

Approche computationnelle de l'analyse conceptuelle

Présentation opérationnelle et approfondissement méthodologique de la détection d'un concept dans des extraits textuels

Francis Lareau

Volume 49, numéro 2, automne 2022

Philosophie québécoise. Histoire et humanités numériques

URI : <https://id.erudit.org/iderudit/1097460ar>

DOI : <https://doi.org/10.7202/1097460ar>

[Aller au sommaire du numéro](#)

Éditeur(s)

Société de philosophie du Québec

ISSN

0316-2923 (imprimé)

1492-1391 (numérique)

[Découvrir la revue](#)

Citer cet article

Lareau, F. (2022). Approche computationnelle de l'analyse conceptuelle : présentation opérationnelle et approfondissement méthodologique de la détection d'un concept dans des extraits textuels. *Philosophiques*, 49(2), 413–431. <https://doi.org/10.7202/1097460ar>

Résumé de l'article

Une tâche importante en philosophie est la lecture et l'analyse de textes pour en dégager les concepts. L'objectif de la présente étude est d'explorer la possibilité d'une assistance computationnelle pour effectuer cette tâche. Une méthode classique est le concordancier, mais celle-ci ne permet pas de distinguer les extraits où le concept n'est pas exprimé de manière canonique. Nous proposons une méthode permettant de reconnaître ces extraits, que nous appliquons à un corpus d'articles de la revue *Philosophiques*. Nous déterminons d'abord les extraits où le concept est exprimé de manière explicite. Ensuite, nous déterminons les extraits les moins susceptibles d'exprimer le concept cible. Enfin, nous utilisons plusieurs classificateurs afin de distinguer les extraits où le concept est exprimé de manière implicite. Les résultats montrent une différence significative entre les classificateurs les plus performants, machines à vecteurs de support et réseaux de neurones, et certains modèles probabilistes classiques.

Approche computationnelle de l'analyse conceptuelle

Présentation opérationnelle et approfondissement méthodologique de la détection d'un concept dans des extraits textuels¹

FRANCIS LAREAU

Département d'informatique, Université du Québec à Montréal (UQÀM)

RÉSUMÉ — Une tâche importante en philosophie est la lecture et l'analyse de textes pour en dégager les concepts. L'objectif de la présente étude est d'explorer la possibilité d'une assistance computationnelle pour effectuer cette tâche. Une méthode classique est le concordancier, mais celle-ci ne permet pas de distinguer les extraits où le concept n'est pas exprimé de manière canonique. Nous proposons une méthode permettant de reconnaître ces extraits, que nous appliquons à un corpus d'articles de la revue *Philosophiques*. Nous déterminons d'abord les extraits où le concept est exprimé de manière explicite. Ensuite, nous déterminons les extraits les moins susceptibles d'exprimer le concept cible. Enfin, nous utilisons plusieurs classifieurs afin de distinguer les extraits où le concept est exprimé de manière implicite. Les résultats montrent une différence significative entre les classifieurs les plus performants, machines à vecteurs de support et réseaux de neurones, et certains modèles probabilistes classiques.

MOTS-CLEFS: concept; lecture et analyse conceptuelle de texte assistées par ordinateur; LACTAO; word2vec; doc2vec; machine à vecteurs de support; réseau de neurones artificiel

ABSTRACT — An important task in philosophy is reading and analyzing texts to identify concepts. The objective of this study is to explore the possibility of a computational support for this task. A classic method is the concordancer, but this approach does not make it possible to distinguish the excerpts where the concept is not expressed in a canonical way. We propose a method to identify these excerpts that we apply to a corpus of articles from the journal *Philosophiques*. We first determine the excerpts where the concept is expressed explicitly. Then, we identify the excerpts least likely to express the target concept. Finally, we train several classifiers to distinguish the excerpts where the concept is implicitly expressed. The results show a significant difference between the most

1. Ce texte est issu d'une conférence présentée lors d'une journée d'étude intitulée « Philosophie québécoise: histoire et humanités numériques », qui s'est déroulée dans le cadre du Congrès annuel de la Société de philosophie du Québec, tenu durant le 88^e congrès de l'ACFAS, à l'Université de Sherbrooke, à l'Université Bishop's et en mode virtuel, du 3 au 7 mai 2021. L'auteur remercie l'auditoire de cette activité, Jean-Claude Simard ainsi que Jean-Guy Meunier pour leurs commentaires sur une version antérieure du manuscrit. L'auteur reconnaît le financement du Fonds de recherche du Québec — Société et culture (FRQSC-276470).

efficient classifiers, support vector machines and neural networks, and some classical probabilistic models.

KEYWORDS: concept; computer assisted conceptual analysis of texts; CACAT; word2vec; doc2vec; support vector machine; svm; artificial neural network

Introduction

En philosophie, une tâche importante est la lecture et l'analyse de textes pour en dégager les concepts. Ce travail d'analyse conceptuelle contribue au progrès de la connaissance dans les différents domaines de la science et de la philosophie. Il peut se concevoir comme une activité visant à éclaircir nos pensées au moyen d'une décomposition de concepts en leurs constituants, c'est-à-dire en déterminant les propriétés ou les catégories importantes ainsi que leurs relations². Or on reproche parfois à l'analyse conceptuelle son éloignement de l'empirie³, mais cette activité peut aussi constituer une entreprise empirique, partiellement du moins, puisque de facto différentes communautés linguistiques utilisent certains termes à certaines fins (ces faits pouvant s'étudier comme tels), et les inférences à leur sujet sont faillibles, de sorte qu'il peut être utile de les analyser⁴. En se fondant sur l'hypothèse de Firth selon laquelle on connaît un mot par ce qui l'accompagne⁵, il est alors possible de connaître un concept exprimé sous forme linguistique en observant l'entourage qui se trouve dans nos usages. Bref, si on peut concevoir l'analyse conceptuelle a priori en visant, par exemple, la reconnaissance des conditions nécessaires et suffisantes définissant des concepts, on peut autrement la concevoir comme visant à découvrir a posteriori les caractéristiques ou les relations importantes se trouvant dans le discours. De plus, si les faits linguistiques portant sur un concept peuvent être attribués à différentes communautés linguistiques, alors on peut aussi envisager différentes analyses en fonction des communautés qui nous intéressent et qui sont coextensives des usages linguistiques. Par exemple, on peut s'intéresser à un concept tel qu'exprimé par une communauté experte ou encore examiner ceux exprimés par une autre communauté qui serait plus naïve en la matière. Enfin, en raison des avancées technologiques spectaculaires de l'intelligence artificielle, l'analyse conceptuelle en philosophie connaît aujourd'hui des développements importants en ce que plusieurs chercheurs tentent de modéliser de manière informatique différentes étapes d'une approche matérielle de l'analyse

2. Étienne Bonnot de Condillac, *Essai sur l'origine des connoissances humaines*, Paris, Houel, 1798; Peter Frederick Strawson, *Analyse et métaphysique*, Paris, Vrin, 1985.

3. Voir Brian Talbot, « Why so negative? Evidence aggregation and armchair philosophy », *Synthese*, vol. 191, 2014, p. 3865-3896.

4. Voir Frank Jackson, *From Metaphysics to Ethics, a Defence of Conceptual Analysis*, Oxford, Clarendon Press, 1998.; Voir aussi Claudine Tiercelin, « La métaphysique et l'analyse conceptuelle », *Revue de métaphysique et de morale*, vol. 4, n° 36, 2002, p. 529-554.

5. John R. Firth, *Papers in Linguistics*, London, Oxford University Press, 1957, p. 11.

conceptuelle⁶. Récemment, la recherche en philosophie montre la pertinence d'une telle approche pour l'étude de divers *corpora*, notamment à propos de Descartes⁷, Darwin⁸, Bergson⁹, Peirce¹⁰ et Evans¹¹. Suivant la présentation des différentes étapes de l'approche computationnelle à l'analyse conceptuelle, nous explorons plus en détail l'étape de détermination des contextes linguistiques dans lesquels un concept est exprimé en appliquant et testant différentes modélisations de cette tâche sur un corpus original issu d'une communauté philosophique particulière. Plus précisément, nous proposons une approche capable de reconnaître non seulement les contextes où un concept est exprimé de manière explicite ou canonique (à l'aide de termes ou de syntagmes reconnus par une communauté linguistique comme exprimant un concept), mais de détecter également les contextes dans lesquels un concept est exprimé de manière implicite ou non canonique et que certains qualifient de pertinents¹² ou de périphériques¹³ à un concept donné.

6. Paul Thagard, *Computational philosophy of science*, MIT Press, 1988; Mathieu Valette, Fløttum Kjersti et François Rastier, « Conceptualisation and Evolution of concepts. The example of French Linguist Gustave Guillaume », dans *Academic Discourse – Multidisciplinary Approaches*, 2003, p. 55-74; Dominic Forest et Jean-Guy Meunier, « Classification et catégorisation automatiques: application à l'analyse thématique des données textuelles », *Actes des 7es Journées internationales d'Analyse statistique des Données Textuelles*, vol. 1, Louvain-la-Neuve, Presses Universitaires de l'Université Catholique de Louvain, 2004, p. 434-444; Sylvain Loiseau, « Thématique et sémantique contextuelle d'un concept philosophique », dans *La linguistique de corpus*, G. Williams, dir., Rennes, Presse Universitaires de Rennes, 2005.

7. Dominic Forest, *Lecture et analyse de textes philosophiques assistées par ordinateur: Application d'une approche classificatoire mathématique à l'analyse thématique du Discours de la méthode et des Méditations métaphysiques de Descartes*, Mémoire de maîtrise en philosophie, UQÀM, 2002.

8. Maxime Sainte-Marie, Jean-François Chartier, Louis Chartrand, Jean Danis et Jean-Guy Meunier, « Nouveaux outils, nouveaux jeux de mots: perspectives de recherche et applications de la LATAO », *Cahiers de l'ISC*, 2012.

9. Jean Danis, *L'Analyse Conceptuelle de Textes Assistée par Ordinateur (LACTAO): une expérimentation appliquée au concept d'évolution dans l'œuvre d'Henri Bergson*, Mémoire en philosophie, UQÀM, 2012.

10. Davide Pulizzotto, Jean-François Chartier, Francis Lareau, Jean-Guy Meunier et Louis Chartrand, « Conceptual Analysis in a computer-assisted framework: mind in Peirce », *Umanistica Digitale*, vol. 2, 2018, p. 185-205.

11. Francis Lareau, *Analyse philosophique en contexte numérique du concept de dualité chez Jonathan St B. T. Evans*, Mémoire en philosophie, UQÀM, 2016; Francis Lareau, Louis Chartrand, Jean-François Chartier et Davide Pulizzotto, « Lecture et analyse conceptuelle de texte assistées par ordinateur (LACTAO) appliquées à des textes de haut niveau théorique: illustration de la méthode et résultats de l'analyse du concept de dualité chez Jonathan St B. T. Evans », *Applied Semiotics/Sémiotique appliquée*, vol. 26, 2018.

12. Louis Chartrand, Jean-Guy Meunier, Davide Pulizzotto, José López González, Jean-François Chartier, Ngoc Tan Le, Francis Lareau et Julian Trujillo Amaya, « CoFiH: A heuristic for concept discovery in computer-assisted conceptual analysis », *JADT 2016: 13^e Journées internationales d'Analyse statistique des Données Textuelles*, 7-10 juin 2016, Nice, vol. 1.

13. Davide Pulizzotto, José López González, Jean-François Chartier, Jean-Guy Meunier, Louis Chartrand, Francis Lareau et Ngoc Tan Le, « Recherche de "périségments" dans un

Les étapes de l'analyse conceptuelle assistée par ordinateur

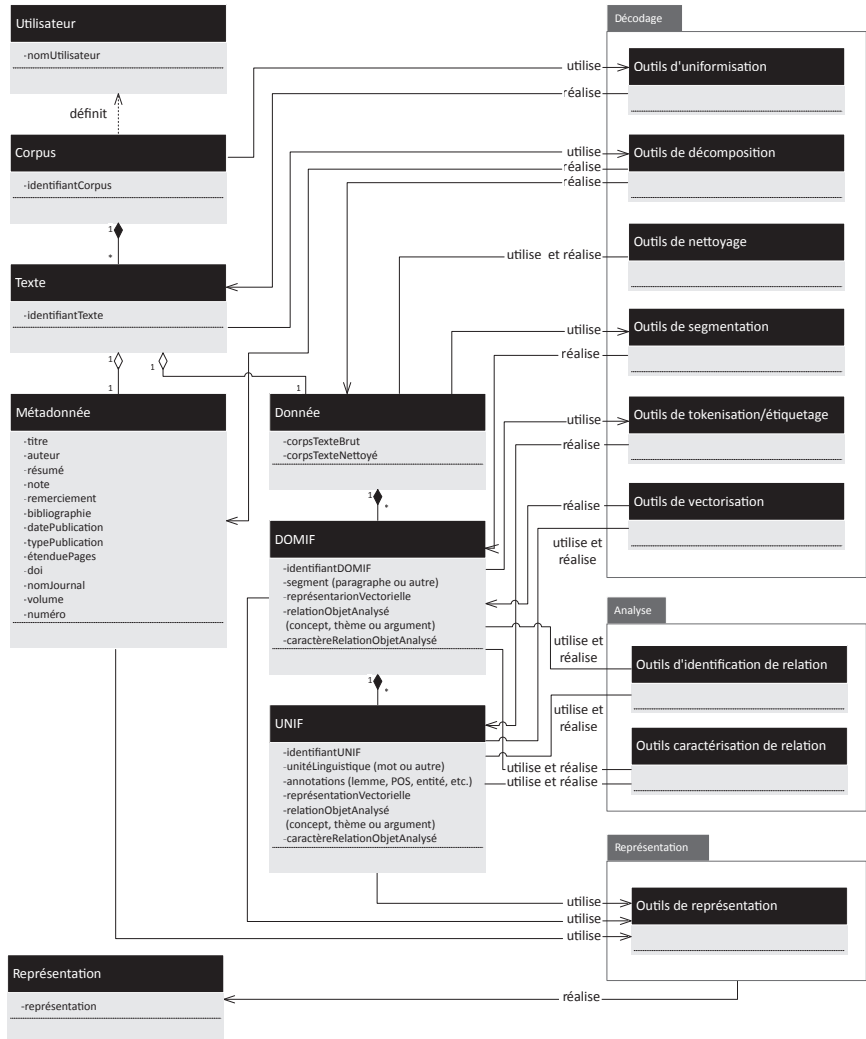
L'assistance informatique d'une tâche telle que l'analyse conceptuelle n'est possible que dans la mesure où celle-ci est décomposable en sous-opérations, dont certaines sont traduisibles en fonctions ou algorithmes exécutables par ordinateur. Elle s'inspire de techniques informatiques issues des domaines de recherche en fouille de texte (*text mining*), en apprentissage machine et en traitement automatique des langues naturelles. Ces techniques mettent en jeu des procédures déductives, inductives ou hybrides et sont utilisées pour atteindre divers objectifs comme l'extraction d'information, la catégorisation, la visualisation, la synthèse ou encore l'analyse conceptuelle. Généralement, les grandes étapes d'une fouille de texte apportant une assistance computationnelle à l'analyse de corpus sont :

- i. la constitution du corpus;
- ii. le décodage du corpus;
- iii. les analyses;
- iv. la représentation des résultats d'analyse;
- v. l'interprétation des résultats d'analyse.

Les étapes d'une analyse conceptuelle sont similaires à celles d'une analyse de corpus assistée par ordinateur, mais les analyses sont spécifiquement conceptuelles. Outre la constitution du corpus et l'interprétation des résultats qui se prêtent plus difficilement à une modélisation computationnelle, nous exemplifions dans le diagramme de classes ci-dessous une décomposition de l'analyse de corpus suivant les trois grandes étapes habituellement modélisées en fouille de texte. Ce diagramme s'inspire du langage de modélisation unifié (*Unified Modeling Language*) qui est une méthode normalisée de visualisation en développement logiciel. Trois types d'objets y sont présentés, c'est-à-dire la classe *utilisateur*, les classes d'objets manipulés par le système et les classes d'objets qui sont en fait des outils permettant de manipuler les objets précédemment énoncés. Ces outils représentent les différents algorithmes en jeu et sont organisés en paquets (*packages*) correspondant aux trois étapes précédemment mentionnées, c'est-à-dire le *décodage* du corpus, l'*analyse* et la *représentation* des résultats. Spécifions que la classe *utilisateur* représente l'expert-e ou chercheur-se utilisant une telle modélisation informatique de l'analyse de corpus.

contexte d'analyse conceptuelle assistée par ordinateur: le concept d'«esprit» chez Peirce», *JEP-TALN-RECITAL*, vol. 2, 2016.

Diagramme de classes d'un système d'analyse de corpus



Lors de la première étape — le *décodage* du corpus —, le *corpus* préalablement constitué est défini par l'*utilisateur* et utilisé par les *outils d'uniformisation* afin de réaliser des *textes* uniformes, lesquels sont tous membres du *corpus* de départ. Chaque *texte* est ensuite utilisé par les *outils de décomposition* des textes pour extraire et réaliser les *données* ainsi que les *métadonnées* associées à chaque *texte*. Les *données* représentent le corps du texte de chaque *texte* tandis que les *métadonnées* représentent les informations à propos du corps du texte, comme le titre, le ou les auteurs, le résumé, les notes, les remerciements, les références, la date de publication, le type de

publication, l'étendue des pages, l'identifiant numérique d'objet (*doi*), le nom de la revue, le volume, le numéro, etc.¹⁴ Les *données* brutes de chaque *texte* sont ensuite utilisées par les *outils de nettoyage* pour réaliser des *données* nettoyées où les corps de texte sont exempts de toute information périphérique ou de bruit numérique. Les *données* nettoyées de chaque *texte* sont ensuite utilisées par les *outils de segmentation* afin de réaliser des segments textuels nommés *domaines d'information* (DOMIFs) et pouvant correspondre à des paragraphes, des phrases ou tout autre partitionnement possible. Chaque DOMIF est associé à un identifiant et utilisé par les *outils de tokenisation et d'étiquetage* afin de décomposer ceux-ci en unités linguistiques atomiques nommées *unités d'information* (UNIFs), lesquelles peuvent correspondre à des caractères, des mots ou des groupes de mots. Les *outils d'étiquetage* vont attribuer un identifiant et réaliser l'étiquetage de chaque UNIF et, au besoin, ces outils vont reconnaître et noter la fonction de l'UNIF dans la phrase (*part of speech* ou POS en anglais) ainsi que son lemme, qui correspond à la forme retenue comme entrée dans un dictionnaire¹⁵. Aussi, les UNIFs sont utilisées par les *outils de vectorisation* afin de réaliser une représentation vectorielle des DOMIFs et des UNIFs. Ainsi, le *corpus* est composé de *textes* qui sont chacun associés à une *métadonnée* ainsi qu'une *donnée* qui est composée de DOMIFs, eux-mêmes composés d'UNIFs. Lors de la deuxième étape — l'*analyse* (conceptuelle dans notre cas) —, les UNIFs sont utilisées par les *outils de détermination de relation* pour découvrir des régularités linguistiques permettant de reconnaître les UNIFs qui sont en relation avec un objet d'analyse (en l'occurrence conceptuel¹⁶). Selon les objectifs de la recherche, il est possible de caractériser plus finement la relation à l'objet d'analyse en découvrant d'autres régularités détectables à l'aide d'*outils de caractérisation de relation*. Lors de la troisième étape — la *représentation* des résultats — les UNIFs et les DOMIFs en relation avec l'objet d'analyse ainsi que les régularités observées et, au besoin, certaines *métadonnées* collectées (l'année de publication ou les auteurs, par exemple) sont utilisées par les *outils de représentation* afin de réaliser une *représentation* des résultats d'analyse, laquelle sera éventuellement interprétée.

14. Selon la structure inhérente aux textes constituant le corpus d'étude et en fonction des objectifs de la recherche, il est possible de définir autrement ce qu'est une donnée et une métadonnée.

15. Spécifions que d'autres formes d'étiquetage sont possibles comme la reconnaissance d'entités nommées, c'est-à-dire la détermination des unités linguistiques référentielles ainsi que de leur type.

16. L'objet d'analyse peut être conceptuel, mais aussi thématique, argumentaire ou autre.

Décodage du corpus

Le corpus sélectionné pour cette étude est composé des 1 476 textes de la revue semestrielle *Philosophiques* (1974-2018), lesquels nous décomposons en 45 818 segments de texte (contenant environ 5 phrases) eux-mêmes décomposables en 233 501 phrases. Suivant la lemmatisation des mots et la détermination de leur fonction dans la phrase, seuls les noms communs, les verbes, les adjectifs et les adverbes sont retenus, pour un total de 3 611 243 lemmes, qu'on peut classer sous 61 802 types distincts.

Analyse conceptuelle (partielle)

Cette étude n'est pas une analyse conceptuelle complète, puisqu'elle n'inclut pas certaines étapes importantes comme la caractérisation fine des relations avec le concept ainsi que leur interprétation. L'étape spécifique qui nous intéresse est la détermination des segments de texte dans lesquels un concept cible est exprimé autrement qu'avec le ou les termes y faisant habituellement référence. L'approche que nous proposons présuppose que notre problématique est un cas particulier du problème de classification de cas indéterminés à partir de cas positifs. Nous nous inspirons de la méthode proposée par Nigam et coll.¹⁷ qui suit les étapes suivantes :

- a) la reconnaissance de cas positifs ;
- b) la reconnaissance de cas négatifs ;
- c) l'apprentissage machine à propos de ces cas positifs et négatifs permettant, à terme,
- d) la détermination des cas auparavant indéterminés comme positifs ou négatifs.

Notre première hypothèse (h1) est que la reconnaissance de certains cas positifs peut s'effectuer à l'aide de l'approche classique du concordancier où on reconnaît l'expression d'un concept cible à l'aide d'un ou de plusieurs termes censés y référer. Notre deuxième hypothèse (h2) est que la reconnaissance de certains cas négatifs peut se faire à l'aide d'un calcul de similarité sur une représentation numérique de nos données, dans la mesure où les cas les plus dissimilaires aux cas reconnus comme positifs peuvent être considérés comme négatifs. Notre troisième hypothèse (h3) est qu'on peut entraîner un modèle de classification sur la base des cas positifs et négatifs préalablement reconnus. Notre dernière hypothèse (h4) est que le modèle de classification préalablement entraîné permet de reconnaître les cas positifs et négatifs parmi les cas indéterminés. Plus précisément, nous testerons nos hypothèses sur différents types de représentations vectorielles de notre corpus, ainsi qu'avec divers modèles de classification.

17. Kamal Nigam, Andrew Kachites McCallum, Sebastian Thrun et Tom Mitchell, « Text classification from labeled and unlabeled documents using EM », *Machine Learning*, vol. 39, 2000, p. 103-134.

Reconnaissance de cas positifs (h1)

Le concept sélectionné pour cette étude est celui de **cognition**. Nous proposons la reconnaissance de cas positifs à l'aide d'une règle exprimée par une expression régulière (*regex*) permettant de subsumer différentes expressions du langage naturel. L'expression régulière « cognit » permet de reconnaître 1 370 segments textuels contenant les termes suivants : affectivo-cognitives, anti-cognitives, anti-cognitivisme, anti-cognitiviste, cognitiens, cognitif-s, cognitif-clef, cognitif-développemental, cognitif-émotionnel, cognitifs-scientifiques, cognitio, cognition-s, cognitione, cognitionem, cognitionis, cognitique, cognitit, cognitive, cognitive-s, cognitive-émotionnelle, cognitivement, cognitiviste, cognitivisme, cognitiviste-s, cognitivité, cognitivo-comportementale-s, logico-cognitif, métacognitif, métacognition, métacognitive, non-cognitivisme, non-cognitiviste-s, néo-cognitivisme, psychocognitif, psychologico-cognitif et sociocognitifs-tives. Notons que le sens des termes reconnus apparaît relié au concept de **cognition**, mais dans le cas où certaines expressions se révéleraient ambiguës, il serait alors nécessaire d'effectuer une opération de désambiguïsation afin de retirer celles sans relation avec le concept cible. Enfin, une fois les cas positifs trouvés, les termes précédemment énoncés sont retirés du vocabulaire, puisque les cas indéterminés en sont exempts, de sorte que la détermination des cas positifs parmi les cas indéterminés ne peut pas se réaliser à l'aide de ces termes.

Reconnaissance de cas négatifs (h2)

La représentation numérique classique des textes est le sac de mots — *bag of words* (BOW) — où chaque domaine d'information¹⁸ (corps du texte ou segment de texte) est représenté par un vecteur dont chacune des valeurs correspond à la fréquence d'une unité d'information¹⁹ (mot, groupe de mots, lemme, groupe de lemmes, caractère, groupe de caractères, etc.) dans le domaine d'information²⁰. Cette représentation peut être manipulée de manière à favoriser certains termes (le TFIDF²¹ en est un exemple) ou à normaliser la représentation. Autrement, Mikolov et coll.²² proposent de transformer les textes en des représentations de types *word2vec* ou *doc2vec*

18. Le domaine d'information (DOMIF) est parfois nommé « document » dans la littérature et, subséquemment, les deux expressions seront utilisées indifféremment.

19. L'unité d'information (UNIF) est parfois nommée « terme » dans la littérature et, subséquemment, les deux expressions seront utilisées indifféremment.

20. Voir Gerard Salton et Michael J. McGill, *Introduction to modern information retrieval*, McGraw-Hill, 1983.

21. *TF* est la fréquence totale d'une unité lexicale donnée dans un document donné, et *IDF* est l'inverse du logarithme de la fréquence documentaire, de sorte que *TFIDF* est la multiplication d'une variable locale (*TF*) et d'une variable globale (*IDF*). Notons que nous utilisons le *TFIDF* sur la représentation BOW.

22. Tomas Mikolov, Kai Chen, Greg Corrado et Jeffrey Dean, « Efficient estimation of word representations in vector space », *arXiv*, 2013, p. 1301-3781.

qu'on nomme parfois des « plongements » (*embeddings*), dans la mesure où on passe des textes à un espace vectoriel continu de dimension inférieure à la représentation classique. Les méthodes utilisées par les auteurs afin de générer la représentation *word2vec* se trouvent en deux variantes, c'est-à-dire le *continuous bag of words* (CBOW) et le *skip-gram* (SG). CBOW est un réseau de neurones artificiels muni d'une couche cachée dont l'entraînement vise à prédire un terme à partir d'une fenêtre de termes adjacents, c'est-à-dire son contexte. Le réseau est amorcé avec des valeurs aléatoires, puis ces valeurs sont modifiées itérativement de manière à réduire la différence entre les valeurs de sortie obtenues et celles désirées. On utilise ensuite ce modèle de prédiction du terme à partir des contextes pour extraire la représentation *word2vec*. Celle-ci est composée de chaque vecteur de valeurs générées par la couche cachée pour chacun des termes, de sorte que l'étendue de ces vecteurs correspond au nombre de neurones de la couche cachée (de 50 à 300 habituellement). SG fonctionne de manière similaire, mais inversée, c'est-à-dire qu'à partir d'un terme, l'entraînement vise à prédire le contexte. Les équivalents de CBOW et de SG pour *doc2vec* se nomment respectivement *distributed memory* (DM) et *distributed bag of words* (DBOW). Le procédé pour générer la représentation *doc2vec* ne se distingue que par l'ajout d'un identifiant de document à chacun des vecteurs d'entrée et de sortie. *Doc2vec* permet ainsi de représenter à la fois les termes et les documents par des vecteurs de caractéristiques (*features*). Ces caractéristiques encapsulent de manière implicite des propriétés syntaxiques ou sémantiques comme la synonymie, l'antonymie, l'analogie, etc. L'exemple classique montrant la puissance de ce type de représentation est celui où l'addition du vecteur correspondant au terme « roi » à celui de « femme », puis la soustraction du vecteur « homme » à la résultante permettent de trouver un vecteur synthétique dont la proximité est maximale avec le vecteur « reine » dans un espace multidimensionnel. Similairement, on peut générer un vecteur synthétique en faisant la moyenne d'un premier ensemble de vecteurs correspondants à des documents afin de reconnaître, parmi un second ensemble de vecteurs, ceux qui sont les plus similaires ou dissimilaires avec les premiers. Nous utiliserons donc la représentation issue de la méthode DBOW (que nous nommerons DBOW) afin de reconnaître les cas négatifs que nous concevons comme les documents les plus éloignés d'un vecteur synthétique faisant la moyenne des vecteurs correspondant aux cas positifs. Plus précisément, la mesure de similarité utilisée est le cosinus. Or, les cas négatifs sont définis par la négative et, ce faisant, ceux-ci sont probablement moins homogènes que les cas positifs. De plus, il semble raisonnable de croire qu'il y a un nombre de cas négatifs plus élevé que le nombre de cas positifs, puisque la cognition n'est qu'un des nombreux sujets de la revue *Philosophiques*. Conséquemment, nous sélectionnons un nombre de cas négatifs plus élevé que celui des cas positifs dans un rapport de 4 pour 1

(5 480 cas négatifs) afin de représenter très approximativement une distribution réaliste des classes positives et négatives dans le corpus²³.

Entraînement et validation des modèles de classification (h3)

Les modèles de classification sont entraînés sur un sous-ensemble aléatoire de nos données textuelles (80 %) pour les deux types de représentation BOW et DBOW. Les données restantes (20 %) sont utilisées afin de valider les modèles. Parmi ceux testés se trouvent les machines à vecteurs de support (*support vector machine* ou SVM). Les SVM sont un ensemble de techniques d'apprentissage supervisé permettant la classification binaire de données multidimensionnelles (ou non). L'avantage des SVM est que leur performance est similaire ou supérieure à celle d'un réseau simple de neurones artificiels, dans la mesure où les SVM permettent de trouver un minimum global comparativement aux réseaux de neurones artificiels qui peuvent s'enliser dans un minimum local²⁴. Les SVM permettent de classer de manière binaire des données de haute dimensionnalité en traçant un hyperplan (un sous-espace affine plat) entre deux sous-ensembles et en maximisant la marge entre l'hyperplan et les supports constitués des données limitrophes. Notons que la marge est dite molle (*soft*) en ce qu'elle est déterminée par une validation croisée maximisant la performance du classifieur et permettant de se protéger du bruit, des données d'entraînement incorrectement classées ou des chevauchements de classes. Dans le cas des problèmes non linéaires (où il n'existe pas d'hyperplan capable de partitionner l'ensemble de données en deux sous-ensembles correspondant aux classes recherchées), il est possible de transformer l'espace vectoriel de manière à rendre cette partition possible. Autrement dit, un problème de partition non linéaire peut devenir un

23. Si le nombre de cas positifs n'est pas équivalent au nombre de cas négatifs, alors les données sont dites déséquilibrées ou mal balancées. Cela est problématique, dans la mesure où plus les données sont déséquilibrées, plus la performance de classification est dégradée et biaisée vers la classe majoritaire. Or, dans Gary M. Weiss et Foster Provost, « Learning when training data are costly: the effect of class distribution on tree induction », *Journal of Artificial Intelligence Research*, vol. 19, 2003, p. 315-354, les auteurs montrent qu'un ratio de 1 pour 1 n'est pas nécessairement une partition optimale pour un problème de classification donné. De plus, la recherche montre que la classification, notamment à l'aide de machines à vecteurs de support, tend à être résiliente lorsque le nombre de données est suffisant pour la complexité du problème et lorsque le ratio de déséquilibre ne dépasse pas 1 pour 10 (Yanmin Sun, Andrew Wong et Mohamed S. Kamel, « Classification of imbalanced data: a review », *International Journal of Pattern Recognition and Artificial Intelligence*, vol. 23, n° 4, 2011, p. 687-719; Akila Somasundaram et U. Srinivasulu Reddy, « Data Imbalance: Effects and Solutions for Classification of Large and Highly Imbalanced Data », *Proceedings of ICRECT-16*, 2016, p. 28-34.).

24. Christopher J. C. Burges, « A Tutorial on Support Vector Machines for Pattern Recognition », *Data Mining and Knowledge Discovery*, vol. 2, 1998, p. 121-167; Matthias Rychetsky, *Algorithms and Architectures for Machine Learning based on Regularized Neural Networks and Support Vector Approaches*, Shaker Verlag, 2001.

problème linéaire lorsqu'on représente les données dans un espace de plus haute dimensionnalité. Cette astuce de « redescription » des données s'effectue à l'aide de méthodes noyaux (*kernel*) comme celle, polynomiale (espace de degré polynomial plus élevé que celui des variables d'origine), celle, sigmoïde ou celle de fonction à base radiale (*radial basis function* ou *RBF*). Entre autres, le SVM de type *RBF* génère de nouvelles caractéristiques (*features*) en mesurant la distance entre chacun des points ainsi qu'un centre (ou un ensemble de points) dans l'espace et prend habituellement la forme d'une fonction radiale gaussienne²⁵. Nous testons un SVM de type linéaire ainsi qu'un autre de type *RBF* que nous comparons à la classification naïve bayésienne de type gaussien²⁶, aux arbres de décision²⁷, aux forêts aléatoires²⁸ et aux réseaux de neurones artificiels multicouches²⁹ (deux couches cachées de 300 neurones chacune)³⁰. Enfin, nous validons nos modèles par la pertinence qui se conçoit et se mesure par la précision et le rappel. La précision correspond aux items correctement attribués à la classe cible par rapport à l'ensemble des items qui lui sont attribués, c'est-à-dire le nombre de vrais positifs divisé par le nombre de vrais positifs et de faux positifs. Le rappel correspond aux items correctement attribués à la classe cible par rapport à l'ensemble des items qui y appartiennent effectivement, c'est-à-dire le nombre de vrais positifs divisé par le nombre de vrais positifs et de faux négatifs. Une mesure combinatoire est la moyenne harmonique (*F-score*) où :

$$F - score = 2 * \frac{\text{précision} * \text{rappel}}{\text{précision} + \text{rappel}}$$

Une mesure similaire que nous utilisons est l'exactitude (*accuracy*), puisqu'elle tient compte des vrais négatifs et, ce faisant, elle permet de synthétiser la performance sur les cas positifs et négatifs :

25. John Shawe-Taylor et Nello Cristianini, *Kernel Methods for Pattern Analysis*, Cambridge University Press, New York, 2004.

26. T.F. Chan, G.H. Golub et R.J. LeVeque, « Updating Formulae and a Pairwise Algorithm for Computing Sample Variances », dans H. Caussinus, P. Ettinger, R. Tomassone, dir., *COMPSTAT 1982 5th Symposium held at Toulouse 1982*, Physica, Heidelberg, 1982.

27. Leo Breiman, Jerome Friedman, R. A. Olshen et Charles J. Stone, *Classification and regression trees*, Pacific Grove, Wadsworth & Brooks, 1984.

28. Leo Breiman, « Random forests », *Machine Learning*, vol. 45, n° 1, 2001, p. 5-32.

29. Geoffrey E. Hinton, « Connectionist learning procedures », *Artificial intelligence*, vol. 40, n° 1, 1989, p. 185-234.

30. Toutes les implémentations des modèles de classification utilisés dans cette étude sont tirées de la librairie *scikit-learn* en langage Python (Fabian Pedregosa, Gael Varoquaux, Alexandre Gramfort, Vincent Michel, Bertrand Thirion, Olivier Grisel, Mathieu Blondel, Peter Prettenhofer, Ron Weiss, Vincent Dubourg, Jake Vanderplas, Alexandre Passos, David Cournapeau, Matthieu Brucher et Edouard Duchesnay, « Scikit-learn: Machine Learning in Python », *JMLR*, vol. 12, 2011, p. 2825-2830.). Voir la documentation disponible au https://scikit-learn.org/dev/supervised_learning.html#supervised-learning.

$$Exactitude = \frac{\text{vrais positifs} + \text{vrais négatifs}}{\text{vrais positifs} + \text{vrais négatifs} + \text{faux positifs} + \text{faux négatifs}}$$

Résultats des modèles de classification

Le tableau suivant montre les mesures d’exactitude résultant de l’application de chacun de nos classifieurs sur les données de validation des représentations *BOW* ainsi que *DBOW*, c’est-à-dire les *SVM* de type linéaire (*SVM-LIN*), les *SVM* de type *RBF* (*SVM-RBF*), les classifieurs naïfs bayésiens de type gaussien (*NBG*), les arbres de décision (*AD*), les forêts aléatoires (*FA*) et les réseaux de neurones artificiels (*RNA*).

	<i>SVM-LIN</i>	<i>SVM-RBF</i>	<i>NBG</i>	<i>AD</i>	<i>FA</i>	<i>RNA</i>
<i>BOW</i>	0,95	0,95	0,89	0,83	0,92	0,95
<i>DBOW</i>	0,97	0,98	0,97	0,83	0,91	0,97

Sur la représentation *BOW*, la meilleure valeur d’exactitude est partagée entre les *SVM* et le *RNA*. Avec la représentation *DBOW*, la performance est optimale lorsqu’on utilise un *SVM* de type *RBF*. À l’exception des arbres de décision, les classifieurs sont également, sinon plus performants, lorsqu’on les applique à la représentation *DBOW*. Bref, le modèle le plus performant est un *SVM* de type *RBF* appliqué à une représentation de type *DBOW*. Afin de nous assurer que ces résultats ne soient pas un artefact de l’application préalable de notre mesure de distance entre les cas positifs et les cas négatifs sur la représentation *DBOW*, nous répétons l’expérience en déterminant le même nombre de cas négatifs à l’aide d’une mesure de distance cosinus sur la représentation *BOW*. Le tableau suivant montre les mesures d’exactitude suivant l’application des mêmes types de classifieurs sur les mêmes représentations *BOW* et *DBOW*.

	<i>SVM-LIN</i>	<i>SVM-RBF</i>	<i>NBG</i>	<i>AD</i>	<i>FA</i>	<i>RNA</i>
<i>BOW</i>	0,89	0,89	0,73	0,80	0,83	0,88
<i>DBOW</i>	0,86	0,89	0,83	0,75	0,83	0,88

À nouveau, on remarque que la performance est meilleure lorsqu’on utilise un *SVM*, mais cette performance est maintenant identique entre les représentations *BOW* et *DBOW* pour le type *RBF*, et celle-ci est meilleure sur *BOW* pour le type linéaire. La performance des arbres de décision est meilleure sur *BOW*. La performance des classifieurs naïfs bayésiens de type gaussien est meilleure sur *DBOW*, et les performances des forêts aléatoires et des *RNA* sont identiques entre *BOW* et *DBOW*. En somme, le type de

représentation sur laquelle est effectuée la mesure de distance initiale semble légèrement favoriser la performance suivant l'application des classifieurs sur le même type de représentation. Malgré tout, la mesure de distance initiale obtenue sur *DBOW* est optimale, dans la mesure où la performance de tous classifieurs est meilleure avec celle-ci comparativement à la performance obtenue à partir d'une mesure de distance initiale prise sur *BOW*.

Reconnaissance des cas positifs parmi les cas indéterminés avec le modèle optimal (h4)

L'évaluation d'un modèle computationnel à l'aide de mesures de pertinence est habituellement considérée comme suffisante si on ne se prononce pas sur la capacité du modèle à généraliser. En humanité numérique, notamment en philosophie, une telle évaluation, bien que nécessaire, semble insuffisante, dans la mesure où une tradition, irréductible à un ensemble de méthodes, se compose aussi d'applications bien réelles. Subséquemment, nous appliquons aux cas indéterminés le modèle optimal préalablement entraîné de manière à vérifier si l'apprentissage du modèle est généralisable (à ces cas) et observer si cet outil fait bien ce qu'il est censé faire. Tel quel, le modèle optimal permet de reconnaître 25 295 nouveaux cas positifs, ce qui semble beaucoup par rapport à un total de 38 968 cas indéterminés. Ce nombre élevé peut s'expliquer par le fait que les cas négatifs de départ ne sont pas reconnus en soi comme tels, mais par opposition aux cas positifs, de sorte qu'ils sont moins homogènes et moins bien représentés que les cas positifs. Une manière de résoudre ce problème est de biaiser le modèle en faveur de la classe problématique. Pour les *SVM*, cela peut s'effectuer en déplaçant l'hyperplan de manière à favoriser une classe plutôt qu'une autre³¹. Une fois l'hyperplan biaisé en faveur des cas négatifs, les nouveaux cas positifs se réduisent à un nombre plus raisonnable de 14 911 segments de texte.

À titre d'exemple, nous présentons les trois cas indéterminés reconnus comme exprimant le concept de **cognition** qui sont les plus similaires (mesure cosinus) aux cas positifs de départ. Notons que ceux-ci proviennent du même article de Jean-Nicolas Kaufmann intitulé « Critique du programme de naturalisation en philosophie de l'esprit »³². Dans cet article, Kaufman met à mal la théorie représentationnaliste de l'esprit (*TRE*) soutenue par

31. Notons qu'un déplacement de l'hyperplan implique une dégradation de la performance du modèle. Plus précisément, certains cas positifs limitrophes correctement reconnus comme tels ne le seront plus après le déplacement de l'hyperplan en faveur des cas négatifs. Par conséquent, on peut s'attendre à une amélioration de la précision du modèle au détriment du rappel pour les cas positifs. Autrement dit, le modèle ne détecte pas tous les cas positifs, car il peine à déterminer correctement les cas positifs limitrophes, mais ce qu'il détecte comme positif est rarement membre de l'ensemble des cas négatifs. À défaut d'être exhaustif, on peut considérer notre modèle comme une heuristique de découverte relativement précise.

32. J. Nicolas Kaufmann, « Critique du programme de naturalisation en philosophie de l'esprit », *Philosophiques*, vol. 35, n° 2, 2008, p. 483-512.

Dretske³³. Le premier cas est une discussion à propos d'une des sept thèses impliquées par la TRE que l'auteur numérote de 1 à 7. Dans cet extrait, on traite de cognition, dans la mesure où on peut la concevoir comme mettant en jeu des états mentaux. On y distingue deux types d'états mentaux (ou cognitions) où l'un est conscient et intentionnel et l'autre, infrapersonnel et possiblement non intentionnel :

La thèse [1], formulée de cette manière générale, ne semble pas faire problème. On peut naturellement se demander si tous les états mentaux ont de l'intentionnalité. On doit probablement distinguer les états mentaux conscients des états mentaux infrapersonnels. Les derniers pourraient être non intentionnels. La thèse [2] comporte des problèmes multiples³⁴.

Dans le deuxième cas, l'auteur expose les thèses 2, 3, 4 et 5 de la TRE. En bref, on définit la composante « intentionnelle » précédemment mentionnée en mettant en relation des états mentaux et des contenus, lesquels peuvent être de nature représentationnelle, symbolique, syntaxique, sémantique ou causalement dépendante du monde extérieur :

[2] L'intentionnalité consiste en une relation (à spécifier) à un contenu.

[3] Le contenu est de nature représentationnelle (TRE).

[4] Le contenu est de nature symbolique (Fodor).

a. Le contenu a une structure syntaxique (TSE)

b. Le contenu est sémantiquement évaluable (croyances vraies/fausses, etc.).

[5] Le contenu intentionnel dépend du monde extérieur : dépendance causale et nomologique (thèse externaliste)³⁵.

Enfin, dans le troisième cas, l'auteur expose un argument de Dretske à propos de la thèse 7 selon laquelle les états mentaux correspondent localement à des caractéristiques physiques intrinsèques. L'argument met en jeu l'hypothèse de l'efficacité causale des états intentionnels, c'est-à-dire que ceux-ci sont capables de causer des effets dans le monde externe à l'esprit ou à son support matériel :

L'efficacité causale dépend des caractéristiques intrinsèques des états intentionnels (thèse [7]). Elle ne peut pas être survenante sur les caractéristiques extrinsèques. Les propriétés externes (ou relationnelles) ne sont pas causalement pertinentes. Il s'agit d'un argument que Dretske considère pour le récuser. Il soutient que cet argument repose sur des confusions³⁶.

Sans surprise, les cas indéterminés reconnus comme positifs qui sont les plus similaires aux cas positifs de départ semblent effectivement porter sur le concept de **cognition**. Maintenant, voyons ce qu'il en est des cas trouvés positifs parmi les cas indéterminés et qui sont les moins susceptibles de

33. Fred Dretske, *Naturalising the Mind*, Cambridge MA, The MIT Press, 1995.

34. Kaufmann, « Critique du programme de naturalisation en philosophie de l'esprit », p. 486.

35. *Ibid.*, p. 485.

36. *Ibid.*, p. 489.

porter sur le concept cible, dans la mesure où ils sont les plus dissimilaires face aux cas positifs de départ. Le premier cas traite de l'influence méconnue de Karl Bühler sur la pensée de Karl R. Popper. La relation au concept de **cognition** est ténue, mais présente en ce qu'on fait mention de « fonction fondamentale de l'esprit ». L'auteure montre que Bühler, qu'on associe à tort à la Gestaltpsychologie, cherche à comprendre les fonctions cognitives permettant l'organisation perceptive et théorique du réel, ce qui poussera Popper à adopter une approche déductiviste mettant en jeu des solutions kantienne :

[...] il y eut une corrélation effective entre l'activité et les réflexions psychologiques menées sous la direction de Bühler et l'utilisation de la perspective kantienne de façon déductiviste. Cependant, [...] définir Karl Bühler comme un psychologue gestaltiste démontre que la connaissance de sa pensée n'est proportionnelle ni à sa célébrité ni à l'importance qu'on lui attribue [et] on ne peut pas ne pas donner raison à Gaetano Kanizsa lorsqu'il affirme que, dans la Gestaltpsychologie [...], sont en vigueur certains « lieux communs courants », qui souvent se révèlent être imprécis, voire faux dans certains cas; le fait qu'il ait adhéré à des solutions de type kantien est précisément l'une de ces simplifications: en effet, le fait que les principes proposés par l'école de Wertheimer aillent dans une direction totalement opposée à ce qu'une telle affirmation devrait signifier (c'est-à-dire la reconnaissance de la fonction fondamentale de l'esprit dans l'organisation perceptive et théorique de la donnée) fut même l'un des arguments pour lesquels elle fut critiquée par de très éminents censeurs, Bühler en tête. De ce point de vue, [...] l'interprétation de Kant par Popper se place au sein d'une conception dont l'inspiration de fond se posait en net contraste avec les positions adoptées par la Gestaltpsychologie de Wertheimer³⁷.

Le deuxième cas est issu d'un compte rendu à propos d'un texte sur la philosophie des sciences. Dans cet extrait, on présente la position des auteurs à propos de la relation entre l'intersubjectivité et la science. On soutient de manière plus ou moins explicite que des procédés individuels rationnels divergeant (de nature cognitive) se confrontent ou s'assemblent pour former des opérations sociales de nature argumentative ou critique, desquelles émerge le phénomène scientifique :

Cette dernière [l'intersubjectivité] peut être argumentative et discutoire, la libre discussion et la critique étant des gages de précision et de rigueur: le sujet qui a d'abord travaillé en solitaire est confronté à des points de vue qui stimulent sa créativité, lui font voir des aspects qui lui ont échappé. Cela suppose l'acceptation de certaines règles rationnelles, bien que l'argumentation n'empêche pas la tricherie ou l'anarchisme. Des compromis entre rationalité scientifique et rationalité politique ou économique sont toujours possibles,

37. Fiorenza Toccafondi, « De Karl Bühler à Karl R. Popper », *Philosophiques*, vol. 26, n° 2, 1999, p. 282.

bien que pas forcément souhaitables. L'intersubjectivité peut aussi s'inspirer de l'idéal husserlien du passage du je au nous sans sacrifice de transparence. C'est « la vigilance de la conscience présente à ses actes, sensible à la diversité phénoménologique de l'expérience, résistant à la robotisation du travail de recherche, insistant pour rester sujet vivant apte à relativiser ses propres tentatives de formalisation du réel naturel » (p. 221), ouverte aux réactions des autres³⁸.

Dans le troisième cas, François Recanati répond à une critique de Michel Seymour³⁹ à propos de son livre intitulé *Literal Meaning*⁴⁰. L'auteur se positionne par rapport à une variante extrême du contextualisme radical qui, en théorie de l'interprétation, se nomme *meaning eliminativism* (ME). Comme son nom l'indique, cette approche n'accorde pas de contenu sémantique aux énoncés, c'est-à-dire que le sens ne provient pas d'un réseau sémantique préalable, mais plutôt de « processus pragmatiques » ou d'une « opération qui constitue le sens ». Dans la mesure où de tels processus ou opérations sont cognitifs, on peut soutenir une relation entre cet extrait et le concept investigué :

Il n'y a guère qu'avec la plus extrême des formes du contextualisme radical, à savoir ME, que paraît difficilement surmontable la tension dont parle Seymour entre le caractère optionnel de la modulation (par opposition à la saturation) et son caractère nécessaire pour obtenir un contenu suffisamment déterminé. ME nie en effet que les mots possèdent autre chose qu'un [] potentiel sémantique que, dans mon livre, j'identifie à un ensemble de situations sources dans lesquelles ou à propos desquelles le mot a été légitimement employé. Ce potentiel sémantique étant fondamentalement différent d'un contenu déterminé, les processus pragmatiques qui conduisent à celui-ci à partir de celui-là sont obligatoires : il ne s'agit plus d'une opération facultative sur un sens primitif engendrant un nouveau sens « modulé », mais d'une opération qui constitue le sens []⁴¹.

Enfin, si les quelques exemples présentés et qui représentent les cas extrêmes sont en nombre insuffisant pour tirer des conclusions définitives, ceux-ci indiquent néanmoins que notre méthode permet la reconnaissance d'extraits portant sur le concept de **cognition** dans les textes sélectionnés de la revue *Philosophiques*. Sans pratiquer une investigation complète, il est possible de pousser un peu plus loin l'analyse conceptuelle afin de monter les possibilités de l'approche computationnelle pour les philo-

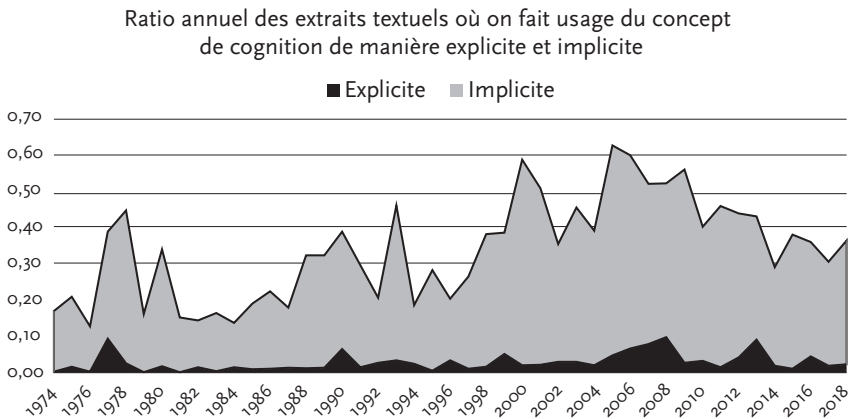
38. Maurice Gagnon, « Compte rendu de [Daniel Andler, Anne Fagot-Largeault et Bertrand Saint-Sernin, *Philosophie des sciences I et II*, Paris, Gallimard, collection Folio Essais, 2002, 1334 pages.] », *Philosophiques*, vol. 31, n° 1, 2004, p. 278.

39. Michel Seymour, « Le contextualisme sémantique en perspective : au sujet de *Literal Meaning*, de François Recanati », *Philosophiques*, vol. 33, n° 1, 2006, p. 249-262.

40. François Recanati, *Literal Meaning*, Cambridge, Cambridge University Press, 2004.

41. François Recanati, « Réponses à mes critiques », *Philosophiques*, vol. 33, n° 1, 2006, p. 283.

sophes. Lorsqu'on compare le nombre de cas positifs de départ (1 370) au nombre de cas positifs trouvés parmi les cas indéterminés (16 281), relativement à l'ensemble de tous les cas possibles (45 818), on saisit que ce qui est publié dans *Philosophiques* traite rarement de la cognition explicitement (~3 %), mais on observe que le concept est souvent présent de manière implicite dans les extraits textuels (~33 %) laissant présager un certain intérêt de la communauté à ce sujet. D'ailleurs, les résultats présentés diachroniquement au tableau suivant montrent que les usages implicites et explicites du concept de **cognition** connaissent une croissance en dents de scie, depuis la fondation de la revue jusqu'aux années 2005-2008. Depuis, on observe une légère, mais constante décroissance. On note une crête précoce autour de la publication d'un numéro intitulé « Philosophie et psychologie » en 1977. On observe un plateau autour de 2006, année où est publié un numéro intitulé « Philosophie et psychopathologie ». L'écart entre le minimum de 1976 (~13 %) et le maximum de 2005 (~62 %) ainsi que la valeur de 2018 (~36 %) indiquent l'ampleur de la croissance et de la décroissance des usages implicites et explicites du concept de cognition par la communauté philosophique publiant dans *Philosophiques*.



Conclusion

Somme toute, les résultats partiels présentés indiquent que l'approche computationnelle proposée fait ce qu'elle est censée faire, c'est-à-dire reconnaître dans un corpus donné des extraits textuels qui sont en relation avec un concept cible. Cette approche ne se prétend ni infaillible ni exhaustive, de sorte qu'elle se présente comme une heuristique de découverte relativement performante et pouvant contribuer au progrès de la connaissance

dans les différents domaines de la philosophie. Elle se démarque d'une approche traditionnelle de l'analyse conceptuelle en ce que sa méthode, fondée sur l'état de l'art en intelligence artificielle, est computationnellement explicite, reproductible, et n'est pas restreinte par l'étendue du corpus à analyser. Sous toute réserve, les résultats présentés indiquent que la relation conceptuelle est inégale et semble proportionnelle au degré de similitude entre les expressions explicites ou canoniques du concept (les cas positifs de départ) et les expressions implicites ou non canoniques du concept (les cas indéterminés reconnus comme positifs par le modèle optimal). En principe, spécifions qu'il est possible pour des locuteurs de la langue d'approximer la force de la relation entre le concept cible et les extraits textuels sélectionnés et, le cas échéant, il serait intéressant de comparer leurs résultats aux nôtres et d'évaluer plus en détail la performance du modèle proposé. Néanmoins, notons que ce dernier est empiriquement validé par une mesure de pertinence ainsi que par une brève évaluation de son application réelle. Bref, notre étude montre qu'il est effectivement possible de modéliser computationnellement l'analyse conceptuelle, partiellement du moins, par décomposition de cette tâche en sous-opérations dont, notamment, la détermination des extraits textuels dans lesquels un concept est exprimé. Cette modélisation peut s'effectuer de différentes manières et, ce faisant, elle implique des choix préalables déterminants pour la qualité des résultats ultérieurs. Parmi ces choix se trouvent celui du type de représentation du corpus sélectionné et celui du modèle de classification. Pour la tâche étudiée, une représentation par plongement telle que *DBOW* apparaît autant sinon plus performante que celle, classique, *BOW*, pour la majorité des modèles de classification. Autant sur *BOW* que sur *DBOW*, les modèles de classification de type *SVM* à noyaux *RGB* performant autant sinon mieux que tous les autres modèles de classification que nous avons comparés, c'est-à-dire les *SVM* de type linéaire, les naïfs bayésiens de type gaussien, les arbres de décision, les forêts aléatoires et les réseaux simples de neurones artificiels. De plus, les *SVM* offrent une solution de déplacement de l'hyperplan qui semble satisfaisante pour le traitement de nos données mal balancées, mais cela présuppose une approximation a priori du nombre de cas positifs à découvrir. Enfin, il serait intéressant de voir si les *SVM* performant aussi bien que d'autres modèles de classification non couverts par cette étude, notamment les réseaux de neurones artificiels plus complexes ou profonds, dans la mesure où ceux-ci semblent donner de bons résultats dans plusieurs domaines connexes⁴². Notons, cependant, que ces derniers nécessitent une grande quantité de données, et notre corpus de 1 476 articles possède une étendue qu'on pourrait considérer comme

42. Voir Nehad Mohamed Ibrahim, « Text Mining using Deep Learning », *International Journal of Scientific & Engineering Research*, vol. 9, n° 9, 2018.

pauvre pour approximer correctement une fonction complexe comme la reconnaissance d'extraits textuels exprimant un concept cible. Spécifions également que les réseaux de neurones artificiels ainsi que les classifications sur une représentation par plongement comme *DBOW* sont opaques, dans la mesure où ils ne permettent pas de reconnaître facilement les caractéristiques saillantes qui sont déterminantes pour la tâche qui nous intéresse.