n'est cependant pas garanti, et dépend notamment de la distribution des données relatives aux différents problèmes traités (Mou et al (2016) [17], Alonso et Plank (2017) [18], Bingel et Søgaard (2017) [19]).

3 Modèle proposé

Nous proposons un modèle ayant pour objectif 1) de déterminer les frontières de composants argumentatifs présents dans un ensemble de dissertations et 2) de déterminer le type de chaque composant argumentatif dans lesdites dissertations. Nous nous inspirons du travail de Hashimoto et al [16] et optons pour un modèle multi-tâches s'affranchissant de la définition de caractéristiques manuellement définies. Plus particulièrement, nous utilisons des techniques issues de l'apprentissage profond et entraînons un modèle capable d'effectuer de l'étiquetage morpho-syntaxique (EMS), du chunking, de la détection de limites de composants argumentatifs et de la classification de com-

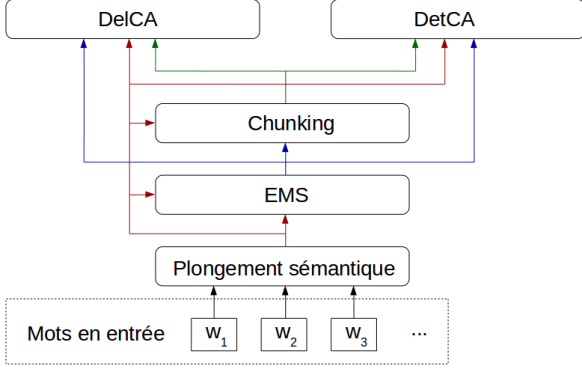


FIGURE 1 – Aperçu de l'architecture utilisée couche par couche. EMS, DelCA et DetCA sont respectivement des acronymes pour Etiquetage Morpho-Syntaxique, Délimitation des Composants Argumentatifs et Détermination du type des Composants Argumentatifs.

posants argumentatifs. Une illustration de l'architecture du modèle est proposée en Figure 1. Les différentes couches utilisées sont présentées ci-dessous.

3.1 Plongement sémantique

Nous utilisons une première couche de plongement sémantique assignant une représentation vectorielle e_t à chaque mot w_t donné en entrée du système. Nous utilisons Glove [20] afin d'obtenir un ensemble de représentations vectorielles entraînées de manière non-supervisée¹. Les représentations vectorielles de mots sont continuellement optimisées au cours de l'entraînement du modèle sur les différentes tâches explicitées ci-dessous. Les mots pour lesquels nous ne disposons pas de représentation vectorielle pré-entraînée sont transformés en un mot spécial $\langle \text{UNK} \rangle$.

3.2 Etiquetage morpho-syntaxique

La seconde couche du modèle correspond à une tâche d'étiquetage morpho-syntaxique (EMS), consistant à assigner pour chaque mot w_t en entrée du système une étiquette morpho-syntaxique (e.g, nom commun, verbe, déterminant, etc.). Nous utilisons un Gated Recurrent Unit (GRU) [21] bi-directionnel afin d'encoder les séquences de mots en entrée du système.

GRU est un réseau de neurones récurrent utilisant un mécanisme de déclenchement sans utilisation de cellule mémoire séparée. A l'instant t , GRU calcule l'état caché h_t de la manière suivante :

$$h_t = (1 - z_t)n_t + z_th_{(t-1)}$$

avec

$$n_t = \tanh(W_n x_t + b_n + r_t(W_{hn} h_{(t-1)} + b_{hn}))$$

1. Le modèle pré-entraîné est issu de <https://nlp.stanford.edu/projects/glove/>

$$r_t = \sigma(W_r x_t + b_r + W_{hr} h_{(t-1)} + b_{hr})$$

$$z_t = \sigma(W_z x_t + b_z + W_{hz} h_{(t-1)} + b_{hz})$$

où x_t représente l'entrée à l'instant t , r_t , z_t et n_t sont respectivement les portes de réinitialisation, d'entrée et de nouveauté, σ représente la fonction sigmoïde, et W et b sont des matrices et vecteurs de paramètres.

En vue d'exploiter le contexte "passé" et "futur" d'un élément d'une séquence de N éléments $[x_1, x_1, \dots, x_N]$, nous pouvons construire un encodage bi-directionnel par concaténation des états cachés obtenus par un encodage séquentiel "à l'endroit" (e.g, à l'instant $t = 1$, l'entrée est x_1 , à l'instant $t = 2$, l'entrée est x_2 , etc.) et un encodage "à l'envers" (e.g, à l'instant $t = 1$, l'entrée est x_N , à l'instant $t = 2$, l'entrée est x_{N-1} , etc.) :

$$\vec{h}_t = \overrightarrow{GRU}(x_t), t \in [1, N]$$

$$\overleftarrow{h}_t = \overleftarrow{GRU}(x_t), t \in [N, 1]$$

$$h_t = [\vec{h}_t; \overleftarrow{h}_t]$$

Nous utilisons les représentations vectorielles des mots constituant l'exemple en cours comme entrée de la couche EMS :

$$\vec{h}_t^{(1)} = \overrightarrow{GRU}(e_t)$$

$$\overleftarrow{h}_t^{(1)} = \overleftarrow{GRU}(e_t)$$

$$h_t^{(1)} = [\vec{h}_t^{(1)}; \overleftarrow{h}_t^{(1)}]$$

Ensuite pour chaque instant t , nous calculons la probabilité d'assigner l'étiquette k au mot w_t de la manière suivante :

$$p(y_t^{(1)} = k | h_t^{(1)}) = \frac{\exp(W_{sm(1)} f c_t^{(1)} + b_{sm(1)})}{\sum_{c_1} \exp(W_{sm(1)} f c_t^{(1)} + b_{sm(1)})} \quad (1)$$

$$f c_t^{(1)} = \text{relu}(W_{fc(1)} h_t^{(1)} + b_{fc(1)}) \quad (2)$$

Avec W et b matrices et vecteurs de paramètres, relu la fonction Unité de Rectification Linéaire [22], et c_1 l'ensemble des classes possibles pour l'étiquette EMS.

3.3 Chunking

Le chunking consiste à assigner une étiquette chunk (chunk nom, chunk verbe, etc.) à chaque mot. Nous calculons les états cachés relatifs au chunking en exploitant ce que le modèle a appris pour la tâche EMS :

$$\vec{h}_t^{(2)} = \overrightarrow{GRU}([e_t; h_t^{(1)}; y_t^{(EMS)}])$$

$$\overleftarrow{h}_t^{(2)} = \overleftarrow{GRU}([e_t; h_t^{(1)}; y_t^{(EMS)}])$$

$$h_t^{(2)} = [\vec{h}_t^{(2)}; \overleftarrow{h}_t^{(2)}]$$

Avec $h_t^{(1)}$ l'état caché obtenu à l'instant t pour la tâche EMS et $y_t^{(EMS)}$ la représentation vectorielle pondérée de

l'étiquette EMS. En suivant Hashimoto et al. [16], $y_t^{(EMS)}$ est défini comme suit :

$$y_t^{(EMS)} = \sum_{j=1}^{card(c_1)} p(y_t^{(1)} = j | h_t^{(1)}) l(j) \quad (3)$$

où $l(j)$ est une représentation vectorielle de la j -ème étiquette EMS. Les représentations vectorielles des étiquettes sont pré-entraînées avec GloVe.

La probabilité d'assigner une étiquette chunk à un mot est ensuite calculée de manière similaire à celle pour les étiquettes EMS (équations (1) et (2)), mais avec un ensemble de paramètres propres à la couche chunking.

3.4 Délimitation des composants argumentatifs (DelCA)

L'objectif de cette tâche est de déterminer, au mot près, les frontières de chaque composant argumentatif au sein d'une dissertation. Nous suivons Stab et Gurevych [2] et traitons cette tâche comme un problème de segmentation de texte supervisée dont les étiquettes suivent un IOB-tagset [23] : le premier mot de chaque composant argumentatif porte l'étiquette "Arg-B", les mots restant dudit composant argumentatif portent l'étiquette "Arg-I", et les mots n'appartenant pas à un composant argumentatif portent l'étiquette "O".

Chaque dissertation est traitée comme une unique séquence de mots que nous encodons de la manière suivante :

$$\begin{aligned} \overrightarrow{h_t^{(3)}} &= \overrightarrow{GRU}([e_t; h_t^{(1)}; y_t^{(EMS)}; h_t^{(2)}; y_t^{(chunk)}]) \\ \overleftarrow{h_t^{(3)}} &= \overleftarrow{GRU}([e_t; h_t^{(1)}; y_t^{(EMS)}; h_t^{(2)}; y_t^{(chunk)}]) \\ h_t^{(3)} &= [\overrightarrow{h_t^{(3)}}; \overleftarrow{h_t^{(3)}}] \end{aligned}$$

où $y_t^{(chunk)}$ est la représentation vectorielle pondérée de l'étiquette chunk, calculée de manière similaire à celle de l'étiquette EMS (équation (3)).

La probabilité d'assigner une étiquette à un mot est ensuite calculée de manière similaire à celle pour les étiquettes EMS, mais avec un ensemble de paramètres propres à la couche DelCA.

3.5 Déterminer le type des composants argumentatifs (DetCA)

L'objectif de cette tâche est de déterminer le type de chaque composant argumentatif parmi prémisses, conclusions intermédiaires et conclusion majeure. Nous traitons cette tâche comme un problème d'étiquetage de segment. Nous considérons qu'un segment peut être la séquence des mots appartenant à un même composant argumentatif ou la séquence des mots appartenant à une même portion

[S1] *The greater our goal is, the more competition we need.*
[S2] *Take Olympic games which is a form of competition for instance, it is hard to imagine how an athlete could win the game without the training of his or her coach, and the help of other professional staffs such as the people who take care of his diet, and those who are in charge of the medical care [S3].* The winner is the athlete but the success belongs to the whole team. Therefore [S4] without the cooperation, there would be no victory of competition [S5].

Consequently, no matter from the view of individual development or the relationship between competition and cooperation we can receive the same conclusion that [S6] **a more cooperative attitudes towards life is more profitable in one's success.**

FIGURE 2 – Un extrait d'une dissertation extrait du corpus. Les passages soulignés par un trait continu constituent des prémisses, ceux soulignés par un trait discontinu constituent des conclusions intermédiaires, et les passages en gras sont des conclusions majeures. Les numéros des segments [S#] sont rajoutés à titre indicatif. Le premier segment correspond à la portion du début du texte jusqu'à la première prémisse. Le second segment correspond à la première prémisse. Le troisième segment correspond à la portion non surlignée entre la première prémisse et la première conclusion intermédiaire, etc.

de texte continue dont les mots n'appartiennent pas à un composant argumentatif. La notion de segment est illustrée en Figure 2.

Nous encodons chaque segment $s_i, i \in [1, L]$ de la manière suivante :

$$\begin{aligned} \overrightarrow{h_{it}} &= \overrightarrow{GRU}([e_{it}; h_{it}^{(1)}; y_{it}^{(EMS)}; h_{it}^{(2)}; y_{it}^{(chunk)}]) \\ \overleftarrow{h_{it}} &= \overleftarrow{GRU}([e_{it}; h_{it}^{(1)}; y_{it}^{(EMS)}; h_{it}^{(2)}; y_{it}^{(chunk)}]) \\ h_{it} &= [\overrightarrow{h_{it}}; \overleftarrow{h_{it}}] \end{aligned}$$

où it représente l'instant t du segment s_i .

Afin que le modèle se concentre davantage sur les marqueurs potentiellement importants (comme "I firmly believe that" ou "we can receive the same conclusion that") nous utilisons un mécanisme d'attention [24], nous permettant de surcroît de synthétiser l'information portée par les états cachés d'un segment en un vecteur de taille fixe :

$$\begin{aligned} u_{it} &= \tanh(W_{att}h_{it} + b_{att}) \\ \alpha_{it} &= \frac{\exp(u_{it}^\top u_{att})}{\sum_t \exp(u_{it}^\top u_{att})} \\ sh_i &= \sum_t \alpha_{it} h_{it} \end{aligned}$$

Avec W_{att} , b_{att} et u_{att} respectivement matrices, biais et vecteurs de paramètres.

Nous encodons ensuite la dissertation à partir des états cachés synthétiques sh_i des segments :

$$\begin{aligned}\overrightarrow{h_j^{(4)}} &= \overrightarrow{GRU}(sh_i), i \in [1, L] \\ \overleftarrow{h_j^{(4)}} &= \overleftarrow{GRU}(sh_i), i \in [L, 1] \\ h_j^{(4)} &= [\overrightarrow{h_j^{(4)}}; \overleftarrow{h_j^{(4)}}]\end{aligned}$$

La probabilité d'assigner une étiquette à un segment est ensuite calculée de manière similaire à celle pour les étiquettes EMS, mais avec un ensemble de paramètres propres à la couche DetCA.

4 Entraînement du modèle

Nous entraînons le modèle en alternant les couches à chaque "epoch" dans l'ordre suivant : EMS, chunking, DelCA, DetCA. Afin d'évaluer la pertinence d'implémenter un modèle multi-tâches, nous avons entraîné une version du modèle en omettant l'optimisation des couches EMS et chunking (nommée "w/o EMS & chunking") et une version du modèle en optimisant l'ensemble des couches (nommée "w/ EMS & chunking"). Les détails de l'entraînement de chaque couche sont explicités ci-dessous.

4.1 Couche EMS

Nous suivons Hashimoto et al. [16] et notons $\theta_{EMS} = (W_{EMS}, b_{EMS}, \theta_e)$ l'ensemble des paramètres intervenant dans la couche EMS. W_{EMS} représente l'ensemble des matrices de paramètres de la couche EMS, b_{EMS} l'ensemble des biais de la couche EMS et θ_e l'ensemble des paramètres de la couche de plongement sémantique des mots. La fonction de coût est définie par :

$$\begin{aligned}J^{(1)} &= - \sum_s \sum_t \log p(y_t^{(1)} = k | h_t^{(1)}) \\ &+ \lambda \|W_{EMS}\|^2 + \delta \|\theta_e - \theta'_e\|^2\end{aligned}$$

Avec $p(y_t^{(1)} = k | h_t^{(1)})$ la probabilité d'assigner la bonne étiquette k au mot w_t de la séquence de mots s , $\lambda \|W_{EMS}\|^2$ est la régularisation L2 et $\delta \|\theta_e - \theta'_e\|^2$ un régularisateur successif. λ et δ sont des hyper-paramètres.

Le régularisateur successif a pour vocation de stabiliser l'entraînement en empêchant θ_e d'être trop modifié spécifiquement par la couche EMS. θ_e étant partagé par l'ensemble des couches du modèle, des modifications trop importantes apportées par l'entraînement de chaque couche empêcheraient le modèle d'apprendre convenablement. θ'_e est l'ensemble des paramètres intervenant dans la couche de vectorisation des mots à l'époque précédente.

4.2 Couche chunking

Nous notons $\theta_{chunk} = (W_{chunk}, b_{chunk}, E_{EMS}, \theta_e)$ l'ensemble des paramètres intervenant dans la couche chunking. W_{chunk} et b_{chunk} sont respectivement les matrices

de paramètres et biais de la couche chunking, incluant ceux de θ_{EMS} . E_{EMS} est l'ensemble des paramètres caractérisant la représentation vectorielle des étiquettes EMS. La fonction de coût est définie de la manière suivante :

$$\begin{aligned}J^{(2)} &= - \sum_s \sum_t \log p(y_t^{(2)} = k | h_t^{(2)}) \\ &+ \lambda \|W_{chunking}\|^2 + \delta \|\theta_{EMS} - \theta'_{EMS}\|^2\end{aligned}$$

Avec $p(y_t^{(2)} = k | h_t^{(2)})$ la probabilité d'assigner la bonne étiquette k au mot w_t de la séquence de mots s . θ'_{EMS} est l'ensemble des paramètres de la couche EMS obtenus avant d'entamer l'"epoch" courante d'entraînement de la couche chunking.

4.3 Couche DelCA

Notons $\theta_{DelCA} = (W_{DelCA}, b_{DelCA}, E_{EMS}, E_{chunk}, \theta_e)$ l'ensemble des paramètres intervenant dans la couche DelCA, avec W_{DelCA} et b_{DelCA} respectivement matrices de paramètres et biais de la couche DelCA, incluant ceux de la couche chunking et EMS. E_{chunk} est l'ensemble des paramètres caractérisant la représentation vectorielle des étiquettes de la couche chunking. La fonction de coût est définie de la manière suivante :

$$\begin{aligned}J^{(3)} &= - \sum_d \sum_t \log p(y_t^{(3)} = k | h_t^{(3)}) \\ &+ \lambda \|W_{DelCA}\|^2 + \delta \|\theta_{chunk} - \theta'_{chunk}\|^2\end{aligned}$$

Avec $p(y_t^{(3)} = k | h_t^{(3)})$ la probabilité d'assigner la bonne étiquette k au mot w_t de la dissertation d . θ'_{chunk} est l'ensemble des paramètres de la couche chunking obtenus avant d'entamer l'"epoch" courante d'entraînement de la couche DelCA.

4.4 Couche DetCA

Notons $\theta_{DetCA} = (W_{DetCA}, b_{DetCA}, E_{EMS}, E_{chunk}, \theta_e)$ l'ensemble des paramètres intervenant dans la couche DetCA, avec W_{DetCA} et b_{DetCA} respectivement matrices de paramètres et biais de la couche DetCA, incluant ceux de la couche chunking et EMS. La fonction de coût est définie de la manière suivante :

$$\begin{aligned}J^{(4)} &= - \sum_d \sum_i \log p(y_i^{(4)} = k | sh_i^{(4)}) \\ &+ \lambda \|W_{DetCA}\|^2 + \delta \|\theta_{chunk} - \theta'_{chunk}\|^2\end{aligned}$$

Avec $p(y_i^{(4)} = k | sh_i^{(4)})$ la probabilité d'assigner la bonne étiquette k au segment s_i de la dissertation d .

5 Expérimentations et résultats

5.1 Hyper-paramètres et données utilisées

Optimisation. Nous entraînons le modèle en alternant les couches, suivant l'ordre suivant : EMS, chunking, DelCA, DetCA. Chaque couche est entraînée pendant une "epoch" avant de passer à la couche suivante. Nous utilisons Adam [25] comme algorithme d'apprentissage, avec

$\beta_1 = 0.9$, $\beta_2 = 0.999$ et $\epsilon = 10^{-8}$. Le coefficient d'apprentissage est commun à toutes les couches et fixé à 10^{-3} au début de l'entraînement, puis multiplié par 0.75 toutes les 10 "epoch". Afin de limiter le problème d'explosion du gradient, nous redimensionnons sa norme avec une stratégie de gradient clipping [26]. Nous suivons [16] et appliquons un gradient clipping de $\min(3.0, \text{profondeur})$, où *profondeur* représente le nombre de GRU impliquées dans la couche entraînée.

Initialisation des paramètres. Afin de faciliter la propagation du gradient lors de l'entraînement, nous utilisons des matrices orthogonales générées aléatoirement comme états initiaux pour les matrices de paramètres des GRU, comme préconisé par Saxe et al. [27]. Les autres matrices de paramètres sont initialisées avec des valeurs issues d'une loi normal $\mathcal{N}(0, \sqrt{2/n_{in}})$, où n_{in} représente le nombre de neurones entrant dans la couche concernée, comme proposé par He et al [28]. Les vecteurs de biais sont initialisés en tant que vecteurs nuls.

Dimensions vectorielles utilisées. La représentation vectorielle utilisée pour les mots en entrée du système et les représentations vectorielles des étiquettes EMS et chunking sont de dimension 50. Les états cachés des GRU sont de dimension 100 pour toutes les couches du modèle.

Régularisation. En suivant [16], nous fixons les coefficients λ à 10^{-6} pour les matrices de paramètres des GRU et 10^{-5} pour les autres matrices de paramètres. Le coefficient de régularisation successif δ est fixé à 10^{-2} pour toutes les couches. Nous utilisons aussi Dropout [29] sur toutes les couches, avec taux de neurones affectés de 0.2.

Données d'entraînement pour les couches EMS et chunking. Nous utilisons le corpus issu de la tâche partagée CoNLL-2000 [30] avec les étiquettes associées pour entraîner les couches EMS et chunking.

Données d'entraînement pour les couches DelCA et DetCA. Nous utilisons le corpus Argument Annotated Essays (version 2) partagé par Stab et Gurevych [2] en suivant le découpage entraînement/test fourni pour l'entraînement des couches DelCA et DetCA.

Arrêt de l'entraînement. Dans un cas d'entraînement uni-tâche, une pratique généralement adoptée est d'arrêter l'entraînement du modèle peu avant le surapprentissage. Dans le cas de notre modèle, il n'est pas évident de déterminer le meilleur moment pour arrêter l'entraînement, puisque le modèle peut surapprendre sur une tâche particulière, mais pas sur les autres. Ainsi, nous arrêtons l'entraînement du modèle lorsqu'il surapprend sur les couches DelCA et DetCA, et reportons les meilleurs résultats obtenus pour chaque tâche avant le surapprentissage de celle-ci.

DetCa simple. Nous nommons DetCa simple la tâche DetCa avec la modification suivante : tous les segments des dissertations correspondant à des composants argumentatifs sont traités comme ne comportant qu'un unique mot spécial <VIDE>. L'hypothèse est que cette transformation

Tâche	w/o EMS & chunking	w/ EMS & chunking
DelCA	0.5934	0.8688
DetCA	0.7464	0.7105
DetCA simple	0.7529	0.7911

TABLE 1 – Macros f1-scores obtenus sur les différentes tâches.

Tâche	F1-score obtenus en [2]	F1-score humain
DelCA	0.867	0.886
DetCA	0.826	0.868

TABLE 2 – F1-scores obtenus sur les tâches DelCA et DetCA par Stab et Gurevych [2] et des agents humains.

forcera le modèle à se concentrer sur le contexte entourant les composants argumentatifs, et l'empêchera donc de se surentraîner en considérant les mots à l'intérieur des composants.

5.2 Résultats obtenus

Les résultats obtenus sur les données de test pour les tâches DelCA, DetCA et DetCA simple sont présentés en Table 1. La colonne "w/o EMS & chunking" fait référence à la version du modèle pour laquelle l'optimisation des couches EMS et chunking a été omise. La colonne "w/ EMS & chunking" fait référence à la version du modèle pour laquelle l'optimisation des couches EMS et chunking a été réalisée. Nous prenons comme référence les performances atteintes par des agents humains² ainsi que les résultats présentés par Stab et Gurevych [2], illustrés en Table 2.

Evaluation générale des performances. Nous obtenons un macro f1-score de 0.8688 sur DelCA avec la version "w/ EMS & chunking". Ces résultats sont obtenus sans définition de caractéristiques manuelles et sont comparables à ceux enregistrés en [2] ; ils atteignent 98,06% de la performance humaine. Concernant la classification des composants argumentatifs, nous obtenons un macro f1-score de 0.7911 avec DetCA simple pour la version "w/ EMS & chunking", ce qui représente 95,8% des performances obtenues en [2] et 91,1% de la performance humaine.

Pertinence de DetCA simple. Selon nous, les mots formant un composant argumentatif ne sont pas réellement caractéristiques de sa classe, et en se concentrant dessus, le modèle peut être amené à modéliser du bruit l'empêchant de généraliser correctement. En revanche, le contexte dans lequel apparaissent les composants semble très important. Par exemple, des mots tels que "we can receive the same conclusion that" semblent indiquer que l'auteur va annoncer une conclusion intermédiaire ou majeure. Cela peut expliquer la différence de performances entre DetCA et DetCA simple, notamment pour la version "w/ EMS

2. La performance humaine correspond à la moyenne des résultats obtenus par des annotateurs humains, tels que présentés en [2]

& chunking", avec respectivement un f1-score de 0.7105 contre 0.7911, soit une amélioration de 11,3%.

Pertinence du modèle multi-tâches. Les macro f1-scores sur les tâches DelCA et DetCA simple sont respectivement de 0.5934 et 0.7529 pour la version "w/ EMS & chunking" et de 0.8688 et 0.7911, soit des améliorations de 46,4% et 5,1%. Ces résultats permettent donc de valider l'intérêt d'entraîner un modèle multi-tâches et incitent à l'ajout de tâches auxiliaires supplémentaires.

6 Travaux à venir et perspectives

Les résultats obtenus sont encourageants et pourraient sûrement être améliorés, notamment avec une recherche plus approfondie d'hyper-paramètres optimaux. La différence de performances entre les versions du modèle "w/ EMS & chunking" et "w/o EMS & chunking" portent à croire qu'implémenter davantage de tâches auxiliaires pourrait être bénéfique. Une piste serait d'introduire une couche modélisant un arbre de dépendances syntaxiques en complément de la couche chunking, comme effectué en [16].

En vue d'implémenter un système complet d'argument mining tel que présenté par Stab et Gurevych [2], nous prévoyons d'implémenter des couches permettant la génération automatique de graphes d'arguments. A cette fin il est nécessaire de déterminer s'il existe un arc entre chaque paire ordonnée de composants argumentatifs, ainsi que d'inférer l'étiquette portée par ledit arc.

7 Conclusion

Cet article a présenté une méthode d'extraction et d'analyse automatique d'arguments à partir de textes bruts. L'utilisation de techniques d'apprentissage profond nous permet de nous affranchir de la définition de caractéristiques manuellement définies. Par ailleurs, l'amélioration des performances de notre système par l'exploitation de paramètres optimisés sur des tâches auxiliaires met en avant l'intérêt de l'utilisation d'un modèle multi-tâches. Nous avons comme perspective la complétion de la chaîne de traitement existante en vue d'obtenir un système capable de synthétiser une dissertation par modélisation automatique d'un graphe d'arguments.

Références

- [1] M. Lippi et P. Torroni, Argumentation mining : State of the art and emerging trends, *ACM Transactions on Internet Technology (TOIT)*, 16(2), p.10, 2016.
- [2] C. Stab et I. Gurevych, Parsing argumentation structures in persuasive essays, *Computational Linguistics*, 43(3), pp.619-659, 2017.
- [3] N. Madnani, M. Heilman, J. Tetreault et M. Chodorow, Identifying high-level organizational elements in argumentative discourse, *Proceedings of the 2012 Conference of the North American Chapter of the Association for Computational Linguistics : Human Language Technologies*, pp. 20-28, Association for Computational Linguistics, 2012.
- [4] R. Levy, Y. Bilu, D. Hershcovich, E. Aharoni et N. Sionim, Context dependent claim detection, *Proceedings of COLING 2014, the 25th International Conference on Computational Linguistics : Technical Papers*, pp. 1489-1500, 2014.
- [5] Y. Ajjour, W.F. Chen, J. Kiesel, H. Wachsmuth et B. Stein, Unit Segmentation of Argumentative Texts, *Proceedings of the 4th Workshop on Argument Mining*, pp. 118-128, 2017.
- [6] T. Goudas, C. Louizos, G. Petasis et V. Karkaletsis, Argument extraction from news, blogs, and social media, *Hellenic Conference on Artificial Intelligence*, pp. 287-299, Springer, Cham, 2014.
- [7] C. Sardianos, I.M. Katakis, G. Petasis et V. Karkaletsis, Argument extraction from news, *Proceedings of the 2nd Workshop on Argumentation Mining*, pp. 56-66, 2015.
- [8] S. Eger, J. Daxenberger et I. Gurevych, Neural End-to-End Learning for Computational Argumentation Mining, *arXiv preprint arXiv :1704.06104*, 2017.
- [9] J. Eckle-Kohler, R. Kluge et I. Gurevych, On the role of discourse markers for discriminating claims and premises in argumentative discourse, *Proceedings of the 2015 Conference on Empirical Methods in Natural Language Processing*, pp. 2236-2242, 2015.
- [10] J. Park et C. Cardie, Identifying appropriate support for propositions in online user comments. *Proceedings of the 1st Workshop on Argumentation Mining*, pp. 29-38, 2014.
- [11] I. Persing et V. Ng, End-to-End Argumentation Mining in Student Essays, *Proceedings of the 2016 Conference of the North American Chapter of the Association for Computational Linguistics : Human Language Technologies*, Association for Computational Linguistics, pages 1384-1394, 2016.
- [12] P. Potash, A. Romanov et A. Rumshisky, Here's My Point : Joint Pointer Architecture for Argument Mining, *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing*, pp. 1364-1373, 2017.
- [13] S. Ruder, An overview of multi-task learning in deep neural networks, *CoRR*, abs/1706.05098, 2017.
- [14] A. Søgaard et Y. Goldberg, Deep multi-task learning with low level tasks supervised at lower layers, *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 2 : Short Papers)*, Vol. 2, pp.231-235, 2016.
- [15] Z. Yang, R. Salakhutdinov et W. Cohen, Multi-task cross-lingual sequence tagging from scratch, *arXiv preprint arXiv :1603.06270*, 2016.

- [16] K. Hashimoto, C. Xiong, Y. Tsuruoka et R. Socher, A joint many-task model : Growing a neural network for multiple nlp tasks, *Empirical Methods in Natural Language Processing (EMNLP)*, 2017.
- [17] L. Mou, Z. Meng, R. Yan, G. Li, Y. Xu, L. Zhang et Z. Jin, How transferable are neural networks in nlp applications ?, *Empirical Methods in Natural Language Processing (EMNLP)*, pp. 479–489, 2016.
- [18] H.M Alonso et B. Plank, When is multitask learning effective ? Semantic sequence prediction under varying data conditions, *15th Conference of the European Chapter of the Association for Computational Linguistics*, 2017.
- [19] J. Bingel et A. Søgaard, Identifying beneficial task relations for multi-task learning in deep neural networks, *arXiv preprint arXiv :1702.08303*, 2017.
- [20] J. Pennington, R. Socher et C. Manning, Glove : Global vectors for word representation. *Proceedings of the 2014 conference on empirical methods in natural language processing (EMNLP)*, pp. 1532-1543, 2014.
- [21] K. Cho, B. Van Merriënboer, C. Gulcehre, D. Bahdanau, F. Bougares, H. Schwenk et Y. Bengio, Learning phrase representations using RNN encoder-decoder for statistical machine translation, *arXiv preprint arXiv :1406.1078*, 2014.
- [22] V. Nair et G.E. Hinton, Rectified linear units improve restricted boltzmann machines, *Proceedings of the 27th international conference on machine learning (ICML-10)*, pp. 807-814, 2010.
- [23] L.A. Ramshaw et M.P. Marcus, Text chunking using transformation-based learning, *Natural language processing using very large corpora*, pp. 157-176, Springer, Dordrecht, 1999.
- [24] D. Bahdanau, K. Cho et Y. Bengio, Neural machine translation by jointly learning to align and translate, *ICLR*, 2015.
- [25] D.P. Kingma et J. Ba, Adam : A method for stochastic optimization, *ICLR*, 2015.
- [26] R. Pascanu, T. Mikolov et Y. Bengio, On the difficulty of training recurrent neural networks, *Proceedings of The 30th International Conference on Machine Learning*, pp. 1310–1318, 2013.
- [27] A.M. Saxe, J.L. McClelland, S. Ganguli, Exact solutions to the nonlinear dynamics of learning in deep linear neural networks, *Proceedings of the International Conference on Learning Representations (ICLR)*, 2014.
- [28] K. He, X. Zhang, S. Ren et J. Sun, Delving Deep into Rectifiers : Surpassing Human-Level Performance on ImageNet Classification, *ICCV*, 2015.
- [29] N. Srivastava, G.E. Hinton, A. Krizhevsky, I. Sutskever et R. Salakhutdinov, Dropout : a simple way to prevent neural networks from overfitting, *Journal of machine learning research*, 15(1) :1929–1958, 2014.
- [30] E.F.T.K. Sang, S. Buchholz, Introduction to the CoNLL-2000 shared task : chunking, *Proceedings of the 2nd Workshop on Learning Language in Logic and the 4th Conference on Computational Natural Language Learning*, Lisbon, Portugal, vol. 7, 2000, pp. 127–132, 2000.

Découverte de cardinalité maximale contextuelle dans les bases de connaissances

E. A. Sidi Aly^{1,2} M. L. Diakité¹ A. Giacometti² B. Markhoff² A. Soulet²

¹ Département Mathématiques et Informatique - Université de Nouakchott Al Aasriya (Mauritanie)

² Laboratoire d'Informatique Fondamentale et Appliquée de Tours - Université de Tours (France)

arbi2fr@yahoo.fr, diakite@una.mr, prenom.nom@univ-tours.fr

3 juillet 2018

Résumé

Les bases de connaissances du web sémantique doivent être enrichies par des informations utiles aux applications de fouille, de recherche d'information, de question-réponse, etc. En effet, leur génération à partir de plateformes collaboratives ou d'intégration de sources diverses produit des manques d'information, d'une part, et des erreurs ou incohérences d'autre part. Heureusement, leur volume important permet d'en induire des contraintes vraisemblables. Tel est l'objet de l'algorithme présenté dans cet article, qui extrait des règles de cardinalité maximale à partir d'une base de connaissances. L'enrichissement de la base par ces nouveaux axiomes permet d'y trouver plus de faits, positifs ou négatifs, ce qui rend plus pertinentes les évaluations de la qualité des règles générées par des algorithmes de fouille classiques. Les expérimentations menées sur une partie de DBpedia et sur l'ensemble d'une base de connaissances numismatiques démontrent la faisabilité de l'approche et la pertinence des contraintes extraites.

Mots Clef

Découverte de cardinalité, base de connaissances.

Abstract

The big semantic web knowledge bases have to be enriched for applications in data mining, information retrieval, question answering, etc. Indeed, their generation from collaborative platforms or integration of various sources leads to lack of information on the one hand, and inconsistencies on the other hand. Fortunately, their volume makes it possible to induce probable constraints. This is the aim of the algorithm presented in this article, which extracts maximum cardinality rules from a knowledge base. Adding these new axioms to the knowledge base allows applications to find more facts, positive or negative, which makes more relevant the evaluations of the quality of the rules generated by traditional datamining algorithms. Experiments conducted on part of DBpedia and on an entire numismatic knowledge base demonstrate the feasibility of the approach and the relevance of the discovered contextual constraints.

Keywords

Cardinality discovery, knowledge base.

1 Introduction

Nous considérons de grandes bases de connaissances du web, construites par des algorithmes de recherche d'information à partir de plateformes collaboratives (e.g., DBpedia [2]) et/ou d'intégration de sources diverses. Pour en désigner les éléments, nous utilisons les termes *concept*, *rôle* et *individu* au sens des logiques de description.

Contexte et motivations En représentation des connaissances les restrictions numériques précisant le nombre d'occurrences d'un rôle sont particulièrement utiles [3]. Parmi elles, les contraintes de cardinalité maximale permettent de savoir quand toutes les assertions sur un rôle donné pour un individu donné existent dans la base. C'est utile pour qualifier les réponses aux requêtes sur une base de connaissances, c'est-à-dire les compléter par des informations précises sur leur qualité en terme de *rappel* par rapport à une réalité [13, 17].

Il est illusoire d'espérer des ajouts manuels de telles contraintes d'intégrité dans de grandes bases de connaissances¹, qui soient correctes et suffisantes. Aussi, des techniques de type *rétro-ingénierie* [14] applicables sur ces grandes bases doivent être considérées, afin de les rechercher systématiquement. Des propositions existent déjà pour trouver des contraintes de clés [1, 11, 15, 16] dans des données RDF. Mais à notre connaissance, il n'y a pas encore de travaux sur l'extraction de contraintes de cardinalité maximale dans les bases de connaissances.

Challenge L'extraction de contraintes de cardinalité à partir des données existantes est connue comme un problème important de la *rétro-ingénierie* des bases de données relationnelles [14, 18]. Par rapport au cadre des bases de données traditionnelles, ce problème est bien plus complexe pour les bases de connaissances du web. Tout d'abord, ces bases de connaissances contiennent généralement des **données incohérentes**, que ce soient

1. [5] présente néanmoins un outil pour le faire sur Wikidata.

des assertions fausses ou des assertions dupliquées. De ce fait, la cardinalité maximale observée pour un rôle donné ne saurait être considérée comme sa cardinalité maximale la plus vraisemblable. Par exemple, il est vraisemblable qu’une personne ait au plus une année de naissance et deux parents. Pourtant dans DBpedia (voir les rôles `dbo:birthYear` et `dbo:parent` dans la table 1), certaines personnes ont 5 années de naissance ou 6 parents ! Ces quelques assertions incohérentes ne doivent pas influencer la caractérisation des cardinalités maximales.

Ensuite, ces bases de connaissances sont souvent **incomplètes pour un rôle donné**. Pour cette raison, la cardinalité la plus observée n’est pas forcément la cardinalité maximale. Typiquement, la plupart des personnes décrites dans DBpedia n’ont qu’un seul parent renseigné (voir le rôle `dbo:parent` dans la table 1). Toutefois, certaines en ont plus et ceci n’est pas une anomalie, il faut en tenir compte : la cardinalité maximale du rôle `dbo:parent` pour une personne ne doit pas être sous-estimée (ici à 1) au vu de l’ensemble des cardinalités observées.

Enfin, des travaux récents sur la détection de contraintes de clefs dans les bases de connaissances [16] ont montré que de nombreuses contraintes intéressantes ne sont **valides que sur une partie** d’une base de connaissances. Par exemple, s’il semble difficile de déterminer une cardinalité maximale pour le nombre de nationalités d’une personne en général, comme certains états limitent le nombre de nationalités à 1 il est possible de détecter cette limite pour les ressortissants de tels états. Il est donc essentiel non seulement de détecter des cardinalités maximales, mais aussi d’identifier *les contextes* dans lesquels de telles contraintes peuvent être détectées.

Contributions Etant donnée une distribution de cardinalités $(n_i)_{i \geq 1}$ observées dans une base de connaissances $\mathcal{K}$ pour un rôle R dans un contexte C , nous commençons par proposer **une méthode de calcul d’une cardinalité maximale vraisemblable**, en calculant une estimation du taux de cohérence *réel* que la cardinalité i soit maximale. Cette estimation, notée τ_i , est calculée en prenant en compte tous les individus pour lesquels le rôle R est complet. Son calcul est détaillé et justifié dans la section 4.2. Pour être statistiquement valide, une version corrigée de cette estimation du taux de cohérence, notée $\tilde{\tau}_i$, est également introduite. Des exemples d’estimations de taux de cohérence, corrigés ou non, sont représentés dans la table 1 pour les rôles `dbo:birthYear`, `dbo:parent` et `dbo:nationality` en considérant le concept `dbo:Person` comme contexte.

Etant donnée une arborescence de concepts constituant les contextes candidats, nous proposons ensuite **un algorithme d’exploration systématique d’un ensemble de contraintes contextuelles** pour les rôles desquels nous recherchons les cardinalités maximales. Cet algorithme, décrit en section 4.3, vise à limiter les calculs en élaguant un maximum des contraintes possibles.

Enfin nous présentons et analysons des résultats expérimentaux obtenus sur une base de connaissances

résultant d’un processus d’intégration de 5 bases de données numismatiques [6].

dbo:Person / dbo:birthYear			
i	n_i	τ_i	$\tilde{\tau}_i$
5	1	1.0	0.0
4	2	0.667	0.0
3	4	0.571	0.0
2	91	0.928	0.775
1	159841	0.999	0.996
dbo:Person / dbo:parent			
i	n_i	τ_i	$\tilde{\tau}_i$
6	1	1.0	0.0
4	9	0.9	0.420
3	75	0.882	0.718
2	9392	0.991	0.975
1	10643	0.529	0.518
dbo:Person / dbo:nationality			
i	n_i	τ_i	$\tilde{\tau}_i$
8	2	1,000	0,000
6	1	0,333	0,000
5	1	0,250	0,000
4	13	0,765	0,397
3	167	0,908	0,796
2	3 263	0,947	0,921
1	123 386	0,973	0,969

TABLE 1 – Distributions de cardinalités de rôles de personnes dans DBpedia (i est la cardinalité ; n_i le nombre d’individus étant i fois sujets du rôle considéré ; τ_i est une estimation fréquentielle du taux de cohérence réel ; $\tilde{\tau}_i$ en est une version corrigée s’appuyant sur la borne de Hoeffding)

2 Etat de l’art

Notre algorithme vise à augmenter la connaissance sur les données contenues dans les grandes bases de connaissances du web, en termes de validité comme en termes de complétude par rapport à la réalité représentée. Il permet d’enrichir la partie schéma (TBox en logiques de description) de ces bases pour mieux utiliser leur partie données (ABox). Plusieurs travaux récents vont dans ce sens [1, 11, 15, 16, 10] et d’autres s’en rapprochent [7, 13, 17] mais ciblent des individus (assertions de la ABox) plutôt que des concepts (assertions de la TBox).

Dans [17], une technique de fouille de textes de Wikipedia pour ajouter des précisions sur le degré de complétude des informations dans Wikidata est décrite. Notre proposition est complémentaire puisque notre algorithme traite les données déjà contenues dans les bases de connaissances. Mais surtout, il ne caractérise pas les rôles par rapport à des *individus* précis mais à des *concepts* définis (au sens des logiques de description). Les auteurs de [7, 13] présentent également des propositions pour déterminer quand est-ce qu’un rôle particulier (comme `dbo:parent`) manque pour un individu particulier (comme *Obama*). Plus générale, notre proposition consiste à calculer les cardinalités maximales vraisemblables des rôles relativement à des concepts

définissant des contextes : elle enrichit donc la partie schéma.

Ce sont des clés au sens des bases de données, donc des contraintes au niveau du schéma, qui sont recherchées dans [1, 11, 15, 16]. L'idée est de trouver des axiomes indiquant que tout individu d'un certain concept doit posséder une valeur unique pour un rôle donné R . Cela constitue donc une cardinalité maximale du rôle R pour le concept C . Egalement très proches de nos travaux, dans [10], les auteurs proposent de déterminer automatiquement quels rôles devraient être obligatoirement renseignés pour un concept donné de la base de connaissances. Pour cela ils comparent la densité du rôle pour les individus de ce concept par rapport à sa densité pour les individus d'autres concepts, qui lui sont liés dans la hiérarchie des concepts. Notre proposition s'appuie sur d'autres critères pour calculer la cardinalité maximale du rôle pour un contexte (notion plus générale que seulement les concepts de la base). Elle peut être adaptée au calcul de la cardinalité minimale, auquel cas elle trouverait, entre autres, quels rôles ont une cardinalité minimale au moins supérieure à 1 pour un concept donné, soit plus d'information que seulement savoir si le rôle devrait exister ou pas.

Ces différentes sortes d'information supplémentaire sur la qualité des données de la base de connaissances, en termes de validité comme en termes de complétude par rapport à la réalité représentée, permettent d'améliorer le fonctionnement des applications qui les utilisent, en réduisant le flou de l'hypothèse du monde ouvert. Ainsi pour améliorer la mesure de qualité de règles issues de processus de fouille dans les bases de connaissances du web sémantique, une *hypothèse de complétude partielle* est définie et utilisée dans [8, 7] : cette règle stipule que si un rôle est renseigné pour un individu, alors les informations concernant ce rôle pour cet individu sont considérées complètes. Si on peut noter que cette hypothèse est contredite par l'observation de DBpedia (voir l'extrait fourni dans la table 1), elle rend tout de même plus précis le calcul de la confiance associée aux résultats de fouille. Ces auteurs ont démontré le besoin pour la fouille de ce qu'ils appellent des *oracles de complétude*, et proposé un certain nombre d'heuristiques pour en définir, comme par exemple la popularité des individus (qui augmente les chances que les faits renseignés sur eux soient complets), etc.

La fouille de données est loin d'être le seul domaine qui bénéficie d'axiomes tels que ceux découverts par notre algorithme, par exemple, s'appuyant sur des travaux de référence en base de données, les auteurs de [4, 12] et plus récemment [9] proposent de caractériser les réponses obtenues par des requêtes, en fonction des informations connues concernant le degré de complétude de la base de connaissances interrogée, par rapport à la réalité représentée.

3 Préliminaires

3.1 Bases de connaissances

Dans ce papier, nous considérons des *bases de connaissances* $\mathcal{K} = (\mathcal{T}, \mathcal{A})$ où $\mathcal{T}$ et $\mathcal{A}$ sont respectivement les TBox et ABox de $\mathcal{K}$. $\mathcal{T}$ désigne un ensemble d'axiomes terminologiques définis à partir des concepts et rôles atomiques de $\mathcal{K}$, alors que $\mathcal{A}$ désigne l'ensemble des assertions ou faits de $\mathcal{K}$. Plus précisément, $\mathcal{A}$ contient des expressions de la forme $C(a)$ et $R(a, b)$ où C est un *concept*, R est un *rôle*, et a, b sont des *individus*.

Dans le cas de la base de connaissances DBpedia, `dbo:Country` et `dbo:Person` sont des exemples de concepts atomiques et `dbo:nationality` est un exemple de rôle atomique de sa TBox. Par ailleurs, `dbo:Country(Mauritania)` et `dbo:nationality(Arby, Mauritania)` sont des exemples de faits ou assertions de sa ABox. Le premier indique que *Mauritania* est un pays, alors que le second indique que *Arby* est de nationalité mauritanienne.

Les logiques de description permettent de définir des axiomes pour enrichir la TBox d'une base de connaissances. Par exemple, la relation d'inclusion $\sqsubseteq$ permet d'indiquer qu'un concept C_1 est inclus dans un concept C_2 , noté $C_1 \sqsubseteq C_2$. Plus précisément, une base de connaissances $\mathcal{K}$ implique l'axiome $C_1 \sqsubseteq C_2$ si pour toute interprétation $\mathcal{I}$ de $\mathcal{K}$, $C_1^{\mathcal{I}} \subseteq C_2^{\mathcal{I}}$. Par exemple, les axiomes $\exists \text{dbo:nationality}.\top \sqsubseteq \text{dbo:Person}$ et $\exists \text{dbo:nationality}^{\neg}.\top \sqsubseteq \text{dbo:Country}$ indiquent respectivement que le domaine du rôle `dbo:nationality` est inclus dans le concept `dbo:Person`, et que le co-domaine du rôle `dbo:nationality` est inclus dans le concept `dbo:Country`.

3.2 Contraintes contextuelles de cardinalité maximale

Soit R un rôle d'une base de connaissances $\mathcal{K} = (\mathcal{T}, \mathcal{A})$. On considère généralement que ce rôle satisfait dans $\mathcal{K}$ une contrainte de cardinalité maximale M si pour tout sujet s , le nombre d'objets o tels que $R(s, o)$ soit présent dans $\mathcal{K}$ (directement présent dans sa ABox $\mathcal{A}$ ou inférable à partir de ses TBox $\mathcal{T}$ et ABox $\mathcal{A}$) est inférieur ou égal à M .

En logique de description, une telle contrainte peut se représenter par un axiome de la forme *sqssubseteq* en utilisant le constructeur de restriction de cardinalité ($\leq MR$). En effet, en terme logique, une base de connaissances $\mathcal{K}$ implique l'axiome $\exists R.\top \sqsubseteq (\leq MR)$ si pour toute interprétation $\mathcal{I}$ de $\mathcal{K}$, $\{x : (\exists y)((x, y) \in R^{\mathcal{I}})\} \subseteq \{x : \#\{y : (x, y) \in R^{\mathcal{I}}\} \leq M\}$ où $\#E$ représente la cardinalité d'un ensemble E .

Plus précisément, dans ce papier, nous cherchons à identifier des contraintes *contextuelles* de cardinalité maximale, à savoir des contraintes qui ne sont pas nécessairement vérifiées par tous les sujets s d'un rôle R , mais par tous les sujets instances d'un concept, qu'il soit atomique ou composé, déjà défini dans $\mathcal{K}$ ou pas. Cette notion est introduite

formellement dans la définition suivante :

Définition 3.1 (Contrainte contextuelle). *Etant donné un rôle R , un concept atomique ou défini C et un entier M , une contrainte contextuelle de cardinalité maximale définie sur R est une expression γ de la forme : $C \sqsubseteq (\leq MR)$.*

Le concept C est appelé le contexte de la contrainte γ . La contrainte γ est satisfaite dans une base de connaissances $\mathcal{K}$ si et seulement si pour toute interprétation $\mathcal{I}$ de $\mathcal{K}$, $C^{\mathcal{I}} \subseteq \{x : \#\{y : (x, y) \in R^{\mathcal{I}}\} \leq M\}$.

Par exemple, la contrainte contextuelle $(\text{dbo:Person}) \sqsubseteq (\leq 5 \text{dbo:nationality})$ indique que toutes les personnes ont au plus 5 nationalités, alors que la contrainte contextuelle $(\text{dbo:Person} \sqcap \exists \text{dbo:nationality}.\{China\}) \sqsubseteq (\leq 1 \text{dbo:nationality})$ indique que toutes les personnes de nationalité chinoise ont au plus une nationalité.

Dans ce travail, on cherche à extraire des contraintes contextuelles de cardinalité maximale qui soient les plus générales possibles.

Définition 3.2 (Contrainte contextuelle minimale). *Soient deux contraintes contextuelles de cardinalité maximale $\gamma_1 : C_1 \sqsubseteq (\leq M_1 R)$ et $\gamma_2 : C_2 \sqsubseteq (\leq M_2 R)$ définies sur R . La contrainte γ_1 est dite plus générale que la contrainte γ_2 si $C_2 \sqsubseteq C_1$ et $M_1 \leq M_2$. Etant donné un ensemble de contraintes Γ définies sur R , une contrainte $\gamma_1 \in \Gamma$ est dite minimale dans Γ s'il n'existe aucune contrainte γ_2 dans Γ plus générale que γ_1 .*

Par exemple, la contrainte contextuelle $(\text{dbo:Person}) \sqsubseteq (\leq 2 \text{dbo:nationality})$ est plus générale que la contrainte contextuelle $(\text{dbo:Person} \sqcap \exists \text{dbo:nationality}.\{China\}) \sqsubseteq (\leq 5 \text{dbo:nationality})$ car $(\text{dbo:Person} \sqcap \exists \text{dbo:nationality}.\{China\}) \sqsubseteq \text{dbo:Person}$ et $2 \leq 5$.

La notion de minimalité a pour objectif de ne pas extraire de contraintes contextuelles qui soient redondantes. Intuitivement, considérons les deux contraintes γ_1 et γ_2 introduites dans la définition précédente, et supposons que γ_1 soit plus générale que γ_2 . Etant donnée une base de connaissances $\mathcal{K}$ dans laquelle les contraintes γ_1 et γ_2 sont satisfaites, soit une instance s de C_2 dans $\mathcal{K}$. D'après γ_2 , nous savons que pour toute interprétation $\mathcal{I}$ de $\mathcal{K}$, $\#\{o : (s, o) \in R^{\mathcal{I}}\} \leq M_2$. Mais comme γ_1 est plus générale que γ_2 , nous savons par définition que $C_2 \sqsubseteq C_1$. Il en découle que s est aussi une instance de C_1 dans $\mathcal{K}$, et d'après γ_1 , que pour tout interprétation $\mathcal{I}$ de $\mathcal{K}$, $\#\{o : (s, o) \in R^{\mathcal{I}}\} \leq M_1$, ce qui est une contrainte plus forte que $\#\{o : (s, o) \in R^{\mathcal{I}}\} \leq M_2$. En effet, par définition de la minimalité, nous savons que $M_1 \leq M_2$. Par rapport à la contrainte γ_1 , la contrainte γ_2 est donc inutile car redondante, i.e. elle ne permet pas de déduire d'information supplémentaire.

Le problème traité dans ce papier est alors le suivant : **étant donné une base de connaissances $\mathcal{K}$, un rôle R et une hiérarchie de concepts $(\mathcal{C}, \sqsubseteq)$, nous cherchons à découvrir l'ensemble des contraintes contextuelles de cardinalité maximale de la forme $C \sqsubseteq (\leq MR)$ avec $C \in \mathcal{C}$, qui soient satisfaites sur $\mathcal{K}$ et minimales dans $\mathcal{C}$.**

En pratique, une base de connaissances telle que DBpedia est très incomplète (par exemple, de nombreuses personnes ont seulement un parent), et elle comporte de nombreuses incohérences (par exemple, des personnes peuvent avoir jusqu'à 5 parents). Pour ces raisons, étant donnée une base de connaissances $\mathcal{K}$, il n'est pas pertinent de chercher à extraire des contraintes de cardinalité qui soient *parfaitement* satisfaites dans $\mathcal{K}$, mais les contraintes :

- *les plus probables et suffisamment probables* par rapport à un seuil donné, de manière à prendre en compte et tolérer les incohérences, et
- *suffisamment certaines* par rapport à un degré de confiance, pour ne pas extraire des contraintes qui soient remises en cause régulièrement par l'ajout de nouveaux faits dans la base de connaissances.

Nous détaillons dans la section suivante comment évaluer la probabilité qu'une contrainte soit satisfaite dans une base de connaissances $\mathcal{K}$ et comment mesurer la certitude que cette contrainte soit réelle.

4 Extraction de contraintes contextuelles de cardinalité maximale

Pour résoudre le problème énoncé précédemment, nous commençons par le reformuler en introduisant la notion de taux de cohérence dans la section 4.1, puis nous décrivons dans la section 4.2 comment détecter une cardinalité maximale pour un rôle R dans un contexte C . Ensuite, étant donné un ensemble de contextes candidats $\mathcal{C}$, nous montrons dans la section 4.3 comment explorer efficacement l'ensemble des contraintes contextuelles possibles.

4.1 Taux de cohérence

Etant donnée une base de connaissances $\mathcal{K} = (\mathcal{T}, \mathcal{A})$, supposons que i soit la cardinalité maximale du rôle R dans le contexte C . Soit s un individu de C dans $\mathcal{K}$, complet pour le rôle R dans $\mathcal{K}$ (dans le sens où tous les faits $R(s, o)$ possibles représentant le monde réel sont dans $\mathcal{A}$ ou inférables). Dans le cas où il existe exactement i faits dans $\mathcal{K}$ de la forme $R(s, o)$, cela renforce l'hypothèse que i soit la cardinalité maximale de R dans le contexte C . Inversement, s'il existe plus de i faits dans $\mathcal{K}$ de la forme $R(s, o)$, cela affaiblit l'hypothèse que i soit la cardinalité maximale de R dans le contexte C . Ainsi dans le tableau 1, pour la classe `dbo:Person`, les individus comportant au moins 3 assertions pour le rôle `dbo:parent` affaiblissent l'hypothèse que la cardinalité maximale soit 2 mais ils restent peu nombreux au regard des 9 392 individus qui ont exactement 2 parents.

En suivant ce raisonnement, nous introduisons la notion de

taux de cohérence pour évaluer si une cardinalité i pour le rôle R dans le contexte C a des chances d'être maximale :

Définition 4.1 (Taux de cohérence). *Etant donnée une base de connaissances $\mathcal{K}$, le taux de cohérence de la cardinalité i pour le rôle R dans le contexte C est le ratio :*

$$\tau_i^{C,R}(\mathcal{K}) = \frac{n_i^{C,R}}{n_{\geq i}^{C,R}}$$

où $n_i^{C,R}$ (resp. $n_{\geq i}^{C,R}$) représente le nombre de sujets s tels que i faits $R(s, o)$ (resp. i faits ou plus) appartiennent à $\mathcal{K}$ dans le contexte C .

Par exemple, dans le tableau 1, $n_{\geq 2}^{\text{dbo:Person}, \text{dbo:parent}}$ est égal à 9477 ($9477 = 9392 + 75 + 9 + 1$). De cette manière, le taux de cohérence $\tau_2^{\text{dbo:Person}, \text{dbo:parent}}(\mathcal{K})$ est de 0,991 (i.e., $9392/9477$). Par la suite, quand le contexte et la relation sont clairs, nous pouvons les omettre dans les notations. Dans ce cas, n_i et τ_i désignent respectivement les termes $n_i^{C,R}$ et $\tau_i^{C,R}$.

Maintenant nous allons formaliser le lien entre le taux de cohérence et la notion de contrainte maximale. Originellement introduit dans [13], $\mathcal{K}^* = (\mathcal{T}^*, \mathcal{A}^*)$ désigne une hypothétique base de connaissances idéale qui contiendrait tous les axiomes et toutes les assertions du monde réel. Comme $\mathcal{K}^*$ est correcte et complète, **le taux de cohérence au sein de $\mathcal{K}^*$, noté $\tau_M^{C,R}(\mathcal{K}^*)$, est égal à 1 si et seulement si $C \sqsubseteq (\leq M R)$ appartient à $\mathcal{T}^*$.**

En pratique, le taux de cohérence mesuré dans une base de connaissances est différent du taux de cohérence réel : $\tau_i(\mathcal{K}) \neq \tau_i(\mathcal{K}^*)$. Par exemple, le taux de cohérence $\tau_2(\mathcal{K})$ pour le rôle `dbo:parent` du tableau 1 est égal à 0,991 alors que le taux de cohérence réel est égal à 1. Plus grave, on a $\tau_6^{\text{dbo:Person}, \text{dbo:parent}}(\mathcal{K}) = 1$! Le taux de cohérence sur $\mathcal{K}$ est donc une estimation peu fiable du taux de cohérence réel sur $\mathcal{K}^*$.

4.2 Détection d'une contrainte

L'estimation $\tau_i(\mathcal{K})$ de $\tau_i(\mathcal{K}^*)$ doit être corrigée pour être statistiquement valide. Pour ce faire, nous proposons d'utiliser l'inégalité de Hoeffding qui a l'avantage d'être vraie pour toute distribution. En terme de probabilité, si X est une variable aléatoire indiquant pour un sujet s tiré aléatoirement, le nombre de faits $R(s, o)$ appartenant à $\mathcal{K}$, alors τ_i est une estimation fréquentielle de la probabilité conditionnelle $P(X = i / X \geq i)$. Etant donné un niveau de confiance $1 - \delta$, l'inégalité de Hoeffding stipule que $\tau_i(\mathcal{K}^*)$ est compris entre $\tau_i(\mathcal{K}) - \epsilon_i$ et $\tau_i(\mathcal{K}) + \epsilon_i$ où $\epsilon_i = \sqrt{\frac{\log(1/\delta)}{2n_{\geq i}}}$. Dans ce contexte, afin de prendre des décisions les plus sûres, nous proposons d'utiliser la borne inférieure de l'intervalle de confiance $[\tau_i - \epsilon_i, \tau_i + \epsilon_i]$. Plus formellement, on a la propriété suivante :

Propriété 4.1 (Minoration). *Etant données une base de connaissances $\mathcal{K}$ et une confiance $1 - \delta$, le taux de*

cohérence réel $\tau_i(\mathcal{K}^)$ de la cardinalité i pour le rôle R dans le contexte C est supérieur à $\tilde{\tau}_i(\mathcal{K})$:*

$$\tau_i(\mathcal{K}^*) \geq \tilde{\tau}_i(\mathcal{K})$$

où $\tilde{\tau}_i(\mathcal{K})$ est le taux de cohérence pessimiste défini par :

$$\tilde{\tau}_i(\mathcal{K}) = \max \left\{ \frac{n_i}{n_{\geq i}} - \sqrt{\frac{\log(1/\delta)}{2n_{\geq i}}}; 0 \right\}$$

Cette propriété nous munit d'un outil efficace pour approximer le taux de cohérence réel. Il survient alors la difficulté de choisir la cardinalité maximale une fois que l'on dispose pour chaque cardinalité i du taux de cohérence pessimiste $\tilde{\tau}_i(\mathcal{K})$ (pour un rôle R dans le contexte C).

Plus précisément, **étant donné un seuil minimal de cohérence $\min_\tau$ et un niveau de confiance $1 - \delta$, nous considérons que M est la cardinalité maximale de R dans le contexte C si et seulement si $\tilde{\tau}_M \geq \min_\tau$ et $M = \arg \max_{i \geq 1} \tilde{\tau}_i(\mathcal{K})$.**

Quelques exemples d'estimations $\tilde{\tau}_i$ et de détection de cardinalités maximales contextuelles sont donnés dans la table 1. Dans les 3 exemples, on a considéré `dbo:Person` comme contexte, et on a cherché à détecter la cardinalité maximale contextuelle de trois rôles : `dbo:birthYear`, `dbo:parent` et `dbo:nationality`. Intuitivement, pour les deux premiers rôles, on souhaiterait détecter des cardinalités maximales respectives de 1 et 2. Pour un niveau de confiance $1 - \delta = 99\%$ et un seuil $\min_\tau = 0.97$, on constate que les cardinalités maximales supposées sont effectivement détectées (cf. lignes en gras dans la table 1). De manière intéressante, avec ces mêmes seuils, aucune cardinalité n'est détectée pour `dbo:nationality`.

4.3 Exploration de l'espace de recherche

Etant donné une base de connaissances $\mathcal{K}$, un rôle R , une arborescence de concepts $(\mathcal{C}, \sqsubseteq)$, un degré de confiance δ et un seuil minimal de cohérence $\min_\tau$, nous cherchons à découvrir l'ensemble des contraintes contextuelles de cardinalité maximale de la forme $C \sqsubseteq (\leq M R)$ avec $C \in \mathcal{C}$, qui soient minimales et suffisamment certaines sur $\mathcal{K}$. En pratique, notons que l'arborescence $(\mathcal{C}, \sqsubseteq)$ peut être une arborescence déjà existante dans la TBox de la base de connaissances, ou une arborescence construite dans une phase préalable de préparation des données (voir la section 5.1).

Dans un tel cadre, il y a potentiellement un très grand nombre de contraintes contextuelles à considérer, évaluer et comparer. Néanmoins, il est possible de réduire la taille de l'espace de recherche à explorer en se basant sur les propriétés 4.2 et 4.3 énoncées ci-après. Tout d'abord, la propriété 4.2 montre qu'une contrainte $C \sqsubseteq (\leq M R)$ ne peut pas être suffisamment certaine si le contexte C contient trop peu d'individus dans $\mathcal{K}$, car alors l'intervalle de confiance du taux de cohérence calculé grâce à l'inégalité de Hoeffding est très large et sa borne inférieure ne peut être supérieure au seuil $\min_\tau$ imposé.

Propriété 4.2 (Nombre minimal d’observations). *Etant donné une base de connaissances $\mathcal{K}$, une contrainte contextuelle de cardinalité maximale $C \sqsubseteq (\leq M R)$ et un seuil $\min_\tau$, le taux de cohérence $\tilde{\tau}_M(\mathcal{K})$ que M soit la cardinalité maximale de R dans C ne peut être supérieur à $\min_\tau$ que si $|C| \geq \frac{\log(1/\delta)}{2(1-\min_\tau)^2}$.*

Par ailleurs, supposons qu’une contrainte γ définie par $C \sqsubseteq (\leq M R)$ avec $M = 1$ ait été détectée comme suffisamment certaine au cours de l’exploration. Alors, d’après la propriété 4.3, il n’est pas nécessaire d’explorer les contraintes γ' définies par $C' \sqsubseteq (\leq M' R)$ où C' est plus spécifique que C . Cette propriété découle directement de la définition 3.2 de la minimalité.

Propriété 4.3 (Contrainte minimale). *Soient une base de connaissances $\mathcal{K}$ et une contrainte contextuelle de cardinalité maximale γ définie par $C \sqsubseteq (\leq M R)$ avec $M = 1$. Toute contrainte γ' définie par $C' \sqsubseteq (\leq M' R)$ avec $C' \sqsubseteq C$ et $M' \geq 1$ ne peut être minimale.*

L’algorithme 1 détaille notre fonction récursive d’exploration, la fonction *C3M* (pour *Contextual Cardinality Constraint Mining*). Cette fonction prend en entrée une base de connaissances, un rôle à explorer, un contexte courant, une cardinalité maximale courante ($M = \infty$ si aucune cardinalité maximale n’a encore pu être détectée), et enfin, des seuils δ et $\min_\tau$. **Le démarrage de l’exploration d’une arborescence de racine $\top$ se fait en exécutant la fonction $C3M(\mathcal{K}, R, \top, \infty, \delta, \min_\tau)$.**

Pour commencer, la fonction *C3M* détermine si le nombre d’individus est suffisant dans le contexte C . Si ce n’est pas le cas, elle arrête l’exploration à la ligne 2 conformément à la propriété 4.2. Sinon, le taux de cohérence $\tilde{\tau}_i$ est calculé pour chaque cardinalité i (lignes 4 à 6) et la ligne 7 retient la cardinalité maximale la plus probable. Si le taux de cohérence correspondant n’est pas supérieur au seuil $\min_\tau$, alors cela signifie qu’aucune cardinalité maximale n’a pu être détectée à ce niveau et i_M est fixé ligne 8 à ∞ . Ensuite, si la cardinalité maximale détectée i_M est strictement inférieure à M (la cardinalité maximale détectée au niveau précédent), alors on dispose d’une nouvelle contrainte minimale de cardinalité maximale i_M et on l’ajoute à Γ , l’ensemble des contraintes recherchées. Finalement, conformément à la propriété 4.3, si i_M est égale à 1, il n’est pas nécessaire de poursuivre l’exploration en parcourant les contextes plus spécifiques de C . Sinon, la fonction *C3M* est appelée récursivement à la ligne 12 pour tous les C' qui sont des sous-concepts directs de C .

Dans notre implémentation de la fonction *C3M*, nous avons appliqué une approche client-serveur où les distributions de cardinalité $n_i^{C,R}$ sont calculées par interrogation en SPARQL d’une base de connaissances localisée sur un serveur. Dans un tel cadre, la complexité de notre méthode en nombre de requêtes sur le serveur est en $\mathcal{O}(|\mathcal{C}|)$ où $|\mathcal{C}|$ représente le nombre de concepts dans l’arborescence $\mathcal{C}$ explorée. Dans le pire des cas, côté client, la complexité en nombre d’opérations est en $\mathcal{O}(|\mathcal{C}| \times i_{max})$ où i_{max}

Algorithm 1 C3M

Input: Une base de connaissances $\mathcal{K}$, un rôle R , un contexte C , un entier M , un niveau de confiance δ et un seuil minimal de support $\min_\tau$

Output: Un ensemble Γ de contraintes contextuelles de cardinalité maximale

```

1:  $\alpha := \frac{\log(1/\delta)}{2(1-\min_\tau)^2}$  et  $n_{\geq 0}^{C,R} := |C|$ 
2: if ( $n_{\geq 0}^{C,R} < \alpha$ ) then return  $\emptyset$ 
3:  $\Gamma := \emptyset$  et  $i_{max} := \arg \max_{i \in \mathbb{N}} \{n_i^{C,R} > 0\}$ 
4: for all  $i \in [1..min\{M, i_{max}\}]$  do
5:    $\tilde{\tau}_i := \max \left\{ \frac{n_i^{C,R}}{n_{\geq i}^{C,R}} - \sqrt{\frac{\log(1/\delta)}{2n_{\geq i}^{C,R}}}; 0 \right\}$ 
6: end for
7:  $i_M := \arg \max_{i \in [1..min\{M, i_{max}\}]} \{\tilde{\tau}_i\}$ 
8: if ( $\tilde{\tau}_{i_M} < \min_\tau$ ) then  $i_M = \infty$ 
9: if ( $i_M < M$ ) then  $\Gamma := \{C \sqsubseteq (\leq i_M R)\}$ 
10: if ( $i_M > 1$ ) then
11:   for all  $C' \in subClassOf(C)$  do
12:      $\Gamma := \Gamma \cup C3M(\mathcal{K}, R, C', i_M, \delta, \min_\tau)$ 
13:   end for
14: end if
15: return  $\Gamma$ 

```

représente l’entier maximal pour lequel il existe au moins un sujet s tel que i_{max} faits $R(s, o)$ appartiennent à la base de connaissances $\mathcal{K}$, i.e. $i_{max} = \arg \max_{i \in \mathbb{N}} \{n_i^{\top, R} > 0\}$.

5 Expérimentations

Outre les requêtes sur DBpedia (dont nous montrons des échantillons de réponses en table 1), qui ont été utilisées pour mettre au point la définition du taux de cohérence, nous avons expérimenté l’algorithme 1 sur un jeu de données mis à notre disposition par les auteurs de [6].

5.1 Données et protocole

Le jeu de données utilisé porte sur le domaine numismatique, il est le résultat d’un processus d’intégration mené dans le cadre du projet européen ARIADNE². Ses auteurs ont utilisé le CIDOC-CRM³ pour intégrer les contenus de 5 ressources construites par des institutions de différents pays européens. Il contient 3 123 998 triplets, dont les définitions de 114 classes et 373 rôles ou propriétés du CIDOC-CRM et d’ARIADNE. Il est stocké et interrogé avec le triplestore Blazegraph (v2.1.4), sur une machine virtuelle sous Linux avec 32 GB de mémoire virtuelle, sur un serveur ayant pour processeur un Dual Intel Xeon E5620 4 coeurs. L’algorithme 1 est implémenté en Java et utilise la bibliothèque de programmation pour RDF Jena⁴. La base porte sur des pièces de monnaies mais, par choix des intégrateurs, il n’existe pas de classe *Coin*. Les individus correspondant à des pièces sont des instances de `E22_Man_Made_Object` caractérisées par certains URIs (ex. `<http://nomisma.org/id/coin>`) comme va-

2. <http://ariadne-infrastructure.eu/>

3. <http://www.cidoc-crm.org/>

4. <http://jena.apache.org>

leur objet de certains rôles (ex. `P2_has_type`). Plusieurs rôles et plusieurs URIs sont utilisés pour cela, aussi nous avons décidé de construire notre propre arborescence d'exploration de la façon suivante :

Au **premier niveau**, notre arborescence contient tous les concepts C_i de la base, soit 114 concepts ($i \in [1..114]$). Tous ces concepts sont des sous-concepts du **concept racine** $\top$ au **niveau zéro**, i.e. pour tout i , nous avons $C_i \sqsubseteq \top$. Au **deuxième niveau** notre arborescence contient tous les concepts C_i^j définis par $C_i^j := C_i \sqcap (\exists R_j. \top)$ où C_i ($i \in [1..114]$) et R_j ($j \in [1..373]$) sont respectivement des concepts et rôles de la base. A ce niveau, 42 522 concepts C_i^j sont ainsi définis. Enfin, au **troisième niveau**, notre arborescence contient tous les concepts $C_i^{j,k}$ définis par $C_i^{j,k} := C_i \sqcap (\exists R_j. \{a_k\})$ où C_i ($i \in [1..114]$) et R_j ($j \in [1..373]$) sont respectivement des classes et rôles de la base, et a_k est un individu du co-domaine de R_j , i.e. $a_k \in (\exists R_j^{-1}. \top)$. Grâce à ce dernier niveau, il est possible de considérer des contextes à la manière de notre exemple jouet où `dbo:Person` $\sqcap$ `dbo:nationality.China`. Notons finalement que pour tout i, j, k , nous avons $C_i^{j,k} \sqsubseteq C_i^j \sqsubseteq C_i$. Globalement, cette arborescence comporte 3 160 357 concepts, donc pour les 373 rôles de la base de connaissances cela représente plus d'un milliard de contraintes contextuelles possibles (exactement 1 178 813 161 contraintes). Néanmoins, comme nous le verrons dans la section suivante, l'utilisation des propriétés 4.2 et 4.3 permet d'élaguer une grande partie de l'espace de recherche.

5.2 Résultats

Tous les résultats présentés dans cette section ont été obtenus avec un **seuil minimal de confiance** $1 - \delta = 0,99\%$ (pour des contraintes les plus certaines possibles) et un **seuil minimal de cohérence** $\min_\tau = 0,95$ (pour des contraintes suffisamment probables). Ce seuil a été défini expérimentalement. Sur des bases de connaissances de plus grande taille comme DBpedia, un seuil plus élevé est préférable. Néanmoins, l'approche est relativement peu sensible aux seuils (i.e., l'ensemble des contraintes trouvées est stable).

Analyse quantitative. Avec ces paramètres, la propriété 4.2 nous indique qu'une contrainte $C \sqsubseteq (\leq M R)$ ne peut être suffisamment certaine si son contexte C contient moins de $\alpha = \frac{\log(1/\delta)}{2(1-\min_\tau)^2} = 922$ instances. Ainsi, l'utilisation de la propriété 4.2 permet de n'explorer que 16 641 contraintes, soit moins de 0,002% des plus de 1 milliard de contraintes possibles. Qui plus est, notre expérience montre que la propriété 4.3 permet de réduire encore de 82,5% la taille de l'espace de recherche à explorer. Au final, avec les seuils choisis notre algorithme ne cherche à détecter une cardinalité maximale que pour 2 909 contextes possibles, avec un temps de calcul complet de moins de 50 minutes.

La table 2 donne une vue globale et quantitative du résultat de l'exploration réalisée. Sur les 2 909 contraintes contex-

M	Niveau dans l'arborescence				Total
	0 $\top$	1 C_i	2 C_i^j	3 $C_i^{j,k}$	
1	60	28	10	222	320
2	3	6	9	90	108
3	0	7	14	92	113
4	1	0	8	20	29
5	1	0	0	16	17
6	0	0	0	8	8
Total	65	41	41	448	595

TABLE 2 – Répartition par niveau et cardinalité maximale M des contraintes minimales détectées

tuelles possibles, notre algorithme a détecté au total 887 contraintes de cardinalité maximale, 595 d'entre elles étant des contraintes minimales. Sur cet exemple, le critère de minimalité permet donc de réduire de près de 67% le nombre de contraintes retournées. On constate que les contraintes les plus nombreuses sont des cardinalités maximales avec $M = 1$, ce qui correspond à des contraintes où pour un rôle donné R , tout sujet s est en relation avec au plus un objet o . Néanmoins de très nombreuses contraintes sont trouvées avec des cardinalités maximales $M \in \{2, 3\}$ (37% des contraintes minimales détectées). On note également que si des contraintes de cardinalités maximales sont détectées dès le niveau 0 (65 contraintes avec un contexte $C \equiv \top$), la recherche de contraintes contextuelles est particulièrement pertinente. Il faut en effet noter que les contraintes les plus nombreuses sont trouvées au niveau 3 (75% des contraintes détectées), sachant que par construction de notre arborescence, c'est à ce niveau que sont caractérisées les pièces de monnaie.

Analyse qualitative. Tout d'abord, dès le niveau 0, notre méthode permet de retrouver des contraintes fonctionnelles attendues, par exemple, pour les 3 rôles du CIDOC-CRM `P1_is_identified_by`, `P52_has_current_owner` et `P50_has_current_keeper`, indiquant que si un sujet décrit dans la base possède plus d'un identifiant, un propriétaire ou un conservateur, alors on peut en déduire que ces identifiants (respectivement, propriétaires et conservateurs) sont identiques. Concernant le rôle `P45_consists_of` du CIDOC-CRM (permettant de décrire les matériaux constitutifs d'un objet), il est intéressant de noter qu'une cardinalité maximale de 2 est détectée dès le niveau 1 pour la classe `E22_Man_Made_Object`. La base de connaissances décrit notamment des médailles constituées d'or et de pierre précieuse (telle l'agate). Pour ce même rôle, une cardinalité maximale de 1 est détectée au niveau 3 pour les pièces de monnaie. Cette information est notamment représentée par la contrainte `E22_Man_Made_Object` $\sqcap$ `EP2_has_type.<http://nomisma.org/id/coin>` $\sqsubseteq (\leq 1 P45_consists_of)$. Cette contrainte est détectée bien qu'à certaines pièces la relation `P45_consists_of` associe deux matériaux ; mais c'est rare (et le plus souvent il

s'agit du même matériau dans deux langues différentes). Un même type de contrainte (avec $M = 1$) est trouvée au niveau 3 pour tous les contextes décrivant des pièces, concernant le rôle `P62_depicts` (ce qui est dépeint sur l'objet). C'est raisonnable car dans le cas d'une pièce de monnaie, on trouve le plus souvent une seule représentation figurative (sur une des deux faces de la pièce), alors qu'une telle contrainte n'est pas valide pour d'autres objets.

Au passage, l'étude de l'ensemble des contraintes extraites par notre méthode a mis en évidence des redondances dans la base, sans doute du fait des choix d'intégration. Dans une phase de post-traitement, la connaissance d'axiomes tel que $\exists P2_has_type.\{\dots coin\} \sqsubseteq \exists Thing_has_type_Concept.\{\dots moneta\}$ pourrait réduire encore le nombre de contraintes extraites.

6 Conclusion

Nos expérimentations démontrent la faisabilité d'une exploration systématique d'une base de connaissances, à la recherche de contraintes contextuelles de cardinalité maximale, grâce à l'algorithme que nous proposons dans cet article : dans le cas étudié, cela prend moins d'une heure pour une base de connaissances contenant plus de 3 millions de triplets, décrits par une centaine de concepts et plus de 300 rôles. Les propriétés utilisées par notre algorithme font que seules 595 contraintes ont été obtenues, ce qui reste analysable manuellement. Cela nous a permis de vérifier que ces contraintes sont bien pertinentes dans le contexte de la base étudiée. De plus, nos expérimentations démontrent l'importance du contexte dans cette découverte de contraintes. Il s'agit à notre connaissance de la première proposition de calcul de contraintes contextuelles de cardinalité maximale dans une base de connaissances du web sémantique. Ces grandes bases de connaissances, reflet d'une intelligence collective, sont générées à partir de l'expertise limitée de nombreux contributeurs et souffrent encore, tantôt de lacunes dans les informations, tantôt d'incohérences. Utiliser leurs contenus courants afin de mieux caractériser les connaissances représentées est donc très utile, comme montré dans l'état de l'art : cela permet aux applications qui exploitent ces grandes bases de connaissances de produire des résultats plus fiables.

Nous avons donc pour perspective d'exploiter les contraintes extraites pour calculer la confiance associée à des règles découvertes dans la base de connaissances ainsi enrichie. Mais avant cela, nous travaillons sur des post-traitements pour réduire encore le nombre de contraintes présentées en résultat. Pour cela, nous explorons le potentiel des raisonnements possibles sur la TBox, en particulier comment les relations de subsomption entre classes peuvent éliminer des redondances dans les ensembles de contraintes extraites.

Références

[1] Atencia, M., David, J., Scharffe, F. : Keys and pseudo-keys detection for web datasets cleansing and interlinking. In :

EKAU. pp. 144–153. Springer (2012)

[2] Auer, S., Bizer, C., Kobilarov, G., Lehmann, J., Cyganiak, R., Ives, Z. : Dbpedia : A nucleus for a web of open data. In : The semantic web, pp. 722–735. Springer (2007)

[3] Baader, F., Sattler, U. : Expressive number restrictions in description logics. *Journal of Logic and Computation* 9(3), 319–350 (1999)

[4] Darari, F., Nutt, W., Pirrò, G., Razniewski, S. : Completeness statements about rdf data sources and their use for query answering. In : ISWC. pp. 66–83. Springer Berlin Heidelberg (2013)

[5] Darari, F., Razniewski, S., Prasojo, R.E., Nutt, W. : Enabling Fine-Grained RDF Data Completeness Assessment. In : Web Engineering. pp. 170–187. Springer International Publishing, Cham (2016)

[6] Felicetti, A., Gerth, P., Meghini, C., Theodoridou, M. : Integrating heterogeneous coin datasets in the context of archaeological research. In : EMF-CRM@ICTPDL. pp. 13–27. CEUR-WS.org (2015)

[7] Galárraga, L., Razniewski, S., Amarilli, A., Suchanek, F.M. : Predicting completeness in knowledge bases. In : WSDM. pp. 375–383. ACM (2017)

[8] Galárraga, L.A., Teflioudi, C., Hose, K., Suchanek, F. : Amie : Association rule mining under incomplete evidence in ontological knowledge bases. In : WWW. pp. 413–422. ACM (2013)

[9] Galárraga, L., Hose, K., Razniewski, S. : Enabling Completeness-aware Querying in SPARQL. In : Proceedings of WebDB. pp. 19–22. ACM (2017)

[10] Lajus, J., Suchanek, F.M. : Are All People Married? Determining Obligatory Attributes in Knowledge Bases. In : WWW (2018)

[11] Pernelle, N., Saïs, F., Symeonidou, D. : An automatic key discovery approach for data linking. *Web Semantics : Science, Services and Agents on the World Wide Web* 23, 16–30 (2013)

[12] Razniewski, S., Korn, F., Nutt, W., Srivastava, D. : Identifying the extent of completeness of query answers over partially complete databases. In : SIGMOD. pp. 561–576. ACM (2015)

[13] Razniewski, S., Suchanek, F., Nutt, W. : But what do we actually know? In : 5th Workshop on Automated Knowledge Base Construction. pp. 40–44 (2016)

[14] Soutou, C. : Relational database reverse engineering : algorithms to extract cardinality constraints. *Data & Knowledge Engineering* 28(2), 161–207 (1998)

[15] Symeonidou, D., Armant, V., Pernelle, N., Saïs, F. : Sakey : Scalable almost key discovery in RDF data. In : ISWC. pp. 33–49. Springer (2014)

[16] Symeonidou, D., Galárraga, L., Pernelle, N., Saïs, F., Suchanek, F. : Vicky : Mining conditional keys on knowledge bases. In : ISWC. pp. 661–677. Springer (2017)

[17] Tanon, T.P., Stepanova, D., Razniewski, S., Mirza, P., Weikum, G. : Completeness-aware rule learning from knowledge graphs. In : ISWC. pp. 507–525. Springer (2017)

[18] Yeh, D., Li, Y., Chu, W. : Extracting entity-relationship diagram from a table-based legacy database. *Journal of Systems and Software* 81(5), 764–771 (2008)

Trois approches pour classifier les données du web des données

Justine Reynaud

Yannick Toussaint

Amedeo Napoli

Université de Lorraine, CNRS, Inria, LORIA, F-54000 Nancy, France

prénom.nom@loria.fr

Résumé

Dans cet article, nous nous intéressons au processus de classification de données relationnelles issues du web des données. Nous disposons d'un ensemble d'objets entre lesquels il existe des relations. Ces objets appartiennent à une ou plusieurs classes. Celles-ci sont définies en extension, et nous cherchons à construire une description en intensification en s'appuyant sur les relations des objets qui les composent. Pour cela, nous employons trois approches : les règles d'association qui s'appuient sur l'Analyse de Concepts Formels (FCA), les redescription et les règles de traduction qui s'appuient sur la Longueur de Description Minimale (MDL). À partir d'expérimentations sur DBpedia, nous discutons les spécificités et la complémentarité de ces trois approches. Nous montrons que les règles d'association sont les plus exhaustives tandis que les règles de traduction ont une meilleure couverture des données. Les redescription pour leur part, sont les règles les plus faciles à appréhender et interpréter.

Mots Clef

Données relationnelles, Analyse de Concepts Formels, Fouille de redescription, DBpedia.

Abstract

In this paper we study a classification process on relational data that can be applied to the web of data. We start with a set of objects and relations between objects, and extensional classes of objects. We then study how to provide a definition to classes, i.e. to build an intensional description of the class, w.r.t. the relations involving class objects. To this end, we propose three different approaches based on Formal Concept Analysis (FCA), redescription mining and Minimum Description Length (MDL). Relying on some experiments on RDF data from DBpedia, where objects correspond to resources, relations to predicates and classes to categories, we compare the capabilities and the complementarity of the three approaches. This research work is a contribution to understanding the connections existing between FCA and other data mining formalisms which are gaining importance in knowledge discovery, namely redescription mining and MDL.

Keywords

Relational data, Formal Concept Analysis, Redescription mining, DBpedia.

1 Introduction

Dans cet article, nous nous intéressons à la possibilité de découvrir des définitions dans les données RDF issues du web des données. Ces définitions peuvent être réutilisées dans la conception de bases de connaissances (KBs) ou dans l'enrichissement de KBs préexistantes. Étant donné l'immense quantité de données du web des données, il s'agit d'un enjeu majeur.

Le problème est le suivant : nous disposons d'un ensemble d'objets connectés par des relations, comme une ABox en logique de descriptions (DL) (Baader *et al.*, 2003). L'objectif est de classifier ces objets en fonction des relations dans lesquelles ils sont impliqués.

Les objets qui partagent des éléments communs appartiennent à une même classe. Ce partage peut être exact — les éléments sont identiques — ou approximatif — les éléments sont similaires. Finalement, nous obtenons un ensemble de classes organisées selon un ordre partiel, et les descriptions associées à ces classes. Ces descriptions sont nécessaires afin de construire les définitions des différentes classes. Les définitions sont considérées comme des conditions nécessaires (NC) et suffisantes (SC) pour classifier de nouveaux objets. Si x est une instance de la classe Rouge alors x a la couleur rouge (NC), et inversement, si x a la couleur rouge alors x est une instance de la classe Rouge (SC).

Pour poursuivre l'analogie avec les DLs, l'idée dans cet article est de construire et d'appliquer des règles d'induction de la forme « si $r(x, y)$ et $y:C$ alors $x:\exists r.C$ ». Cela signifie que, étant données y une instance de la classe C et r une relation telle que $r(x, y)$, alors x appartient à une classe, disons D , dont la description comprend $\exists r.C$. C'est-à-dire que les instances de D sont reliées à au moins une instance de C par la relation r . Nous utilisons ce type de règles pour construire les définitions des classes. Ces définitions sont de la forme $C_i \equiv e_1 \sqcap e_2 \sqcap \dots \sqcap e_n$ où e_j est une expression de la forme $\exists r.C_j$. Notre travail se divise en trois tâches principales, à savoir (i) préparer les données, (ii) découvrir les définitions, (iii) évaluer la qualité de ces

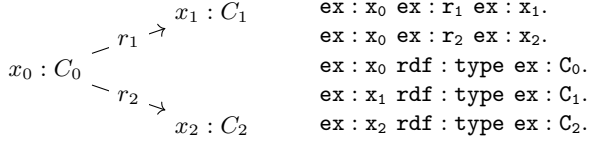


FIGURE 1 – Exemple de données relationnelles et les triplets associés.

définitions. Nous nous appuyons ici sur trois approches, les règles d'association, la fouille de redescriptions et la découverte de règles de traduction.

Cet article est dans la continuité de travaux de recherche sur la découverte de définitions dans le web des données. Son originalité est de comparer trois approches qui ne s'appuient pas sur les mêmes principes mais qui s'avèrent être complémentaires au vu des résultats de nos expérimentations. À notre connaissance, il s'agit du premier article où une telle comparaison est proposée à la fois à un niveau théorique et à un niveau expérimental.

L'article est organisé comme suit : dans la section 2, nous présentons la nature des données sur lesquelles nous travaillons et les processus de classification dans le web des données. Dans la section 3, nous détaillons les trois approches de classification et leurs applications. La section 5 présente les expériences qui ont été menées ainsi que l'évaluation des règles obtenues. Enfin, la section 6 présente une discussion ainsi que les pistes de recherches futures avant la conclusion section 7.

2 Représentation des données

2.1 Web des Données

Dans cette section, nous présentons les données issues du Web des Données (LOD) que nous considérons. Le LOD peut être vu comme un ensemble de KBs interconnectées. Une KB est composée de deux éléments : une TBox qui définit son *schema* et une ABox qui introduit les individus et leurs relations. L'unité de base d'une KB est le triplet RDF, noté $\langle s, p, o \rangle$, qui encode une assertion sous la forme sujet-prédicat-objet. Les différentes composantes d'un triplet peuvent être des *ressources* U identifiées de manière unique, des *littéraux* L (chaîne de caractères, numérique, date, ...) ou des *nœuds anonymes* B (assimilables à une variable existentiellement quantifiée), de telle sorte que $\langle s, p, o \rangle \in (B \cup U) \times U \times (B \cup U \cup L)$. Dans cet article, nous nous restreignons aux triplets composés uniquement de ressources, c'est-à-dire tels que $\langle s, p, o \rangle \in U \times U \times U$. Les ressources peuvent faire référence à n'importe quel objet ou abstraction et sont identifiées par une URI (*Uniform Resource Identifier*). Une URI est une adresse composée de deux parties. La première partie est l'espace de nom (*namespace*) qui indique de quelle KB viennent les ressources. La seconde partie donne un nom à la ressource au sein de cette KB.

La figure 1, présente un exemple de données relationnelles

et l'ensemble de triplets correspondant. Le préfixe `ex:` correspond à un *namespace* créé pour notre exemple, tandis que le préfixe `rdf:` correspond à un *namespace* pré-existant. `rdf:type` est la relation d'instanciation. Ainsi le triplet $\langle \text{ex}:x_0 \text{ rdf:type ex}:C_0 \rangle$ indique que x_0 est une instance de la classe C_0 .

Le LOD peut être interrogé grâce aux requêtes SPARQL. Par exemple, la requête `SELECT ?x WHERE { ?x rdf:type ex:C0 }` retourne toutes les instances de C_0 . Si l'on prend les données de la figure 1, seul `ex:x0` est retourné.

2.2 Analyse de Concepts Formels

Nous utilisons ici l'Analyse de Concepts Formels (FCA) de Ganter et Wille (1999) pour présenter et comparer les différentes approches. Étant donné un ensemble G d'entités¹, un ensemble M d'attributs et une relation binaire $I \subseteq G \times M$, (G, M, I) est un contexte formel. L'expression gIm s'interprète comme « l'entité g possède l'attribut m ». Les correspondances de Galois (notées $'$) pour un ensemble d'entités $X \subseteq G$ et un ensemble d'attributs $Y \subseteq M$ sont définies comme suit :

$$X' = \{m \in M \mid \forall x \in X, xIm\} \quad \text{et} \\ Y' = \{g \in G \mid \forall y \in Y, gIy\}.$$

À partir des données RDF, nous construisons un contexte formel dont les entités sont les sujets des triplets. Les attributs sont les paires (prédicat, objet) issues des triplets. Nous distinguons deux types d'attributs : des descriptions et des classes. Le premier ensemble d'attributs, dénoté $\mathcal{C}$, est composée des paires de la forme $(\text{rdf:type}, C)$, qui font d'un sujet une instance de la classe C . Ce sont ces classes que nous cherchons à définir. Le second ensemble, dénoté $\mathcal{D}$ est composé de paires (p, o) telles que $p \neq \text{rdf:type}$. Les attributs du contexte sont donc $M = \mathcal{C} \cup \mathcal{D}$. Si l'on prend l'exemple de la figure 1, la figure ?? présente le contexte associé.

		Attributs				
		Descriptions		Classes		
		$\exists r1:C_1$	$\exists r2:C_2$	C_0	C_1	C_2
Objets	x_0	×	×	×		
	x_1				×	
	x_2					×

FIGURE 2 – Contexte formel associé aux données représentées figure 1.

À partir de là, il s'agit de trouver un ensemble de catégories et un ensemble de descriptions de manière à ce que leurs correspondances soient les mêmes. Sur la figure ?? par exemple, on a $\{C_0\}' = x_0$ et $\{\exists r1:C_1, \exists r2:C_2\}' = x_0$. On peut alors construire la définition suivante :

$$C_0 \equiv \exists r1:C_1 \sqcap \exists r2:C_2$$

1. En FCA, G est l'ensemble des *objets*, renommés *entités* afin de ne pas confondre avec les objets en RDF.

Comme les données peuvent être incomplètes, il se peut qu'on ne retrouve pas une égalité entre une classe C_i et sa description $\exists x:C_j$. Autrement dit $\{C_i\}' \neq \{\exists x:C_j\}'$. Il nous faut donc des approches qui tolèrent une forme d'approximation. C'est ce que permettent les trois algorithmes utilisés, et nous allons les détailler plus précisément dans la section suivante.

3 Algorithmes de fouille de règles

3.1 Règles d'association

Le but de la fouille de règles d'association (Agrawal *et al.*, 1993; Klemettinen *et al.*, 1994) est de trouver des dépendances entre les attributs. Une règle d'association entre deux ensembles d'attributs X et Y , notée $X \rightarrow Y$, signifie « si un objet a tous les attributs de X , alors il a tous les attributs de Y ». À cette règle est associée une *confiance*, qui peut être vue comme une probabilité conditionnelle :

$$\text{conf}(X \rightarrow Y) = \frac{|X' \cap Y'|}{|X'|},$$

où $'$ correspond à la correspondance de Galois. La confiance est une mesure de qualité des règles d'association. Une règle d'association est valide si sa confiance est supérieure à un seuil θ défini par l'utilisateur. La confiance n'est pas symétrique : $X \rightarrow Y$ peut être valide sans que $Y \rightarrow X$ le soit. Si la confiance vaut 1, on dit qu'il s'agit d'une *implication* et l'on note $X \Rightarrow Y$. Si on a également $Y \Rightarrow X$, alors X et Y forment une définition, notée $X \equiv Y$.

Nous nous intéressons ici à des définitions. Nous considérons donc conjointement $X \rightarrow Y$ et sa réciproque $Y \rightarrow X$ et cherchons à estimer à quel point ces règles s'approchent d'une implication. Pour cela, nous introduisons la notion de *quasi-définition* qui est à la définition ce que la règle d'association est à l'implication.

Définition 1 (Quasi-définition). Étant donné deux ensembles d'attributs X et Y , une quasi-définition $X \leftrightarrow Y$ est valide si, pour un seuil θ donné,

$$\min(\text{conf}(X \rightarrow Y), \text{conf}(Y \rightarrow X)) \geq \theta.$$

L'algorithme **Eclat** (Zaki, 2000) est l'un des nombreux algorithmes de fouille de règles d'association existant. Il énumère de manière exhaustive toutes les règles d'association valides pour un seuil donné. Nous utilisons cet algorithme pour notre comparaison.

Une règle d'association $r_0 : x_1, \dots, x_n \rightarrow y_1, \dots, y_m$ peut se décomposer sous la forme de plusieurs règles d'association $r_i : x_1, \dots, x_n \rightarrow y_i$ pour $i \in \{1, \dots, m\}$. Si r_0 est valide, l'ensemble des r_i est valide. Pour que les règles obtenues définissent des classes comme nous le souhaitons, il faut que les règles d'association soient de la forme «Classes $\rightarrow$ Descriptions» ou «Descriptions $\rightarrow$ Classes». C'est à dire que la règle $X \rightarrow Y$ est telle que $X \subseteq \mathcal{C}, Y \subseteq \mathcal{D}$ ou $X \subseteq \mathcal{D}, Y \subseteq \mathcal{C}$. Il nous faut donc

utiliser un post-traitement afin de s'assurer que la règle ne contient que des catégories d'un côté et que des descriptions de l'autre. Pour cela, seules les règles $X \rightarrow Y$ qui satisfont l'une des deux contraintes alternatives suivantes sont conservées : (i) l'antécédent ne contient que des classes ($X \subseteq \mathcal{C}$) et il y a au moins une description dans la conséquence ($Y \cap \mathcal{D} \neq \emptyset$); (ii) l'antécédent ne contient que des descriptions ($X \subseteq \mathcal{D}$) et il y a au moins une classe dans la conséquence ($Y \cap \mathcal{C} \neq \emptyset$). Les règles conservées sont décomposées de manière à ne garder que des classes (resp. des descriptions) dans la conséquence si l'antécédent ne contient que des descriptions (resp. des classes).

Par exemple, $\exists r_1:C_1, C_0 \rightarrow \exists r_2:C_2$ n'est pas conservée parce que l'antécédent contient à la fois une catégorie et une description. En revanche, la règle $\exists r_1:C_1 \rightarrow \exists r_2:C_2, C_0$ peut être décomposée en deux règles $r_1: \exists r_1:C_1 \rightarrow \exists r_2:C_2$ et $r_2: \exists r_1:C_1 \rightarrow C_0$. La règle r_2 est conservée. Si sa réciproque est valide, la quasi-définition obtenue est $C_0 \leftrightarrow \exists r_1:C_1$.

3.2 Redescriptions

La fouille de redescriptions, introduite par Ramakrishnan *et al.* (2004), a pour but de trouver des caractérisations multiples d'un même ensemble d'entités. À la différence des règles d'association, les redescriptions s'appuient sur une séparation de l'ensemble des attributs en *vues*. Une vue est un sous-ensemble d'attributs. L'ensemble des vues forme une partition des attributs. Nous utilisons ici deux vues, qui correspondent aux deux types d'attributs — classes et descriptions.

La similarité entre deux ensembles d'attributs, provenant de deux vues différentes, est mesurée avec le coefficient de Jaccard :

$$\text{jacc}(A, B) = \frac{|A' \cap B'|}{|A' \cup B'|},$$

où $'$ correspond à la correspondance de Galois. Une paire d'ensembles d'attributs (A, B) est une redescription si le coefficient de Jaccard $\text{jacc}(A, B)$ est supérieur à un seuil donné. Le coefficient de Jaccard est symétrique, contrairement à la confiance. Une redescription dont le coefficient de Jaccard vaut 1 est une définition. Toutes les redescriptions sont nécessairement des quasi-définitions. En effet,

$$\min(\text{conf}(A \rightarrow B), \text{conf}(B \rightarrow A)) \geq \text{jacc}(A, B).$$

L'algorithme **ReReMi** (Galbrun et Miettinen, 2012) est utilisé ici. En plus des données binaires, **ReReMi** permet de tenir compte de données numériques et catégorielles. Les redescriptions obtenues peuvent être des fonctions booléennes contenant des disjonctions et des négations d'attributs. Afin de comparer les trois algorithmes, **ReReMi** est ici restreint à des conjonctions d'attributs, raison pour laquelle nous ne parlons pas de fonctions booléennes.

Dans un premier temps, **ReReMi** cherche des paires d'attributs — un attribut de chaque vue — susceptibles de prendre part à une définition, c'est-à-dire celles qui ont le coefficient de Jaccard le plus élevé. Ces paires sont ensuite

étendues tour à tour, en ajoutant à chaque fois un attribut à l'un des côtés de la redescription, jusqu'à ce que le nombre d'attributs maximum soit atteint ou que le coefficient de Jaccard n'augmente plus. Les redescriptions ainsi obtenues qui satisfont les critères de sélection sont retournées à l'utilisateur.

3.3 Règles de traduction

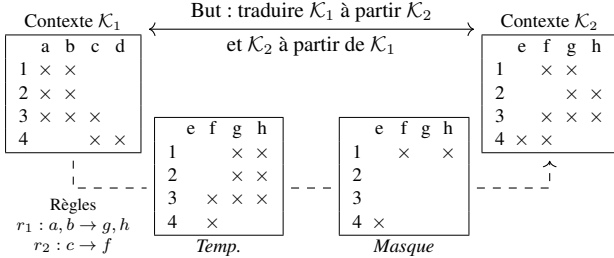


FIGURE 3 – Intuition du fonctionnement de Translator. Ici, seule la traduction de $\mathcal{K}_1$ vers $\mathcal{K}_2$ est représentée. À chaque étape, le but est de trouver la règle qui va minimiser le nombre de croix dans le masque.

L'algorithme Translator (van Leeuwen et Galbrun, 2015) s'appuie également sur une séparation du contexte en deux vues, et cherche un ensemble d'associations entre ces deux vues. Ces associations ont la même forme que des règles d'association, la différence se situant au niveau de la constitution de l'ensemble de règles.

Cet ensemble doit être compact et représentatif. D'une part, il doit couvrir la majorité des données. D'autre part, les règles doivent être aussi petites que possible en terme d'attributs. Afin de trouver un équilibre entre ces deux contraintes, Translator s'appuie sur le concept de Longueur de Description Minimum (MDL). Étant donné un contexte $\mathcal{K}$ et $X \subseteq M$ un ensemble d'attributs, la *longueur* de X est

$$L(X) = - \sum_{x \in X} \log_2 P(x | \mathcal{K}) \quad \text{où } P(x | \mathcal{K}) = \frac{|x'|}{|G|}.$$

Cette longueur correspond au nombre minimal de bits requis pour encoder X .

Pour construire l'ensemble des règles, les auteurs font l'analogie avec la notion de traduction. Une règle est une traduction d'une vue vers une autre. L'idée sous-jacente est la suivante : on souhaite construire un ensemble de règles qui permettent, connaissant une des deux vues, de reconstruire la seconde et *vice versa*. L'idée générale est représentée figure 2. Les erreurs entre le contexte cible et le contexte reconstruit sont corrigées à l'aide d'un masque. La taille de ce masque indique donc le nombre d'erreurs introduites.

L'ensemble de règles est construit de manière itérative, en prenant, à chaque étape, la règle $X \rightarrow Y$ qui maximise Δ

$$\Delta(X \rightarrow Y) = \underbrace{L(Mask^-) - L(Mask^+)}_{\text{Gain d'information}} - \underbrace{L(X \rightarrow Y)}_{\text{Longueur de la règle}}$$

où $Mask^+$ correspond aux éléments ajoutés au masque (erreurs introduites par la règle) et $Mask^-$ correspond aux éléments retirés du masque (erreurs corrigées par la règle). Des règles sont ajoutées tant que Δ est positif.

L'algorithme Translator est le seul dont la mesure de qualité tient compte des règles déjà extraites. Le masque étant mis à jour à chaque étape, la qualité d'une règle dépend des règles déjà extraites. Ainsi, Translator favorise un *bon ensemble de règles* plutôt qu'un *ensemble de bonnes règles*. Chaque règle apporte une information qui n'est pas contenue dans les autres, ce qui limite la redondance d'information.

4 Travaux connexes

La FCA est un outil de conceptualisation qui permet de structurer une ontologie par une approche *bottom-up*. Sertkaya (2011) fait une synthèse des différentes contributions qui s'intéressent au lien entre la FCA et les ontologies. La plupart des ontologies sont construites avec une approche *top-down*, c'est à dire en construisant d'abord le schéma. L'une des façons de tirer parti de la FCA est donc de consuire des connaissances qui vont enrichir une ontologie préexistante. C'est ce que nous faisons dans notre approche, en partant de triplets pour trouver des définitions. Dans (Ferré et Cellier, 2016), les auteurs proposent une extension de la FCA pour les graphes conceptuels. Contrairement aux graphes RDF qui reposent sur des relations binaires, les graphes conceptuels permettent de tenir compte de relation n -aires. Les concepts construits correspondent à une association entre une requête SPARQL et l'ensemble des solutions. Dans notre approche, nous cherchons à mettre en relation deux *graph patterns* et nous ne connaissons pas *a priori* la requête SPARQL qui correspond à la description d'une catégorie.

Notre travail poursuit un travail initié dans (Alam *et al.*, 2015). Les auteurs s'appuient sur la fouille de règles d'association pour fournir un espace de navigation sur des données RDF. Pour cela, les auteurs recherchent des implications et les ordonnent en fonction de la confiance de leur réciproque. Dans notre approche, nous introduisons la séparation des attributs et nous comparons les règles d'association avec les redescriptions et les règles de traduction.

L'algorithme AMIE, proposé par Galárraga *et al.* (2013), est une référence en matière de fouille de règles sur le web des données. Ces règles sont de la forme $B_1, B_2, \dots, B_{n-1} \rightarrow B_n$ où B_i correspond à une relation entre deux variables, que l'on peut exprimer sous forme de triplet $\langle x \text{ } r \text{ } y \rangle$. Par exemple, $manufacturer(x, Samsung), successor(y, x) \rightarrow manufacturer(y, Samsung)$ est une règle qui peut être trouvée par AMIE. Les auteurs ajoutent une condition, à savoir que toutes les variables doivent apparaître deux fois dans la règle. Si l'on considère les triplets sous forme de graphe, il en résulte qu'une règle est un motif cyclique entre des variables (indépendamment de l'orientation des arcs). Dans notre cas, les motifs recherchés

caractérisent une seule ressource identifiée (motif en étoile) et l'association est bidirectionnelle. On trouve par exemple la règle $Samsung_Mobile_Phone(x) \leftrightarrow manufacturer(x, Samsung)$.

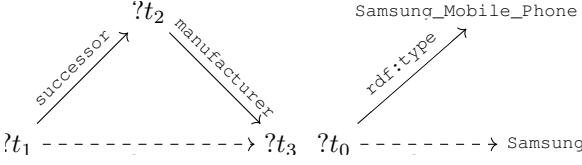


FIGURE 4 – Motifs recherchés dans la base de connaissances par AMIE et par notre approche. Les traits pleins concernent la partie gauche de la règle, et les pointillés la partie droite. Les $?t_i$ sont des variables existentielles.

5 Expérimentations

Nous avons réalisé nos expérimentations sur des données issues de *DBpedia*, qui est une KB construite à partir de *Wikipédia*. Nous cherchons à définir des *catégories*, c'est-à-dire les ressources qui sont dans le co-domaine de la relation `dct:subject`. Pour cela, nous extrayons un sous-ensemble de *DBpedia* à l'aide d'une requête SPARQL avant de transformer les triplets obtenus en contexte comme mentionné en Section 2.2. Enfin, nous utilisons les trois algorithmes présentés pour comparer et évaluer les résultats obtenus.

TABLE 1 – Statistiques sur les jeux de données extraits.

	Triplets	Objets	Attributs	
			Cat.	Descr.
Turing_Award	2 642	65	503	857
Smartphones	8 418	598	359	1 730
Sports_cars	9 047	604	435	2 295
French_films	121 496	6 039	6 028	19 459

5.1 Méthodologie

Nous avons extrait quatre sous-ensembles de triplets de *DBpedia*, de différents domaines². Afin d'isoler un sous-ensemble de triplets, nous avons récupéré tous les triplets dont le sujet est lié à une catégorie donnée. Pour cela, nous avons utilisé la requête SPARQL suivante :

```
SELECT DISTINCT * WHERE {
  ?s ?p ?o .
  ?s dct:subject dbc:X .
  ?p a owl:ObjectProperty .
}
```

Les quatre domaines utilisés (X) sont *French_films*, *Turing_Award_laureates*,

2. Les jeux de données ainsi que les résultats obtenus sont disponibles à l'adresse <https://gitlab.inria.fr/jreynaud/DefinitionMiningComparison>.

Smartphones et *Sports_cars*. Le triplet $?p$ a `owl:objectProperty` assure que $?o$ n'est pas un littéral.

Les triplets extraits sont divisés en deux groupes. L'ensemble des triplets qui ont pour prédicat `dct:subject` correspond aux attributs de classe (C). L'ensemble des triplets dont le prédicat n'est pas `dct:subject` correspond aux attributs de description (D).

5.2 Résultats et évaluation

TABLE 2 – Évaluation des résultats pour chaque jeu de données. $|C_i|$ (resp. $|D_i|$) désigne le nombre moyen de catégories (resp. de descriptions) dans une règle.

Turing_Award_laureates			
X	Eclat	ReReMi	Translator
$ B_{cand}^X $	47	12	11
$ B_{def}^X $	30	9	9
Précision	.64	.75	.85
$ C_i - D_i $	2 — 4	1 — 1	3 — 5
Sports_cars			
X	Eclat	ReReMi	Translator
$ B_{cand}^X $	132	52	31
$ B_{def}^X $	95	30	23
Précision	.72	.68	.74
$ C_i - D_i $	2.8 — 4.5	1.3 — 1.4	2.6 — 4.1
Smartphones			
X	Eclat	ReReMi	Translator
$ B_{cand}^X $	810	98	41
$ B_{def}^X $	521	57	31
Précision	.64	.58	.76
$ C_i - D_i $	4.3 — 7.8	1.6 — 1.8	3.1 — 3.1
French_films			
X	Eclat	ReReMi	Translator
$ B_{cand}^X $	546	36	93
$ B_{def}^X $	371	12	89
Précision	.68	.33	.96
$ C_i - D_i $	2.8 — 4.4	1.2 — 1.1	2.3 — 4.2

Chaque algorithme retourne un ensemble ordonné de quasi-définitions de la forme $C_0, \dots, C_n \leftrightarrow D_0, \dots, D_m$. Chacune des quasi-définitions est évaluée manuellement par trois personnes jouant le rôle d'experts. Étant donnée une quasi-définition $C_0, \dots, C_n \leftrightarrow D_0, \dots, D_m$ d'un jeu de données J , l'évaluateur répond à la question « Étant donné que nous parlons de J , est-il correct de dire que *faire partie à la fois de C_0 , de ... et de C_n et avoir les propriétés $D_0, \dots, D_m$ est équivalent ?* »

L'évaluation finale est l'évaluation majoritaire des trois experts. Si la règle est acceptée, elle rejoint l'ensemble des

définitions obtenues, dont 20 exemples sont présentés figure 4. Les experts ont été du même avis dans 95.4% des cas. La règle R5 de la figure 4 est une règle qui n’a pas fait l’unanimité.

Les algorithmes sont comparés sur la base des définitions extraites et des catégories définies. La figure 5 présente un diagramme de Venn pour chaque corpus, qui représente le nombre de définitions extraites par chaque algorithme. Par exemple, pour le corpus *Turing_Award_laureates*, il y a 22 définitions trouvées uniquement par Eclat, et 8 trouvées à la fois par Eclat et Translator. Au total, Eclat a extrait 30 définitions. La figure 6 présente également un diagramme de Venn pour chaque corpus, qui correspond au nombre de catégories définies par chaque algorithme. Une catégorie est considérée comme définie dès lors qu’elle apparaît dans au moins une définition. Il peut donc y avoir une ou plusieurs catégories considérées comme définies pour une seule définition. C’est notamment le cas des règles R2, R8 – R10, R12, R14, R17 – R20 de la figure 4.

6 Discussion

Ci-dessous, $\mathcal{B}_{cand}[X]$ désigne l’ensemble de toutes les quasi-définitions extraites par l’algorithme X et $\mathcal{B}_{def}[X]$ l’ensemble des définitions de X , c’est-à-dire, l’ensemble des quasi-définitions de $\mathcal{B}_{cand}[X]$ évaluées vraies par les experts. L’ensemble $\mathcal{B}_{cand}$ désigne la totalité des quasi-définitions, indépendamment de l’algorithme qui les a extraites. De la même façon, $\mathcal{B}_{def}$ désigne l’ensemble de toutes les définitions extraites.

6.1 Précision et rappel

La précision d’un algorithme X est $\frac{|\mathcal{B}_{def}^X|}{|\mathcal{B}_{cand}^X|}$. La précision de ReReMi a une très forte variabilité (entre 33 et 75%) et est globalement la plus faible, tout particulièrement sur le jeu de données *French_films*. La précision d’Eclat est stable (entre 64 et 72%). La meilleure précision est obtenue par Translator : quel que soit le jeu de données, elle toujours supérieure à 74%.

Le rappel, défini comme $\frac{|\mathcal{B}_{def}^X|}{|\mathcal{B}_{def}|}$ ne peut pas être utilisé comme une mesure de performance. En effet, certaines règles se recouvrent (ont des attributs en commun des deux côtés). C’est le cas des règles R6 à R10 (figure 4) qui définissent toutes la catégorie *McLaren_vehicles*. Tandis que Translator n’extrait qu’une seule règle (R8), ReReMi extrait 2 règles (R6 et R7) et Eclat en extrait 9 (seules R8 à R10 sont reportées ici).

La table ?? indique le nombre de catégories pour chaque jeu de données. Si l’on se fie à ces valeurs, le rappel est très faible : dans le corpus *Turing_Award_laureates* par exemple, pour 503 catégories dans les données de départ, seules 19 font partie d’une règle. Cela s’explique principalement par la très faible densité des contextes ; une grande quantité de catégories ne concerne qu’une seule entité du jeu de données. Si l’on ne compte que les catégories qui

ont un support d’au moins 3, elles ne sont plus que 105, et 47 catégories seulement ont un support d’au moins 5. Nous considérons donc le rappel non pas par rapport au nombre de catégories présentes dans le jeu de données, mais par rapport aux catégories retrouvées par l’ensemble des algorithmes, comme présenté figure 6.

6.2 Forme et interprétation des règles

D’après la figure 6, 70% des catégories définies par Eclat ou Translator sont définies par les deux algorithmes. Translator extrait considérablement moins de règles que Eclat (jusqu’à 16 fois moins pour le corpus *Smartphones*). Cela s’explique par la façon dont sont générées les règles d’association. Si dans les données, la règle $A \rightarrow B$ a le même support que la règle $A \rightarrow B, C$, alors seule la règle $A \rightarrow B, C$ est conservée. En revanche si le support de $A \rightarrow B$ est strictement supérieur au support de $A \rightarrow B, C$, Eclat génère deux règles. En conséquence, on obtient avec Eclat des règles qui ne diffèrent que d’un attribut, comme R9 et R10 par exemple, alors que Translator ne génère qu’une seule règle, R8 dans l’exemple.

ReReMi n’extrait aucune définition en commun avec Eclat et Translator. Cette différence s’explique par l’heuristique employée par ReReMi. Si C est une catégorie, et que D_1 et D_2 sont deux descriptions telles que $C' = D_1' = D_2'$, alors ReReMi privilégie deux définitions $C \equiv D_1$ et $C \equiv D_2$ tandis que Eclat ne génère qu’une seule définition $C \equiv D_1 \sqcap D_2$. C’est par exemple le cas des définitions R6, R7 générées par ReReMi et R8, générée par Eclat (figure 4). Si $C' = D_1'$ et $D_1' \subset D_2'$, ReReMi génère la définition $C \equiv D_1$, tandis que Eclat génère la définition $C \equiv D_1 \sqcap D_2$. Ce cas a également été retrouvé dans nos résultats, comme le montrent les définitions R12 et R13.

La taille des définitions générées par ReReMi est inférieure à celle des définitions de Translator et Eclat. Cela s’explique par l’heuristique de ReReMi déaillée dans le paragraphe précédent. ReReMi a en moyenne entre 1 et 2 attributs de chaque côté de la définition tandis que pour Eclat et Translator, il y a en moyenne 3 catégories et 4 descriptions par définition.

Les conjonctions dans les définitions extraites par ReReMi n’ont pas le même sens que celles extraites par Eclat et Translator. Par exemple, si l’on considère la définition R15, l’attribut (a Device) peut être enlevé sans altérer la définition : parce que toutes les entités considérées sont des instances de Device. À l’inverse, dans la définition R14, aucun attribut ne peut être retiré sans que la définition ne devienne fausse : tous les attributs font partie de la condition nécessaire. Dans notre approche, il nous semble plus pertinent de n’intégrer que des attributs qui contribuent pleinement à la définition dont ils font partie. Ainsi, la définition R14 nous paraît préférable à la définition R15, parce qu’elle est apparue plus facile à interpréter.

Turing_Award_laureates		
R	Harvard_University_alumni ↔ (almaMater Harvard_University)	R1
ET	Harvard_University_alumni, Turing_Award_laureates ↔ (a Agent), (a Person), (a Scientist), (almaMater Harvard_University)	R2
E	Turing_Award_laureates ↔ (a Agent), (a Person), (award Turing_Award)	R3
ET	Turing_Award_laureates ↔ (a Agent), (a Person), (a Scientist), (award Turing_Award)	R4
E	Modern_cryptographers ↔ (field Cryptography)	R5
Sports_cars		
R	McLaren_vehicles ↔ (manufacturer McLaren_Automotive)	R6
R	McLaren_vehicles ↔ (assembly Surrey)	R7
ET	McLaren_vehicles, Sports_cars ↔ (a Automobile), (a MeanOfTransportation), (assembly Woking), (assembly Surrey), (assembly England), (bodyStyle Coupé), (manufacturer McLaren_Automotive)	R8
E	McLaren_vehicles, Sports_cars ↔ (a Automobile), (a MeanOfTransportation), (assembly England), (assembly Surrey), (bodyStyle Coupé)	R9
E	McLaren_vehicles, Sports_cars ↔ (a Automobile), (a MeanOfTransportation), (assembly Surrey), (bodyStyle Coupé)	R10
Smartphones		
ET	Firefox_OS_devices, Open-source_mobile_phones, Smartphones, Touchscreen_mobile_phones ↔ (a Device), (operatingSystem Firefox_OS)	R11
R	Nokia_mobile_phones ↔ (manufacturer Nokia)	R12
ET	Nokia_mobile_phones, Smartphones ↔ (a Device), (manufacturer Nokia)	R13
R	Samsung_Galaxy ↔ (manufacturer Samsung_Electronics), (operatingSystem Android_(operating_system))	R14
ET	Samsung_Galaxy, Samsung_mobile_phones, Smartphones ↔ (a Device), (manufacturer Samsung_Electronics), (operatingSystem Android_(operating_system))	R15
French_films		
R	Pathé_films ↔ (distributor Pathé)	R16
R	Films_directed_by_Georges_Méliès ↔ (director Georges_Méliès)	R17
ET	Films_directed_by_Georges_Méliès, French_films, French_silent_short_films ↔ (a Film), (a Wikidata :Q11424), (a Work), (director Georges_Méliès)	R18
ET	Films_directed_by_Jean_Rollin, French_films ↔ (a Film), (a Wikidata :Q11424), (a Work), (director Jean_Rollin)	R19
ET	Film_scores_by_Gabriel_Yared, French_films ↔ (a Film), (a Wikidata :Q11424), (a Work), (musicComposer Gabriel_Yared)	R20

FIGURE 5 – Exemples de quasi-définitions obtenues par Eclat, ReReMi et Translator pour chaque corpus. Lorsque le préfixe d'un attribut de droite n'est pas spécifié, il s'agit de dbo. Le préfixe des catégories (dbc) n'est pas spécifié.

7 Conclusion

Dans cet article, nous nous sommes intéressés à trois algorithmes pour trouver des définitions de classes dans le web des données. Nous avons vu que chaque algorithme avait ses spécificités et nous avons vérifié empiriquement que les résultats extraits reflètent ces spécificités. Nous avons aussi montré que malgré des approches très différentes, les algorithmes Eclat et Translator extraient de nombreuses règles communes. À l'inverse, ReReMi, malgré une mesure de qualité semblable à Eclat, extrait des règles plus courtes. L'intérêt de ces algorithmes dépend de l'objectif de l'utilisateur. Dans le cas de notre expérimentation, Eclat est l'algorithme qui a qualifié le plus de catégories, au prix d'un nombre de règles extraites très important. Translator extrait significativement moins de règles avec un nombre de catégories définies légèrement inférieur. Enfin, ReReMi, malgré un nombre de catégories définies plus faible, offre des définitions plus simples en n'incluant pas les attributs qui ne participent pas pleinement à la règle. Par la suite, plusieurs directions de recherche sont envisagées, au niveau de la préparation des données, de la fouille et de l'évaluation.

Comme mentionné précédemment, les contextes construits

à partir des données RDF sont très creux. Dans le cas d'Eclat et de ReReMi, un seuil de support étant utilisé, l'existence d'attributs à très faible support n'a pas d'impact sur les règles retrouvées. Pour Translator en revanche, les règles doivent permettre de reconstruire intégralement le contexte. Dans ce cas, il peut être pertinent de simplifier le contexte avant de procéder à la fouille de règles. Cela peut notamment passer par la suppression des attributs à très faible support, ou au contraire, des attributs ayant un support quasi maximal, et donc qui concernent toutes les entités. Nous pourrions également travailler sur le type des données en entrée, en intégrant notamment des valeurs numériques, ce qui nous permettra de prendre en compte des triplets laissés de côté pour cette expérimentation, tels que les informations de date, d'âge, de taille, ...

Pour améliorer le processus de fouille, il est possible d'ajouter des contraintes sur les règles obtenues afin de restreindre l'espace de recherche, mais aussi d'obtenir des règles plus faciles à interpréter. Dans notre cas, n'autoriser qu'une catégorie par règle serait une contrainte pertinente. Nous pouvons également étendre l'expressivité des règles recherchées en permettant par exemple des opérateurs logiques tels que les disjonctions ou les négations. Si ReReMi permet déjà d'obtenir de telles formules lo-

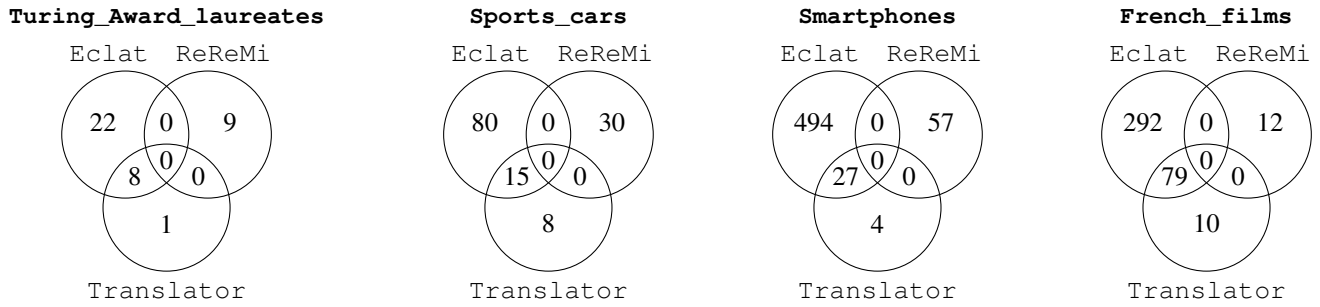


FIGURE 6 – Règles extraites par Eclat, ReReMi et Translator pour chaque jeu de données.

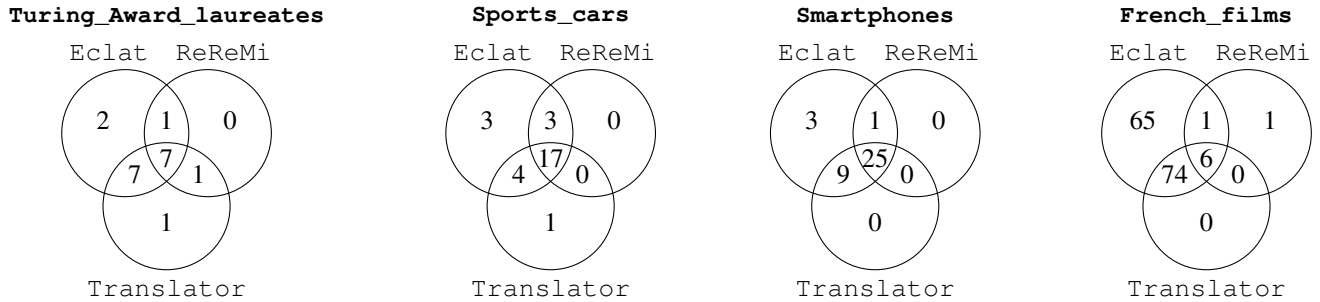


FIGURE 7 – Catégories définies par Eclat, ReReMi et Translator pour chaque jeu de données.

giques, et que des travaux dans ce sens existent pour les règles d'association, ce n'est pas le cas pour les règles de traduction. Cela implique de revoir le calcul de la longueur d'une règle.

Concernant l'évaluation, il nous faudra un moyen de comparer plus formellement les règles entre elles, peut-être *via* l'introduction d'une notion de redondance, qui permettrait de définir une base de règles.

Remerciements

Ce travail est financé avec le support de la Direction Générale de l'Armement et de la Région Lorraine. Merci à Esther Galbrun pour ses conseils ainsi que Quentin Brabant et Jérémie Nevin pour leur aide à l'évaluation des règles.

Références

- R. AGRAWAL, T. IMIELIŃSKI et A. SWAMI : Mining association rules between sets of items in large databases. *In ACM SIGMOD Rec.*, vol. 22, p. 207–216. ACM, 1993.
- M. ALAM, A. BUZMAKOV, V. CODOCEDO et A. NAPOLI : Mining definitions from rdf annotations using formal concept analysis. *In IJCAI*, p. 823–829, 2015.
- F. BAADER, D. CALVANESE, D. MCGUINNESS, D. NARDI et P. PATEL-SCHNEIDER, éd. *The Description Logic Handbook*. Cambridge University Press, Cambridge, UK, 2003.
- S. FERRÉ et P. CELLIER : Graph-FCA in Practice. *In Proceedings of 22nd ICCS*, p. 107–121, 2016.
- L. A. GALÁRRAGA, C. TEFLIOUDI, K. HOSE et F. M. SUCHANEK : AMIE : association rule mining under incomplete evidence in ontological knowledge bases. *In WWW'13*, p. 413–422, 2013.
- E. GALBRUN et P. MIETTINEN : From Black and White to Full Color : Extending Redescription Mining Outside the Boolean World. *Statistical Analysis and Data Mining*, 5(4):284–303, 2012.
- B. GANTER et R. WILLE : *Formal concept analysis - mathematical foundations*. Springer, 1999.
- M. KLEMETTINEN, H. MANNILA, P. RONKAINEN, H. TOIVONEN et A. I. VERKAMO : Finding interesting rules from large sets of discovered association rules. *In CIKM'94*, p. 401–407, 1994.
- N. RAMAKRISHNAN, D. KUMAR, B. MISHRA, M. POTTS et R. F. HELM : Turning CARTwheels : an Alternating Algorithm for Mining Redescriptions. *In KDD'04*, p. 266–275, 2004.
- B. SERTKAYA : A survey on how description logic ontologies benefit from formal concept analysis. *CoRR*, abs/1107.2822, 2011.
- M. van LEEUWEN et E. GALBRUN : Association Discovery in Two-View Data. *TKDE*, 27(12):3190–3202, déc. 2015.
- M. J. ZAKI : Scalable algorithms for association mining. *TKDE*, 12(3):372–390, May 2000.

Diagnostic automatique de l'état dépressif

S. Cholet

H. Paugam-Moisy

Laboratoire de Mathématiques Informatique et Applications (LAMIA - EA 4540)

Université des Antilles, Campus de Fouillole - Guadeloupe

Stephane.Cholet@univ-antilles.fr

Résumé

Les troubles psychosociaux sont un problème de santé publique majeur, pouvant avoir des conséquences graves sur le court ou le long terme, tant sur le plan professionnel que personnel ou familial. Le diagnostic de ces troubles doit être établi par un professionnel. Toutefois, l'IA (l'Intelligence Artificielle) peut apporter une contribution en fournissant au praticien une aide au diagnostic, et au patient un suivi permanent rapide et peu coûteux. Nous proposons une approche vers une méthode de diagnostic automatique de l'état dépressif à partir d'observations du visage en temps réel, au moyen d'une simple webcam. A partir de vidéos du challenge AVEC'2014, nous avons entraîné un classifieur neuronal à extraire des prototypes de visages selon différentes valeurs du score de dépression de Beck (BDI-II).

Abstract

Psychosocial disorders are a major public health problem that can have serious consequences in the short or long term, on a professional, personal or family level. The diagnosis of these disorders must be made by a professional. However, AI (Artificial Intelligence) can make a contribution by providing the practitioner with diagnostic assistance, and the patient with rapid and inexpensive ongoing follow-up. We propose an approach towards an automatic diagnosis of the depressive state based on real-time facial observations, using a simple webcam. From videos of the AVEC'2014 challenge, we trained a neural classifier to extract prototypes of faces according to different values of Beck's depression score (BDI-II).

Mots Clefs

Informatique affective; visages; classifieur incrémental; réseaux de neurones; apprentissage de prototypes.

1 Introduction

1.1 Contexte

Les troubles psychosociaux, et singulièrement les troubles dépressifs, sont une maladie touchant plus de 300 millions de personnes dans le monde. Ces troubles mentaux se caractérisent par une tristesse, une perte d'intérêt ou de plaisir, des sentiments de culpabilité ou de dévalorisation de

soi, un sommeil ou un appétit perturbé, une certaine fatigue et des problèmes de concentration [1]. La maladie se décline en plusieurs termes, souvent liés au contexte (par exemple, la dépression post-partum, après la grossesse; ou la dépression saisonnière, liée au manque de lumière l'hiver) et à la durée des symptômes, qui doivent persister au moins deux semaines pour caractériser une dépression [2]. Elle peut durer de quelques semaines à plusieurs mois voire années. Les conséquences sur l'individu atteint peuvent être multiples et de gravité variable. Parmi celles-ci, on peut citer l'isolement, l'absentéisme au travail, voire même les mutilations ou le suicide. L'importance de venir en aide aux personnes touchées est plébiscitée, et ce à différentes échelles. Dans les entreprises, de plus en plus de mesures sont prises afin d'assurer le bien-être des employés et de réduire ainsi les facteurs de risque liés à la dépression. A moins de consulter un spécialiste, les malades ne sont pas toujours en mesure de réaliser qu'ils sont atteints d'un trouble qui, dans une grande majorité des cas, peut se guérir grâce à un suivi psychologique et/ou à la prescription de médicaments adaptés [3].

1.2 Travaux antérieurs

Ces dernières années ont vu le nombre de travaux relatifs à l'analyse automatique du comportement émotionnel humain progresser de manière significative [4]. Plusieurs tentatives pour modéliser les émotions humaines ont été proposées, dont certaines sont très largement utilisées : une modélisation soit continue (le *circumplex* de Russell [5]), soit discrète (les émotions de base de Ekman [6] : tristesse, joie, colère, peur, dégoût et surprise). L'usage de la vidéo s'est petit à petit imposé comme source de données de choix pour l'analyse émotionnelle, bien qu'historiquement, des procédés plus invasifs aient été préférés, comme l'électrocardiogramme ou la conductance cutanée. Deux cadres d'études se distinguent. Le premier concerne la reconnaissance des émotions et le second, sur lequel nous nous focalisons ici, la prédiction des états dépressifs.

Encourageant les travaux dans ce domaine, des challenges internationaux tels que AVEC [7] ou FERA [8] invitent les chercheurs à confronter leurs méthodes sur une base de données commune. Wen *et al.* [9] ont utilisé des descripteurs visuels dynamiques (LPQ-TOP) associés à une régression par vecteurs supports (SVR) pour diagnostiquer l'état dépressif, avec une erreur RMSE de 8.17 sur le cor-

pus du challenge AVEC'2014. Zhu *et al.* [10] obtiennent une erreur de 9.55, en associant des images de flux optique aux images statiques des visages dans des réseaux de neurones profonds. D'autres approches tiennent compte de la modalité auditive, utilisée notamment pour l'apport d'informations de contexte importantes. Ainsi, Williamson *et al.* [11] ont combiné des descripteurs faciaux (sélection d'unités d'action) et auditifs (durée des phonèmes et analyses des fréquences, notamment) dans des mixtures de modèles gaussiens et obtiennent une erreur de 8.50 sur le corpus AVEC'2014. Gong *et al.* [12] tirent profit des trois modalités visuelle, auditive et contextuelle (retranscription d'interviews) au moyen de régresseurs courants (forêts aléatoires, descente de gradient stochastique et SVM, machine à vecteurs supports) pour une erreur de 4.99 sur le corpus DAIC-WOZ. Si certains travaux accordent une majeure partie de leur effort à la sélection des descripteurs, d'autres optent pour des méthodes connues pour leur capacité à extraire l'information directement depuis les images ou les bandes audios à disposition. Le corpus AVEC'2014 est étiqueté en termes de scores BDI-II et DAIC-WOZ en termes de scores PHQ-8 (*Patient Health Questionnaire*, ver. 8), qui sont deux méthodes d'évaluation de l'état dépressif (voir partie 2).

1.3 Contribution

Le développement d'un système qui, suite à la collaboration d'experts du domaine de la psychiatrie, pourra fournir un diagnostic automatique de l'état dépressif est un axe de bataille offert par l'Intelligence Artificielle pour prévenir l'apparition de tels troubles. En ce sens, l'effort proposé ici est double. Le premier est un outil orienté vers l'usage individuel, permettant à l'utilisateur d'évaluer la sévérité de la dépression dont il souffre. De cette manière, il pourra décider de la suite à donner à l'évaluation en allant consulter un spécialiste. Le second effort s'oriente vers l'usage du système par les experts, notamment pour sa capacité à se spécialiser sur un individu et à augmenter sa précision au fil des entretiens. Ainsi, il disposera d'une aide au diagnostic adaptée à chaque patient.

Le système est basé sur un classifieur neuronal, adapté au traitement de vidéos enregistrées ou capturées en direct. Le traitement produit en sortie un score dépressif, en termes du test *Beck Depression Inventory II* (BDI-II). Dans la partie 2, l'on présentera les données utilisées avant de décrire le classifieur utilisé dans la partie 3. La partie 4 pose les conditions expérimentales retenues dans le cadre de cette étude, et la partie 5 présente les résultats obtenus. Enfin, une conclusion et une ouverture sur des perspectives composera la partie 6.

2 Données

2.1 Corpus AVEC'2014

L'*AudioVisual Emotion Challenge* (AVEC) [7] est un concours international invitant les chercheurs à confronter leurs méthodes et à comparer les performances obtenues

sur un même jeu de données.



FIGURE 1 – Image extraite d'une vidéo de AVEC'2014.

Le jeu de données de l'édition 2014 [7] se présente sous la forme de 100 vidéos (tâche *Freeform*) où un individu est en interaction avec un avatar et répond à une question d'ordre général (e.g. comment vous sentez-vous ? pouvez-vous raconter un souvenir d'enfance ?) et de 100 autres (tâche *Northwind*) où l'individu lit un passage écrit, en langue allemande. Ces 200 vidéos sont réparties en deux sous-ensembles : la partition dite de développement (ensemble de motifs pour tester la *généralisation*) et la partition d'apprentissage (ensemble des *exemples* pour construire le modèle). Les vidéos sont constituées de *frames*, en nombre variable (les vidéos n'ayant pas toutes la même durée), enregistrées à raison de 30 par seconde, et contiennent des informations visuelles et auditives. Chaque vidéo est annotée d'un score, celui obtenu au test BDI-II (voir partie 2.2). Une troisième partition, dite partition de test, ne comprend que des vidéos (100 éléments), sans annotations. Les performances prises en compte par les organisateurs du challenge pour départager les participants sont calculées sur cette dernière partition.

Pour l'étude présentée dans cet article, nous retenons les données visuelles des 200 vidéos des tâches *Freeform* et *Northwind*. Cela représente un jeu de données de 291 155 images semblables à celles de la Figure 1.

2.2 Beck Depression Inventory II

Le test d'évaluation de l'état dépressif *Beck Depression Inventory* (BDI) [13] a été créé par Aaron T. Beck., père de la thérapie cognitive, en 1961. Il a subi plusieurs modifications, visant à l'améliorer. En 1996, sa version II (BDI-II) est un test auto-administré, comptant 21 questions.

Le score obtenu peut prendre une valeur de 0 à 63 ; il donne une indication sur la sévérité de la dépression dont souffre le patient, tel que précisé dans la Table 1. A l'époque de

TABLE 1 – Interprétation du score au test BDI-II

Score obtenu	Sévérité de la dépression
0-13	Minimale
14-19	Moyenne
20-28	Modérée
29-63	Sévère

l'apparition du test, il va à contrecourant des pratiques, en

se focalisant sur la perception qu'a le patient de son propre état, plutôt que sur les enjeux psychologiques motivant son comportement et ses réactions à un environnement donné [14] (que l'on appelle, dans la littérature, la psychodynamique). Le test se fonde sur des années de collectes de données, de collaborations entre psychiatres, d'entretiens docteur-malade et de révisions [15]. Le test BDI-II bénéficie d'une corrélation positive avec l'échelle *Hamilton Rating Scale* (HRS) [13], qui est un test administré par un professionnel en psychiatrie. Il est important de noter que dans le processus d'évaluation, le test BDI-II ne tient pas compte de l'expression faciale ou verbale du sujet. L'étude présentée ici démontre qu'il existe bien une forte corrélation entre l'expression faciale et l'état dépressif puisque le classifieur construit à partir des visages extraits des vidéos permet de prédire de manière fiable la sévérité dépressive.

2.3 Extraction des descripteurs et changement de repère

Afin de classifier les vidéos selon leur score BDI-II, on extrait, pour chaque image, un ensemble de 68 points faciaux d'intérêt (voir Figure 2). Cela nécessite, en amont, la détection et le redimensionnement des visages. L'extracteur de points d'intérêts utilise le modèle de Kazemi et Sullivan [16]. Le détecteur de visages (entraîné sur l'ensemble i-BUG 300-W, voir Sagonas *et al.* [17]) implémente un classifieur linéaire sur une pyramide d'images dans des fenêtres temporelles, ainsi que sur des histogrammes de gradients orientés. L'outil Dlib [18] a été utilisé pour mettre en œuvre l'extraction.

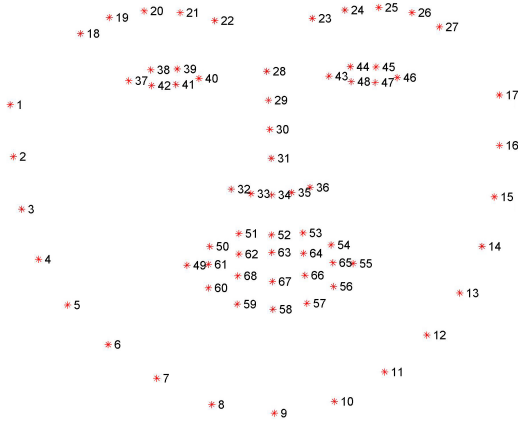


FIGURE 2 – Points faciaux d'intérêt - Modèle MultiPIE.

L'alignement des visages est également une étape importante [19], qui permet notamment de limiter le biais introduit par des facteurs tels que la distance à la webcam ou la morphologie faciale du sujet, mais aussi d'homogénéiser les données avant la classification. Afin d'aligner les yeux, de centrer les visages dans l'image et d'homogénéiser leur taille, les points subissent une translation de vecteur $\vec{T}$, ainsi qu'une rotation d'angle θ (facteur d'échelle : S). Cette transformation isométrique est un cas particulier

de similitude calculée pour chaque visage. La matrice de transformation M est donnée par l'équation 1. Ces étapes sont réalisées au moyen de la librairie OpenCV [20].

$$M = \begin{bmatrix} s_x \cos(\theta) & \sin(\theta) & t_x \\ -\sin(\theta) & s_y \cos(\theta) & t_y \end{bmatrix} \quad (1)$$

3 Modèle à base de prototypes

Le choix du classifieur incrémental pour prédire l'état dépressif a été motivé à la fois par les inspirations biologiques sous-jacentes, comme démontré par Grossberg dès la fin des années 80 (voir [21] pour une synthèse ce sujet) et par le récent regain d'intérêt pour les modèles à base de prototypes : Biehl, Hammer et Villmann [22] affirment en 2016 que de tels systèmes sont très intéressants pour l'analyse de données complexes et de grande dimension.

3.1 Le classifieur incrémental

Le modèle ART (*Adaptive Resonance Theory*) de Grossberg [23] est un système de classification neuronal capable de s'adapter aux entrées dites significatives, tout en restant stable face aux entrées non-significatives. Ainsi, si l'on présente au système un exemple proche d'une représentation qu'il connaît, il la modifiera en conséquence. En revanche, si on lui présente un exemple inconnu, une nouvelle représentation sera créée pour le prendre en compte.

Le classifieur incrémental utilisé ici est inspiré du modèle ART et suit le même principe. Il a été proposé par Azcarraga [24] puis modifié par Puzenat [25], qui l'utilisait pour la reconnaissance de formes manuscrites. Il s'agit d'un réseau de neurones dont la couche d'entrée est, classiquement, adaptée à la dimension de l'espace des données. La seconde couche est constituée de "neurones-distance", les prototypes, qui sont totalement connectés aux neurones d'entrée. Ainsi, à chaque présentation d'un exemple, celui-ci est comparé à tous les prototypes en mémoire. Dans la troisième couche, chaque neurone est connecté à un seul et unique prototype (voir Figure 3); aucun apprentissage n'est effectué par la couche de sortie.

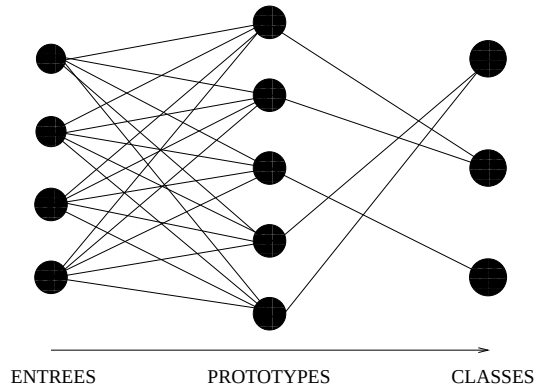


FIGURE 3 – Architecture du classifieur incrémental.

Selon le protocole précisé ci-dessous (partie 3.2), de nouveaux prototypes sont créés pendant le processus d'appren-

tissage. Néanmoins le fait que leur nombre, qui dépend de la taille de la base d'apprentissage ainsi que du nombre de classes, reste faible par rapport au nombre d'exemples garantira que le classifieur aura extrait des données une information synthétique. Les prototypes sont représentés dans le même espace que celui des données, ce qui rend aisée leur interprétation, ainsi que leur appréhension par des experts dans le cadre d'un travail transversal [22].

3.2 Apprentissage

La phase d'apprentissage consiste en une seule passe de la base d'exemples et elle suit un algorithme par compétition. Initialement, le premier exemple est recopié comme unique prototype et on l'associe en sortie à la classe de l'exemple. Par la suite, pour chaque nouvel exemple présenté X , on cherche le prototype $P_{meilleur}$ qui en est le plus proche, au sens de la mesure choisie (voir discussion ci-dessous). *A priori*, si la classe de $P_{meilleur}$ est celle de l'exemple, le prototype est gagnant et sa connexion avec l'exemple est modifiée afin de l'en rapprocher. Sinon, un nouveau prototype est créé à l'image de l'exemple et est associé en sortie à la classe de l'exemple.

Une exception à l'adaptation de $P_{meilleur}$ relève d'une condition plus subtile : on cherche, parmi les prototypes les plus proches de l'exemple, le premier prototype dont la classe est différente de celle de $P_{meilleur}$, et on le nomme P_{second} . S'il y a risque de confusion, i.e. dans une zone de l'espace d'entrée où des motifs proches doivent être associés à des classes distinctes, alors un nouveau prototype est créé.

Le sens de proximité se réfère ici à une mesure de distance ou de similarité entre un exemple et un prototype. La mesure est choisie en accord avec le problème à traiter, ce qui rend les classifieurs incrémentaux flexibles et adaptables. On peut utiliser une distance de Mahalanobis, qui accorde un poids moins important aux composantes les plus dispersées, ou une distance de Minkowski (équation 2) qui permet un calcul plus rapide.

$$d_{mink_p} = \left(\sum_{i=1}^n |x_i - y_i|^p \right)^{1/p} \quad (2)$$

Pour $p = 2$, on retrouve la distance euclidienne qui est particulièrement adaptée aux situations où les descripteurs sont des coordonnées, comme c'est le cas pour les points faciaux d'intérêt. Après avoir vérifié que des valeurs $p > 2$ ne produisaient pas de résultats significativement meilleurs, nous avons opté pour la distance euclidienne et nous la noterons d .

3.3 Hyperparamètres de contrôle

L'algorithme d'apprentissage du classifieur incrémental propose trois hyperparamètres de contrôle afin de répondre au dilemme stabilité-plasticité (*stability-plasticity dilemma*), c'est-à-dire de tenir compte des nouveaux éléments à apprendre sans oublier ceux déjà mémorisés.

Seuil d'influence. Le seuil d'influence s_{inf} du modèle définit la valeur à laquelle doit être inférieure la distance entre un exemple et son meilleur prototype, tel que décrit par l'équation 3. Si cette condition n'est pas vérifiée, on crée un nouveau prototype.

$$d(P_{meilleur}, X) \leq s_{inf} \quad (3)$$

Seuil de confusion. Le seuil de confusion s_{conf} aide à lever les ambiguïtés dans les zones frontières entre classes distinctes et s'utilise comme décrit par l'équation 4. Si cette condition n'est pas vérifiée, on crée un nouveau prototype.

$$d(P_{meilleur}, X) - d(P_{second}, X) > s_{conf} \quad (4)$$

Coefficient de rapprochement. Lorsqu'aucun nouveau prototype n'est créé, l'adaptation de $P_{meilleur}$ à l'exemple X est contrôlée par le coefficient de rapprochement α . L'équation 5 décrit la modification des poids du neurone prototype.

Pour i tel que $P_i = P_{meilleur}$,

$$\forall j, w_{ji} \leftarrow w_{ji} + \alpha(x_j - w_{ji}) \quad (5)$$

Plus ce coefficient α est grand, plus la représentation modélisée par $P_{meilleur}$ se rapproche de l'exemple et plus on accroît la création de prototypes. *A contrario*, pour un coefficient petit, $P_{meilleur}$ sera peu modifié et le nombre de prototypes restera limité. On note que, pour $\alpha = 0.5$, on calcule le barycentre entre les deux entités.

3.4 Généralisation

La généralisation respecte les mêmes contraintes que l'apprentissage, mais il n'y a plus de création ni de modification de prototypes. Pour chaque nouveau motif n'ayant pas participé à l'apprentissage du modèle :

1. **Présenter** un motif X
2. **Rechercher** $P_{meilleur}$ tel que $d(P_{meilleur}, X)$ soit minimale
3. **Rechercher** P_{second} tel que la classe de P_{second} soit différente de celle de $P_{meilleur}$
4. **Si** $P_{meilleur}$ est trop éloigné de X ou s'il y a risque de confusion :

$$\begin{cases} d(P_{meilleur}, X) > s_{inf} \\ \text{ou} \\ d(P_{meilleur}, X) - d(P_{second}, X) \leq s_{conf} \end{cases}$$

Alors rejeter X (non-réponse)

Sinon Si la classe de $P_{meilleur}$ est celle de X ,

Alors X est reconnu (bonne réponse)

Sinon X n'est pas reconnu (mauvaise réponse)

3.5 Non-réponses

Le classifieur incrémental est en mesure de produire, en sortie, trois types réponses : une "non-réponse", une

"bonne-réponse" ou une "mauvaise réponse". En particulier, une non-réponse est rendue lorsque le meilleur prototype $P_{meilleur}$ de l'exemple d'entrée est trop éloigné de ce dernier, ou lorsque la distance entre $P_{meilleur}$ et P_{second} est trop faible au regard de l'exemple. Cette réponse, en plus de se rapprocher du diagnostic que ferait un humain, peut être considérée comme un indicateur de fiabilité du score dépressif calculé pour un sujet donné (voir 5.3). Dans le cas où le système produirait un grand nombre de non-réponses, la classification pourrait être jugée peu fiable. Le cas échéant, le système peut être spécialisé sur l'individu, via un réapprentissage du modèle, sur décision d'un expert. Cette possibilité d'obtenir une "non-réponse" est une spécificité précieuse de ce type de classifieur, le rendant plus proche d'un diagnostic humain.

4 Conditions expérimentales

Le classifieur incrémental est entraîné pour associer à chaque image (cf. 2.1) le score BDI-II de la vidéo dont elle a été extraite. Afin de réduire les risques liés au sur-apprentissage, et pour disposer d'un plus grand nombre de classes représentées, nous avons mélangé les partitions de développement et d'apprentissage dont nous disposions. De plus, les exemples ont été stratifiés afin que chaque classe soit toujours représentée en quantité raisonnable dans les ensembles d'apprentissage et de généralisation.

A priori, il conviendrait d'apprendre un modèle en régression pour lire en sortie la valeur du score. Cependant les différentes valeurs sont en nombre limité (seulement 41 présentes dans les vidéos étudiées, parmi les 64 valeurs possibles en théorie) et chacune sera considérée comme une classe. Il est important de noter que le système ne sera pas en mesure de discriminer, en généralisation, une classe inexistante dans les données d'apprentissage. De plus, compte tenu de la Table 1, on pourra *a posteriori* regrouper les scores numériques dans des intervalles pour qualifier de manière descriptive la sévérité de la dépression.

4.1 Stratégies de test

Nous retenons trois stratégies pour les expériences :

Classique : construction d'un modèle sur une base d'apprentissage puis estimation de la performance en généralisation sur une base disjointe ;

Validation croisée : une partition $S = \cup_{m=1}^M S_m$ de la base de données S étant réalisée, apprentissage de M modèles, chacun sur $S = \cup_{k \neq m} S_k$, avec estimation de sa performance en généralisation sur S_m ;

Flux continu : un modèle ayant été appris, utilisation en temps-réel pour prédire l'état dépressif d'un individu placé devant une webcam.

La stratégie "classique" a été mise en œuvre en premier, afin d'étudier le comportement du modèle et de valider les choix de prétraitements et d'extraction des descripteurs (présentés en 2.3). Au fil de ces expérimentations, la taille de la base de généralisation a été fixée à 3/10e des don-

nées, distinctes des 7/10e ayant servi à entraîner le modèle, comme le récapitule la Table 2.

TABLE 2 – Composition des ensembles de données pour la stratégie classique

	Freeform	Northwind	Total
Apprentissage	113 876	89 933	203 809
Généralisation	48 945	38 401	87 346
Total	162 821	128 334	291 155

Les meilleurs hyperparamètres pour le modèle ont été déterminés au moyen d'une recherche en grille (*grid search*) pour tenir compte des interactions entre hyperparamètres. Les résultats présentés ci-dessous ont été obtenus avec un seuil d'influence de 70, un seuil de confusion de 0.1 et un coefficient de rapprochement de 0.1.

4.2 Mesures de performances

Afin d'évaluer les performances de notre approche, nous retenons quatre indicateurs, dont deux estiment un taux de succès en classification et deux autres mesurent une erreur :

- Le taux de succès, en termes de **score** BDI-II
- Le taux de succès au sens des **intervalles** de sévérité de dépression (cf. Table 1)
- L'erreur quadratique moyenne (RMSE)
- L'erreur absolue moyenne (MAE)

Ces deux derniers indicateurs sont bien adaptés aux cas de la classification multi-classe, particulièrement lorsque les classes sont hétérogènes en nombre de données. Pour les taux de succès, il est important de noter que ces indicateurs seront calculés en tenant compte des images bien classées, et non des vidéos dans leur ensemble.

5 Résultats

5.1 Entraînement et validation croisée

Les meilleures performances obtenues dans le cadre de la stratégie "classique" sont présentées dans la Table 3 pour les taux de succès et la Table 4 pour les indicateurs d'erreur. Notons que ces résultats, en particulier les RMSE, ne sont pas directement comparables avec ceux du challenge AVEC'2014 cités en 1.2 dans la mesure où nous n'avons pas accès à la partition de test réservée aux organisateurs du challenge, et où nous n'avons pas choisi le même partitionnement des données pour nos essais.

TABLE 3 – Taux de succès pour la stratégie classique

	Freeform		Northwind	
	Bon score	Bon intervalle	Bon score	Bon intervalle
App.	92.43%	95.25%	92.29%	95.40%
Gén.	89.78%	93.70%	91.01%	94.52%

Cette stratégie oblige à construire le modèle en n'utilisant qu'une partie des données (70% ici). En revanche, la stratégie "validation croisée" permet, après estimation moyenne

TABLE 4 – Les erreurs pour la stratégie classique

	Freeform		Northwind	
	RMSE	MAE	RMSE	MAE
App.	4.07	0.83	3.79	0.8
Gén.	4.67	1.11	4.05	0.92

de la performance en généralisation sur M modèles, de construire un $M + 1^{eme}$ modèle qui apprend sur toutes les données à disposition. La Table 5 donne les performances pour une validation croisée avec $M = 10$, où l'algorithme apprend sur 262 040 exemples et généralise sur les 29 115 restants. En effet, les bases *Freeform* et *Northwind* ont été mélangées puisque la stratégie "classique" a démontré la similarité de leur comportement.

TABLE 5 – Performances en validation croisée

	Bon score	Bon intervalle	RMSE	MAE
Moyenne	90.73%	94.51%	4.30	0.97

On note un gain de performance d'environ 4 % en passant du nombre de bien classés par score BDI-II au nombre de bien classés par intervalle de sévérité dépressive. Cette amélioration confirme l'existence d'une continuité entre états dépressifs de sévérité proche, et témoigne de la capacité du système à la saisir. Le nombre de prototypes est de l'ordre de 15% à 18% du nombre d'exemples.

5.2 Comparaison avec d'autres classifieurs

TABLE 6 – Comparaison des performances des classifieurs de la littérature

	Bon score	Temps de classification (en sec., pour un ex.)
SVM	73.23 %	0.009
MLP	66.98 %	0.001
Random Forest	94.87 %	0.118
C. Incrémental	90.25%	0.023

Les performances du classifieur incrémental sont comparées aux performances de classifieurs de la littérature dans la Table 6. Les taux de succès ont été obtenus en généralisation sur 30% des données via la stratégie classique exposée en 4.1, après un apprentissage sur 70% de la base complète. Les temps de réponse des classifieurs à un nouveau motif présenté ont aussi été mesurés. Le CI n'est pas le plus performant en termes de taux de succès, mais présente le meilleur compromis entre qualité et rapidité de la réponse. Ce point est essentiel dans le cadre d'un outil d'aide au diagnostic puisque le système doit pouvoir donner une estimation en flux continu.

5.3 Flux continu

À l'issue de la validation croisée, le classifieur a été entraîné sur l'ensemble des 291 155 images disponibles. Les performances à considérer sont celles obtenues en moyenne (voir Table 5). Il peut désormais être utilisé en prédiction pour fournir une estimation automatique de l'état dépressif d'un individu faisant face à une simple webcam [26].

La fréquence des images est de 30 par seconde lors de la capture. Cependant, l'expression dépressive s'évaluant sur la durée, il n'est pas nécessaire de traiter toutes les images produites. On fixe un nombre d'images n à traiter par secondes (par exemple, $n = 10$) ainsi qu'une durée d'enregistrement. Les images sont prétraitées et les descripteurs extraits comme décrit dans la partie 2.3. Chacune est alors comparée aux prototypes par l'algorithme de généralisation (cf. partie 3.4).

La sortie du système est une valeur d'état dépressif du sujet filmé estimée par le score BDI-II majoritaire sur une période p donnée en secondes (par exemple : $p = 20$). Notons au passage que cette procédure permet d'effacer au fur et à mesure les données personnelles qui n'auront été enregistrées que temporairement. Comme suggéré dans la partie 3.5, les non-réponses pourront être à terme exploitées comme indicateur de fiabilité du classifieur.

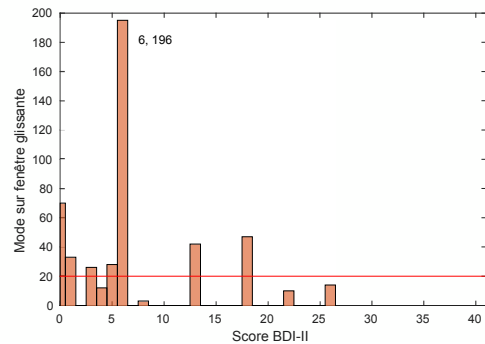


FIGURE 4 – Distribution des scores proposés par la classification en flux continu

Il est important de noter que, dans ce dernier contexte, il ne nous est pas encore possible d'évaluer de réelles performances. Par exemple, la pertinence des scores BDI-II fournis par le système ne pourra être validée que lorsqu'un protocole expérimental sera mis en place, en collaboration avec un expert humain (voir partie 6). Néanmoins, la faisabilité du traitement *on-line* a été établie par l'un des auteurs de cet article : en filmant son propre visage, il a obtenu un score stable de 6 sur une période d'une vingtaine de secondes, avec un nombre de non-réponses de 20 sur 500 matérialisé par la ligne horizontale sur la Figure 4.

6 Conclusion et discussion

Nous avons proposé un classifieur incrémental à base de prototypes afin de déterminer l'état dépressif d'un individu à partir d'une vidéo. Le prétraitement des images permet de réduire fortement le biais introduit par les différences d'échelle et les spécificités morphologiques des sujets. La classification rapide autorise, sous couvert de validation par un expert, le développement d'un module de classification en temps-réel de l'état dépressif, en capturant le flux vidéo directement via une webcam.

Le système pourra facilement être utilisé par un praticien comme outil d'aide au diagnostic et de suivi de patient, ce dernier pouvant lui-même effectuer des évaluations de son état à l'aide d'un matériel peu coûteux. Si cela s'avère nécessaire, l'outil pourra être ré-étalonné (phase d'apprentissage complémentaire) pour mieux s'adapter à un patient précis. Sur le plan technique, notons cependant que l'accroissement du nombre de prototypes aura pour effet de ralentir le traitement. Pour cela, nous proposons en perspective l'étude d'une procédure d'élagage, visant à réduire le nombre de prototypes. Ce type de procédure va de paire avec tout système incrémental, et est en phase active de développement, raison pour laquelle elle n'est pas présentée dans cet article.

Notons enfin que la plupart des travaux sur la dépression utilisent à la fois les modalités visuelle et auditive, à l'instar de Yu *et al.* [27]. La prochaine étape de ce travail consistera à entraîner, de manière indépendante et sur le même modèle, un classifieur permettant de prédire l'état dépressif à partir des données audio uniquement. La mise en commun des deux modèles pourra ensuite se faire au moyen d'un modèle de mémoire associative multimodale qui réalise la fusion des données à l'aide d'une *Bidirective Associative Memory* (BAM). Le modèle complet a déjà été développé [28], sur la base d'une modélisation cognitive, et il a démontré l'amélioration des performances par la prise en compte de plusieurs modalités [29].

Remerciements

Le travail décrit dans cet article a été réalisé en Python. En ce sens, ses auteurs souhaitent remercier les contributeurs de Numpy [30], Scipy [31] et Scikit-learn [32].

Références

- [1] (2018) Depression. Organization World Health. Accessed 2018-04-02. [Online]. Available : <http://www.who.int/mediacentre/factsheets/fs369/fr/>
- [2] (2018) Depression. Health National Institute of Human. Accessed 2018-04-02. [Online]. Available : <https://www.nimh.nih.gov/health/topics/depression/>
- [3] R. J. DeRubeis, G. J. Siegle, and S. D. Hollon, "Cognitive therapy vs. medications for depression : Treatment outcomes and neural mechanisms," *Nat. Rev. Neurosci.*, no. 10, pp. 788–796, oct.
- [4] Z. Zeng, M. Pantic, G. I. Roisman, and T. S. Huang, "A survey of affect recognition methods : Audio, visual, and spontaneous expressions," *IEEE Trans. Pattern Anal. Mach. Intell.*, vol. 31, no. 1, pp. 39–58, 2009.
- [5] J. Russell, "A circumplex model of affect," *J. Pers. Soc. Psychol.*, vol. 39, no. 6, pp. 1161–1178, 1980.
- [6] P. Ekman, "Differential Communication Of Affect By Head And Body Cues.pdf," *J. Pers. Soc. Psychol.*, vol. 2, no. 5, pp. 726–735, 1965.
- [7] F. Ringeval, M. Pantic, B. Schuller, M. Valstar, J. Gratch, R. Cowie, S. Scherer, S. Mozgai, N. Cummins, and M. Schmitt, "Avec 2017 - Real-life Depression, and Affect Recognition Workshop and Challenge," *Proc. 7th Annu. Work. Audio/Visual Emot. Chall. - AVEC '17*, pp. 3–9.
- [8] M. F. Valstar, E. Sanchez-Lozano, J. F. Cohn, L. A. Jeni, J. M. Girard, Z. Zhang, L. Yin, and M. Pantic, "FERA 2017 - Addressing Head Pose in the Third Facial Expression Recognition and Analysis Challenge," *Autom. Face Gesture Recognit. (FG 2017)*, pp. 839–847, 2017.
- [9] L. Wen, X. Li, G. Guo, and Y. Zhu, "Automated depression diagnosis based on facial dynamic analysis and sparse coding," *IEEE Trans. Inf. Forensics Secur.*, vol. 10, no. 7, pp. 1432–1441, 2015.
- [10] Y. Zhu, Y. Shang, Z. Shao, and G. Guo, "Automated Depression Diagnosis based on Deep Networks to Encode Facial Appearance and Dynamics," *IEEE Trans. Affect. Comput.*, no. X, pp. 1–1.
- [11] J. R. Williamson, T. F. Quatieri, B. S. Helfer, G. Ciccarelli, and D. D. Mehta, "Vocal and Facial Biomarkers of Depression based on Motor Incoordination and Timing," *Proc. 4th Int. Work. Audio/Visual Emot. Chall. - AVEC '14*, pp. 65–72.
- [12] Y. Gong and C. Poellabauer, "Topic Modeling Based Multi-modal Depression Detection," *Proc. 7th Annu. Work. Audio/Visual Emot. Chall. - AVEC '17*, pp. 69–76.
- [13] A. T. Beck, R. A. Steer, and G. K. Brown, "Beck depression inventory-II," *San Antonio*, vol. 78, no. 2, pp. 490–498, 1996.
- [14] A. T. Beck, *Depression : Causes and Treatment*. University of Pennsylvania Press, 1972.
- [15] L. R. Aiken, *Psychological Testing and Assessment, 4th edition*. Allyn & Bacon, 1982.
- [16] V. Kazemi and J. Sullivan, "One millisecond face alignment with an ensemble of regression trees," in *Proc. IEEE Comput. Soc. Conf. Comput. Vis. Pattern Recognit.*, 2014, pp. 1867–1874.

- [17] C. Sagonas, E. Antonakos, G. Tzimiropoulos, S. Zafeiriou, and M. Pantic, “300 Faces In-The-Wild Challenge : database and results,” *Image Vis. Comput.*, pp. 3–18.
- [18] D. E. King, “Dlib-ml : A Machine Learning Toolkit,” *J. Mach. Learn. Res.*, pp. 1755–1758.
- [19] A. Rosebrock. (2017) Face Alignment using OpenCV and Python. Accessed 2018-05-24. [Online]. Available : <https://www.pyimagesearch.com/2017/05/22/face-alignment-with-opencv-and-python/>
- [20] G. Bradski, “The OpenCV Library,” *Dr Dobbs J. Softw. Tools*, pp. 120–125.
- [21] G. A. Carpenter and S. Grossberg, “Adaptive Resonance Theory,” *Handb. brain theory neural networks*, pp. 87–90, 2003.
- [22] M. Biehl, B. Hammer, and T. Villmann, “Prototype-based models in machine learning,” *Wiley Interdiscip. Rev. Cogn. Sci.*, vol. 7, no. 2, pp. 92–111, 2016.
- [23] S. Grossberg, “Adaptive Resonance Theory : How a brain learns to consciously attend, learn, and recognize a changing world,” *Neural Networks*, pp. 1–47.
- [24] A. P. Azcarraga, “Modèles neuronaux pour la classification incrémentale de formes visuelles,” Ph.D. dissertation, Genoble INPG.
- [25] D. Puzenat, “Parallélisme et modularité des modèles connexionnistes,” p. 176.
- [26] S. Cholet and H. Paugam-Moisy, “Démonstration du diagnostic automatique de l ’ état dépressif,” in *Conférence Natl. en Intell. Artif.* Nancy : Plateforme Intelligence Artificielle, 2018, p. To Appear.
- [27] S. Yu, S. Scherer, D. Devault, J. Gratch, G. Stratou, L. P. Morency, and J. Cassel, “Multimodal prediction of psychosocial disorders : Learning verbal and non-verbal commonalities in adjacent pairs,” in *Proc. 17th Work. Semant. Pragmat. Dialogue*, 2013, pp. 160–169.
- [28] E. Reynaud, A. Crépet, H. Paugam-Moisy, and D. Puzenat, “A computational model for binding sensory modalities,” in *Abstr. Conscious. Cogn.* Academic Press, 2000, ch. 9, pp. 97–88.
- [29] H. Paugam-Moisy and E. Reynaud, “Multi-network system for sensory integration,” *Int. Jt. Conf. Neural Networks, Vols 1-4, Proc.*, vol. 1-4, no. February 2001, pp. 2343–2348, 2001.
- [30] T. E. Oliphant, *A Guide to Numpy*. Trelgol Publishing, 2006.
- [31] E. Jones, T. Oliphant, P. Peterson *et al.*, “SciPy : Open source scientific tools for Python,” 2001–. [Online]. Available : <http://www.scipy.org/>
- [32] F. Pedregosa, G. Varoquaux, A. Gramfort, V. Michel, B. Thirion, O. Grisel, M. Blondel, P. Prettenhofer, R. Weiss, V. Dubourg, J. Vanderplas, A. Passos,
- D. Cournapeau, M. Brucher, M. Perrot, and E. Duchesnay, “Scikit-learn : Machine learning in Python,” *Journal of Machine Learning Research*, vol. 12, pp. 2825–2830, 2011.

Analyse d'opinions multi-aspects pour la recommandation fine de restaurants

Isabelle Tellier¹

Hamid Hammouche²

Didier Cholvy²

Jean-Baptiste Tanguy²

¹ Lattice, CNRS (UMR8094), ENS Paris, Université Sorbonne Nouvelle - Paris 3

² advanced decision, 243 B Boulevard Pereire, 75017 Paris

isabelle.tellier@sorbonne-nouvelle.fr, hamid.hammouche@advanceddecision.fr

didier.cholvy@advanceddecision.fr, jean-baptiste.tanguy@advanceddecision.fr

Résumé

Dans cet article, nous présentons la tâche d'analyse d'opinions multi-aspects et son application pour la recommandation dans un cadre industriel. Les textes analysés concernent des avis en français sur des restaurants, postés sur Internet. Ces textes ont été manuellement annotés conformément à une ontologie du domaine, pour en extraire les différentes facettes de l'expérience utilisateur décrite et qualifier leur polarité. Plusieurs expériences ont ensuite été menées pour apprendre automatiquement à reproduire ce type d'annotation sur de nouveaux textes. L'utilisation de deux CRF successifs ont permis d'atteindre les meilleurs résultats.

Mots Clés

traitement automatique des langues, analyse multi-aspects, fouille d'opinion, CRF

Abstract

In this paper, we present the task of multi-aspect opinion mining and its application for recommendation in an industrial context. The texts used are comments in French about restaurants, posted on the Internet. These texts have been manually labeled according to a domain ontology, so as to extract the various aspects of the user experience and attach a polarity to them. Several Machine Learning experiments have been conducted to reproduce this annotation on new texts. The best results have been obtained by using two consecutive CRFs.

Keywords

natural language processing, multi-aspect analysis, opinion mining, CRF

1 Introduction

L'objectif final de nos travaux est de proposer un service industriel de recommandation personnalisée de produits touristiques. Ce service s'appuiera notamment sur l'analyse automatique de commentaires utilisateurs postés sur Internet concernant ces produits, en y recherchant les "expériences utilisateurs" fines qui y sont décrites. Cette dé-

marche correspond à la tâche d'*extraction automatique d'opinions multi-aspects* dans des textes. Dans un premier temps, nous nous sommes concentrés sur le domaine de la restauration. Dans ce cadre, les différents "aspects" pertinents pouvant apparaître dans un commentaire sont par exemple l'ambiance, le cadre, la cuisine, le service, etc. Chacun d'eux décrit une facette de l'expérience utilisateur et peut être qualifié positivement ou négativement.

Nous avons constitué un corpus original en français de commentaires et l'avons manuellement annoté en distinguant 12 aspects distincts et en associant une polarité (au minimum : positive, négative ou neutre) à chaque occurrence de ces aspects. Nous avons ensuite utilisé ce corpus pour procéder à des expériences d'apprentissage automatique visant à reproduire cette annotation sur de nouveaux textes. Comme cela est traditionnellement fait dans le domaine, les deux étapes d'identification des aspects et de qualification de leur polarité ont été traitées séquentiellement : la première par un CRF, la deuxième en mettant en concurrence un SVM et un nouveau CRF. Sur nos données, c'est la deuxième approche qui s'est avérée la plus performante.

Dans la suite de l'article, nous présentons tout d'abord rapidement l'état de l'art du domaine. Nous décrivons ensuite le corpus que nous avons constitué, puis les expériences que nous avons menées et leurs résultats.

2 Etat de l'art

L'analyse d'opinion multi-aspects est une sous-tâche bien identifiée du Traitement Automatique des Langues [4, 6, 5], elle a notamment fait l'objet d'une compétition SemEval en 2014 [9]¹, 2015 [8]² et 2016 [7]³. Elle est traditionnellement traitée en deux étapes distinctes : identification des "aspects" d'abord puis qualification de leur polarité. Pour l'identification des aspects, quelques approches à base de patrons linguistiques manuellement définis ont été testées, mais ce sont bien sûr les méthodes d'apprentissage automatique supervisé qui ont depuis quelques années pris

1. <http://alt.qcri.org/semeval2014/task4/>

2. <http://alt.qcri.org/semeval2015/task12/>

3. <http://alt.qcri.org/semeval2016/task5/>

le dessus. Comme un "aspect" est un empan de texte (généralement supposé contigu) pouvant couvrir plusieurs mots, ce sont les techniques d'annotation de séquences qui sont les plus adaptées, avec un codage BIO (Begin/In/Out) pour bien en délimiter les frontières. Les CRF [12], et plus récemment les réseaux de neurones convolutifs [10] ont montré leur efficacité en la matière. Associer une polarité à un aspect est une tâche de classification simple, qui a été traitée notamment par de la régression logistique [3], des SVM [1] ou des réseaux neuronaux [11].

3 Constitution du corpus annoté

Lors de SemEval 2016, une partie des données proposées étaient en français et portaient sur des restaurants [2]. Les données d'entraînement et de test de ce corpus comprennent en tout 457 textes correspondant à 2365 phrases. Notre corpus a été développé indépendamment⁴. Il contient 8381 avis issus de différentes sources. Ces avis portent sur des restaurants parisiens, ils correspondant à 46213 phrases : un avis a en moyenne 6 phrases, 75 unités (mots, émoticônes...) dont 51 distincts. Les annotations des deux corpus ne sont pas exactement les mêmes :

- le corpus de SemEval est annoté au niveau des phrases, une propriété "target" permet de citer la portion de texte correspondant à l'aspect annoté : cette propriété est parfois vide, parfois dupliquée dans les annotations. Notre corpus est annoté au niveau des unités linguistiques, chaque unité ne peut donc recevoir qu'une seule étiquette.
- les étiquettes du corpus SemEval sont des paires Entité-Attribut (6 entités, auxquelles peuvent être attachés certains attributs : 12 paires possibles en tout). Notre jeu d'étiquettes, inspiré d'ontologies du domaine, est aussi composé de 12 unités, mais elles ne coïncident que partiellement (certaines sont plus orientées "avis" qu'"aspect" proprement dit) : Ambiance, Astuce, Cadre, Boisson, Cuisine-Générale, Nourriture, Emplacement-Géographique, Public-Cible, Qualité-Prix, Service, Opinion générale et Recommandation.
- la polarité dans SemEval peut valoir "positif", "négatif" ou "neutre", nous avons ajouté les modalités "très positif" et "très négatif".

Pour toutes ces raisons, les résultats de nos expériences seront difficilement comparables avec ceux de la littérature.

Pour préparer nos expériences, des pré-traitements linguistiques ont été appliqués sur nos textes : segmentation, étiquetage POS (Treetagger pour l'oral et pour l'écrit, SEM⁵, Talismane⁶), lemmatisation (Talismane), chunking et entités nommées (SEM), analyse syntaxique (Talismane). Des lexiques spécialisés (termes de cuisine par exemple) ainsi que des lexiques d'opinion ont également été constitués.

4. il ne peut malheureusement pas être diffusé librement pour des raisons de confidentialité industrielle

5. <http://www.lattice.cnrs.fr/sites/itellier/SEM.html>

6. <http://redac.univ-tlse2.fr/applications/talismane.html>

4 Expériences

Pour ces expériences, nous avons systématiquement utilisé le protocole de la validation croisée à 5 plis.

Conformément à la littérature, nous avons d'abord cherché à retrouver les "aspects", indépendamment de leur polarité. Nous avons pour cela utilisé des CRF, tels qu'implémentés dans Wapiti⁷. Les textes ont été traités phrase par phrase, mais en gardant en indice leur position dans le texte. Les "patrons" (templates) permettant de générer les fonctions caractéristiques les plus efficaces comprenaient la prise en compte des tokens courant, précédant et suivant, d'expressions régulières pour les chiffres et les monnaies, des lexiques. Les meilleurs résultats obtenus sont mesurés en micro-moyennes des précisions (P), rappels (R) et F1-mesures (F1) à différents niveaux de granularité :

- pour chaque étiquette distincte : P=73,52, R=72,76, F1=73,13
- pour chaque "segment de texte" (commençant par un B d'un certain type, se prolongeant éventuellement par des I de même type) : P=57,59, R=54,23, F1=55,85
- par segment sans tenir compte de leur type : P=70,72, R=66,27, F1=68,41.

Pour identifier les polarités (classification parmi les 5 classes possibles), nous avons testé des SVM et, de nouveau, des CRF. Nous avons d'abord réalisé des expériences en utilisant les segments de texte "Gold". Pour les SVM⁸, les segments de textes ont été transformés en vecteurs d'indices comprenant divers nombres d'occurrences et proportions (des différentes catégories de mots lexicaux, mots appartenant aux lexiques d'opinion, ponctuations, mots en majuscules, indices de négation). Le meilleur SVM testé (noyau RBF, stratégie "un contre tous", vote pondéré par la précision du modèle associé à l'étiquette proposée) a atteint P=39,56, R=43,3, F1=41,19. Ce SVM n'a pas directement été confronté au vocabulaire des segments, ce qui explique sans doute ses faibles performances. Quant au CRF, il n'a pas à réaliser de segmentation : c'est donc en fait plutôt un modèle de "maximum d'entropie" qui a été appris. Les attributs utilisés se sont réduits aux tokens, aux lemmes, aux POS et à l'appartenance aux lexiques d'opinion. Les résultats ont alors atteint : P=R=F1=74,79, et même P=R=F1=80,83 quand on se ramène à 3 polarités au lieu des 5 initiales. Mis bout à bout, les meilleurs modèles de deux tâches ont la performance suivante (par bloc de segments typés) : P=35,18, R=35,27, F1=35,23.

5 Conclusion

Nous avons reporté dans cet article des expériences préliminaires visant à extraire des informations fines de commentaires en français. La tâche est difficile, nous l'avons évaluée de la façon la plus exigeante possible. Les CRF ont pour l'instant été privilégiés, pour leur utilisation facile de ressources linguistiques externes. Nous comptons par la suite tester des alternatives neuronales à nos modèles.

7. <https://wapiti.limsi.fr/>

8. avec la librairie Python sklearn

Références

- [1] T. Álvarez López, J. Juncal-Martínez, M. Fernández-Gavilanes, E. Costa-Montenegro, and F. J. González-Castaño. Gti at semeval-2016 task 5 : Svm and crf for aspect detection and unsupervised aspect-based sentiment analysis. In *Proceedings of the 10th International Workshop on Semantic Evaluation (SemEval-2016)*, pages 306–311, San Diego, California, June 2016. Association for Computational Linguistics.
- [2] M. Apidianaki, X. Tannier, and C. Richart. Datasets for aspect-based sentiment analysis in french. In *Proceedings of the Tenth International Conference on Language Resources and Evaluation LREC 2016, Portorož, Slovenia, May 23-28, 2016.*, 2016.
- [3] H. Hamdan. Sentisys at semeval-2016 task 5 : Opinion target extraction and sentiment polarity detection. In *Proceedings of the 10th International Workshop on Semantic Evaluation, SemEval@NAACL-HLT 2016, San Diego, CA, USA, June 16-17, 2016*, pages 350–355, 2016.
- [4] M. Hu and B. Liu. Mining and summarizing customer reviews. In *Proceedings of the Tenth ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, Seattle, Washington, USA, August 22-25, 2004*, pages 168–177, 2004.
- [5] B. Liu. *Sentiment Analysis and Opinion Mining*. Synthesis Lectures on Human Language Technologies. Morgan & Claypool Publishers, 2012.
- [6] B. Pang and L. Lee. Opinion mining and sentiment analysis. *Foundations and Trends in Information Retrieval*, 2(1-2) :1–135, 2007.
- [7] M. Pontiki, D. Galanis, H. Papageorgiou, I. Androutsopoulos, S. Manandhar, M. Al-Smadi, M. Al-Ayyoub, Y. Zhao, B. Qin, O. D. Clercq, V. Hoste, M. Apidianaki, X. Tannier, N. V. Loukachevitch, E. Kotelnikov, N. Bel, S. M. J. Zafra, and G. Eryigit. Semeval-2016 task 5 : Aspect based sentiment analysis. In *Proceedings of the 10th International Workshop on Semantic Evaluation, SemEval@NAACL-HLT 2016, San Diego, CA, USA, June 16-17, 2016*, pages 19–30, 2016.
- [8] M. Pontiki, D. Galanis, H. Papageorgiou, S. Manandhar, and I. Androutsopoulos. Semeval-2015 task 12 : Aspect based sentiment analysis. In *Proceedings of the 9th International Workshop on Semantic Evaluation, SemEval@NAACL-HLT 2015, Denver, Colorado, USA, June 4-5, 2015*, pages 486–495, 2015.
- [9] M. Pontiki, D. Galanis, J. Pavlopoulos, H. Papageorgiou, I. Androutsopoulos, and S. Manandhar. Semeval-2014 task 4 : Aspect based sentiment analysis. In *Proceedings of the 8th International Workshop on Semantic Evaluation, SemEval@COLING 2014, Dublin, Ireland, August 23-24, 2014.*, pages 27–35, 2014.
- [10] S. Poria, E. Cambria, and A. F. Gelbukh. Aspect extraction for opinion mining with a deep convolutional neural network. *Knowl.-Based Syst.*, 108 :42–49, 2016.
- [11] D. Tang, B. Qin, and T. Liu. Aspect level sentiment classification with deep memory network. In *Proceedings of the 2016 Conference on Empirical Methods in Natural Language Processing, EMNLP 2016, Austin, Texas, USA, November 1-4, 2016*, pages 214–224, 2016.
- [12] Z. Toh and W. Wang. DLIREC : aspect term extraction and term polarity classification system. In *Proceedings of the 8th International Workshop on Semantic Evaluation, SemEval@COLING 2014, Dublin, Ireland, August 23-24, 2014.*, pages 235–240, 2014.

A De Novo Robust Clustering Approach for Amplicon-Based Sequence Data

Alexandre Bazin

Didier Debroas

Engelbert Mephu Nguifo

contact@alexandrebazin.com

1 Introduction

Studying the structure of the communities in an ecosystem is central in environmental microbiology [8, 14]. The biosphere’s diversity can be determined by amplifying and sequencing specific phylogenetic markers (e.g. 16S rRNA). From there, these amplicons need to be clusterized in “species” named Operational Taxonomic Units (OTUs) [4, 9, 11, 15]. As the volume of sequences has drastically increased in recent times, new clustering tools have emerged to treat the data in reasonable time. The currently used algorithms are, from the point of view of algorithmic complexity, the fastest available that do not produce random results. However, due to their simplicity, the reliability of the results are often discussed. These tools being essentially black boxes, their sensitivity to the sequence order, clustering threshold and structure of the data makes it that the users have no way of knowing whether better Operational Taxonomic Units (OTUs) could have been obtained with different parameters or even whether they correctly represent the data. In these circumstances, there is no choice but to blindly trust them.

Distance-based greedy clustering algorithms such as the ones implemented in OTUclust [1], VSEARCH [13], CD-HIT [10] or USEARCH [5] all share the same base naive algorithm. While more sophisticated algorithms [3, 6, 12, 7, 2] could produce better results quality-wise, their runtime would render them unusable on millions of sequences. As the quality of the OTUs is important, we have to find a way to improve it without increasing the runtime. The different available implementations use a variety of heuristics to counterbalance the simplicity of the algorithm but, to the best of our knowledge, no approach has tried to add a measure of uncertainty to the process. This is why, in order to help increase the quality and trustworthiness of the clustering, we propose to add uncertainty to this simple algorithm through the use of fuzzy clustering.

2 Adding uncertainty to clustering

Distance-based greedy clustering algorithms, such as the one in VSEARCH, produce a number of OTUs and assign each sequence to one of them. The OTU to which a sequence is said to belong to is usually the first one to be encountered that is sufficiently close, i.e. within the specified threshold. This makes a sequence belong only to a single OTU and OTUs either completely include or exclude a sequence. While these would not be problems were the clustering optimal, the need for fast algorithms gives rise to results that are not always trustworthy. The OTUs being presented as absolute, the end user has no choice, should consider them correct and cannot know whether the algorithm has encountered ambiguity. We believe that being less strict in the way the OTUs partition sequences would help produce better results from the end user’s point of view. To help increase the quality of the clustering and maximize the information that can be gathered from the data, we propose to add uncertainty to the clustering by means of fuzzy sets.

Using fuzzy OTUs allows us to discern the difference between sequences close to the OTU and sequences extremely far. Using the parameters t_1 and t_2 , we can tune the “detection radius” around OTUs to gather information that would normally be discarded by the clustering algorithm.

3 Evaluating fuzzy OTUs

An ideal OTU would contain only sequences with a membership value of 1, meaning a group of sequences has been perfectly regrouped with a good threshold and no sequence lies ambiguously on the border. More realistically, a good OTU would contain many sequences with high membership values and little sequences with low values. A bad OTU with the majority of its sequences having low membership values could mean that the algorithm has chosen as a center a sequence on the border of a group or, even worse, between two distinct groups.

We can quickly evaluate the quality of an OTU with this repartition. If we suppose that each sequence lowers the quality of the OTU depending on its membership value, we can use the following formula :

$$Quality(OTU) = 1 - \sum_{i=1}^9 \omega_i \times \frac{\# \text{ sequences with membership value } i \times 0.1}{\# \text{ sequences in the OTU}}$$

with ω_i being the “cost” of having a sequence with membership value $i \times 0.1$.

A problem arises with singletons that always have perfect quality but these can safely be treated separately.

A sequence can belong to multiple OTUs due to fuzzy membership. However, in the end, we want each sequence to be assigned to a single OTU. Hence, we have to choose one of the possible OTUs. We have two types of values left from the clustering process : membership and quality. The first one is based on the distance between the OTU and the sequence and the second one is used to recognize bad OTUs. Choosing the OTU with the best membership value is akin to running VSEARCH. Choosing the OTU with the best quality tends to create bigger OTUs that absorb distant sequences. To better compromise, we can use a linear combination of both values :

$$\alpha \times quality + \beta \times membership$$

Increasing the importance of the quality reduces the number of OTUs containing sequences. When α is low, the “best” OTUs quality-wise absorb very close sequences that would have been attributed to other OTUs. When α gets too high, the best OTUs start absorbing all the sequences around them, effectively acting like an increase of the distance threshold.

4 Experimental Results

4.1 Data

We used our algorithm on a dataset containing 5977 sequences of length between 900 and 3081 for an average of 1442 and taxonomies extracted from the SILVA database. We used a threshold of 0.97 (97% similarity) for determining new OTUs and a threshold of 0.95 for fuzzy membership. For the choice of the OTU for each sequence, we present the results of three strategies : best quality ($\alpha = 1$ and $\beta = 0$), compromise ($\alpha = 0.5$ and $\beta = 0.5$) and distance ($\alpha = 0$ and $\beta = 1$). The comparison with VSEARCH is done using identical parameters when applicable.

The program, dataset and corresponding taxonomy are available on <http://projets.isima.fr/sclust/Expe.html>.

4.2 Relevant Metrics

To measure the effects of introducing uncertainty to the clustering, we consider the computation time, memory usage, number of OTUs, singletons and pairs and average distance in the taxonomy tree between sequences in the same cluster.

4.3 Analysis

Results show that the choice strategy affects every metric relevant to the quality of the clustering : number of OTUs, singletons and pairs, average misclassification. The fuzzy approach uses slightly more memory than VSEARCH but all choice strategies are similar on this metric. When using the default *-maxaccepts* and *-maxrejects* values, computation time is lower for VSEARCH. However, when using higher values for these parameters – and thus more precise clustering – the computation time is the same for both approaches. We observe that increasing the importance of the quality in the OTU choice strategy lowers the final number of OTUs. This is due to the fact that some OTUs are initially created centered on isolated sequences near good OTUs. That isolation lowers their quality and the good OTUs absorb their sequences.

Using the quality also lowers the number of singletons and increases the number of pairs. This most likely means that singletons were created close to either good clusters or one another. The fuzzy approach allows the algorithm to merge those sequences that were slightly too far from the center with their corresponding OTU. The increase in the number of pairs appears to be due to the merging of singletons lying too close to one another. The average taxonomy distance in OTUs is shown to vary wildly. Using only the quality to choose OTUs increases this number as the “best” OTUs attract all the sequences in their fuzzy surroundings. This causes some sequences belonging to different species to be classified together. However, using a compromise between quality and distance lowers this metric as the best clusters only absorb sequences that are sufficiently close to them and should probably be together while rejecting the sequences that are too different.

acknowledgements

This work was supported by the European Union’s “*Fonds Européen de Développement Régional (FEDER)*” program and the Auvergne-Rhone-Alpes region.

Method	Time (min)	Memory	#OTUs	#Singletons	#Doubletons	Distance
Fuzzy (best quality)	1:06	652744	3461	2581	442	0.75
Fuzzy (compromise)	1:06	651980	3596	2776	413	0.54
Fuzzy (distance)	1:06	683772	3631	2837	395	0.59
VSEARCH	0:21	632832	3716	2935	388	0.57

Table 1: Results of the clustering.

References

- [1] Davide Albanese, Paolo Fontana, Carlotta De Filippo, Duccio Cavalieri, and Claudio Donati. Micca: a complete and accurate software for taxonomic profiling of metagenomic data. *Scientific reports*, 5:9743, 2015.
- [2] Violaine Antoine, Benjamin Quost, Marie-Hélène Masson, and Thierry Denoeux. CECM: constrained evidential c-means algorithm. *Computational Statistics & Data Analysis*, 56(4):894–914, 2012.
- [3] Violaine Antoine, Benjamin Quost, Marie-Hélène Masson, and Thierry Denoeux. CEVCLUS: evidential clustering with instance-level constraints for relational data. *Soft Comput.*, 18(7):1321–1335, 2014.
- [4] Wei Chen, Clarence K Zhang, Yongmei Cheng, Shaowu Zhang, and Hongyu Zhao. A comparison of methods for clustering 16s rna sequences into otus. *PloS one*, 8(8):e70837, 2013.
- [5] Robert C Edgar. Search and clustering orders of magnitude faster than blast. *Bioinformatics*, 26(19):2460–2461, 2010.
- [6] Isak Gath and Amir B. Geva. Unsupervised optimal fuzzy clustering. *IEEE Transactions on pattern analysis and machine intelligence*, 11(7):773–780, 1989.
- [7] Sarra Ben Hariz, Zied Elouedi, and Khaled Melouli. Clustering approach using belief function theory. In *International Conference on Artificial Intelligence: Methodology, Systems, and Applications*, pages 162–171. Springer, 2006.
- [8] Mylène Hugoni, Najwa Taib, Didier Debroas, Isabelle Domaizon, Isabelle Jouan Dufournel, Gisèle Bronner, Ian Salter, Hélène Agogué, Isabelle Mary, and Pierre E Galand. Structure of the rare archaeal biosphere and seasonal dynamics of active ecotypes in surface coastal waters. *Proceedings of the National Academy of Sciences*, 110(15):6004–6009, 2013.
- [9] Weizhong Li, Limin Fu, Beifang Niu, Sitao Wu, and John Wooley. Ultrafast clustering algorithms for metagenomic sequence analysis. *Briefings in bioinformatics*, page bbs035, 2012.
- [10] Weizhong Li and Adam Godzik. Cd-hit: a fast program for clustering and comparing large sets of protein or nucleotide sequences. *Bioinformatics*, 22(13):1658–1659, 2006.
- [11] Frédéric Mahé, Torbjørn Rognes, Christopher Quince, Colomban de Vargas, and Micah Dunthorn. Swarm: robust and fast clustering method for amplicon-based studies. *PeerJ*, 2:e593, 2014.
- [12] Airel Pérez-Suárez, José F Martínez-Trinidad, Jesús A Carrasco-Ochoa, and José E Medina-Pagola. Oclustr: A new graph-based algorithm for overlapping clustering. *Neurocomputing*, 121:234–247, 2013.
- [13] Torbjørn Rognes, Tomáš Flouri, Ben Nichols, Christopher Quince, and Frédéric Mahé. Vsearch: a versatile open source tool for metagenomics. *PeerJ*, 4:e2584, 2016.
- [14] Simon Roux, Michaël Faubladier, Antoine Mahul, Nils Paulhe, Aurélien Bernard, Didier Debroas, and François Enault. Metavir: a web server dedicated to virome analysis. *Bioinformatics*, 27(21):3074–3075, 2011.
- [15] Sarah L Westcott and Patrick D Schloss. De novo clustering methods outperform reference-based methods for assigning 16s rna gene sequences to operational taxonomic units. *PeerJ*, 3:e1487, 2015.

Démonstration du diagnostic automatique de l'état dépressif

S. Cholet

H. Paugam-Moisy

Laboratoire de Mathématiques Informatique et Applications (LAMIA - EA 4540)

Université des Antilles, Campus de Fouillole - Guadeloupe

Stephane.Cholet@univ-antilles.fr

Résumé

Les troubles psychosociaux sont un problème de santé publique majeur, pouvant avoir des conséquences graves sur le court ou le long terme, tant sur le plan professionnel que personnel ou familial. Le diagnostic de ces troubles doit être établi par un professionnel. Toutefois, l'IA (l'Intelligence Artificielle) peut apporter une contribution en fournissant au praticien une aide au diagnostic, et au patient un suivi permanent rapide et peu coûteux. Nous proposons un outil d'aide au diagnostic automatique de l'état dépressif à partir d'observations du visage en temps réel, au moyen d'une simple webcam. A partir de vidéos du challenge AVEC'2014, nous avons entraîné un classifieur neuronal à extraire des prototypes de visages selon différentes valeurs du score de dépression de Beck (BDI-II).

1 Introduction

La démonstration est associée à l'article "Diagnostic automatique de l'état dépressif" présenté à la conférence CNIA'2018. Des visages dont les points d'intérêt sont extraits à partir de vidéos sont associés à un score mesurant l'état dépressif du sujet. Le système permettant cette association est un classifieur neuronal incrémental dont l'apprentissage a été réalisé sur les données du challenge AVEC'2014 [1]. Le présent article décrit comment ce classifieur peut être mis en œuvre pour estimer, en temps réel, l'état dépressif d'un sujet filmé par une webcam.

Le système d'estimation de l'état dépressif est composé de quatre phases qui s'enchaînent comme décrit Figure 1 : la capture de la vidéo est suivie de l'extraction des descripteurs, du traitement par le classifieur et de la présentation du diagnostic. On se focalise ici sur la description et le fonctionnement des procédés mis en œuvre pour la classification. L'attention du lecteur est attirée sur le fait que cette chaîne de traitement est encore à l'état de prototype, en phase active de développement et en attente de validation, en particulier par des experts psychiatres.

2 Chaîne des traitements

2.1 Capture vidéo

La capture vidéo est réalisée au moyen d'une webcam connectée à l'ordinateur utilisé pour la classification. Afin

d'optimiser le traitement du flux, les images sont capturées par un thread dédié et bufferisées dans une file d'attente. De plus, la taille des images est réduite à 300x300 pixels avant leur traitement. La bufferisation des images assure la continuité de la capture pendant leur traitement. La fréquence de capture est de 25 images par secondes.

2.2 Prétraitement des images

Dans la mesure où l'état dépressif s'évalue sur la durée, on ne traite que n images par seconde (malgré une fréquence de capture fixe). Ce paramètre est notamment utilisé pour optimiser la fluidité du traitement global. Chaque image fait l'objet d'une recherche de visage et d'estimation de la position des points d'intérêts via la méthode de Kazemi et Sullivan implémentée dans dLib.

L'alignement des visages consiste en une série de transformations géométriques ayant les objectifs suivants :

- que les visages aient la même taille ;
- que les visages soient centrés dans l'image ;
- que les yeux soient alignés horizontalement.

Cette phase, capitale, assure aux données l'homogénéité nécessaire pour la classification. Afin de réduire la durée de l'alignement, ce dernier est réalisé directement sur les points extraits et non sur les images

La sortie du prétraitement est, pour chaque image, un vecteur de 136 composantes comprenant les abscisses et ordonnées des 68 points d'intérêt extraits.

2.3 Classification

Les données sont classifiées au moyen d'un classifieur neuronal incrémental à base de prototypes. C'est un réseau de neurones à trois couches, la première recevant les entrées, la seconde étant constituée de "neurones-distance", i.e. des prototypes, qui sont totalement connectés aux neurones d'entrée. A chaque présentation d'un exemple, ce dernier est comparé à tous les prototypes en mémoire. Dans la troisième couche, plusieurs prototypes sont associés à un neurone de sortie, chacun représentant un score dépressif BDI-II.

La démonstration dont est l'objet ce papier n'utilise qu'en phase de généralisation un classifieur déjà entraîné, donc il n'y a pas création de nouveaux prototypes. Les règles de fonctionnement du classifieur permettent de sortir, ou non, une classe de score dépressif. La possibilité d'obtenir une

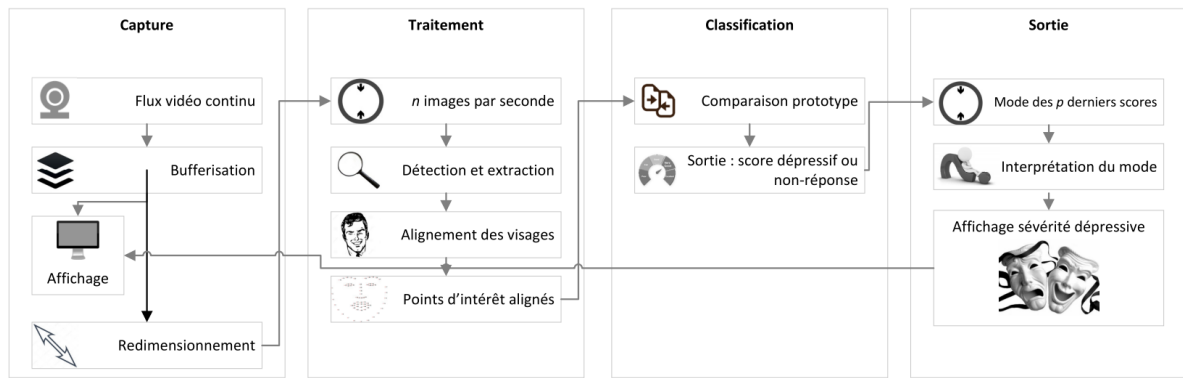


FIGURE 1 – Schéma d'ensemble du système d'estimation de l'état dépressif

TABLE 1 – Interprétation du score au test BDI-II

Score obtenu	Sévérité de la dépression
0-13	Minimale
14-19	Moyenne
20-28	Modérée
29-63	Sévère

"non-réponse" est une spécificité précieuse de ce classifieur qui le rend plus proche d'un diagnostic humain. Lors de la classification en temps réel, les non-réponses peuvent permettre de moduler la décision. Ainsi, le nombre de non-réponses rendues peut permettre le calcul d'une mesure de fiabilité de l'interprétation du score dépressif.

2.4 Diagnostic

La sortie du système est l'interprétation du score (voir Table 1) le plus représenté lors de la classification des images sur les p scores les plus récents. Le paramètre p est à ajuster en accord avec la durée de l'interaction. Les résultats en généralisation sur le corpus AVEC'2014 atteignent un taux de succès de 94% pour la classification de l'intervalle dépressif, contre 90% pour la classification du score BDI-II, d'où le choix de cette sortie qui se veut plus précise et interprétable en l'état.

3 Scénario

On présente ici un cas d'utilisation simple du système de classification de l'état dépressif, où l'utilisateur lit un passage affiché à l'écran pendant qu'il est filmé par une webcam.

Le nombre d'images traitées par seconde est fixé à $n = 10$, et la durée de l'enregistrement à 20 secondes. On fixe à $p = 20$ le nombre de scores pris en comptes dans l'affichage du résultat.

Quelques secondes après le démarrage de la procédure, l'interprétation de l'état dépressif est affichée. On présente également un indicateur de détection du visage, faisant savoir à l'utilisateur s'il doit modifier sa posture face à la

caméra. La Figure 2 présente un aperçu de la sortie du système.

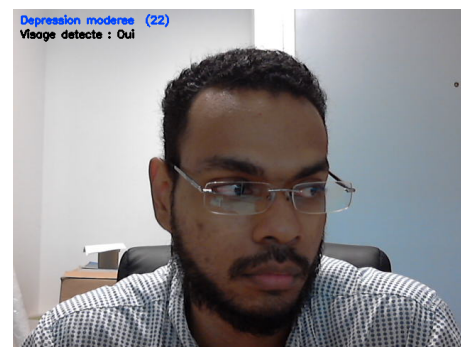


FIGURE 2 – Exemple capturé depuis la sortie du système de classification

4 Conclusion

Le système proposé permet de classifier en temps réel l'état dépressif d'un sujet humain au moyen d'un classifieur neuronal incrémental. Une chaîne de traitement rapide du flux vidéo est mise en œuvre afin de produire une interprétation de l'état dépressif, au regard du test BDI-II. Les résultats obtenus sont encourageant pour la poursuite du développement de l'outil. Toutefois, il est important de considérer avec précaution le résultat. D'une part, le système n'a pas encore pu bénéficier de l'expertise d'un professionnel de la psychiatrie. D'autre part, dans sa définition, le score BDI-II est évalué sur des sujets exprimant des symptômes dépressifs depuis au moins deux semaines, d'où l'importance de fixer un cadre d'utilisation du système en amont.

Références

- [1] M. Valstar, B. Schuller, K. Smith, T. Almaev, F. Eyben, J. Krajewski, R. Cowie, and M. Pantic, "Avec 2014 : 3d dimensional affect and depression recognition challenge," *Proc. 4th Int. Work. Audio/Visual Emot. Chall. - AVEC '14*, pp. 3–10, 2014.

ONTOREV : un moteur de révision d'ontologies OWL

Thinh Dong, Chan Le Duc, Myriam Lamolle

LIASD - EA4383, IUT de Montreuil, Université Paris8, France
{dong, leduc, lamolle}@iut.univ-paris8.fr

English abstract We propose a Web-based application for revising an OWL ontology. The main functionality of this application consists in changing as less as possible an initial ontology when adding a new piece of knowledge such as an axiom, an assertion. To be able to ensure such a minimal change in preserving consistency of the resulting ontology, we use a set of models generated by a tableau algorithm to characterize the semantics of an ontology, and define a semantic distance between ontologies. Our Web-based application provides also other services such as visualizing an OWL ontology in a succinct syntax, inputting an OWL axiom/assertion.

Introduction. L'intelligence collective s'appuie sur des connaissances partagées. Afin que ces connaissances soient exploitables via des systèmes informatiques, elles doivent être représentées et manipulées en utilisant un mécanisme de raisonnement automatique. En outre, les connaissances qui reflètent la compréhension d'un domaine d'application évoluent au cours du temps. Ceci nécessite un mécanisme assurant à tout moment la cohérence sémantique de la totalité des connaissances du domaine représenté ; modifier une seule information modélisée implique de réviser/vérifier toute la sémantique portée par cette connaissance. Pour répondre à ces deux problématiques, nous proposons dans cet article une approche pour des modèles ontologiques consistant en un mécanisme de révision permettant la représentation et la gestion de l'évolution des connaissances.

Après une étude des approches existantes pour réviser des ontologies, nous avons proposé un nouvel algorithme tableau pour des ontologies d'expressivité *SHIQ* (incluse dans OWL-DL) qui garantit la priorité aux nouvelles connaissances, la cohérence de la nouvelle ontologie et les changements minimaux. Nous avons implémenté ce nouvel algorithme dans le prototype ONTOREV mis en ligne par des services Web afin de favoriser le travail collaboratif et la gestion de l'évolution de la modélisation collective d'une ontologie.

Fondement formel. Comme mentionné précédemment, les ontologies étudiées dans cet article sont exprimées en la logique de description *SHIQ* (Baader *et al.*, 2003) qui autorise l'expression de concepts, atomiques ou complexes, composés de constructeurs logiques booléens, de rôles (relations binaires) transitifs et inverses, de restrictions existentielle, universelle et numérique. Une ontologie consiste en un ensemble de connaissances dont chacune peut être soit une *assertion* de la forme $C(a)$ ou $R(a, b)$ où C est un concept, R un rôle, a, b des individus, soit un *axiome* de la forme $C \sqsubseteq D$ ou $R \sqsubseteq S$ où C, D sont des concepts et R, S des rôles. Comme toutes les logiques de description, la sémantique de *SHIQ* est fondée sur la théorie des modèles, *i.e.* elle utilise un domaine non-vide Δ et une fonction d'interprétation $\mathcal{I}$ pour associer à chaque

concept C un ensemble $C^{\mathcal{I}} \subseteq \Delta$, à chaque rôle R un ensemble $R^{\mathcal{I}} \subseteq \Delta \times \Delta$. Le couple $\mathcal{I} = (\Delta, \cdot^{\mathcal{I}})$ est un modèle de l'ontologie (*i.e.* elle est cohérente) si $\mathcal{I}$ vérifie toute contrainte sémantique (ou connaissance) contenue dans l'ontologie ; par exemple $C^{\mathcal{I}} \subseteq D^{\mathcal{I}}$ pour tout axiome $C \sqsubseteq D$.

Soient $\mathcal{O}$ et $\mathcal{O}'$ deux ontologies en $\mathcal{SHIQ}$. La révision de $\mathcal{O}$ par $\mathcal{O}'$ consiste à construire une ontologie $\mathcal{O}^*$ conservant toutes les connaissances de $\mathcal{O}'$ et le plus possible les connaissances de $\mathcal{O}$. Nous avons adopté une approche fondée sur les modèles pour la construction de $\mathcal{O}^*$. Ceci veut dire que la sémantique d'une ontologie $\mathcal{O}$ est caractérisée par un ensemble fini de *modèles représentatifs*, noté $\text{FM}(\mathcal{O})$, qui est construit par l'algorithme de tableau. Par conséquent, la révision de $\mathcal{O}$ par $\mathcal{O}'$ revient à construire un ensemble de modèles, noté $\text{FM}(\mathcal{O} \oplus \mathcal{O}')$, qui inclut $\text{FM}(\mathcal{O}')$ et la partie la plus grande de $\text{FM}(\mathcal{O})$ compatible avec $\text{FM}(\mathcal{O}')$. En particulier, si $\mathcal{O} \cup \mathcal{O}'$ est cohérente alors $\text{FM}(\mathcal{O} \oplus \mathcal{O}') = \text{FM}(\mathcal{O}) \cup \text{FM}(\mathcal{O}')$, et $\mathcal{O}^* = \mathcal{O} \cup \mathcal{O}'$, *i.e.* la révision devient simplement l'ajout de $\mathcal{O}'$ dans $\mathcal{O}$. Les lecteurs intéressés peuvent trouver le détail de cet algorithme de révision dans l'article (Dong *et al.*, 2017).

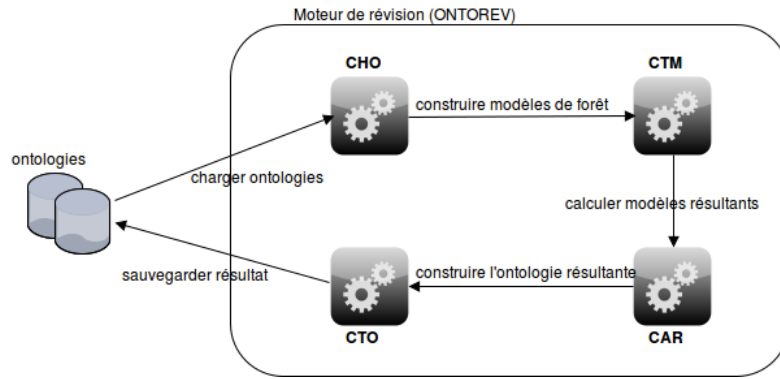


Fig. 1: Architecture d'ONTOREV

Architecture du noyau. L'algorithme de révision présenté dans la section précédente a été implémenté dans un prototype, appelé ONTOREV, et évalué avec plusieurs ontologies du monde réel. L'architecture d'ONTOREV est décrite dans la figure 1. Elle est composée de quatre modules principaux, à savoir, CHO qui charge une ontologie $\mathcal{O}$ à réviser et une autre ontologie $\mathcal{O}'$ contenant les assertions/axiomes à ajouter ; CTM qui construit les modèles de forêt de $\mathcal{O}$ et $\mathcal{O}'$; CAR qui calcule $\text{FM}(\mathcal{O} \oplus \mathcal{O}')$ représentant les modèles de l'ontologie résultante à partir de $\text{FM}(\mathcal{O})$ et $\text{FM}(\mathcal{O}')$; enfin, CTO qui construit l'ontologie résultante $\mathcal{O}^*$ à partir de $\text{FM}(\mathcal{O} \oplus \mathcal{O}')$ obtenu.

Le module CHO, qui dépend fortement de l'OWLAPI¹, doit effectuer également certains pré-traitements comme l'absorption, la suppression de tautologies, etc. Le module CTM implémente un algorithme tableau spécifique (Dong *et al.*, 2017) permettant de construire l'ensemble des modèles représentatifs $\text{FM}(\mathcal{O})$ d'une ontologie $\mathcal{O}$.

¹ <http://owlapi.sourceforge.net/>

L'ensemble $FM(\mathcal{O})$ peut être très grand s'il y a un nombre important de *non-déterminismes intrinsèques* (e.g. ceux qui ne proviennent pas directement du concept $\neg C \sqcup D$ pour un axiome $C \sqsubseteq D$). Le module CAR implémente l'algorithme calculant une distance sémantique entre deux modèles de forêt et plusieurs techniques d'optimisation pour réduire la cardinalité de $FM(\mathcal{O} \oplus \mathcal{O}')$. Enfin, le module CTO génère l'ontologie résultante en *SHIQ* la plus "petite" sémantiquement qui prend pour modèles tous les éléments dans $FM(\mathcal{O} \oplus \mathcal{O}')$. Remarquons que l'existence d'une ontologie en *SHIQ* prenant *exactement* les éléments dans $FM(\mathcal{O} \oplus \mathcal{O}')$ pour modèles n'est pas nécessaire.

Exemple. Pour illustrer le processus de révision, nous considérons une ontologie, notée $\mathcal{O}$, contenant les axiomes et assertions suivants :

- (1) $\text{Professor} \sqsubseteq \text{Researcher} \sqcup \text{Expert}$ (les professeurs sont des chercheurs ou experts),
- (2) $\text{Professor} \sqsubseteq \exists \text{supervises. Student}$ (un professeur supervise au moins un étudiant),
- (3) $\text{Professor} \sqsubseteq (\geq 2 \text{ teaches. Course})$ (un professeur enseigne au moins deux cours),
- (4) $\text{Professor}(\text{Alex})$ (Alex est un professeur).

On souhaite maintenant réviser $\mathcal{O}$ en prenant en compte le fait que les chercheurs et experts ne supervisent aucun étudiant. On note $\mathcal{O}'$ une deuxième ontologie contenant deux nouveaux axiomes δ_1 et δ_2 qui expriment les nouvelles connaissances :

$\delta_1 : \text{Researcher} \sqsubseteq \forall \text{supervises. } (\neg \text{Student})$ et $\delta_2 : \text{Expert} \sqsubseteq \forall \text{supervises. } (\neg \text{Student})$

Si l'on ajoute $\mathcal{O}'$ dans $\mathcal{O}$ directement, l'union $\mathcal{O} \cup \mathcal{O}'$ sera incohérente car selon (4) Alex est un Professor et qui doit superviser un Student selon (2). Cependant, Alex doit être soit un Researcher, soit un Expert selon (1) ; ce qui contredit δ_1 ou δ_2 . Ceci nécessite la révision de $\mathcal{O}$ afin que δ_1 et δ_2 soient prises en compte et les connaissances de $\mathcal{O}$ soient conservées le plus possibles. Pour ce faire, ONTOREV construit deux ensembles de modèles $FM(\mathcal{O})$ et $FM(\mathcal{O}')$. Ensuite, il extrait de $FM(\mathcal{O}')$ l'ensemble de modèles, noté $FM(\mathcal{O} \oplus \mathcal{O}')$, qui sont les plus proches (par rapport à la distance sémantique) des modèles dans $FM(\mathcal{O})$. Enfin, ONTOREV génère à partir de $FM(\mathcal{O} \oplus \mathcal{O}')$ l'ontologie résultante, notée $\mathcal{O}^*$, contenant les axiomes et assertion suivants :

$\text{Expert} \sqsubseteq \forall \text{supervises. } (\neg \text{Student})$,

$\text{Researcher} \sqsubseteq \forall \text{supervises. } (\neg \text{Student})$, $\text{Professor}(\text{Alex})$,

$\top \sqsubseteq (\text{Professor} \sqcap \neg \text{Researcher} \sqcap \neg \text{Expert} \sqcap \neg \text{Course} \sqcap \neg \text{Student} \sqcap \exists \text{supervises. Student} \sqcap (\geq 2 \text{ teaches. Course}) \sqcup (\text{Professor} \sqcap \text{Researcher} \sqcap \neg \text{Expert} \sqcap \neg \text{Course} \sqcap \neg \text{Student} \sqcap \exists \text{supervises. Student} \sqcap (\geq 2 \text{ teaches. Course}) \sqcup (\text{Student} \sqcap \forall \text{supervises. } (\neg \text{Student}) \sqcap (\leq 1 \text{ teaches. Course}) \sqcap \neg \text{Course} \sqcap \neg \text{Professor} \sqcap \neg \text{Researcher} \sqcap \neg \text{Expert}) \sqcup (\text{Course} \sqcap \forall \text{supervises. } (\neg \text{Student}) \sqcap (\leq 1 \text{ teaches. Course}) \sqcap \neg \text{Student} \sqcap \neg \text{Professor} \sqcap \neg \text{Researcher} \sqcap \neg \text{Expert})$,

Remarquons que $\mathcal{O}^*$ peut inclure des axiomes de taille importante car ils proviennent des modèles trouvés dans $FM(\mathcal{O} \oplus \mathcal{O}')$. Comme toutes les approches fondées sur les modèles, les axiomes de l'ontologie résultante construite par notre méthode décrivent les connaissances plutôt au niveau des modèles qu'au niveau de notre propre conception du domaine. Par conséquent, ils sont moins "parlant" que ceux des ontologies initiales. Pour évaluer ONTOREV, nous avons réalisé plusieurs tests avec différentes ontologies du monde réel. Les lecteurs intéressés peuvent trouver les résultats de ces tests dans l'article (Dong *et al.*, 2017).

Interface Web d'ONTOREV. Nous avons intégré ONTOREV dans une application Web² qui permet d'éditer, de réviser et d'interroger une ontologie OWL. Les avantages de cette interface sont que (i) tous les calculs lourds d'ONTOREV sont effectués sur un serveur puissant, (ii) l'affichage du contenu de l'ontologie de travail et la syntaxe de saisie d'un nouvel axiome/assertion sont facilités car exprimable en syntaxe Manchester (très utilisée par la communauté), (iii) les requêtes d'utilisateur sont analysées de façon centralisée. Cela permet d'optimiser la recherche des réponses en utilisant un mécanisme de cache ou d'apprentissage, et (iv) les requêtes courants sur une ontologie OWL tels que la vérification de la cohérence, la déduction d'un nouvel axiome/assertion y sont naturellement intégrés.

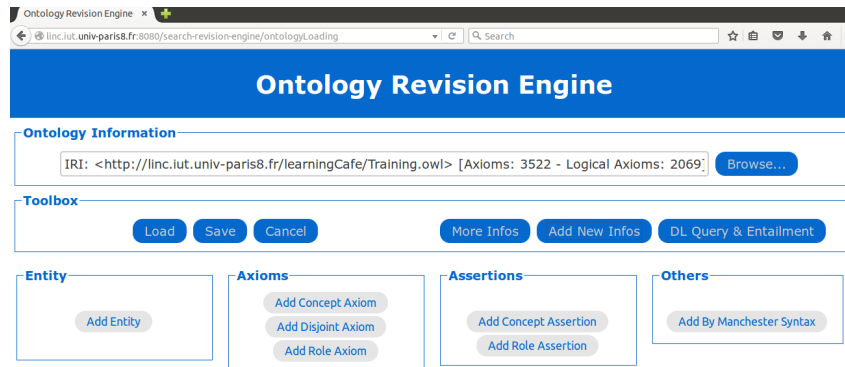


Fig. 2: Interface Web d'ONTOREV

Toutes les fonctionnalités accessibles via un navigateur Web sont aussi disponibles comme Web services du type REST. La page d'accueil (cf. figure 2) est découpée en deux zones :

(1) **Ontology Information** : Un utilisateur peut charger une ontologie en tapant l'URI de l'ontologie ou en choisissant (**Browse...**) une ontologie en local. Notons que si le bouton **Load** situé dans la zone **Toolbox** est cliqué mais que l'utilisateur n'a pas tapé de lien ni choisi une ontologie en local, l'ontologie *Training.owl* de la plate-forme *LearningCafé* se charge par défaut.

(2) **Toolbox** : contient les boutons permettant de charger une ontologie (**Load**), sauvegarder le résultat (**Save**), annuler le traitement (**Cancel**). De plus, l'utilisateur peut cliquer sur **More Infos** pour afficher les entités, les axiomes/assertions de l'ontologie de travail ; et **Add New Infos** lui permet de mettre à jour l'ontologie. Enfin, il peut aussi faire des requêtes sur l'ontologie en utilisant **DL Query & Entailment**. Cette fonctionnalité utilise les services du raisonneur HermiT (Shearer *et al.*, 2008) à l'heure actuelle.

Quand l'utilisateur clique sur **Add New Infos**, une nouvelle zone s'ouvre (cf. figure 2). Cette zone comporte plusieurs boutons **Entity**, **Axioms**, etc. permettant de saisir un

² <http://linc.iut.univ-paris8.fr:8080/search-revision-engine/>

nouvel axiome ou une nouvelle assertion à ajouter dans l'ontologie de travail. Ces boutons peuvent déclencher ONTOREV si l'axiome/assertion à ajouter n'est pas compatible avec ceux déjà existant dans l'ontologie de travail.

Conclusion et perspective. Nous avons présenté le prototype ONTOREV, implémenté en Java pour la révision d'ontologie OWL. Le service de révision d'ontologie basé sur ONTOREV a été intégré dans une interface Web permettant de faciliter l'accès au service et de profiter de la puissance de calcul d'un serveur. Comme perspective à court terme, nous visons à proposer une approche de révision hybride fondée non seulement sur la sémantique (modèles) mais aussi sur la syntaxe. Une telle approche améliorerait la lisibilité des axiomes de l'ontologie résultante.

References

- [Baader *et al.*, 2003] Franz Baader, Diego Calvanese, Deborah L. McGuinness, Daniele Nardi, and Peter F. Patel-Schneider, editors. *The Description Logic Handbook: Theory, Implementation, and Applications*. Cambridge University Press, 2003.
- [Dong *et al.*, 2017] Tinh Dong, Chan Le Duc, and Myriam Lamolle. Tableau-based revision for expressive description logics with individuals. *J. Web Sem.*, 45:63–79, 2017.
- [Shearer *et al.*, 2008] Rob Shearer, Boris Motik, and Ian Horrocks. Hermit: A Highly-Efficient OWL Reasoner. In Alan Ruttenberg, Ulrike Sattler, and Cathy Dolbear, editors, *Proc. of the 5th Int. Workshop on OWL: Experiences and Directions (OWLED 2008 EU)*, Karlsruhe, Germany, October 26–27 2008.

ABClass: A multiple instance learning approach for sequence data

Manel Zoghliami^{1,2} Sabeur Aridhi³ Mondher Maddouri⁴ Engelbert Mephu Nguifo¹

¹University of Clermont Auvergne, LIMOS, PB 10125, 63173 Clermont Ferrand, France

² University of Tunis El Manar, Faculty of sciences of Tunis, LIPAH, 1060 Tunis, Tunisia

³ University of Lorraine, CNRS, Inria, LORIA, F-54000 Nancy, France

⁴ University of Jeddah, College Of Business, PB 80327, 21589 Jeddah, KSA

manel.zoghliami@gmail.com

In Multiple Instance Learning (MIL) problem for sequence data, the learning data consist of a set of bags where each bag contains a set of instances/sequences. In some real world applications such as bioinformatics comparing a random couple of sequences makes no sense. In fact, each instance of each bag may have structural and/or temporal relation with other instances in other bags. Thus, the classification task should take into account the relation between semantically related instances across bags.

In this work, we present ABClass, a novel MIL approach for sequence data classification. Each sequence is represented by one vector of attributes extracted from the set of related instances. For each sequence of the unknown bag, a discriminative classifier is applied in order to compute a partial classification result. Then, an aggregation method is applied in order to generate the final result. We applied ABClass to solve the problem of bacterial Ionizing Radiation Resistance (IRR) prediction. The experimental results were satisfactory.

Availability: ABClass tool and the used dataset can be found in the following link :

<http://homepages.loria.fr/SAridhi/software/MIL/> . ABClass runs on a Windows or a Linux platform (tested on Ubuntu distribution) that contains a java JRE.