



Actes de la 25e Conférence Nationale en Intelligence Artificielle

Éric Gaussier

► To cite this version:

Éric Gaussier. Actes de la 25e Conférence Nationale en Intelligence Artificielle. Plate-Forme Intelligence Artificielle, Association Française pour l'Intelligence Artificielle, 2022. hal-04564545

HAL Id: hal-04564545

<https://ut3-toulouseinp.hal.science/hal-04564545v1>

Submitted on 30 Apr 2024

HAL is a multi-disciplinary open access archive for the deposit and dissemination of scientific research documents, whether they are published or not. The documents may come from teaching and research institutions in France or abroad, or from public or private research centers.

L'archive ouverte pluridisciplinaire **HAL**, est destinée au dépôt et à la diffusion de documents scientifiques de niveau recherche, publiés ou non, émanant des établissements d'enseignement et de recherche français ou étrangers, des laboratoires publics ou privés.



Distributed under a Creative Commons Attribution - NonCommercial - NoDerivatives 4.0 International License



AfIA

Association française
pour l'Intelligence Artificielle

CNIA

Conférence Nationale en Intelligence Artificielle

PFIA 2022



Table des matières

Éric GAUSSIER

Éditorial	5
-----------------	---

Comité de programme.....	6
--------------------------	---

Sara Maqrot, Stéphanie Roussel, Gauthier Picard, Cédric Pralet

Techniques d'allocation de lots avec des préférences conflictuelles représentées par des graphes acycliques dirigés pondérés	7
---	---

G. Prévost, S. Cardon, T. Cazenave, C. Guettier, É. Jacopin

La planification SAS sous forme de tri topologique	17
--	----

Adrien Metge, Nicolas Maille, Benoît Le Blanc

Décision assistée par un système intelligent : de l'utilisation de l'IA à l'intention de l'opérateur dans la replanification	27
---	----

Alexandre Bettinger, Armelle Brun, Anne Boyer

Influence indépendante de l'exploration et de l'exploitation dans les métaheuristiques : application aux systèmes de recommandation	37
--	----

Éric Pardoux, Louis Devillaine

Vers une éthique processuelle de l'IA.....	47
--	----

Anaëlle Martin

L'intelligence artificielle peut-elle être saisie par le droit de l'Union européenne ?	57
--	----

Valentin Fouillard, Nicolas Sabouret, Safouan Taha, Frédéric Boulanger

Le problème du décor revisité : un modèle logique pour le diagnostic d'erreurs humaines	67
---	----

Martin, M. Donain, E. Fromont, T. Guns, L. Roze, A. Termier

Optimisation du Positionnement de Voitures en Autopartage basée sur la Prédiction de leur Utilité.....	76
---	----

R. Leclercq, G. Magnaval

Intelligence Artificielle & Génie Civil : enjeux et cas d'usages.....	86
---	----

H. Abdine, C. Xypolopoulos, M. Kamal Eddine, M. Vazirgiannis

Évaluation et Production de Plongements de Mots à Partir de Contenus Web Français à Grande Échelle	93
---	----

Olivier Gracianne, Anaïs Halftermeyer, Thi-Bich-Hanh Dao

Des clusters de tweets aux tags de descriptions : présentation d'un évènement par la caractérisation de ses manifestations.....	101
--	-----

Yanzhu Guo, Christos Xypolopoulos, Michalis Vazirgiannis

HowCOVID-19 is Changing Our Language : Detecting Semantic Shift in Twitter Word Embeddings	111
---	-----

T. Zgheib, H. Borges, V. Feldman, T. Guntz, C. Di Loreto, O. Desmaison, F. Corduant

Détection d'objets en temps réel : Entraînement de réseaux de neurones convolutifs sur images réelles et synthétiques.....	122
---	-----

Y. Zegaoui, P. Borianne, M. Chaumont, G. Subsol A. Seriai, M. Derras

Détection d'objets urbains dans des nuages de points LiDAR	130
--	-----

H. Farhat, L. Daniel, M. Benguigui, A. Girard

Modèle et jeu de données pour la détection multi-spectrale de feux de forêt à bord de satellites	140
---	-----

Éditorial

Conférence Nationale en Intelligence Artificielle

La Conférence Nationale en Intelligence Artificielle (CNIA), soutenue par le Conseil d'Administration de l'AFIA, s'adresse à l'ensemble de la communauté de recherche en IA. CNIA se veut un lieu privilégié pour faire connaître les dernières avancées en IA. Elle se veut aussi un forum destiné à renforcer les liens et les interactions entre les différentes sous-disciplines de l'IA et les disciplines faisant appel à l'IA. À ce titre, CNIA encourage les soumissions à la frontière entre sous-branches de l'IA, ainsi que les soumissions à la frontière de l'IA et d'autres disciplines.

Alors que l'IA se trouve aujourd'hui au cœur de nombreux développements, il est important d'avoir un forum qui réunisse l'ensemble des acteurs intéressés de près ou de loin par l'IA. L'objectif de CNIA est d'aborder à la fois les problématiques de recherche, les enjeux technologiques et les enjeux sociétaux liés à l'utilisation de l'IA, à travers l'ensemble des disciplines de l'IA :

- recherche heuristique et résolution de problèmes,
- incertitude et intelligence artificielle,
- logique, satisfiabilité et satisfaction de contraintes,
- apprentissage automatique,
- extraction, ingénierie et gestion des connaissances,
- représentation des connaissances et raisonnement,
- planification, contrôle,
- aide à la décision,
- causalité, agents autonomes et systèmes multi-agents,
- reconnaissance des formes et vision par ordinateur,
- traitement automatique des langues naturelles et de la parole, recherche d'information,
- interactions avec l'humain,
- perception et robotique,
- IA et web,
- environnements informatiques d'apprentissage humain et apprentissage à distance,
- IA responsable, IA de confiance (incluant explicabilité, certification, équité, ...),
- éthique de IA,
- droit et IA,
- IA et société,
- IA & X (X= santé, environnement, énergie, transport, défense, agriculture, matériaux, ...),
- ...

CNIA invite également :

- les soumissions d'articles présentant un panorama ou une synthèse d'un domaine récent qui ne soit ni trop étroit ni trop large ;
- les soumissions d'articles prospectifs présentant des idées et visions qui incitent la communauté à poursuivre de nouvelles voies de recherche (nouveaux problèmes, nouveaux domaines d'application, nouvelles méthodologies). Les auteurs doivent pouvoir convaincre l'audience que le sujet est intéressant et prometteur, et doivent le relier à l'état de l'art.

PFIA est une plateforme qui accueille d'autres conférences spécialisées sur certains de ces thèmes. Toute contribution relevant d'un de ces thèmes peut être soumise soit à la conférence spécialisée, soit à CNIA, mais pas aux deux.

Enfin, CNIA a accueilli « France@International » dédiée aux présentations d'avancées récentes de la communauté française acceptées dans les conférences majeures et générales de l'IA que sont AAAI 2022 et IJCAI-ECAI 2022.

Éric GAUSSIER

Comité de programme

Président

- Éric Gaussier, Université Grenoble Alpes.

Membres

- Isabelle Bloch, Sorbonne Université ;
- Olivier Boissier, École des Mines de Saint-Étienne ;
- Grégory Bonnet, Université de Caen Normandie ;
- Elise Bonzon, Université Paris Cité ;
- Robert Bossy, INRAE Centre de Jouy en Josas, MaIAGE ;
- Sylvie Coste-Marquis, Université d'Artois ;
- Benjamin Dalmas, École des Mines de Saint-Étienne ;
- Yves Demazeau, CNRS ;
- Arnaud Doniec, IMT Lille Douai ;
- Jean-Gabriel Ganascia, Sorbonne Université ;
- Gregor Goessler, Inria ;
- Salima Hassas, Université Lyon Claude Bernard ;
- Nathalie Hernandez, Université Toulouse Jean Jaurès ;
- Nicolas Lachiche, Université de Strasbourg ;
- Florence Le Ber, École Nationale du Génie de l'Eau et de l'Environnement de Strasbourg ;
- Marie-Jeanne Lesot, Sorbonne Université ;
- Engelbert Mephu Nguifo, Université Clermont Auvergne ;
- Fabrice Muhlenbach, Université de Saint-Étienne ;
- Marie-Christine Rousset, Université Grenoble Alpes ;
- Catherine Roussey, INRAE Centre de Clermont, TSCF ;
- Olivier Simonin, INSA de Lyon ;
- Catherine Tessier, ONERA ;
- Laurent Vercouter, INSA Rouen Normandie ;
- Bruno Zanuttini, Université de Caen Normandie.

Techniques d'allocation de lots avec des préférences conflictuelles représentées par des graphes acycliques dirigés pondérés

Sara Maqrot, Stéphanie Roussel, Gauthier Picard, Cédric Pralet

ONERA/DTIS, Université de Toulouse

prenom.nom@onera.fr

Résumé

Nous introduisons des techniques d'allocation de ressources pour un problème où (i) les agents expriment des demandes d'obtention de lots d'objets sous forme de graphes acycliques dirigés compacts pondérés (chaque chemin dans ces graphes est un lot dont l'évaluation est la somme des poids des arêtes traversées), et (ii) les agents ne demandent pas exactement les mêmes articles mais les demandes peuvent porter sur des objets conflictuels, qui ne peuvent pas être tous deux assignés. Ce cadre est motivé par des applications réelles telles que l'allocation de créneaux d'observation de la Terre, la virtualisation de fonctions réseaux ou la recherche de chemins multi-agents. Nous étudions plusieurs techniques d'allocation et analysons leurs performances sur un problème d'allocation de portions d'orbite.

Mots-clés

Allocation de ressources, approche utilitariste, partage équitable, graphes.

Abstract

We introduce resource allocation techniques for a problem where (i) the agents express requests for obtaining item bundles as compact edge-weighted directed acyclic graphs (each path in such graphs is a bundle whose valuation is the sum of the weights of the traversed edges), and (ii) the agents do not bid on the exact same items but may bid on conflicting items, that cannot be both assigned. This setting is motivated by real applications such as Earth observation slot allocation, virtual network functions, or multi-agent path finding. We study several allocation techniques and analyze their performances on an orbit slot ownership allocation problem.

Keywords

Resource allocation, utilitarianism, fair division, graphs.

1 Introduction

Nous considérons un problème d'allocation de lots d'articles indivisibles contraint par le chaînage d'articles (pour allouer à chaque agent une chaîne d'articles successifs) et des articles conflictuels. La première contrainte est capturée par un graphe acyclique dirigé (DAG) avec des arêtes

pondérées, un nœud source et un nœud puits. Ce graphe représente tous les lots (*i.e.* chemins) d'articles valides pour un agent, où la qualité d'un lot est représentée par les poids additifs des arêtes. La deuxième contrainte stipule que chaque agent doit obtenir un chemin sans conflit dans son graphe. On trouve ces problèmes d'allocations dans des domaines d'application tels que la virtualisation des fonctions de réseau (NFV), où les utilisateurs demandent l'allocation de graphes dirigés de services dans une infrastructure réseau partagée [20], ou dans l'observation de la Terre à l'aide d'une constellation de satellites, où les utilisateurs demandent l'exclusivité sur des portions d'orbite (sans chevauchement avec les portions d'autres utilisateurs) pour mettre en œuvre leurs demandes d'observation périodiques [11, 17]. Dans de tels contextes, outre les pondérations additives des arêtes, d'autres critères peuvent être pris en compte pour guider le processus d'allocation, en particulier lorsque les utilisateurs de la constellation sont des parties prenantes qui s'attendent à ce que les allocations soient équitables ou proportionnelles à leur investissement.

Travaux connexes. La littérature contient quelques travaux relatifs à l'allocation de biens structurés sous forme de graphes. Dans la division équitable de graphe, l'objectif est de diviser un graphe d'éléments entre plusieurs agents, avec des utilités additives attachées aux nœuds [1, 8]. Ces travaux fournissent des propriétés intéressantes pour trouver des allocations sans envie (*envy-free*) ou Pareto-optimales, de manière efficace dans certaines structures de graphes spécifiques, comme par exemple les chemins, les arbres, les étoiles. Cependant, dans notre problème, (i) les agents ne sont pas en compétition pour le même ensemble d'éléments, (ii) le graphe est dirigé pour composer des chemins d'un article de début à un article de fin, (iii) même en faisant correspondre notre problème à une division de graphe et en regroupant les éléments conflictuels en éléments composites, il est hautement improbable que le graphe résultant soit acyclique. Ici, les graphes sont utilisés pour exprimer les préférences, et non les biens à allouer. En bref, notre travail ne s'inscrit pas dans les cadres existants de partage équitable de graphes, ou ne peut pas bénéficier des résultats théoriques sur les graphes en forme de chemin ou d'étoile.

D'autres travaux connexes concernent les enchères de chemins [9, 6, 21], où les agents font des offres pour des chemins dans un graphe où chaque arête appartient à un

agent. L'objectif est d'attribuer des chemins aux agents par le biais d'enchères, et éventuellement de préserver une certaine confidentialité pour les propriétaires des arêtes. Dans le cas d'une fonction objectif utilitaire pour le problème de détermination du vainqueur, sans confidentialité des prix, cela entre dans le cadre de Vickrey-Clarke-Groves (VCG), et garantit donc des mécanismes efficaces et sans risque de confusion. Mais, là encore, les agents font des offres sur le même ensemble de nœuds et d'arêtes.

Dans le domaine du transport, les recherches sur des structures très similaires, à savoir les réseaux de flux, fournissent des techniques pour un flux maximal équitable dans les réseaux multi-sources et multi-puits [12]. Bien que les techniques utilisées soient très similaires aux nôtres (programmation linéaire), l'objectif de débit maximal est très différent de la maximisation de l'utilité du chemin. Par ailleurs, [7] a travaillé sur des problèmes de plus court chemin multiple basés sur des techniques de déconfliction. Bien que le problème présente des caractéristiques similaires, encore une fois les agents évoluent sur les mêmes graphes, et l'objectif est centré sur la minimisation de la longueur du chemin et la minimisation des chemins conflictuels, sans considérations d'équité.

Dans les jeux de congestion, les agents se voient attribuer des chemins de manière à minimiser les délais dus au croisement des chemins. Si des agents se voient attribuer les mêmes nœuds, alors le retard lié à leurs chemins augmente : [14, 16]. Dans notre travail, nous ne considérons pas le retard mais des incompatibilités. Bien que ces dernières puissent être modélisées par des fonctions non linéaires $\{0, \infty\}$, certaines allocations de chemins sont irréalisables dans notre problème, contrairement aux jeux de congestion. En outre, l'utilisation de méthodes de résolution de jeux de congestion comme dans [16] peut aboutir à des équilibres de Nash injustes, en raison de nombreux chemins irréalisables.

Plus généralement, une autre approche classique de l'allocation équitable de biens indivisibles est le *round-robin*, qui est presque sans envie [3]. C'est notamment une technique privilégiée pour allouer des fonctions réseau virtuelles dans les infrastructures de virtualisation de fonctions réseau [19], ou pour ordonnancer des tâches. Nous comparons nos techniques à cette approche.

Contributions. Cet article introduit et présente un modèle pour de tels scénarios, sous le prisme de l'optimalité et de l'équité. Le modèle capture toute application dans lesquelles les agents expriment leurs préférences sous forme de DAGs avec arêtes pondérées et dans lesquelles il existe des conflits entre certains nœuds de ces graphes. Nous montrons que ce problème d'allocation est NP-difficile. Nous présentons et évaluons plusieurs algorithmes sur des données provenant de constellations de satellites et de demandes simulées, selon trois critères : l'optimalité utilitaire, le temps de calcul et l'équité.

2 Modèle du problème

Nous étudions des problèmes d'allocation où les évaluations des agents sur les lots d'articles sont représentées comme des DAGs avec arêtes pondérées, avec la présence de certains conflits entre les nœuds de ces graphes. Une allocation consiste à choisir un chemin complet dans chaque graphe, de sorte que les chemins sélectionnés ne soient pas en conflit les uns avec les autres.

2.1 Définitions et notations

Définition 1. Un problème d'allocation de multiple chemins conflictuels dans des graphes acycliques dirigés avec arêtes pondérées (PADAG) est un tuple $\langle \mathcal{A}, \mathcal{G}, \mu, \mathcal{C} \rangle$, où :

- $\mathcal{A} = \{1, \dots, n\}$ est un ensemble d'agents ;
- $\mathcal{G} = \{g_1, \dots, g_m\}$ est un ensemble de DAGs pondérés, où chaque $g \in \mathcal{G}$ est un tuple $\langle V_g, E_g, u_g \rangle$ qui représente des préférences sur certains éléments dans V_g , avec des connexions entre les éléments dans $E_g \subset V_g \times V_g$, pondéré en utilisant la fonction d'utilité $u_g : V_g \times V_g \rightarrow \mathbb{R}$; nous supposons également que V_g contient deux nœuds spécifiques, la source s_g et le puits t_g , et que E_g contient une arête de s_g à t_g pondérée par l'utilité 0 (afin de traiter les cas où aucun lot d'éléments ne peut être sélectionné dans g) ;
- $\mu : \mathcal{G} \rightarrow \mathcal{A}$ fait correspondre chaque graphe g dans $\mathcal{G}$ à son propriétaire a dans $\mathcal{A}$;
- $\mathcal{C} \subset \{(v, v') | (v, v') \in V_g \times V_{g'}, g, g' \in \mathcal{G}^2, \mu(g) \neq \mu(g')\}$ est un ensemble de conflits entre des paires d'éléments de deux graphes distincts de $\mathcal{G}$ provenant de deux agents distincts.

Pour chaque graphe g et chaque ensemble d'arêtes $X \subseteq E_g$, l'utilité de X pour g est définie par $u_g(X) = \sum_{e \in X} u_g(e)$, ce qui signifie que les évaluations des arêtes sont considérées comme additives dans cet article. Par conséquent, chaque chemin de s_g à t_g dans un graphe g est évalué en additionnant les utilités des arêtes traversées, et chaque DAG représente de manière compacte un ensemble d'évaluations pour des paquets d'éléments, comme dans les enchères combinatoires. De plus, nous désignons par $\mathcal{G}_a = \mu^{-1}(a)$ l'ensemble des graphes possédés par l'agent a , et l'utilité d'un ensemble d'arêtes X pour l'agent a est définie par $u_a(X) = \sum_{g \in \mu^{-1}(a)} u_g(X \cap E_g)$.

Définition 2. Une allocation est une fonction π qui associe, à chaque graphe $g \in \mathcal{G}$, un chemin $\pi(g)$ de s_g à t_g dans g . Formellement, $\pi(g)$ peut être représentée comme un ensemble de nœuds dans V_g . En effet, comme on manipule des DAGs, il est facile de reconstruire les arêtes successivement traversées par le chemin à partir de cet ensemble. Par extension, l'allocation pour l'agent a est donnée par $\pi(a) = \bigcup_{g \in \mu^{-1}(a)} \pi(g)$.

Par convention, nous désignons par $u(\pi(g)) = u_{\mu(g)}(\pi(g))$ (resp. $u(\pi(a)) = u_a(\pi(a))$), l'utilité du graphe g

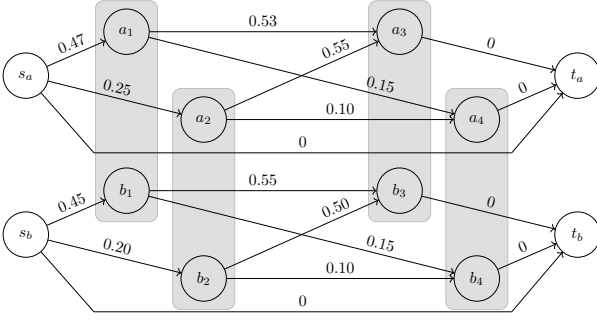


FIGURE 1 – Exemple d'évaluations (ou préférences) de lots par des agents représentées sous forme de DAGs. Les conflits sont représentés par des hypernœuds gris.

(resp. de l'agent a) pour l'allocation π . Enfin, l'utilité globale obtenue avec l'allocation π est donnée par $u(\pi) = \sum_{g \in \mathcal{G}} u(\pi(g))$ (ou de manière équivalente $u(\pi) = \sum_{a \in \mathcal{A}} u(\pi(a))$).

Définition 3. Une allocation π est *valide* si pour chaque paire de graphes distincts g et g' il n'y a pas de conflit entre les nœuds dans les chemins résultants, c'est-à-dire $(\pi(g) \times \pi(g')) \cap \mathcal{C} = \emptyset$.

Exemple 1. La figure 1 illustre un PADAG. Les meilleurs chemins pour les agents a et b sont respectivement $\{s_a, a_1, a_3, t_a\}$ et $\{s_b, b_1, b_3, t_b\}$, tous deux évalués à 1. En raison de conflits de nœuds (par exemple entre a_1 et b_1), ces deux chemins ne peuvent pas être assignés en même temps. Une allocation valide pourrait être $\pi_{\text{ex}} = \{a \mapsto \{s_a, a_2, a_4, t_a\}, b \mapsto \{s_b, b_1, b_3, t_b\}\}$ avec une utilité globale $u(\pi_{\text{ex}}) = u(\pi_{\text{ex}}(a)) + u(\pi_{\text{ex}}(b)) = 0.35 + 1.0 = 1.35$.

Les problèmes que nous considérons dans cet article sont (i) *comment calculer une allocation (utilitaire) optimale* π qui maximise $u(\pi)$, et (ii) *comment calculer une allocation équitable optimale* π , par le biais d'une optimisation lexicmin, c.a.d. maximisant lexicographiquement le vecteur d'utilité ordonné $(\Lambda_1, \dots, \Lambda_n)$ contenant une composante par agent. Plus précisément, ce vecteur contient les utilités de tous les agents. Formellement, il existe une bijection entre les utilités $u(\pi(i))$ des agents i de $\mathcal{A}$ et les éléments Λ_j du vecteur. De plus, le vecteur est ordonné, i.e. $\forall i \leq j, \Lambda_i \leq \Lambda_j$.

2.2 Analyse de complexité

Les deux propositions suivantes fournissent des résultats de complexité sur les deux problèmes étudiés, relatifs à l'optimisation utilitaire et à l'optimisation lexicmin respectivement.

Proposition 1. *Déterminer s'il existe une allocation valide π telle que l'évaluation utilitaire $u(\pi)$ est supérieure ou égale à une valeur donnée est NP-complet.*

Démonstration. Tout d'abord, le problème est NP puisque $u(\pi)$ est calculable en temps polynomial. Ensuite, il existe une réduction polynomiale de 3-SAT (qui est NP-complet) à notre problème. Fondamentalement, dans une

formule 3-SAT, chaque clause sur des variables propositionnelles x, y, z peut être représentée comme un DAG pondéré g où (1) l'ensemble des nœuds est $V_g = \{x, \neg x, y, \neg y, z, \neg z, s_g, t_g\}$, (2) l'ensemble des chemins de s_g à t_g dans g correspond à l'ensemble des valeurs de vérité pour x, y, z qui satisfont la clause (représentation du diagramme de décision), (3) le poids de chaque arête est fixé à 0 sauf pour les arêtes $s_g \rightarrow n$ où $n \neq t_g$ qui ont un poids de $1/m$, avec m le nombre de clauses dans la formule 3-SAT. Enfin, pour chaque variable propositionnelle x , nous pouvons ajouter un conflit (n, n') pour chaque paire de nœuds étiquetés par les littéraux x et $\neg x$ dans deux graphes distincts. Ensuite, comme un chemin est sélectionné dans chaque graphe et comme il y a m graphes, déterminer s'il existe une allocation valide π telle que $u(\pi) \geq 1$ est équivalent à trouver une solution qui satisfait toutes les clauses, d'où le résultat de NP-complétude étant donné que toutes les opérations utilisées dans la transformation sont polynomiales. $\square$

Proposition 2. *Il est NP-complet de décider s'il existe une allocation valide dont l'évaluation lexicmin est supérieure ou égale à un vecteur d'utilité donné. La proposition est valable même s'il existe un graphe unique par agent.*

Démonstration. Dans le cas général, il suffit de considérer un problème impliquant un agent unique possédant tous les graphes, et d'utiliser le résultat de la proposition précédente. S'il y a un seul graphe par agent, il suffit d'utiliser exactement le même codage 3-SAT que précédemment mais de remplacer les poids $1/m$ par des poids 1. Il est alors possible de montrer qu'il existe une allocation valide dont l'évaluation lexicmin est supérieure ou égale à $(1, 1, \dots, 1)$ s'il existe une solution au problème 3-SAT. De plus, l'évaluation lexicmin d'une allocation π peut être calculée en temps polynomial, d'où le résultat de NP-complétude. $\square$

2.3 À propos des conflits

Dans les PADAGs, les agents sont en compétition pour acquérir les nœuds reliés par l'ensemble $\mathcal{C}$. Les PADAGs reviennent à allouer chaque clique maximale K_i dans le graphe de conflit $(\bigcup_{g \in \mathcal{G}} V_g, \mathcal{C})$ à un agent. L'évaluation d'un agent pour un lot $\mathcal{B} = \{K_1, \dots, K_p\}$ de cliques maximales peut être définie comme l'utilité du meilleur chemin parmi tous les graphes de l'agent traversant un nœud de chaque K_i , mais ne traversant aucun nœud d'une clique en dehors de $\mathcal{B}$.

Pour résoudre ce problème, il est possible de s'appuyer sur des enchères combinatoires, où chaque agent enchère sur l'ensemble des cliques maximales qu'il souhaite (lots de cliques), et un commissaire-priseur détermine le gagnant d'une manière utilitaire ou équitable, par exemple. Cependant, il est à noter que (i) trouver toutes les cliques maximales d'un graphe est $\mathcal{O}(3^{\frac{n}{2}})$ dans le pire des cas [2], et (ii) trouver une allocation maximisant le bien-être dans les enchères combinatoires est NP-difficile, en général [5]. De plus, l'évaluation de tels lots nécessite que les enchérisseurs connaissent les conflits, ce qui peut ne pas être une information publique. Par conséquent, nous considérons dans la

suite que les agents expriment leurs demandes (ensuite traduites en DAGs), sans savoir avec quels agents leurs offres sont en conflit. Seul le coordinateur (par exemple, l'opérateur de réseau ou l'opérateur de constellation) détermine ces conflits.

2.4 À propos des préférences et des comportements stratégiques

Alors que les approches leximin et utilitariste (dans un mécanisme VCG) sont à l'abri des stratégies (*strategyproof*), le fait que les agents définissent directement leur évaluation peut être problématique : un agent peut définir un graphe avec un seul chemin, le seul lot qu'il accepte. En utilisant une approche leximin, il est très probable qu'on lui attribue ce seul lot, pour éviter une utilité de 0 pour le pire agent. Cette situation pourrait être évitée dans certains contextes. Par exemple, les agents devraient seulement être autorisés à demander des articles, et ne pas exprimer directement les graphes. Le coordinateur pourrait ainsi générer des graphes connexes, qui ont de nombreux chemins possibles, en raison de la configuration du système. De plus, dans de nombreux contextes opérationnels, plusieurs utilisateurs peuvent avoir certaines priorités (par exemple, les agences de défense dans le cadre de l'observation par des satellites), ce qui pourrait nécessiter que certaines demandes aient des poids plus élevés et de désactiver la normalisation, ou d'adopter une approche d'allocation itérative, par niveau de priorités. Ainsi, dans la suite du document, nous ne considérerons que les agents ayant la même priorité et qui ne définissent pas les graphes par eux-mêmes.

3 Schémas d'allocation de chemins

Nous proposons ici plusieurs schémas d'allocation pour les PADAGs. Certains d'entre eux sont basés sur la programmation linéaire en nombres entiers (ILP) et la programmation linéaire en nombres entiers mixtes (MILP). Nous commençons par introduire les variables de décision et les contraintes pour ces modèles.

Pour tout DAG $g = \langle V_g, E_g, u_g \rangle$, nous définissons des variables binaires $x_e \in \{0, 1\}$, pour tout $e \in E_g$, indiquant si l'arête e est sélectionnée dans le chemin définissant le lot solution. Nous utilisons également des variables binaires auxiliaires β_v indiquant si le nœud v est sélectionné dans le chemin de solution $\pi(g)$, c'est-à-dire que $\beta_v = 1$ si $v \in \pi(g)$, 0 sinon. Pour tout nœud v dans V_g , nous désignons par $\text{In}(v)$ (resp. $\text{Out}(v)$) son ensemble d'arêtes entrantes (resp. sortantes). Dans tous les modèles ILP introduits par la suite, nous imposons les contraintes (1)–(3) pour définir tous les chemins possibles, les contraintes (4)–(5) pour tenir compte des conflits de sélection d'éléments, et la contrainte (6) pour s'assurer que les sources et les puits sont sélectionnés.

$$\sum_{e \in \text{In}(v)} x_e = \sum_{e \in \text{Out}(v)} x_e, \quad \forall g \in \mathcal{G}, \forall v \in V_g \setminus \{s_g, t_g\} \quad (1)$$

$$\sum_{e \in \text{Out}(s_g)} x_e = 1, \quad \forall g \in \mathcal{G} \quad (2)$$

$$\sum_{e \in \text{In}(t_g)} x_e = 1, \quad \forall g \in \mathcal{G} \quad (3)$$

$$\sum_{e \in \text{In}(v)} x_e = \beta_v, \quad \forall g \in \mathcal{G}, \forall v \in V_g \setminus \{s_g, t_g\} \quad (4)$$

$$\sum_{v \in c} \beta_v \leq 1, \quad \forall c \in \mathcal{C} \quad (5)$$

$$\beta_{s_g} = \beta_{t_g} = 1, \quad \forall g \in \mathcal{G} \quad (6)$$

3.1 Allocation utilitariste (util)

L'approche classiquement utilisée en allocation de ressources est l'approche utilitaire. Elle consiste à trouver l'allocation qui maximise la somme des utilités de tous les chemins sélectionnés. Cela correspond à la résolution du programme linéaire en nombres entiers $P_{\text{util}}(\langle \mathcal{A}, \mathcal{G}, \mu, \mathcal{C} \rangle)$ donné ci-dessous :

$$\begin{aligned} \max \quad & \sum_{a \in \mathcal{A}} \sum_{g \in \mathcal{G}_a} \sum_{e \in E_g} u_g(e) \cdot x_e \\ \text{t.q.} \quad & (1), (2), (3), (4), (5), (6) \end{aligned} \quad (7)$$

L'allocation résultante π est décodée à partir des variables β_v . Formellement, pour tout $g \in \mathcal{G}$, $\pi(g) = \{v \in V_g \mid \beta_v = 1\}$.

Exemple 2. Dans la figure 1, l'allocation utilitaire optimale est $\pi_{\text{util}} = \{a \mapsto \{s_a, a_2, a_3, t_a\}, b \mapsto \{s_b, b_1, b_4, t_b\}\}$, avec une utilité globale $u(\pi_{\text{util}}) = u(\pi_{\text{util}}(a)) + u(\pi_{\text{util}}(b)) = 0.80 + 0.60 = 1.40$.

3.2 Allocation leximin (lex)

Au-delà de l'utilitarisme, une façon de mettre en œuvre une allocation équitable et Pareto-optimale est de considérer la règle *leximin* qui sélectionne, parmi toutes les allocations possibles, une allocation conduisant aux meilleurs profils d'utilité par rapport à l'ordre leximin [13]. Plus précisément, soit $z = [z_1, \dots, z_n]$ le vecteur d'utilité où chaque composante $z_a \in [0, Z_a]$ représente l'utilité pour l'agent $a \in \mathcal{A}$. Z_a désigne ici la meilleure valeur d'utilité pour l'utilisateur a considéré seul, c'est-à-dire pour le problème mono-agent où le meilleur chemin peut être choisi pour chaque graphe $g \in \mathcal{G}_a$. En optimisation selon le leximin, l'objectif est de maximiser lexicographiquement le vecteur $\Lambda = [\Lambda_1, \dots, \Lambda_n]$ obtenu après avoir ordonné $[z_1, \dots, z_n]$ suivant un ordre croissant.

Cette règle leximin peut être mise en œuvre par une séquence de programmes linéaires [10]. Nous adaptons ici une telle procédure au cas spécifique des PADAGs. Supposons que nous ayons déjà optimisé sur les premières $K - 1$ composantes $[\Lambda_1, \dots, \Lambda_{K-1}]$ de Λ , pour $K \in [1..n]$. Ensuite, on peut utiliser le programme présenté par la suite pour optimiser la $K^{\text{ième}}$ composante Λ_K du profil leximin de z . Dans ce modèle, λ représente l'utilité obtenue au niveau K dans Λ , avec $\lambda \in [\Lambda_{K-1}, \max_{a \in \mathcal{A}} Z_a]$ et par convention $\Lambda_0 = 0$. y_{ak} est une variable binaire égale à

1 si l'agent $a \in \mathcal{A}$ joue le rôle de l'agent associé au niveau $k \in [1..K-1]$ dans $[\Lambda_1, \dots, \Lambda_{K-1}]$, 0 sinon. L'optimisation de Λ_K peut être effectuée à l'aide du programme $P_{\text{lex}}(\langle \mathcal{A}, \mathcal{G}, \mu, \mathcal{C} \rangle, K, [\Lambda_1, \dots, \Lambda_{K-1}])$ donné ci-dessous :

$$\max \quad \lambda \quad (8)$$

$$\text{t.q.} \quad (1), (2), (3), (4), (5), (6)$$

$$z_a = \sum_{g \in \mathcal{G}_a} \sum_{e \in E_g} u_g(e) \cdot x_e, \quad \forall a \in \mathcal{A} \quad (9)$$

$$\sum_{a \in \mathcal{A}} y_{ak} = 1, \quad \forall k \in [1..K-1] \quad (10)$$

$$\sum_{k \in [1..K-1]} y_{ak} \leq 1, \quad \forall a \in \mathcal{A} \quad (11)$$

$$\lambda \leq z_a + M \sum_{k \in [1..K-1]} y_{ak}, \quad \forall a \in \mathcal{A} \quad (12)$$

$$z_a \geq \sum_{k \in [1..K-1]} \Lambda_k \cdot y_{ak}, \quad \forall a \in \mathcal{A} \quad (13)$$

Dans la contrainte (12), $M = \max_{a \in \mathcal{A}} Z_a$ est utilisé pour ignorer les agents associés à des niveaux strictement inférieurs à K lors de l'optimisation de λ (formulation big-M). La contrainte (13) garantit que l'utilité obtenue pour l'agent associé au niveau $k \in [1..K-1]$ ne soit pas être inférieure à Λ_k . Pour mettre en œuvre la règle du leximin, il suffit alors de résoudre une séquence de problèmes P_{lex} pour $K \in \mathcal{A}$ afin d'optimiser la valeur de chaque composante du profil d'utilité.

Algorithme 1 : Algorithme leximin (lex)

Données : Un problème PADAG $\langle \mathcal{A}, \mathcal{G}, \mu, \mathcal{C} \rangle$

Résultat : Une allocation leximin-optimale π

```

1 pour  $K = 1$  à  $|\mathcal{A}|$  faire
2    $(\lambda^*, \text{sol}) \leftarrow$ 
     résoudre  $P_{\text{lex}}(\langle \mathcal{A}, \mathcal{G}, \mu, \mathcal{C} \rangle, K, [\Lambda_1, \dots, \Lambda_{K-1}])$ 
3    $\Lambda_K \leftarrow \lambda^*$ 
4 pour  $g \in \mathcal{G}$  faire
5    $\pi(g) \leftarrow \{v \in V_g \mid \text{sol}(\beta_v) = 1\}$ 
6 retourner  $\pi$ 
    
```

Exemple 3. L'allocation leximin optimale pour l'exemple de la figure 1 est $\pi_{\text{lex}} = \{a \mapsto \{s_a, a_1, a_4, t_a\}, b \mapsto \{s_b, b_2, b_3, t_b\}\}$, avec une utilité globale $u(\pi_{\text{lex}}) = u(\pi_{\text{lex}}(a)) + u(\pi_{\text{lex}}(b)) = 0.62 + 0.70 = 1.32$ et le vecteur d'utilité $(0.62, 0.70)$.

3.3 Allocation leximin approchée (a-lex)

Le modèle précédent met en œuvre une règle leximin exacte et assure donc l'équité de l'allocation résultante, mais il peut difficilement passer à l'échelle avec l'augmentation des tests classiques de répartition équitable (proportionnalité, absence d'envie...) est impossible puisque les agents n'enchérissent pas sur le même ensemble d'articles. Nous proposons ainsi une version approximative du calcul du leximin, basée sur un schéma maximin itéré. Fondamentalement, cette approche considère à chaque étape une utilité

minimale $\Delta_a \geq 0$ pour certains agents et maximise la pire utilité parmi les agents restants, pour lesquels nous supposons arbitrairement $\Delta_a = -1$. Le problème à résoudre, appelé $P_{\text{a-lex}}(\langle \mathcal{A}, \mathcal{G}, \mu, \mathcal{C} \rangle, \Delta)$, est le suivant :

$$\max \quad \delta \quad (14)$$

$$\text{t.q.} \quad (1), (2), (3), (4), (5), (6)$$

$$\delta \leq \sum_{g \in \mathcal{G}_a} \sum_{e \in E_g} u_g(e) x_e, \quad \forall a \in \mathcal{A} \mid \Delta_a = -1 \quad (15)$$

$$\sum_{g \in \mathcal{G}_a} \sum_{e \in E_g} u_g(e) x_e \geq \Delta_a, \quad \forall a \in \mathcal{A} \mid \Delta_a \neq -1 \quad (16)$$

La méthode de résolution consiste alors à optimiser $P_{\text{a-lex}}$ de manière itérative, comme pour leximin. Comme l'indique l'algorithme 2, à chaque itération (une par agent), $P_{\text{a-lex}}$ est résolu, un pire agent $\hat{a}$ est déterminé, et son utilité minimale $\Delta_{\hat{a}}$ est fixée.

Algorithme 2 : Algorithme leximin approché (a-lex)

Données : Un problème PADAG $\langle \mathcal{A}, \mathcal{G}, \mu, \mathcal{C} \rangle$

Résultat : Une allocation maximin-optimale π

```

1  $\Delta \leftarrow [-1, \dots, -1]$ 
2 pour  $K = 1$  à  $|\mathcal{A}|$  faire
3    $(\delta^*, \text{sol}) \leftarrow$  résoudre  $P_{\text{a-lex}}(\langle \mathcal{A}, \mathcal{G}, \mu, \mathcal{C} \rangle, \Delta)$ 
4    $S \leftarrow \underset{a \in \mathcal{A} \mid \Delta_a = -1}{\text{argmin}} \sum_{g \in \mathcal{G}_a} \sum_{e \in E_g} u_g(e) \text{sol}(x_e)$ 
5    $\hat{a} \leftarrow$  choisir un agent  $a \in S$ 
6    $\Delta_{\hat{a}} \leftarrow \delta^*$ 
7 pour  $g \in \mathcal{G}$  faire
8    $\pi(g) \leftarrow \{v \in V_g \mid \text{sol}(\beta_v) = 1\}$ 
9 retourner  $\pi$ 
    
```

La principale différence avec P_{lex} est qu'à chaque itération, dans $P_{\text{a-lex}}$, la position d'un agent dans l'ordre est déterminée une fois pour toutes, alors que dans P_{lex} l'ordre peut être révisé à chaque itération. De plus, si une égalité se produit à la ligne 5 pour déterminer le pire agent (cas $|S| > 1$), on peut se fier à une certaine heuristique ou à un choix arbitraire. Ainsi, $P_{\text{a-lex}}$ est une approximation de P_{lex} qui contient moins de variables et de contraintes.

Exemple 4. L'allocation leximin approchée pour l'exemple de la figure 1 est $\pi_{\text{a-lex}} = \{a \mapsto \{s_a, a_1, a_4, t_a\}, b \mapsto \{s_b, b_2, b_3, t_b\}\}$, avec le vecteur d'utilité $(u(\pi_{\text{a-lex}}(a)), u(\pi_{\text{a-lex}}(b))) = (0.62, 0.70)$ et une utilité globale $u(\pi_{\text{a-lex}}) = 0.62 + 0.70 = 1.32$. C'est la même chose que π_{lex} puisqu'il n'y a que deux agents et aucune égalité entre les pires utilités.

3.4 Allocation gloutonne (greedy)

Pour les décisions très rapides, le maximin itéré peut encore être trop lent. Dans de tels cas, une approche gloutonne peut fournir des allocations valides très rapidement. L'idée principale de l'allocation de chemins par approche gloutonne est d'itérer sur l'ensemble des graphes. À chaque étape, un

graphe g^* qui a le meilleur chemin d'utilité est sélectionné, ce chemin est choisi comme $\pi(g^*)$, et tous les nœuds des autres graphes qui sont en conflit avec les nœuds de $\pi(g^*)$ sont désactivés. Le graphe g^* est alors supprimé, et le processus continue jusqu'à ce qu'il n'y ait plus de graphe à considérer. Ce processus garantit le respect des contraintes (1), (2), (3), (4), (5), (6).

Ce processus est décrit dans l'algorithme 3.

Algorithme 3 : Algorithme glouton (greedy)

Données : Un problème PADAG $\langle \mathcal{A}, \mathcal{G}, \mu, \mathcal{C} \rangle$

Résultat : Une allocation π

```

1 tant que  $\mathcal{G} \neq \emptyset$  faire
2   déterminer  $g^*$  le graphe d'utilité maximale avec le
   chemin  $p$ 
3    $\pi(g^*) \leftarrow p$ 
4   pour  $g \in \mathcal{G}$  faire
5      $V_g \leftarrow \{v \in V_g \mid \forall w \in \pi(g^*), \{v, w\} \notin \mathcal{C}\}$ 
6      $E_g \leftarrow \{(v, w) \in E_g \mid v \in V_g, w \in V_g\}$ 
7    $\mathcal{G} \leftarrow \mathcal{G} \setminus \{g^*\}$ 
8 retourner  $\pi$ 

```

Déterminer le meilleur chemin dans un DAG g est linéaire en temps : $\mathcal{O}(|E_g| + |V_g|)$ [4]. De toute évidence, la méthode gloutonne est équivalente à la méthode utilitaire lorsqu'il n'y a pas de conflit entre les graphes. En effet, greedy retournera le meilleur chemin pour chaque graphe, ce qui est la meilleure solution utilitaire dans un tel contexte. De plus, cette approche aboutit à un équilibre de Nash où aucun agent ne peut améliorer son utilité sans un impact négatif sur les autres agents. Ceci est équivalent à la procédure *Nashify* de [16] dans le contexte des jeux de congestion, avec un seul tour. Nous verrons dans les expériences que cet équilibre est loin d'être équitable.

Exemple 5. L'allocation gloutonne pour l'exemple de la figure 1 est $\pi_{\text{greedy}} = \{a \mapsto \{s_a, a_1, a_3, t_a\}, b \mapsto \{s_b, b_2, b_4, t_b\}\}$, avec une utilité globale $u(\pi_{\text{greedy}}) = u(\pi_{\text{greedy}}(a)) + u(\pi_{\text{greedy}}(b)) = 1.0 + 0.3 = 1.3$ et le vecteur d'utilité $(1.0, 0.3)$.

3.5 Allocations *round-robin* (p-rr et n-rr)

Une approche rapide pour l'allocation équitable de biens indivisibles est le *round-robin*. Elle consiste à faire en sorte que chaque agent choisisse à tour de rôle (dans un ordre fixe prédéfini) un élément (en fonction de ses préférences) jusqu'à ce qu'il n'y ait plus d'élément à allouer. Comme greedy, il est polynomial en nombre d'agents et d'articles. Dans notre cas, nous pouvons considérer deux types d'articles à allouer : les chemins (approches notée p-rr) ou les nœuds (approche notée n-rr). Dans le cas des chemins, chaque agent choisit à son tour le meilleur chemin possible, compte tenu des nœuds déjà alloués (pour éviter les conflits). Ce processus fonctionne de manière similaire à celui de greedy, mais alterne entre les utilisateurs pour équilibrer les utilités. Dans le cas des nœuds, chaque agent

construit de manière incrémentale le chemin associé à chacun de ses graphes, en choisissant tour à tour son meilleur prochain nœud réalisable jusqu'à ce qu'il atteigne le puits ou qu'il n'y ait plus de nœud réalisable à choisir (un chemin sans issue). Dans ce dernier cas, l'agent se voit attribuer le chemin source-puits d'utilité 0 et perd les nœuds précédemment choisis. Dans les deux approches, les contraintes (1), (2), (3), (4), (5), (6) sont respectées puisque les chemins considérés sont tous faisables. Notons que p-rr résulte en un équilibre de Nash où chaque agent s'est vu allouer le meilleur chemin compte tenu des autres allocations. Ce n'est pas le cas pour n-rr, puisque certains nœuds laissés par un agent tombant dans une impasse peuvent avoir empêché d'autres agents de trouver une meilleure solution.

Exemple 6. L'allocation *round-robin* par chemin $\pi_{\text{p-rr}}$ pour l'exemple de la figure 1 est équivalente à π_{greedy} , puisque a choisit $\{s_a, a_1, a_3, t_a\}$ et ensuite b choisit $\{s_b, b_2, b_4, t_b\}$. L'allocation *round-robin* par nœud $\pi_{\text{n-rr}}$ est également équivalente à π_{greedy} car a choisit d'abord a_1 , puis b choisit b_2 (seule option réalisable), puis a choisit a_3 (meilleure option), et enfin b choisit b_4 (seule option réalisable).

4 Evaluation expérimentale

Dans cette section, nous évaluons les performances des méthodes d'allocation proposées pour résoudre un problème d'allocation de portions d'orbite¹ codé en PADAGs. Nous présentons le dispositif expérimental et analysons quelques résultats obtenus sur des instances synthétiques.

4.1 Scénario d'évaluation

Comme décrit dans [18], les constellations de satellites d'observation de la Terre soulèvent de nombreux défis. Nous abordons ici celui dans lequel des utilisateurs peuvent demander a priori l'exclusivité sur des portions d'orbite, afin de disposer de temps satellite pour réaliser des prises de vue sans avoir à passer systématiquement par l'arbitrage d'un opérateur central. L'objectif de notre étude est d'allouer des portions d'orbite aux utilisateurs, sachant que chacun a des requêtes d'observation définies par un point d'intérêt (POI) à observer avec une fréquence de revisite donnée, par exemple observer la ville de Saint-Étienne toutes les 2 heures pendant plusieurs jours. Étant donné que plusieurs satellites peuvent capturer le même point sur la Terre autour des plots temporels définis, plusieurs lots sont spécifiés par chaque utilisateur, qui se valorisent différemment en fonction de la qualité de la séquence des portions d'orbite, par exemple la proximité des plots temporels ou les angles d'acquisition. En outre, plusieurs utilisateurs peuvent être intéressés par des points d'intérêt très proches, ce qui entraîne un chevauchement des portions d'orbite qui ne peuvent être attribués simultanément aux utilisateurs correspondants.

¹. qui est un problème assez récent identifié mais non formalisé dans [17], à ne pas confondre avec les problèmes d'ordonnancement d'orbite [11].

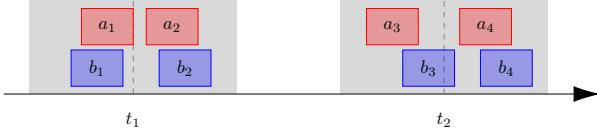


FIGURE 2 – Un problème d'allocation de portions d'orbites impliquant 2 agents (a en rouge, b en bleu) demandant des portions d'orbite autour de deux plots temporels (t_1 et t_2), avec des fenêtres de tolérance autour de chaque plot (en gris) et deux portions candidates pour chaque plot ($a_1, \dots, a_4, b_1, \dots, b_4$).

4.2 Cadre expérimental

Nous considérons une constellation en orbite basse (500 km d'altitude) composée de n_p plans orbitaux régulièrement espacés ayant une inclinaison de 60 degrés, avec $n_p \in \{2, 4, 8, 16\}$ et 2 satellites régulièrement espacés sur chaque plan orbital. Pour générer des instances PADAG, nous générons aléatoirement des requêtes pour 4 agents souhaitant obtenir la propriété de portions d'orbite afin de réaliser des acquisitions au sol répétitives de points d'intérêt (POI) appartenant à la même zone. Les POIs sont choisis aléatoirement dans un sous-ensemble extrait de [15]. Tous les agents ont le même modèle pour une requête r : obtenir une observation tous les jours à $8:00 + \delta_r$, $12:00 + \delta_r$, et $16:00 + \delta_r$, avec une tolérance de 1 heure autour de chaque plot temporel, et un décalage aléatoire uniforme $\delta_r \in [-2, 2]$ pour tous les plots temporels de la même requête. Pour chaque POI et chaque plot temporel, les portions d'orbite candidates sont déterminées grâce à une librairie logicielle de mécanique spatiale, en partant de l'hypothèse qu'un satellite est pertinent pour un POI dès que son élévation au-dessus de l'horizon est supérieure à 15 degrés. Les portions d'orbite incompatibles sont celles qui se chevauchent alors qu'elles appartiennent au même satellite.

Nous formulons ensuite ces requêtes et les portions d'orbite candidates sous forme de PADAGs. Chaque requête est représentée par un graphe, dans lequel les nœuds (à l'exception de la source et du puits) sont des portions d'orbite pour capturer un POI à un moment donné, et les arêtes relient deux portions d'orbite consécutives pour répondre à la requête. Par exemple, la figure 1 représente un PADAG pour deux requêtes de deux utilisateurs (a et b), avec deux plots temporels. Chaque utilisateur a deux portions d'orbite candidates par plot temporel, illustré en figure 2. Pour simplifier, nous ne considérons que les utilités attachées aux portions d'orbite, sans prendre en considération les transitions entre elles. Nous étudions une fonction d'utilité *linéaire*, qui est linéaire en fonction de la distance entre le milieu τ de la portion allouée et le plot temporel demandé (utilité linéairement décroissante de 1 lorsque τ est exactement sur le plot à 0 lorsque τ atteint les limites de la fenêtre de tolérance). Nous normalisons chaque utilité par rapport à l'utilité maximale qui peut être obtenue pour chaque graphe individuellement. Nous considérons 2 requêtes par agent, et un horizon de 365 jours, ce qui donne des DAGs ayant 1095 couches. Ce paramètre donne des DAGs avec les propriétés suivantes en moyenne :

n_p	2	4	8	16
largeur	3.08	5.41	10.05	19.38
conflits	26798.80	45636.06	82971.20	158180.20
durée (s)	603.28	600.10	599.87	598.75

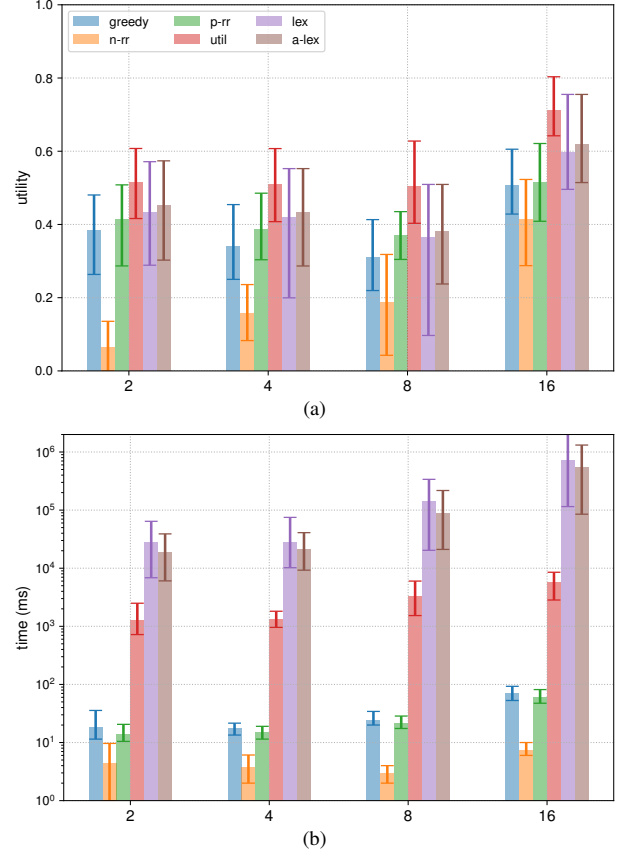


FIGURE 3 – Utilité globale moyenne (en haut) et temps de calcul (en bas) obtenus pour chaque algorithme et chaque taille de constellation.

En bref, la largeur du DAG (nombre de nœuds par couche) et le nombre de conflits augmentent proportionnellement avec le nombre de satellites. La durée des portions d'orbite est d'environ 10 minutes.

Les solveurs sont codés en Java 1.8 et exécutés sur un CPU Intel(R) Xeon(R) E5-2660 v3 @ 2.60GHz à 20 cœurs, 62GB RAM, Ubuntu 18.04.5 LTS. util, lex et a-lex utilisent l'API Java de IBM CPLEX 20.1 (avec un temps limite de 10 minutes). Nous avons exécuté sur 30 instances de PADAGs générés aléatoirement et tracé la moyenne avec une confiance de $[0.05, 0.95]$.

4.3 Utilité

La figure 3a montre l'utilité globale normalisée moyenne pour chaque algorithme et chaque taille de constellation (exprimée en nombre de plans orbitaux). L'utilité globale normalisée est l'utilité moyenne du graphe, donc entre 0 et 1. De toute évidence, util fournit l'allocation utilitaire optimale. En deuxième position, a-lex fournit de bonnes allo-

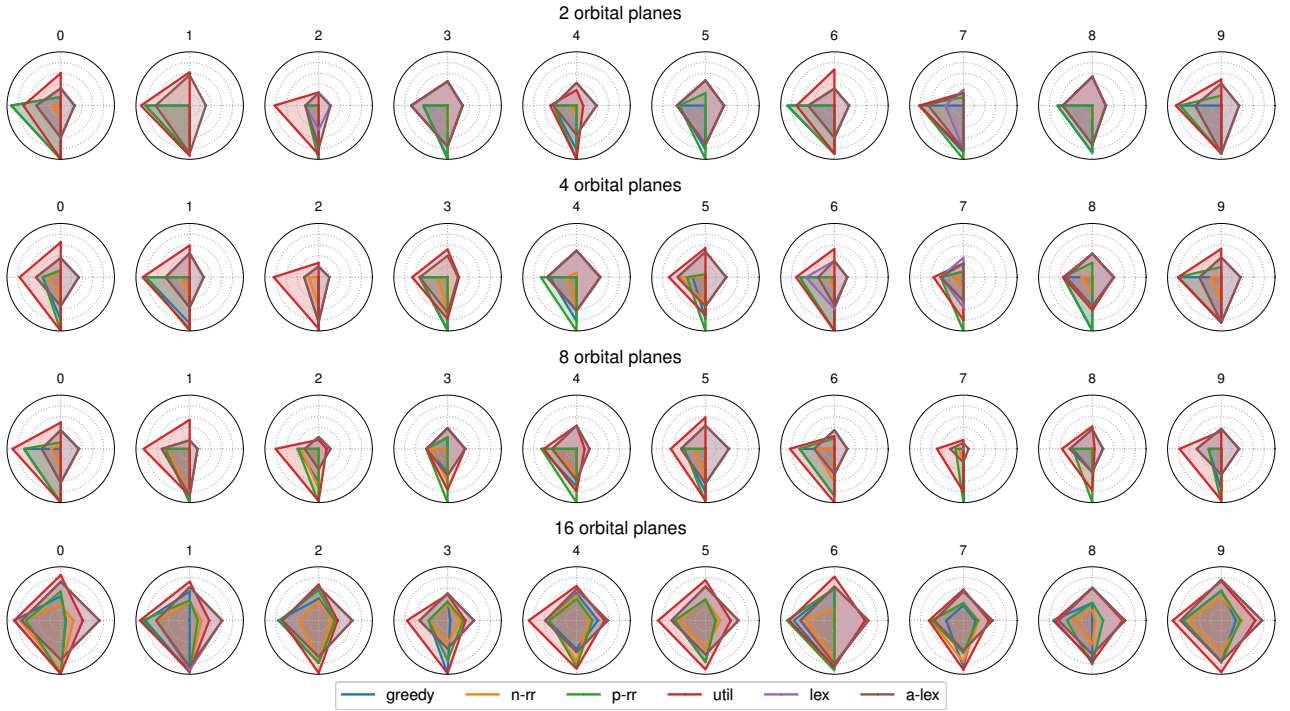


FIGURE 4 – Profils d'utilité (par ordre lexicmin) pour les 10 premières instances (sur 30) pour chaque taille de constellation et chaque algorithme (sud : meilleure utilité sur tous les agents ; ouest : deuxième meilleure utilité ; nord : troisième meilleure utilité ; est : pire utilité).

cations à presque 85% de la valeur optimale en moyenne. Il est intéressant de noter que lex a des performances équivalentes à a-lex (écart inférieur à 5% en moyenne). En effet, a-lex fournit des allocations légèrement meilleures du point de vue utilitaire, au détriment de l'approximation de l'équité. Les approches *round-robin* diffèrent vraiment en termes d'utilité. Bien que p-rr fournit des allocations à presque 71% de l'optimal, n-rr donne des allocations de faible utilité, à presque 10% de l'optimal sur des constellations plus petites. D'ailleurs, dans de telles configurations, il y a peu de chemins réalisables pour chaque demande. Ainsi, pour la plupart des demandes, la construction incrémentale myope des chemins aboutit à des impasses, et donc à des allocations d'une utilité de 0; alors qu'en considérant un autre ordre d'agents, les allocations auraient pu être meilleures. Enfin, greedy se comporte légèrement moins bien que p-rr, à près de 68% de la valeur optimale en moyenne. Sur des instances plus grandes, où de nombreux chemins existent pour répondre à chaque requête, greedy se comporte encore mieux. Plus généralement, les problèmes de constellation de grande taille sont plus faciles à résoudre du point de vue utilitaire par les algorithmes non optimaux, puisqu'il existe plus d'options pour éviter les conflits malgré leur nombre élevé.

4.4 Équité

Pour analyser l'équité des allocations résultantes, la Figure 4 fournit les profils d'utilité obtenus pour les 10 premières instances pour chaque algorithme. greedy, par principe, cherche à allouer d'abord les chemins à plus forte uti-

lité, ce qui entraîne des allocations injustes où seuls les 2 ou 3 premiers utilisateurs sont servis, alors que plus souvent le quatrième utilisateur n'a aucune requête satisfaite. Les approches *round-robin* sont plus équitables que l'approche greedy et satisfont souvent les requêtes d'un plus grand nombre d'utilisateurs. util donne des profils avec la plus grande surface, mais la plupart du temps le quatrième utilisateur est négligé. Enfin, lex et a-lex se comportent presque identiquement (leurs profils sont superposés), montrant que l'approximation a-lex est suffisante pour produire des allocations équitables. Notons qu'avec des constellations plus grandes, puisqu'il y a de plus en plus d'options pour servir les utilisateurs, tous les algorithmes ont tendance à donner des allocations plus équitables. Néanmoins, lex et a-lex sont les meilleurs choix ici, et les concurrents *round-robin* produisent des allocations insatisfaisantes.

4.5 Temps de calcul

La figure 3b montre le temps de calcul moyen en millisecondes pour chaque algorithme et chaque taille de constellation. Comme prévu (par conception), greedy, p-rr et n-rr sont les plus rapides. n-rr qui n'effectue même pas les opérations de chemin le plus court, est de loin le plus rapide, mais donne lieu à des allocations très mauvaises et peu équitables. greedy et p-rr sont très rapides, mais sont basés sur de multiples recherches du chemin maximum dans les DAGs. p-rr fournit encore rapidement d'assez bonnes allocations utilitaires et équitables. util, basé sur la résolution d'un seul programme linéaire en nombres entiers, est 100 fois plus lent que les algorithmes les plus rapides. Ensuite,

lex et a-lex sont jusqu'à deux ordres de grandeur plus lents que util sur les plus grandes constellations. Ceci est dû aux multiples appels ($|A|$) au solveur MILP sur les grands problèmes. a-lex est 2 à 3 fois plus rapide que lex puisqu'il résout des MILP plus petits, tout en donnant des allocations aussi justes que celles de lex.

5 Conclusion

Dans cet article, nous avons proposé le modèle PADAG, un nouveau problème d'allocation où les agents expriment leurs préférences sur des lots d'articles sous forme de DAGs pondérés par les arêtes. Nous avons introduit et analysé plusieurs méthodes de résolution (utilitaire, leximin, leximin approché, glouton) contre les allocations *round-robin*, du point de vue de l'utilitarisme et de l'équité. Nous avons évalué ces méthodes sur de grandes instances de problèmes d'allocation de créneaux orbitaux, générées aléatoirement, avec plus de 1000 couches. Sur de grandes constellations, nous observons que la méthode leximin approché constitue un bon compromis entre l'optimalité utilitaire, l'optimalité leximin et le temps de calcul, par rapport à toutes les autres méthodes de solution. Elle est même équivalente au leximin exact sur la plupart des instances, tout en divisant les temps de calcul par un facteur de 2 à 3.

Nous identifions plusieurs pistes pour de futures investigations. Tout d'abord, comme le leximin approximatif est un compromis entre l'équité et les temps de calcul, nous aimerions étudier d'autres méthodes de résolution, encore plus réactives, notamment pour résoudre des PADAGs plus grands et des topologies spécifiques. En effet, comme pour la division des chemins et les jeux de congestion, des techniques dédiées pourraient être conçues pour des topologies données (étoiles, chaînes, etc.). Deuxièmement, puisque les PADAGs sont fortement contraints par les conflits, nous cherchons à explorer des heuristiques *min-conflict* pour améliorer nos algorithmes. Enfin, nous pensons que les PADAGs ont un grand potentiel pour être utilisés dans une variété de domaines, et nous souhaitons donc évaluer les techniques proposées sur des problèmes provenant d'autres applications, comme le domaine NFV où les chaînes de fonctions sont modélisées comme des graphes, et les incompatibilités contrôlent l'accès aux nœuds ou le domaine de la recherche de chemins dans un cadre multi-agent (les préférences de chemin sont modélisées comme des graphes, et les incompatibilités modélisent les contraintes pour éviter que deux agents occupent la même position en même temps).

Références

- [1] Sylvain Bouveret, Katarína Cechlárová, Edith Elkind, Ayumi Igarashi, and Dominik Peters. Fair division of a graph. In Carles Sierra, editor, *Proceedings of the Twenty-Sixth International Joint Conference on Artificial Intelligence, IJCAI 2017, Melbourne, Australia, August 19-25, 2017*, pages 135–141. ijcai.org, 2017.
- [2] Coen Bron and Joep Kerbosch. Algorithm 457 : Finding all cliques of an undirected graph. *Commun. ACM*, 16(9) :575–577, September 1973.
- [3] Ioannis Caragiannis, David Kurokawa, Hervé Moulin, Ariel D. Procaccia, Nisarg Shah, and Junxing Wang. The unreasonable fairness of maximum nash welfare. In *Proceedings of the 2016 ACM Conference on Economics and Computation, EC '16*, page 305–322, New York, NY, USA, 2016. Association for Computing Machinery.
- [4] Thomas H. Cormen, Charles E. Leiserson, Ronald L. Rivest, and Clifford Stein. *Introduction to Algorithms*. The MIT Press, 2nd edition, 2001.
- [5] Peter Cramton, Yoav Shoham, and Richard Steinberg. *Combinatorial Auctions*. MIT Press, 2005.
- [6] Ye Du, Rahul Sami, and Yaoyun Shi. Path auctions with multiple edge ownership. *Theor. Comput. Sci.*, 411(1) :293–300, 2010.
- [7] Michael S. Hughes, Brian J. Lunday, Jeffrey D. Weir, and Kenneth M. Hopkinson. The multiple shortest path problem with path deconfliction. *Eur. J. Oper. Res.*, 292(3) :818–829, 2021.
- [8] Ayumi Igarashi and Dominik Peters. Pareto-optimal allocation of indivisible goods with connectivity constraints. In *The Thirty-Third AAAI Conference on Artificial Intelligence, AAAI 2019*, pages 2045–2052. AAAI Press, 2019.
- [9] Nicole Immorlica, David R. Karger, Evdokia Nikolova, and Rahul Sami. First-price path auctions. In John Riedl, Michael J. Kearns, and Michael K. Reiter, editors, *Proceedings 6th ACM Conference on Electronic Commerce (EC-2005), Vancouver, BC, Canada, June 5-8, 2005*, pages 203–212. ACM, 2005.
- [10] David Kurokawa, Ariel D. Procaccia, and Nisarg Shah. Leximin allocations in the real world. *ACM Transactions on Economics and Computation*, 6(3–4), 2018.
- [11] M. Lemaître, G. Verfaillie, H. Fargier, J. Lang, N. Bataille, and J.-M. Lachiver. Equitable allocation of earth observing satellites resources. In *5th ONERA-DLR Aerospace Symposium (ODAS'03)*, 2003.
- [12] Nimrod Megiddo. Optimal flows in networks with multiple sources and sinks. *Math. Program.*, 7(1) :97–107, 1974.
- [13] Hervé Moulin. *Fair division and collective welfare*. MIT Press, 2003.
- [14] Noam Nisan, Tim Roughgarden, Eva Tardos, and Vijay V. Vazirani. *Algorithmic Game Theory*. Cambridge University Press, USA, 2007.
- [15] OpenStreetMap. Openstreetmap points of interest (on french territory). <https://www.data.gouv.fr/fr/datasets/points-dinterets-openstreetmap/>, 2021. Accessed : 30-08-2021.

- [16] Panagiota N. Panagopoulou and Paul G. Spirakis. Algorithms for pure nash equilibria in weighted congestion games. *ACM J. Exp. Algorithmics*, 11 :2.7–es, feb 2007.
- [17] Gauthier Picard. Auction-based and Distributed Optimization Approaches for Scheduling Observations in Satellite Constellations with Exclusive Orbit Portions. In *International Workshop on Planning and Scheduling for Space (IWSPSS'21)*, 2021.
- [18] Gauthier Picard, Clément Caron, Jean-Loup Farges, Jonathan Guerra, Cédric Pralet, and Stéphanie Rousel. Autonomous Agents and Multiagent Systems Challenges in Earth Observation Satellite Constellations. In U. Endriss, A. Nowé, F. Dignum, and A. Lomuscio, editors, *Proceedings of the 20th International Conference on Autonomous Agents and MultiAgent Systems, AAMAS '21*, page 39–44, Richland, SC, 2021. International Foundation for Autonomous Agents and Multiagent Systems.
- [19] Jordi Ferrer Riera, Eduard Escalona, Josep Batallé, Eduard Grasa, and Joan A. García-Espín. Virtual network function scheduling : Concept and challenges. In *2014 International Conference on Smart Communications in Network Technologies (SaCoNeT)*, pages 1–5, 2014.
- [20] Song Yang, Fan Li, Stojan Trajanovski, Ramin Yahyapour, and Xiaoming Fu. Recent advances of resource allocation in network function virtualization. *IEEE Transactions on Parallel and Distributed Systems*, 32(2) :295–314, 2021.
- [21] Lei Zhang, Haibin Chen, Jun Wu, Chong-Jun Wang, and Junyuan Xie. False-name-proof mechanisms for path auctions in social networks. In Gal A. Kaminka, Maria Fox, Paolo Bouquet, Eyke Hüllermeier, Virginia Dignum, Frank Dignum, and Frank van Harmelen, editors, *ECAI 2016 - 22nd European Conference on Artificial Intelligence, 29 August-2 September 2016, The Hague, The Netherlands - Including Prestigious Applications of Artificial Intelligence (PAIS 2016)*, volume 285 of *Frontiers in Artificial Intelligence and Applications*, pages 1485–1492. IOS Press, 2016.

La planification SAS sous forme de tri topologique

G. Prévost¹, S. Cardon¹, T. Cazenave², C. Guettier³, É. Jacopin¹

¹ Centre de Recherche de l'Académie Militaire de Saint-Cyr Coëtquidan, CReC

² Université Paris-Dauphine - PSL, LAMSADE

³ SAFRAN Electronics and Defense

Résumé

Nous abordons la question de comment fournir en temps-réel un plan d'actions pour plusieurs millions de personnages non joueurs (PNJ) dans les mondes virtuels. Suite à une étude approfondie des plans générés par des jeux vidéo commerciaux utilisant la planification, nous présentons une nouvelle classe de problèmes au domaine de planification Structure d'Action Simplifiée (SAS) applicable à ces jeux vidéo. Nous présentons également un nouvel algorithme de complexité linéaire pour cette classe de problèmes qui, contrairement aux autres planificateurs SAS, nous permet effectivement de considérer la gestion de millions de PNJs par image.

Mots-clés

Intelligence Artificielle, planification, temps-réel, actions, SAS, tri topologique.

Abstract

We address the question of how to provide an action plan in real-time to several million non-player characters (NPCs) in virtual worlds. After a close study of the plans generated in commercial video games using planning, we present a new class of problems in the Simplified Action Structure (SAS) planning framework applicable to these video games. We also present a new linear-time algorithm for this class of problems which, unlike previous SAS planners, effectively allows us to consider the management of millions of NPCs per frame.

Keywords

Artificial Intelligence, planning, real-time, actions, SAS, topological sort.

1 Introduction

F.E.A.R., un jeu de tir à la première personne (FPS) sorti en 2005, a été le premier jeu vidéo à utiliser la planification pour générer les comportements des personnages en temps réel [21, 22]. Le succès de F.E.A.R. [20] a été tel qu'il a conduit à une large diffusion de l'utilisation de la planification dans les FPS [14, 17, 24, 27] ; cette diffusion a notamment été facilitée par la publication l'année suivante d'un kit de développement (SDK) contenant le code du planificateur [19]. Aujourd'hui, non seulement le succès de F.E.A.R. est toujours reconnu [16], mais les plus grandes produc-

tions, touchant des millions de joueurs, n'hésitent pas à utiliser la planification et à la faire connaître [13, 8, 12], et ce malgré les contraintes de temps réel de plus en plus exigeantes.

En 2005, un moteur de jeu affichait environ 30 images par seconde, soit 33 ms pour toute la logique du jeu : entre les graphismes, la physique et les mécaniques de gameplay, cela laissait moins d'une milliseconde au planificateur pour générer des plans pour les quelques personnages qui l'appelaient [26] ; dix ans plus tard, le passage à 60 images par seconde n'a fait que réduire le budget de traitement disponible pour la planification. Aujourd'hui, pour satisfaire le budget alloué au planificateur, les studios limitent explicitement le nombre d'appels à quelques dizaines tout au plus [6, 13, 12] malgré l'augmentation des performances matérielles. De plus, les situations de jeu sont conçues de manière à ce que le nombre de personnages soit également limité, réduisant ainsi le nombre d'appels au planificateur. Cependant, la dynamique du marché du jeu vidéo pousse les productions à simuler des univers de plus en plus vastes avec, pour l'instant, des dizaines de milliers de personnages en vue [15] et demain des millions, en particuliers pour les jeux massivement multijoueurs. Les résultats de [5] suggèrent que les GPU sont une solution potentielle pour de tels univers dans le cloud gaming ; mais qu'en est-il des PC ou des consoles de jeu pour lesquels les GPU sont dédiés au graphisme ?

La complexité temporelle des problèmes de planification dépend d'une part des restrictions du langage pour représenter le problème de planification [4] et d'autre part de la partie de l'entrée qui est fixée [11]. À partir du code dans le SDK de F.E.A.R., nous pouvons observer que les états sont des vecteurs de valeurs discrètes et qu'à certains moments, comme les combats ou les tâches routinières, les actions sont fixées et seuls les états initiaux et finaux font partie de l'entrée du problème de planification. De plus, la grande majorité des actions sont unaires [10], c'est-à-dire qu'elles ne changent la valeur que d'une seule variable d'état. Enfin, les actions sont post-unicques par type (attaque, défense, inspection, ...), c'est-à-dire qu'un type d'actions est le seul à modifier une variable d'état donnée ; ceci permet d'insérer dans le plan un opérateur de défense générique, par exemple, et de retarder le choix du type de défense au moment de l'exécution du plan. Les réponses à un questionnaire rempli par plusieurs développeurs d'IA de jeux corro-

borent que notre approche est toujours d'actualité : la modélisation des problèmes de planification dans les jeux commerciaux correspond au domaine SAS avec actions unaires et post-uniques. Ces jeux appartiennent ainsi à la classe de problèmes SAS-PU [3] qui est NP-Hard.

L'analyse des données de planification récoltées en jeu par [18] nous montre que les plans dans les jeux de tir à la première personne ont deux caractéristiques spécifiques au domaine : (1) ils sont totalement ordonnés, (2) ils n'ont qu'une seule occurrence d'un type d'action donné : par exemple, une seule action pour menacer, recharger, se mettre à couvert, esquiver, etc. Par conséquent, dans la section suivante, nous présentons deux nouvelles restrictions SAS correspondant à "Totalement ordonné" (T) et à "isomorphisme de type" (T_1) que nous combinons en T_1 ; nous présentons également un algorithme de complexité en temps linéaire pour la classe de problèmes SAS-PUT₁. Dans la section suivante, nous présentons et discutons des tests visant à illustrer les propriétés linéaires de notre algorithme qui est capable de générer suffisamment de plans pour des millions de personnages tout en respectant le budget de traitement que le moteur de jeu alloue au planificateur à chaque image.

2 SAS-PUT₁ en temps linéaire

2.1 Contexte

Nous utilisons les notations de [1] tout au long de ce document. La structure d'action simplifiée (SAS) et sa version étendue (SAS⁺) représentent des *états* avec un ensemble $\mathcal{M}$ de m variables dont chacune peut prendre au plus n valeurs discrètes ou être *indéfinie*¹ ; nous notons $\mathcal{D}_v$ l'ensemble des valeurs de la variable d'état v . Trois ensembles de variables d'état sont utilisés dans SAS pour représenter les conditions d'une *action* : (1) les *post-conditions* (post) définissent de nouvelles valeurs pour certaines variables d'état, (2) les *pré-conditions* (pre) définissent les valeurs des variables d'état avant leur modification par les post-conditions, et (3) les *prevail-conditions* (prv) sont des pré-conditions qui ne seront pas modifiées par les post-conditions. Une action est *applicable* dans un état si ses pré- et prevail-conditions sont toutes satisfaites dans cet état ; *appliquer* une action dans un état change les valeurs de toutes les variables d'état définies dans les pré-conditions par celles définies dans les post-conditions, les prevail-conditions quant à elles restent inchangées. SAS diffère de SAS⁺ par l'ajout de deux restrictions pertinentes pour notre domaine d'application : une action ne peut changer une variable d'état que d'une valeur définie à une autre valeur définie (S6), et aucune variable d'état *initial* ou *but* ne peut être indéfinie (S7).

Un *plan* est une séquence d'actions telle que tout état résultant de l'application d'une action de la séquence est cohérent avec l'action suivante de la séquence ; un plan résout l'*instance d'un problème de planification*, constitué des états initial et but (s_0, s_*) et d'un ensemble de types

d'action, si : (1) toute action du plan est une instance distincte d'un *type d'action* du problème de planification, (2) la première action est applicable dans s_0 , et (3) l'application de la dernière action de la séquence conduit à s_* . Nous introduisons une *première nouvelle restriction* qui exige que le nombre d'instances d'actions distinctes du même type apparaissant dans un plan soit au maximum $k \in \mathbb{N}$ (T_k).

Trois restrictions supplémentaires sont pertinentes pour notre étude : (*Post-unicité*) deux types d'action distincts ne peuvent pas changer la même variable d'état à la même valeur (P), (*Unaire*) chaque type d'action change exactement la valeur d'une variable d'état (U), et (*Single-valuedness*) si plusieurs types d'action ont la même prevail-condition de définie alors la valeur sur cette prevail-condition doit être la même pour tous ces types d'action (S). Comme conséquence à la fois du **théorème 4.4** [1, p. 76], qui stipule que toute solution de plan minimal d'une instance de problème SAS⁺-PUS contient au plus deux actions de chaque type, et de notre nouvelle restriction (T_k), toute instance de problème SAS⁺-PUS est également une instance de problème SAS⁺-PUST₂.

Nous imposons finalement comme *seconde nouvelle restriction* que les plans soient totalement ordonnés (T), et comme ils correspondent à notre domaine d'application, nous cherchons à résoudre les instances de problème SAS-PUT₁ que nous notons SAS-PUT₁. Nous observons que le problème qui consiste à allumer de façon répétée une série de lumières dans un tunnel [1, pp. 16-17] puis à éteindre ces lumières, appartient à la classe de problèmes SAS-PUST₁ qui est un sous-ensemble de la classe de problèmes SAS-PUT₁ ; à notre connaissance, ces classes de problèmes ne sont pas nouvelles mais il n'était pas nécessaire d'explicitement les restrictions (T_k) et (T) qui y conduisent. La restriction (T) également n'est pas totalement nouvelle [3] mais elle n'a jamais été explicitement énoncée. Le problème du tunnel n'est autre que le problème MultiPrv_2_Cycle de la section 3. Enfin, comme nos plans solutions sont T_1 , nous ne ferons plus la différence en une action et un type d'action par la suite.

2.2 Exemple d'un problème SAS-PUT₁

Avant de présenter un problème concret de la classe SAS-PUT₁, l'Éleveur de Chevaux (Tab. 1), nous présentons des notions et notations utiles pour la suite. Premièrement, du fait de la restriction (P) et (U), chaque action a est identifiable par la paire (v_i, p) , où p est la postcondition de a sur v_i : $v_i \in \mathcal{M}, p \in \mathcal{D}_{v_i}$ tel que $post(a)[v_i] = p$. Nous dénotons dorénavant l'*identifiant* de chaque action (PU) avec le format $a_{v_i}^p$. De plus, $post(a_{v_i}^p) \equiv post(a_{v_i}^p)[v_i]$ et $pre(a_{v_i}^p) \equiv pre(a_{v_i}^p)[v_i]$. Ces identifiants nous permettent de définir trois ensembles de *prédécesseurs* pour une action $a_{v_i}^p$: (1) $\mathcal{N}_{pre}(a_{v_i}^p) = \{a_{v_j}^q \mid v_j = v_i \wedge q = pre(a_{v_i}^p)\}$, l'ensemble des prédécesseurs de $a_{v_i}^p$ par *dépendance post-pré*. Du fait de (P), cet ensemble est un singleton. (2) $\mathcal{N}_{prv}(a_{v_i}^p) = \{a_{v_j}^q \mid v_j \neq v_i \wedge q = prv(a_{v_i}^p)[v_j]\}$, l'ensemble des prédécesseurs de $a_{v_i}^p$ par *dépendance post-prv*. (3) $\mathcal{N}_m(a_{v_i}^p) = \{a_{v_j}^q \mid v_i \neq v_j \wedge a_{v_j}^q \notin \mathcal{N}_{prv}(a_{v_i}^p) \wedge pre(a_{v_i}^p) = prv(a_{v_j}^q)[v_i]\}$, l'ensemble des actions dont la

1. Les valeurs indéfinies sont notées u pour *undefined*.

prevail-condition sur v_i est *menacée* par $a_{v_i}^p$. Les ensembles $\mathcal{N}_{pre}$ et $\mathcal{N}_{prv}$ sont intuitifs à comprendre, $\mathcal{N}_m$ en revanche découle de la deuxième condition² du *Modal Truth Criterion* (MTC) développé par David Chapman [7, p.340]. Soit $a_{v_i}^p, a_{v_j}^q \in \mathcal{A}$, l'égalité $pre(a_{v_i}^p) = prv(a_{v_j}^q)[v_i]$ signifie que l'action $a_{v_i}^p$ peut “clobber”³, i.e. *menace*, la prevail-condition de $a_{v_j}^q$ sur v_i car, par définition de la pré-condition, $pre(a_{v_i}^p)$ est la valeur de v_i qui sera changée, ou “clobbered”, par $a_{v_i}^p$. Par conséquent, pour respecter la deuxième condition du MTC et sachant que les plans solutions sont T_1 , l'action $a_{v_j}^q$ doit être ordonnée avant $a_{v_i}^p$ dans le plan solution. $a_{v_j}^q$ devient donc un prédécesseur de $a_{v_i}^p$, d'où l'ensemble $\mathcal{N}_m(a_{v_i}^p)$. On dénote $\mathcal{N}(a_{v_i}^p) = \mathcal{N}_{pre}(a_{v_i}^p) \cup \mathcal{N}_{prv}(a_{v_i}^p) \cup \mathcal{N}_m(a_{v_i}^p)$ l'ensemble des prédécesseurs (*Neighbors*) de $a_{v_i}^p$. C'est un ensemble partiellement ordonné d'actions.

La post-unicité implique que $\mathcal{N}_{pre}(a_{v_i}^p)$ est un singleton, i.e. $a_{v_i}^p$ n'a qu'un seul prédécesseur par dépendance post-pré. En revanche, ce prédécesseur de $a_{v_i}^p$ n'a pas nécessairement que $a_{v_i}^p$ comme *successeur* via la dépendance post-pré. En effet, la post-unicité n'implique pas la pré-unicité comme on peut le voir avec les actions *Stocker botte de foin* et *Remplir mangeoire* de l'Éleveur de Chevaux (Tab. 1) qui partagent la même pré-condition malgré la post-unicité du problème. Elles ont toutes les deux comme unique prédécesseur *Prendre botte de foin* mais il en résulte que *Prendre botte de foin* a deux successeurs via la dépendance post-pré. Cette implication a son importance pour la recherche des actions menaçantes car une prevail-condition menacée est définie par l'égalité $pre(a_{v_i}^p) = prv(a_{v_j}^q)[v_i]$ ($a_{v_i}^p, a_{v_j}^q \in \mathcal{A}$) mais l'action menaçante $a_{v_i}^p$ n'est pas identifiable par la pair $(v_i, pre(a_{v_i}^p))$ en temps constant. Pour pallier ce problème, on introduit deux accesseurs par action : *Next* et *NextInCycle*, qui pointent vers un et un seul successeur si elles sont définies⁴. L'accesseur *NextInCycle* est définissable lors d'un pré-processing. En effet, plusieurs actions dépendantes entre elles par dépendances post-pré peuvent boucler entre elles, auxquels cas on peut facilement montrer par post-unicité que ce *cycle* est unique⁵. Comme le cycle est unique, s'il existe, alors chaque action du cycle n'a qu'un seul prédécesseur et qu'un seul successeur dans ce cycle. S'il n'y a pas de cycle, $NextInCycle = \emptyset$. Pour tout $v_i \in \mathcal{M}$, on note $Cycle[v_i]$ l'ensemble des actions appartenant à ce cycle. L'accesseur *Next*, quant à lui, est définie dynamiquement et une seule fois par la procédure 1 *BuildChain* qui retourne une *chaîne d'actions*, i.e. une séquence d'actions totalement ordonnée par dépendance post-pré. Une action menaçante est donc identifiable en temps constant une fois la chaîne d'actions construite (avec *Next*) ou si elle appartient à un cycle (avec *NextInCycle*). Soit $a_{v_i}^x, a_{v_j}^q \in \mathcal{A}$ telles que $Next(a_{v_i}^x)$ pointe vers une action

$\mathcal{A}$	pre	post	prv	Description
$a_{v_0}^0$	$v_0 = 1$	$v_0 = 0$	$\langle u, u, u \rangle$	Stocker botte de foin
$a_{v_0}^1$	$v_0 = 0$	$v_0 = 1$	$\langle u, 0, u \rangle$	Prendre botte de foin
$a_{v_0}^2$	$v_0 = 1$	$v_0 = 2$	$\langle u, u, u \rangle$	Remplir mangeoire
$a_{v_1}^0$	$v_1 = 1$	$v_1 = 0$	$\langle u, u, u \rangle$	Poser seau
$a_{v_1}^1$	$v_1 = 0$	$v_1 = 1$	$\langle 0, u, u \rangle$	Prendre seau
$a_{v_2}^1$	$v_2 = 0$	$v_2 = 1$	$\langle u, 1, u \rangle$	Remplir seau d'eau
$a_{v_2}^2$	$v_2 = 1$	$v_2 = 2$	$\langle u, 1, u \rangle$	Remplir abreuvoir

TABLE 1 – Actions de l'Éleveur de Chevaux dont le but est de nourrir les chevaux : $s_* = \langle 2, 0, 2 \rangle$ où v_0 représente la *Botte de foin* avec $\mathcal{D}_{v_0} = \{0 : stock, 1 : enMain, 2 : dansMangeoire\}$, v_1 représente un *Seau* avec $\mathcal{D}_{v_1} = \{0 : auSol, 1 : enMain\}$, et v_2 représente de l'*Eau* avec $\mathcal{D}_{v_2} = \{0 : dansFontaine, 1 : dansSeau, 2 : dansAbreuvoir\}$.

qui menace $a_{v_j}^q$, on a $pre(Next(a_{v_i}^x)) = prv(a_{v_j}^q)[v_i]$, et donc $post(a_{v_i}^x) = prv(a_{v_j}^q)[v_i]$. Ce qui implique que si $a_{v_j}^q \in \mathcal{N}_m(Next(a_{v_i}^x))$, alors $a_{v_i}^x \in \mathcal{N}_{prv}(a_{v_j}^q)$.

Tab. 1 présente le problème de l'Éleveur de Chevaux, un problème de classe SAS-PUT₁. Toutes les actions sont post-unicues et unaires car chaque action ne modifie qu'une seule variable d'état et il n'y en a pas deux modifiant une même variable d'état à la même valeur. Chaque action est donc identifiable avec le format $a_{v_i}^p$ et ces identifiants se trouvent dans la colonne $\mathcal{A}$. Ce problème n'est pas (S) car les actions *Remplir Seau d'Eau* ($a_{v_2}^1$) et *Remplir abreuvoir* ($a_{v_2}^2$) ont une prevail-condition différente que *Prendre botte de foin* ($a_{v_0}^1$) sur la variable d'état *Seau* (v_1) : $prv(a_{v_2}^2)[v_1] = prv(a_{v_1}^1)[v_1] \neq prv(a_{v_0}^1)[v_1]$. Le problème aurait été (S) si : $prv(a_{v_0}^1)[v_1] = u$ ou $prv(a_{v_0}^1)[v_1] = 1$, par exemple. Ce problème est T_1 car, pour chaque instance (s_0, s_*) du problème, le plan solution minimal, i.e. le plan solution possédant le moins d'actions, s'il existe, entre s_0 et s_* , ne contient pas deux fois la même action. Considérons maintenant l'action *Remplir abreuvoir* identifiée par $a_{v_2}^2$ pour illustrer $\mathcal{N}(a_{v_2}^2)$. On a $\mathcal{N}_{pre}(a_{v_2}^2) = \{a_{v_2}^1\}$ et $\mathcal{N}_{prv}(a_{v_2}^2) = \{a_{v_1}^1\}$. En fonction de l'instance (s_0, s_*) , *Remplir abreuvoir* peut avoir sa prevail-condition sur v_1 menacée par l'action *Poser seau*, auquel cas $a_{v_2}^2 \in \mathcal{N}_m(a_{v_1}^0)$ afin que *Remplir abreuvoir* soit ordonnée avant *Poser seau* dans le plan solution minimal. Enfin, considérons les actions $a_{v_0}^0$, $a_{v_0}^1$ et $a_{v_0}^2$ pour illustrer le comportement des dépendances post-pré et comment sont définis *Next* et *NextInCycle*. On constate que *Stocker botte de foin* ($a_{v_0}^0$) et *Prendre botte de foin* ($a_{v_0}^1$) bouclent entre elles. On a $Cycle[v_0] = \{a_{v_0}^0, a_{v_0}^1\}$, $a_{v_0}^1 \in \mathcal{N}_{pre}(a_{v_0}^0) \wedge NextInCycle(a_{v_0}^0) = a_{v_0}^1$, et $a_{v_0}^0 \in \mathcal{N}_{pre}(a_{v_0}^1) \wedge NextInCycle(a_{v_0}^1) = a_{v_0}^0$. En revanche, *Remplir mangeoire* ($a_{v_0}^2$) n'est pas dans le cycle, on a $a_{v_0}^1 \in \mathcal{N}_{pre}(a_{v_0}^2) \wedge NextInCycle(a_{v_0}^2) = \emptyset$. $a_{v_0}^1$ a deux successeurs possibles en fonction de l'instance du problème, on peut avoir $Next(a_{v_0}^1) = a_{v_0}^2$ ou $Next(a_{v_0}^1) = NextInCycle(a_{v_0}^1) = a_{v_0}^0$.

2. La première condition du MTC est satisfaite par $\mathcal{N}_{pre}$ et $\mathcal{N}_{prv}$.

3. *Clobber* est un terme utilisé et défini par D. Chapman.

4. *Next* et *NextInCycle* peuvent pointer vers le même successeur.

5. Les actions (PU) (noeuds) et leurs dépendances post-pré (arcs) forment un graphe dirigé faiblement connecté. Par contraposition, s'il y a deux cycles dans le graphe alors il y a deux actions avec la même post-condition.

Procédure 1 BuildChain($v_i, s, g, \mathcal{D}, E_{\mathcal{D}}, \mathcal{A}$)

Input : $v_i \in \mathcal{M}$, $s, g \in \mathcal{D}_{v_i}$ deux variables du domaine de v_i ; $\mathcal{D}$, l'ensemble des actions jaunes; $E_{\mathcal{D}}$, l'ensemble des ordres entre les actions de $\mathcal{D}$; $\mathcal{A}$, l'ensemble des actions du problème.

Parameters : x, y , deux valeurs de $\mathcal{D}_{v_i}$.

Output : Chaque action a de la chaîne est jaune, ajouté à $\mathcal{D}$ et $(\mathcal{N}_{pre}(a), a)$ est ajouté à $E_{\mathcal{D}}$.

```

1:  $x \leftarrow \emptyset$ ;  $y \leftarrow \emptyset$ 
2: if  $a_{v_i}^g \notin \mathcal{A}$  then fail
3: end if
4:  $\text{Color}(a_{v_i}^g) \leftarrow \text{yellow}$ ;  $\mathcal{D} \leftarrow \mathcal{D} \cup \{a_{v_i}^g\}$ ;
5:  $y \leftarrow \text{pre}(a_{v_i}^g)$ ;
6:  $E_{\mathcal{D}} \leftarrow E_{\mathcal{D}} \cup \{(a_{v_i}^y, a_{v_i}^g)\}$ 
7: if  $\text{Next}(a_{v_i}^y) = \emptyset$  then
8:    $\text{Next}(a_{v_i}^y) \leftarrow a_{v_i}^g$ 
9: end if
10: while  $y \neq s$  do
11:   if  $a_{v_i}^y \notin \mathcal{A}$  then fail
12:   end if
13:   if  $\text{Color}(a_{v_i}^y) = \text{yellow}$  then fail
14:   end if
15:    $\text{Color}(a_{v_i}^y) \leftarrow \text{yellow}$ ;  $\mathcal{D} \leftarrow \mathcal{D} \cup \{a_{v_i}^y\}$ 
16:    $x \leftarrow y$ ;  $y \leftarrow \text{pre}(a_{v_i}^y)$ 
17:    $E_{\mathcal{D}} \leftarrow E_{\mathcal{D}} \cup \{(a_{v_i}^x, a_{v_i}^y)\}$ 
18:   if  $\text{Next}(a_{v_i}^x) = \emptyset$  then
19:      $\text{Next}(a_{v_i}^x) \leftarrow a_{v_i}^y$ 
20:   end if
21: end while
```

2.3 Planification réduite à un tri topologique

SAS⁺-PUS [2] est une classe abordable de problèmes de planification SAS⁺ pour lesquels le meilleur algorithme [1, pp. 84-90], que nous notons désormais $\mathcal{P}$, s'exécute en $O(m^2n)$ (cf. **Théorème 4.11** [1, p. 89]). Dans une première phase, $\mathcal{P}$ itère sur chacune des m variables d'état du but pour construire des chaînes d'au plus n actions vers l'état initial grâce à la restriction (P). Dans une deuxième phase, $\mathcal{P}$ ordonne les couples de $O(mn)$ actions apparaissant dans des chaînes distinctes grâce à la restriction (S) pour produire un ensemble partiellement ordonné d'actions; pour réaliser cette deuxième phase, $\mathcal{P}$ itère d'abord sur $O(mn)$ actions et ensuite sur $O(m)$ variables. Dans une troisième et dernière phase, si l'ensemble partiellement ordonné d'actions ne contient aucun cycle, il est retourné comme solution.

Notre algorithme, que nous notons désormais $\mathbb{P}$, suit les trois phases de $\mathcal{P}$ tout en simplifiant la complexité temporelle de la deuxième phase. Le point clé de l'algorithme $\mathbb{P}$ pour obtenir une complexité temporelle linéaire est d'utiliser les dépendances post-prv entre actions pour vérifier les prevail-conditions plutôt que de parcourir l'ensemble des variables d'état. Les prevail-conditions sont des états *partiellement* définis et, couplées aux restrictions (P) et (U), elles représentent également un ensemble partiellement ordonné d'actions ($\mathcal{N}_{prv}$, cf. sous-section 2.2), que $\mathbb{P}$ utilise

Procédure 2 DFSTopo($a_{v_i}^p, \mathcal{D}, E_{\mathcal{D}}, s_0, \mathbb{P}$)

Input : $a_{v_i}^p$, une action jaune identifiée; s_0 , l'état initial; $\mathbb{P}$, le plan solution.

Output : $a_{v_i}^p$ est coloré en vert une fois que tous ses prédécesseurs ont été topologiquement triés; Elle est insérée en queue de la liste $\mathbb{P}$.

```

1:  $\text{Color}(a_{v_i}^p) \leftarrow \text{blue}$ 
2: for  $a_{v_j}^q \in \mathcal{N}(a_{v_i}^p)$  do
3:   if  $a_{v_j}^q \notin \mathcal{N}_{pre}(a_{v_i}^p) \vee a_{v_i}^p$  n'est pas la 1re action à modifier  $s_0[v_i]$  then
4:     if  $\text{Color}(a_{v_j}^q) = \text{blue}$  then fail {Cycle détecté.}
5:   end if
6:   if  $\text{Color}(a_{v_j}^q) = \text{yellow}$  then
7:     DFSTopo( $a_{v_j}^q, \mathcal{D}, E_{\mathcal{D}}, s_0, \mathbb{P}$ )
8:   end if
9: end if
10: end for
11:  $\text{Color}(a_{v_i}^p) \leftarrow \text{green}$ ;  $\mathbb{P} \leftarrow \mathbb{P} + \{a_{v_i}^p\}$ 
```

lors des phases 2 et 3.

Soit $\mathcal{A}$ l'ensemble des actions d'un problème de planification et (s_0, s_*) une instance de ce problème, $\mathbb{P}$ planifie en arrière en utilisant l'identification des actions $a_{v_i}^p \in \mathcal{A}$ et leurs ensembles de voisinage $\mathcal{N}_{pre}(a_{v_i}^p)$ et $\mathcal{N}_{prv}(a_{v_i}^p)$ afin de trouver un plan solution minimal et totalement ordonné entre s_0 et s_* . Un pré-processing est ainsi nécessaire pour (1) construire une table de *Hashing* des actions $a_{v_i}^p$ de $\mathcal{A}$ en fonction de (v_i, p) et pour (2) construire leurs ensembles $\mathcal{N}_{pre}(a_{v_i}^p)$ et $\mathcal{N}_{prv}(a_{v_i}^p)$. Pour toute instance d'un problème de planification, (1) permet à $\mathbb{P}$ d'accéder à n'importe quelle action en un temps constant ($O(1)$) tandis que (2) permet de réduire la complexité temporelle de $\mathbb{P}$ par rapport à $\mathcal{P}$ lors de la phase 2. Toutes les actions du problème ne sont pas utiles à la résolution d'une instance, on note ainsi $\mathcal{D}$ l'ensemble des actions nécessaires à la résolution. Il en découle que l'ensemble des prédécesseurs d'une action a ne sont pas nécessairement tous utiles. En effet, l'état initial peut satisfaire la pré-condition de a ainsi que certaines de ses prevail-conditions, auxquels cas les éventuels prédécesseurs concernés ne sont pas utiles. Si l'Éleveur de Chevaux ne porte initialement pas de *botte de foin* ($s_0[v_0] = 0$) et que son objectif est d'en porter une ($s_*[v_0] = 1$), alors l'action *Prendre botte de foin* est nécessaire et réalisable dans s_0 , et l'action qui la précède, *Stocker botte de foin*, n'est pas utile. Ainsi, la résolution d'une instance (s_0, s_*) par $\mathbb{P}$ consiste à trouver l'ensemble $\mathcal{D}$ et à construire un ensemble suffisant des ordres entre les actions de $\mathcal{D}$, on dénote $E_{\mathcal{D}}$ cet ensemble d'ordres. Cette recherche et construction des ordres s'effectue lors des phases 1 et 2. La troisième et dernière phase est un tri topologique du plan partiellement ordonné $\langle \mathcal{D}, E_{\mathcal{D}} \rangle$ afin de retourner un plan totalement ordonné (restriction (T)). Enfin, on introduit 4 *couleurs* pour les actions : (blanc) la couleur initiale lorsque $\mathbb{P}$ est appelé, (jaune) l'action est utile pour résoudre le problème, (bleu) l'action est en cours de tri, (vert) l'action est triée topologiquement et insérée dans le plan solu-

Procédure 3 $\mathbb{P}(\mathcal{M}, \mathcal{A}, s_0, s_*)$

Input : $\mathcal{M}; \mathcal{A}; s_0, s_*$: états initial et final totalement définis.

Parameters : $\mathcal{D}$, l'ensemble des actions jaunes; $E_{\mathcal{D}}$, l'ensemble des ordres entre les actions jaunes;

Output : P : un plan d'action totalement ordonné qui relie s_0 à s_* ; produit un échec si l'instance n'est pas solvable.

```

1:  $P \leftarrow \emptyset; \mathcal{D} \leftarrow \emptyset; E_{\mathcal{D}} \leftarrow \emptyset$ 
2: for  $v_i \in \mathcal{M}$  do {Phase 1}
3:   if  $s_0[v_i] \neq s_*[v_i]$  then
4:     BuildChain( $v_i, s_0[v_i], s_*[v_i], \mathcal{D}, E_{\mathcal{D}}, \mathcal{A}$ )
5:   end if
6: end for
7: if  $\mathcal{D} = \emptyset$  then return  $\emptyset$  { $s_0$  et  $s_*$  sont égaux.}
8: end if
9: for  $a_{v_i}^p \in \mathcal{D}$  do {Phase 2}
10:  for  $a_{v_j}^q \in \mathcal{N}_{prv}(a_{v_i}^p)$  do
11:    if  $q \neq s_0[v_j]$  then
12:      if  $a_{v_j}^q \notin \mathcal{A}$  then fail
13:    end if
14:    if Color( $a_{v_j}^q$ ) = white then
15:      BuildChain( $v_j, s_0[v_j], q, \mathcal{D}, E_{\mathcal{D}}, \mathcal{A}$ )
16:    end if
17:     $E_{\mathcal{D}} \leftarrow E_{\mathcal{D}} \cup \{(a_{v_j}^q, a_{v_i}^p)\}$ 
18:  end if
19:  if  $q \neq s_*[v_j]$  then
20:    if Next( $a_{v_j}^q$ ) =  $\emptyset$  then
21:      BuildChain( $v_j, q, s_0[v_j], \mathcal{D}, E_{\mathcal{D}}, \mathcal{A}$ )
22:    end if
23:    if  $prv(\text{Next}(a_{v_j}^q))[v_i] \neq p$  then
24:       $E_{\mathcal{D}} \leftarrow E_{\mathcal{D}} \cup \{(a_{v_i}^p, \text{Next}(a_{v_j}^q))\}$ 
25:    end if
26:  end if
27:  if  $q = s_0[v_j] \wedge q \in \text{Cycle}[v_j]$  then
28:    if  $a_{v_i}^p$  est s'ordonne avant modification de  $s_0[v_j]$  then {cf. section 2.3}
29:       $E_{\mathcal{D}} \leftarrow E_{\mathcal{D}} \cup \{(a_{v_i}^p, \text{NextInCycle}(a_{v_j}^{s_0[v_j]}))\}$ 
30:    else { $a_{v_i}^p$  s'ordonne après.}
31:       $E_{\mathcal{D}} \leftarrow E_{\mathcal{D}} \cup \{(a_{v_j}^{s_0[v_j]}, a_{v_i}^p)\}$ 
32:    end if
33:  end if
34: end for
35: end for
36: for  $a \in \mathcal{D}$  do {Phase 3}
37:  if Color( $a$ ) = yellow then
38:    DFSTopo( $a, \mathcal{D}, E_{\mathcal{D}}, s_0, P$ )
39:  end if
40: end for
41: return  $P$ 

```

tion. Les actions de $\mathcal{D}$ sont toutes colorées en jaune. Blanc et jaune sont utilisées dans les 3 phases alors bleu et vert ne sont utilisées que dans la phase 3.

(Phase 1, 3.2 à 3.6) $\mathbb{P}$ va d'abord chercher les différences entre s_0 et s_* pour y construire des chaînes d'actions. Ces chaînes sont construites par la procédure 1 *BuildChain* via

les dépendances post-pré des actions (1.5 et 1.13). Dynamiquement, cette procédure colore en jaune, ajoute dans $\mathcal{D}$ (1.4 et 1.13) et définit la variable *Next* (1.7 et 1.16) des actions parcourues. Également, pour chaque action parcourue, l'ordre avec son prédécesseur est ajouté dans $E_{\mathcal{D}}$ (1.6 et 1.15).

(Phase 2, 3.9 à 3.35), $\mathbb{P}$ vérifie les prevail-conditions de toutes les actions de $\mathcal{D}$ (3.9) en parcourant les prédécesseurs de l'ensemble $\mathcal{N}_{prv}$ (3.10). Soient $a_{v_i}^p, a_{v_j}^q \in \mathcal{A}$ telles que $a_{v_i}^p$ est l'action en cours de vérification et $a_{v_j}^q$ un prédécesseur par dépendance post-prv. La prevail-condition sur v_j de $a_{v_i}^p$, $prv(a_{v_i}^p)[v_j]$, est satisfaite soit par $s_0[v_j]$ soit par $a_{v_j}^q \in \mathcal{N}_{prv}(a_{v_i}^p)$ (restriction (P)), sinon l'instance de problème n'est pas solvable. Si le prédécesseur $a_{v_j}^q$ est jaune, cela signifie que la chaîne d'actions v_j entre $s_0[v_j]$ et $s_*[v_j]$ est déjà construite, alors $\mathbb{P}$ se contente d'ordonner l'action en cours de vérification $a_{v_i}^p$ avec son prédécesseur $a_{v_j}^q$ (3.17). Si $a_{v_j}^q$ n'est pas jaune, elle est blanche (3.14) et la chaîne d'actions v_j entre $s_0[v_j]$ et $prv(a_{v_i}^p)[v_j]$ est manquante. Elle est donc construite (3.15), puis l'ordre entre $a_{v_i}^p$ et $a_{v_j}^q$ est construit et ajouté dans $E_{\mathcal{D}}$ (3.17). De 3.19 à 3.26, $\mathbb{P}$ traite les actions menaçantes, ou "clobberer". Comme expliqué dans la sous-section 2.2, si $a_{v_j}^q$ satisfait la prevail-condition de $a_{v_i}^p$ sur v_j et si $q \neq s_*[v_j]$, alors le successeur de $a_{v_j}^q$ ($\text{Next}(a_{v_j}^q)$), s'il existe, menace la prevail-condition sur v_j de $a_{v_i}^p$. Pour cela $\text{Next}(a_{v_j}^q)$ doit être définie (3.20). $\text{Next}(a_{v_j}^q) = \emptyset$ équivaut à dire que le successeur de $a_{v_j}^q$, s'il existe, est blanc. La chaîne d'actions v_j entre $prv(a_{v_i}^p)[v_j]$ et $s_0[v_j]$ est donc manquante et doit être construite (3.21). (3.23) $\text{Next}(a_{v_j}^q)$ est nécessairement défini, et si $\text{Next}(a_{v_j}^q) \notin \mathcal{N}_{prv}(a_{v_i}^p)$, i.e. $\text{Next}(a_{v_j}^q)$ n'est pas un prédécesseur par dépendance post-prv de $a_{v_i}^p$, alors $a_{v_i}^p \in \mathcal{N}_m(\text{Next}(a_{v_j}^q))$, i.e. $a_{v_i}^p$ est un prédécesseur de $\text{Next}(a_{v_j}^q)$ pour respecter le MTC (cf. sous-section 2.2), d'où la construction de l'ordre (3.24). De 3.27 à 3.33, $\mathbb{P}$ traite des cas qui ne pouvaient exister avec la restriction (S). En assouplissant (S), il est dorénavant possible d'avoir des prevail-conditions pour une même variable d'état v_j ayant des valeurs différentes (cf. Tab. 1, $prv(a_{v_0}^1)[v_1] \neq prv(a_{v_2}^1)[v_1]$). En particuliers, il est possible qu'une action ait sa prevail-condition en v_j satisfaite par $s_0[v_j] \in \mathcal{D}_{v_j}$ et qu'une autre action nécessaire à la résolution du problème ait sa prevail-condition sur v_j satisfaite par une autre valeur $x \in \mathcal{D}_{v_j}$. Cette situation est problématique si $a_{v_j}^{s_0[v_j]}$ et $a_{v_j}^x$ appartiennent à $\text{Cycle}[v_j]$. Dans ce cas la valeur $s_0[v_j]$ peut apparaître deux fois, malgré la restriction T_1 , une fois avec l'état initial $s_0[v_j]$ et une fois après l'exécution de l'action $a_{v_j}^{s_0[v_j]}$. Il faut donc déterminer si l'action ayant $s_0[v_j]$ comme prevail-condition s'ordonne avant la modification de v_j (3.29) ou si elle s'ordonne après $a_{v_j}^{s_0[v_j]}$ qui rétablit $s_0[v_j]$ (3.31). On peut déterminer (3.29) ou (3.31) lors du pré-traitement en cherchant si deux actions qui ont une prevail-condition définie différemment sur v_j sont reliées par un chemin dirigé qui ne passe pas par leur dépendance post-prv en v_j . Soit $a_{v_i}^p, a_{v_j}^q \in \mathcal{A}$, si $prv(a_{v_i}^p)[v_k] = s_0[v_k]$, $prv(a_{v_j}^q)[v_k] = x \neq s_0[v_k]$ et il existe un chemin dirigé entre $a_{v_i}^p$ et $a_{v_j}^q$

ne passant pas par les actions affectant v_k où $a_{v_i}^p$ est avant $a_{v_j}^q$, alors $a_{v_i}^p$ doit être exécutée avant la modification de l'état initial en v_j (3.29), sinon, si $a_{v_i}^p$ est après $a_{v_j}^q$ ou $a_{v_i}^p$ et $a_{v_j}^q$ ne sont pas reliées, on est dans le cas (3.31). Prenons l'instance $(s_0 = \langle 0, 0, 0 \rangle, s_* = \langle 2, 0, 2 \rangle)$ pour l'Éleveur de Chevaux, on a $prv(a_{v_0}^1)[v_1] = s_0[v_1]$ et $prv(a_{v_2}^1)[v_1] = prv(a_{v_2}^2)[v_1] = 1 \neq s_0[v_1]$. Les actions $a_{v_1}^0$ et $a_{v_1}^1$ seront utiles pour résoudre l'instance, elles seront colorées en jaune lors de la phase 2 car lors de la phase 1, $s_0[v_1] = s_*[v_1]$ et donc la condition (3.3) n'est pas respectée. Sans la partie 3.27 à 3.33 de l'algorithme, l'action $a_{v_0}^1$ n'a pas d'ordre spécifique relatif à sa prevail-condition en v_1 car elle est satisfaite par $s_0[v_1]$. Or, elle est menacée par $a_{v_1}^1$, mais également rétablie par $a_{v_1}^0$. Sans ordre spécifique, il existe des tris topologiques où $a_{v_0}^1$ est ordonnée après $a_{v_1}^1$ et avant $a_{v_1}^0$, ce qui est incompatible avec sa prevail-condition sur v_1 . Il n'existe pas de chemin dirigé entre $a_{v_0}^1$ et $a_{v_2}^1$, ni entre $a_{v_0}^1$ et $a_{v_2}^2$, sans passer par les actions affectant v_1 , $a_{v_0}^1$ s'ordonne donc après $a_{v_1}^0$ (3.30). (Phase 3, 3.36 à 3.40) Enfin, $\mathbb{P}$ trie topologiquement avec la Procédure 2 *DFSTopo* toutes les actions de $\mathcal{D}$ en fonction des ordres de l'ensemble $E_{\mathcal{D}}$. Concernant la condition d'arrêt (2.3) " $a_{v_i}^p \in \mathcal{A}$ est la première action à modifier $s_0[v_i]$ ", la condition $pre(a_{v_i}^p) = s_0[v_i]$ n'est pas suffisante. Prenons une instance de l'Éleveur de Chevaux où $s_0[v_0] = 1$ et $s_*[v_0] = 2$, il y a deux chaînes d'actions possible entre $s_0[v_0]$ et $s_*[v_0]$: $\langle a_{v_0}^2 \rangle$ et $\langle a_{v_0}^0, a_{v_0}^1, a_{v_0}^2 \rangle$. Si l'instance complète du problème est $(s_0 = \langle 1, 0, 0 \rangle, s_* = \langle 2, 0, 0 \rangle)$, alors $\langle a_{v_0}^2 \rangle$ est le plan solution minimal et $a_{v_0}^2$ est la première action à modifier $s_0[v_0]$. En revanche, si l'instance complète du problème est $(s_0 = \langle 1, 0, 0 \rangle, s_* = \langle 2, 0, 2 \rangle)$, $\langle a_{v_0}^0, a_{v_0}^1, a_{v_0}^2 \rangle$ fera partie du plan solution minimal et la première action à modifier $s_0[v_0]$ sera $a_{v_0}^0$.

Théorème 1. $\mathbb{P}$ est correct et complet.

Démonstration. (Esquisse) $\mathbb{P}$ est correct car il satisfait les deux conditions du MTC lors de la construction de $\mathcal{D}$ et de $E_{\mathcal{D}}$ (Phases 1 et 2) (cf. sous-section 2.2 et la description ci-dessus). La première condition du MTC est satisfaite par (1.6, 1.17, 3.4, 3.15, 3.17, 3.21), la deuxième condition du MTC est satisfaite par (3.24, 3.29, 3.31). Il en découle un tri topologique correct lors de la phase 3, i.e. un plan solution minimal et totalement ordonné dont l'application successive des actions est possible et aboutit finalement à l'état but de l'instance du problème.

(Esquisse) $\mathbb{P}$ est complet car peu importe l'instance du problème il retourne une solution. Pour savoir si une instance est solvable, pour tout $v_i \in \mathcal{M}$ il doit exister une chaîne d'actions entre $s_0[v_i]$ et $s_*[v_i]$. Cette existence se prouve à l'aide de la post-unicité. Si au moins une n'existe pas, $\mathbb{P}$ retourne un échec (1.2, 1.11, 1.13, 3.12). Si les chaînes d'actions existent mais que l'instance n'est pas solvable avec un plan T_1 , $\mathbb{P}$ retourne un échec (1.13, 2.4). Enfin si les chaînes d'actions existent mais que l'instance n'est pas solvable car il y a des cycles via les dépendances post-prv ou via les relations entre actions menacées et menaçantes, alors $\mathbb{P}$ retourne un échec (2.4). Dans tous les autres cas, $\mathbb{P}$ termine et retourne une solution. Pour la terminaison de

$\mathbb{P}$, aucune boucle n'est infinie : (1.10) a $O(n)$ itérations, (2.2) a $O(m)$ itérations, (3.2) a exactement m itérations, (3.9) a $O(\mathcal{A})$ itérations, (3.10) a $O(m)$ itérations et (3.36) a $O(\mathcal{A})$ itérations. Pour finir, la récursivité de *DFSTopo* se termine soit sur un échec, soit sur la condition d'arrêt (2.3). $\square$

Théorème 2. $\mathbb{P}$ a une complexité temporelle de $O(|\mathcal{A}| + |E_{\mathcal{A}}|)$ dans le pire des cas.

Démonstration. Nous avons $E_{\mathcal{A}} = \{(b, a) | a \in \mathcal{A} \wedge b \in \mathcal{N}_{pre}(a) \cup \mathcal{N}_{prv}(a) \cup \mathcal{N}_m(a)\}$ avec l'ensemble $\mathcal{N}_{prv}$ qui permet à $\mathbb{P}$ de ne naviguer qu'à travers les prevail-conditions définies des actions de $\mathcal{D}$. La phase 1 a au plus $|\mathcal{A}| = O(mn)$ étapes : m étapes via la boucle for (3.2) fois n étapes via la procédure *BuildChain* (3.4). En raison de la restriction (T_1) , $\mathcal{D}$ ne peut pas être supérieur à $\mathcal{A}$, donc la boucle for (3.9) nécessite au plus $|\mathcal{A}| = O(mn)$ étapes. Les deux boucles for (3.9) et (3.10) nécessitent $O(|\mathcal{A}| + |E_{\mathcal{A},prv}|)$ étapes avec $E_{\mathcal{A},prv} = \{(b, a) | a \in \mathcal{A} \wedge b \in \mathcal{N}_{prv}(a)\}$. En effet, lors de cette étape, les deux boucles for servent à vérifier les prevail-conditions de l'ensemble des actions de $\mathcal{D}$: il y a donc $|\mathcal{D}|$ étapes pour parcourir toutes les actions de $\mathcal{D}$, plus $|E_{\mathcal{D},prv}|$ étapes pour visiter tous les prédécesseurs de toutes les actions de $\mathcal{D}$. Les procédures *BuildChain* (3.15) et (3.21) peuvent être déclenchées ssi l'action en cours d'évaluation possède une prevail-condition qui n'est satisfaite ni par l'état initial, ni par une action de $\mathcal{D}$. Dans ce cas, les deux procédures cherchent des actions qui sont dans un ensemble $\mathcal{T} \subseteq \mathcal{A} \setminus \mathcal{D}$. Les procédures *BuildChain* rajoutent donc $|\mathcal{T}|$ étapes pour l'ensemble de leurs appels. Ces actions de $\mathcal{T}$ sont ajoutées à l'ensemble $\mathcal{D}$ par la procédure *BuildChain* car il faut dorénavant vérifier leurs prevail-conditions, il y a donc de nouveau $|\mathcal{T}|$ étapes de plus via la boucle for (3.9) ainsi que $|E_{\mathcal{T},prv}|$ de plus pour la boucle for (3.10). Finalement, la phase 2 a une complexité de $O(|Dset| + |E_{\mathcal{D},prv}| + 2 \cdot |\mathcal{T}| + |E_{\mathcal{T},prv}|) \equiv O(|Dset| + |E_{\mathcal{D},prv}| + 2 \cdot |\mathcal{A} \setminus \mathcal{D}| + |E_{\mathcal{A} \setminus \mathcal{D},prv}|) \equiv O(|\mathcal{A}| + |E_{\mathcal{A},prv}|)$. Toutes les instructions de (3.27) à (3.33) se font en $O(1)$ grâce à un pré-traitement décrit dans le paragraphe (Phase 2) de la sous-section 2.3. Enfin, la phase 3 trie topologiquement les actions jaunes de l'ensemble $\mathcal{D}$ à l'aide de l'ensemble des ordres $E_{\mathcal{D}}$ en $O(|\mathcal{D}| + |E_{\mathcal{D}}|)$; dans le pire des cas, cela équivaut à $O(|\mathcal{A}| + |E_{\mathcal{A}}|)$ qui finalement domine l'ensemble et prouve la complexité linéaire en temps de $\mathbb{P}$. $\square$

Théorème 3. $\mathbb{P}$ nécessite au plus $O(m^2n)$ d'espace.

Démonstration. Une action (PU) occupe $O(m)$ d'espace : les pré- et post-conditions peuvent être réduites à une variable chacune, plus une variable pour l'indice de la variable d'état affectée, et l'ensemble $\mathcal{N}_{pre}$ est un singleton en raison des restrictions (P) et (U) ; les prevail-conditions, au contraire, sont des listes de m éléments, et l'ensemble $\mathcal{N}_{prv}$ a au plus m éléments. La table de *Hashing* des actions prend $O(mn)$ d'espace. Enfin, l'ensemble $E_{\mathcal{D}}$ peut avoir au plus (m^2n) éléments : on a $|\mathcal{D}| \leq |\mathcal{A}| \leq mn$ en raison de la restriction (T_1) , il y a donc $O(mn)$ actions, et chaque action

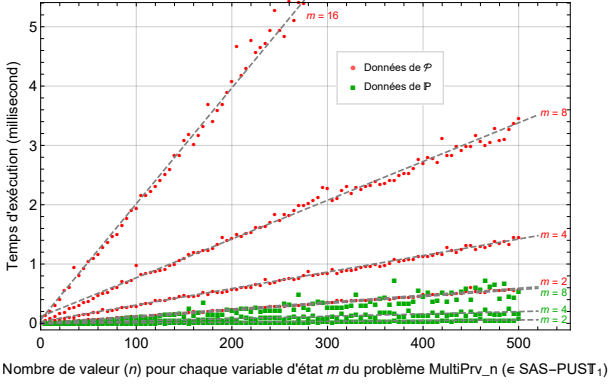


FIGURE 1 – Les temps d’exécution de $\mathbb{P}$ et $\mathcal{P}$ sont linéaires en fonction du nombre de valeurs (n) par variable d’état (m) bien que le problème MultiPrv_n_Cycle est conçu pour générer $O(m^2)$ ordres entre les mn actions. m est constant lors de cette expérience.

peut avoir $O(m)$ prédécesseurs : au plus $(m-1)$ via $\mathcal{N}_{prv}$ et au plus $(m-1)$ via $\mathcal{N}_m$. D’où $O(E_D) = O(m^2n)$. $\square$

3 Analyse comparative de $\mathbb{P}$ vs $\mathcal{P}$

Nous décrivons ici les trois problèmes que nous utilisons dans cet article pour illustrer les complexités temporelles de $\mathbb{P}$ et $\mathcal{P}$. Dans les figures 1, 2 et 3 les carrés verts, resp. les disques rouges, représentent les instances de problèmes résolus par $\mathbb{P}$, resp. $\mathcal{P}$. Les tests ont été effectués avec la configuration suivante : CPU AMD Ryzen 2700X (8-Core) (3,7GHz), 32Gb de RAM et Windows 10 (64 bits); les deux planificateurs sont écrits en C++14 avec les paramètres par défaut de Microsoft Visual Studio 2019.

MultiPrv_n_Cycle est conçu pour vérifier que le temps d’exécution des deux planificateurs est linéaire en (n) (cf. Figure 1); il génère $O(m^2)$ d’ordres entre mn actions avec $0 \leq p < n$; $\forall v_i \in \mathcal{M}$, nous avons :

- $pre(a_{v_i}^p) = p-1$ et $post(a_{v_i}^p) = p$, si $p > 0$
- $pre(a_{v_i}^p) = n-1$ et $post(a_{v_i}^p) = 0$, si $p = 0$
- $prv(a_{v_i}^p)[v_j] = \lfloor n/2 \rfloor$, pour $i < j \leq m$,
- $prv(a_{v_i}^p)[v_j] = u$, pour $1 \leq j \leq i$.

Pour l’expérience, les états initial et final sont de la forme :

- $\forall v_i \in \mathcal{M}, s_0[v_i] = 0$,
- $\forall v_i \in \mathcal{M} \setminus \{v_0\}, s_*[v_i] = 0 \wedge s_*[v_0] = n-1$.

[1] indique avoir implémenté $\mathcal{P}$ avec LISP et rapporte des temps d’exécution superlinéaires en (n), contrairement aux résultats théoriques (cf. **Théorème 4.10**, p. 89); car, entre autres, LISP ne fournit pas de contrôle sur la gestion de la mémoire, il a été suspecté que ce soit la raison des résultats pratiques altérés. Nous avons implémenté $\mathbb{P}$ et $\mathcal{P}$ en C++ sans optimisation spécifique, mais nous avons effectivement dû gérer soigneusement la mémoire pour obtenir des temps d’exécution linéaires en (n) pour les deux algorithmes $\mathcal{P}$ et $\mathbb{P}$.

MultiPrv_2_Cycle est un cas spécifique de MultiPrv_n_Cycle avec $n = 2$; ce n’est rien de plus que

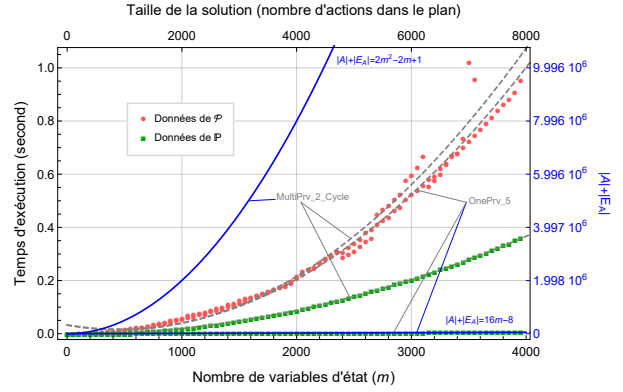


FIGURE 2 – Lorsque le nombre d’ordres est linéaire (resp. quadratique) par rapport au nombre de variables d’état (m), les temps d’exécution de $\mathbb{P}$ sont linéaires (resp. quadratiques) par rapport au nombre de variables d’état (m), ce qui illustre la complexité temporelle de $O(|A| + |E_A|)$, alors que les temps d’exécution de $\mathcal{P}$ sont toujours quadratiques par rapport au nombre de variables d’état (m) malgré les différences dans le nombre d’ordres entre MultiPrv_2_Cycle et OnePrv_5.

le problème du tunnel que nous avons présenté dans la section 2.1. Comme il génère $O(m^2)$ d’ordres entre $2m$ actions, nous pouvons facilement étendre ce problème pour montrer que, dans le pire des cas, le temps d’exécution de $\mathbb{P}$ est quadratique en fonction du nombre de variables m (cf. Figure 2).

OnePrv_5 est conçu pour montrer que le temps d’exécution de $\mathbb{P}$ peut être linéaire avec le nombre de variables (m) alors que celui de $\mathcal{P}$ est quadratique (cf. Figure 2) malgré le nombre $O(m)$ d’ordres d’action à traiter. OnePrv_5 génère $O(m)$ ordres entre $4m$ actions avec $0 < p < n = 5$; $\forall v_i \in \mathcal{M}$, nous avons :

- $pre(a_{v_i}^p) = p-1$ et $post(a_{v_i}^p) = p$,
- $prv(a_{v_i}^p)[v_{i+1}] = \lfloor (n-5)/2 \rfloor = 2$, avec $v_i \neq v_m$,
- $prv(a_{v_i}^p)[v_j] = u$, $\forall v_j \in \mathcal{M} \setminus \{v_{i+1}\}$.

Pour l’expérience, les états initial et final sont de la forme :

- $\forall v_i \in \mathcal{M}, s_0[v_i] = 0$,
- $\forall v_i \in \mathcal{M}, s_*[v_i] = 5$.

Western est conçu pour évaluer l’évolutivité de $\mathbb{P}$ et $\mathcal{P}$ pour notre domaine d’application des jeux vidéo. Il met en œuvre les routines quotidiennes des PNJs [9] du très populaire [25] jeu vidéo commercial Red Dead Redemption 2 [23]. Il est aussi proche que possible du jeu et est de classe SAS-PUT₁ que $\mathcal{P}$ ne peut gérer, avec $m = 29$, $n = 5$, $|A| = 58$, et 7 classes de PNJ avec 9 objectifs spécifiques. L’Éleveur de Chevaux décrit Table 1 est l’une de ces classes de PNJ.

Nous avons conçu une version SAS-PUT₁ de Western, qui est une version restreinte, que $\mathcal{P}$ et $\mathbb{P}$ peuvent traiter, avec $m = 19$, $n = 2$, $|A| = 37$, et 5 classes NPC avec 6 objectifs spécifiques. La figure 3 montre que $\mathbb{P}$ est trente mille fois⁶ plus rapide que $\mathcal{P}$ et peut générer des plans pour environ deux millions de NPCs en 1 milliseconde : la planification

6. (Fig. 3) Pour 1ms, $\mathcal{P}$ produit un plan pour 70 NPCs (abscisse rouge) contre 2,4 millions pour $\mathbb{P}$ (abscisse verte).

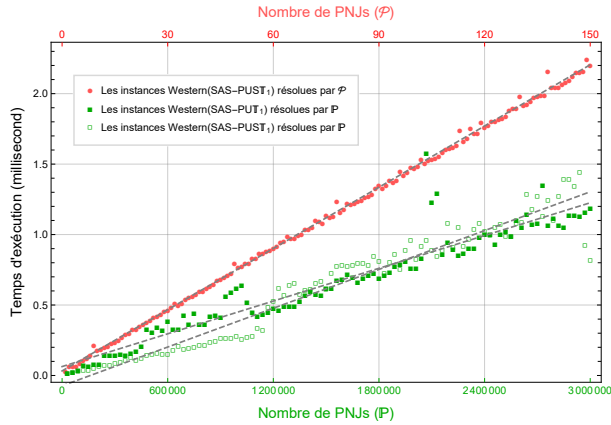


FIGURE 3 – L’abscisse rouge (haut) est l’abscisse du graphe rouge ($\mathbb{P}$) alors que l’abscisse verte (bas) est l’abscisse des deux graphes verts ($\mathcal{P}$). Les plans de solution des instances du problème Western sont longs de 3 à 10 actions : l’impact du nombre d’actions et du nombre de leurs ordonnancements est négligeable sur la génération de chaque plan. Par conséquent, les temps d’exécution de $\mathbb{P}$ et $\mathcal{P}$ sont linéaires par rapport au nombre de PNJs. Cependant, $\mathbb{P}$ construit chaque plan beaucoup plus rapidement que $\mathcal{P}$ et est donc capable de générer des plans pour deux millions de PNJs en moins d’une milliseconde.

avec $\mathbb{P}$ peut être utilisée pour contrôler de grandes villes dans les jeux vidéo commerciaux [15].

4 Discussion

Comme nous l’avons mentionné dans la sous-section 2.1, en conséquence du **Théorème 4.4** [1, p. 76] $\mathcal{P}$ résout les instances des classes de problèmes $\text{SAS}^+ \text{-PUT}_2$. Par conséquent, nous pourrions nous attendre à ce que $\mathbb{P}$ soit environ deux fois plus rapide que $\mathcal{P}$ sur des problèmes spécifiquement conçus en évitant de considérer l’insertion de deux actions de même type plutôt qu’une. $\mathbb{P}$ fait évidemment moins de travail que $\mathcal{P}$, mais quel genre de travail ? La restriction (T_1) n’explique certainement pas à elle seule pourquoi $\mathbb{P}$ est plus rapide de plusieurs ordres de grandeur que $\mathcal{P}$. La figure 2 montre que $\mathcal{P}$ ne profite pas de la croissance linéaire des ordonnancements entre action car il itère sur les variables d’état et non sur les ordres d’actions comme le fait $\mathbb{P}$. Dans la phase 2 de l’algorithme 3, $\mathbb{P}$ ne vérifie que les prevail-conditions définies grâce au système de voisinage $\mathcal{N}$, réduisant ainsi considérablement la charge travail de $\mathbb{P}$ par rapport à $\mathcal{P}$. Ceci est la clé pour expliquer la très grande efficacité d’exécution de $\mathbb{P}$ pour tous les problèmes que nous avons testés. Nous avons finalement observé que le tri topologique de nos graphes est déterministe comme c’est le cas pour le planificateur de [10] qui utilise des actions unaires sans Single-Valuedness mais avec des variables d’état binaires⁷). Les explications clés mineures concernent l’allocation et le remplissage des structures de données autant que possible à l’avance ; cependant, une gestion rigoureuse de la mémoire n’explique que la régularité

7. Les variables binaires sont également une des restrictions (B) du cadre de planification SAS^+ (cf. **Définition 3.9** [1, p. 62]).

des courbes (à la fois $\mathbb{P}$ et $\mathcal{P}$) dans les figures 1, 2, et 3.

Contrairement aux restrictions SAS^+ , la restriction (T_1) ne limite pas l’entrée de l’algorithme de planification mais sa sortie. Dans le domaine d’application des jeux vidéo tels que Western les plans font moins de 10 actions : il est assez simple de vérifier, même à la main, qu’une instance du problème satisfait la restriction T_1 ; c’est beaucoup plus complexe lorsque les plans contiennent des milliers d’actions. En l’état de cet article, la seule solution est d’exécuter $\mathbb{P}$ en tant que test d’appartenance à la classe du problème. Les temps d’exécution de $\mathbb{P}$ ne sont heureusement pas un obstacle à cette fin.

La restriction (T) limite également la sortie de l’algorithme de planification ; cependant, la recherche de plans totalement ordonnés est toujours possible. Un plan est une solution à un problème de planification lorsque son application transforme l’état initial en l’état but du problème : chaque action est appliquée une par une aux situations successives et donc l’application d’un plan correspond à un ordre total sur ses actions. Par conséquent, la restriction (T) peut être appliquée à toute classe de problèmes SAS^+ , qui inclut toutes les classes de problèmes de SAS. Il serait logique de limiter l’utilisation de la restriction (T) aux classes de problèmes telles que MultiPrv_2_Cycle dans lesquelles toutes les solutions sont des plans d’actions totalement ordonnés. Cependant, il est également logique d’utiliser la restriction (T) pour spécifier que notre objectif est de générer des plans d’actions totalement ordonnés, comme c’est actuellement le cas dans les jeux vidéo commerciaux où les PNJ exécutent une tâche après l’autre.

Le budget de traitement est probablement la seule contrainte qui empêcherait un moteur de jeu d’accéder à l’état actuel du jeu et donc de lire la valeur exacte des variables d’état du jeu. Par conséquent, les états initiaux et finaux totalement définis constituent une restriction réaliste dans notre domaine d’application des jeux vidéo. Les moteurs de jeu imposent toutefois un taux de mise à jour pour les variables d’état du jeu, ce qui peut faire que des valeurs périmées fassent partie des problèmes de planification. Cette forme rudimentaire d’incertitude [1, p. 64] devrait définitivement faire partie des futurs travaux sur la planification qui gère de grands mondes virtuels.

Une utilisation plus subtile des valeurs indéfinies est en lien avec la restriction (S2) du domaine SAS^+ , mais qui est assouplie dans le domaine SAS (cf. **Définition 3.2** [1, p. 52]). La restriction (S2) permet à un type d’action de définir une variable d’état qui était auparavant indéfinie : la postcondition de cette action peut définir une variable d’état qui est indéfinie dans la précondition. Il s’agit d’une solution rudimentaire pour gérer le taux de mise à jour des variables d’état du jeu : nous pouvons concevoir un type d’action dont le but est de s’assurer qu’une variable d’état donnée a une valeur ; lors de l’exécution de cette action pendant le jeu, le moteur de jeu attend la prochaine mise à jour de cette variable d’état. Cependant, nous n’avons pas approfondi la restriction (S2) car notre objectif principal dorénavant est de concevoir $\mathbb{P}^+$ pour résoudre les instances de problème de la classe $\text{SAS}^+ \text{-PUT}_1$.

5 Conclusion

Notre objectif était de générer des plans SAS-PU totalement ordonnés pour des millions de PNJs en temps réel et de sorte que deux actions n'aient pas le même type. À cette fin, nous avons apporté les contributions majeures suivantes au domaine SAS :

- Nous avons défini deux nouvelles restrictions : (1) la restriction (T) qui limite les solutions aux plans totalement ordonnés, et (2) la restriction (T_1) qui limite le nombre de types d'action à un dans toute solution ; nous avons noté T_1 la combinaison de ces deux restrictions.
- Nous avons conçu un algorithme, que nous avons noté $\mathbb{P}$, pour résoudre les instances du problème SAS-PUT₁, en assouplissant ainsi la restriction Single-Valuedness (S).
- Nous avons conçu plusieurs problèmes SAS-PUT₁ pour tester diverses caractéristiques et en particulier diverses complexités temporelles puisque notre objectif est la planification en temps réel ; en particulier, MultiPrv_n_Cycle qui généralise l'exemple du tunnel. Nous avons également conçu un domaine Western réaliste par rapport à notre domaine d'application des jeux vidéo.
- La complexité temporelle de notre algorithme dans le pire des cas est linéaire par rapport au nombre d'actions et à leur ordre entre elles ; les temps d'exécution de notre implémentation de $\mathbb{P}$ sont très rapides pour le domaine Western : $\mathbb{P}$ fournit un plan à deux millions de PNJ en environ une milliseconde, atteignant ainsi notre objectif.

Les travaux futurs prendront d'abord en compte les états initiaux et finaux partiels ; nous souhaitons ensuite aborder le problème SAS⁺-PUT₂ : est-il possible de concevoir un algorithme aussi efficace que $\mathbb{P}$ pour cette classe de problèmes ?

Références

- [1] Christer Bäckström. *Computational Complexity of Reasoning about Plans*. PhD thesis, Department of Computer and Information Science, Linköping University, september 1992.
- [2] Christer Bäckström. Equivalence and tractability results for SAS⁺ planning. In *Proceedings of 3rd Conference on the Principles of Knowledge Representation and Reasoning*, pages 126–137. Morgan Kaufmann, october 1992.
- [3] Christer Bäckström and Bernhard Nebel. Complexity results for SAS planning. *Computational Intelligence*, 11(4) :625–655, 1995.
- [4] Tom Bylander. Complexity results for planning. In *Proceedings of 12th IJCAI*, pages 274–279, August 1991.
- [5] Stéphane Cardon and Éric Jacopin. Binary GPU-planning for thousands of NPCs. In *IEEE Conference on Games*, pages 678–681. IEEE Press, August 2020.
- [6] Alex Champandard, Tim Verweij, and Remco Straatman. Killzone 2 multiplayer bots. In *Paris Game AI Conference*. AIGameDev, <https://www.guerrilla-games.com/read/killzone-2-multiplayer-bots> (accessed on December 1st, 2021), June 2009.
- [7] David Chapman. Planning for conjunctive goals. *Artificial intelligence*, 32(3) :333–377, 1987.
- [8] Chris Conway. GOAP in Tomb Raider. GDC AI Summit, March 2015.
- [9] DefendTheHouse. NPC Daily Life in Read Dead Redemption 2. <https://www.youtube.com/watch?v=MrUJgppMn4>, November 2018. Accessed on january 10th, 2022.
- [10] Carmel Domshlak and Ronen Brafman. Structure and complexity in planning with unary operators. In *Proceedings of the 6th International Conference on Artificial Intelligence Planning Systems*, pages 34–43. AAAI Press, 2002.
- [11] Kutluhan Erol, Dana Nau, and V.S. Subrahmanian. Complexity, decidability and undecidability results for domain-independent planning. *Artificial Intelligence*, 76(1-2) :75–88, 1995.
- [12] Simon Girard. Postmortem : AI action planning on Assassin's Creed Odyssey and Immortals Fenyx Rising. <https://www.gamedeveloper.com/programming/postmortem-AI-action-planning-on-Assassins-Creed-Odyssey--and-Immortals-FenyxRising->, November 2021. Accessed on december 1st 2021.
- [13] Peter Higley. GOAP at monolith productions. GDC AI Summit, March 2015.
- [14] Daniel Hillburn. *Simulating Behavior Trees – A Behavior Tree/Hybrid Planner Approach*, volume 1, chapter 8, pages 99–111. CRC Press, 2013.
- [15] Sean Hollister. The matrix awakens didn't blow my mind, but it convinced me next-gen gaming is nigh. <https://www.theverge.com/22828860/the-matrix-awakens-ps5-xbox-series-x-free-next-gen>, december 2021. Accessed on January 12th, 2022.
- [16] Samuel Horti. Why F.E.A.R.'s AI is still the best in first-person shooters – Flank, cover and run away. <https://www.rockpapershotgun.com/why-fears-ai-is-still-the-best-in-first-person-shooters>, April 2017. Accessed on December 3rd, 2021.
- [17] Troy Humphreys. *Exploring HTN Planners through Examples*, volume 1, chapter 12, pages 149–167. CRC Press, 2013.
- [18] Éric Jacopin. Game AI planning analytics : The case of three first-person shooters. In *Proceedings of the 10th AIIDE*, pages 119–124. AAAI Press, 2014.
- [19] Monolith Productions. F.E.A.R. public tools, June 2006.
- [20] Jason Ocampo. F.E.A.R. review. <https://www.gamespot.com/reviews/fear-review/1900-6169771/>, April 2007. Accessed on December 3rd, 2021.

- [21] Jeff Orkin. Agent architecture considerations for real-time planning in games. In *Proceedings of the 1st AIIDE*, pages 105–110, 2005.
- [22] Jeff Orkin. Three States and a Plan : The A.I. of F.E.A.R. In *Proceedings of the Game Developer Conference*, page 17 pages, 2006.
- [23] Rockstar Studios. Red Dead Redemption 2, November 2018.
- [24] Remco Straatman, Tim Verweij, Alex Champandard, Robert Morcus, and Hylke Kleve. *Hierarchical AI for Multiplayer Bots in Killzone 3*, chapter 29, pages 377–390. CRC Press, 2013.
- [25] Take Two Interactive. SEC Filing. <https://ir.take2games.com/node/27706/html>, July 2021. Accessed on January 13th, 2022.
- [26] Michiel van der Leeuw. The PS3’s SPU in the real world – a Killzone 2 case study. Game Developer Conference, March 2009.
- [27] William van der Sterren. *Hierarchical Plan-Space Planning for Multi-Unit Combat Maneuvers*, volume 1, chapter 13, pages 169–183. CRC Press, 2013.

Décision assistée par un système intelligent : de l'utilisation de l'IA à l'intention de l'opérateur dans la replanification.

Adrien Metge^{1,2}, Nicolas Maille¹, Benoît Le Blanc²

¹ ONERA–The French Aerospace Lab, Salon-de-Provence, France

² IMS UMR CNRS 5218, ENSC-Bordeaux INP, Bordeaux, France

adrien.metge@onera.fr, nicolas.maille@onera.fr,
benoit.leblanc@ensc.fr

Résumé

Prendre en compte le processus décisionnel des utilisateurs humains pourrait permettre à des systèmes d'assistance décisionnelle d'être plus adaptés. Cet article présente une étude qualitative où huit participants ont incarné un opérateur aérien chargé d'effectuer deux missions nécessitant une replanification de plan de vol de drone assisté par une IA. Des entretiens d'autoconfrontation ont permis d'établir un lien entre l'évolution des critères du compromis de plan et le processus décisionnel en cours, puis de développer un suivi automatique de ce processus. Nous en tirons des exemples de situations dans lesquelles le système d'IA peut adapter différemment ses recommandations, selon que l'opérateur est identifié comme impliqué dans une phase d'exploration ou bien une phase d'exploitation des plans de vol.

Mots-clés

Interaction avec l'humain, aide à la décision, planification assistée

Abstract

Taking into account the decision-making process of human users could allow decision support systems to be more adapted. This paper presents a qualitative study where eight participants embodied an air operator tasked with two missions requiring AI-assisted drone flight plan replanning. Self-confrontation interviews were used to link the evolution of plan compromise criteria to the ongoing decision process, and then to develop an automatic monitoring of this process. We derive examples of situations in which the AI system can adapt its recommendations differently, depending on whether the operator is identified as being involved in a flight plan exploration phase or exploitation phase.

Keywords

Human-autonomy teaming, decision support, assisted planning

Introduction

Si la multiplication des systèmes basés sur l'Intelligence Artificielle (IA) devrait profondément modifier les fonctionnalités et la façon dont nous utilisons les objets du quotidien, une question importante concerne la façon dont ils seront intégrés dans des applications à haut risque. Il est attendu que l'utilisation d'algorithmes basés sur l'IA permette de mieux appréhender des situations complexes et d'étendre les capacités et la fiabilité de ces systèmes. Dans le domaine aéronautique particulièrement, l'arrivée de drones à haut niveau d'autonomie nécessite l'introduction d'une IA embarquée aux capacités étendues, mais aussi potentiellement plus opaque pour les opérateurs qui en ont la charge. La séparation géographique entre le vecteur aérien et la station au sol où se trouvent les opérateurs conduit à augmenter le niveau d'autonomie de l'aéronef pour faire face à certains aléas comme les évolutions rapides de la situation ou la perte de communication [2]. Face à cette robotisation et numérisation croissante du champ de bataille, l'approche techno-centrée héritée des paradigmes historiques montre ses limites. La saturation informationnelle et la surcharge cognitive induites par le foisonnement des interfaces poussent à reconsidérer la place de l'humain dès les phases initiales de conception des systèmes [4]. Les opérations actuelles et futures se dérouleront de plus en plus dans des environnements de mission hyperconnectés qui exigeront que les processus de commandement et de contrôle soient rapides et agiles, et où le rôle des opérateurs devrait évoluer vers une fonction de supervision de systèmes intelligents dont les capacités de calcul dépassent déjà les capacités humaines en ampleur et en vitesse [11].

Si cette évolution ouvre de nouveaux contextes opérationnels permettant une meilleure efficacité des interventions militaires, il est essentiel que les décisions d'action prises restent entièrement sous la responsabilité des opérateurs humains, car ce sont eux qui en porteront la responsabilité

politique et juridique devant la société. Il existe actuellement un consensus international sur le principe du contrôle humain des systèmes intelligents, qui se traduit dans les textes par « des garanties et des mécanismes, tels que l'attribution de la capacité de décision finale aux humains, qui sont appropriés au contexte et à l'état de l'art », la possibilité de décider de ne pas utiliser un système intelligent afin de conserver des niveaux de jugement humains, ou encore la possibilité que la décision de l'humain prime sur celle calculée par le système [18]. En outre, lorsque des systèmes hautement automatisés fonctionnent dans des environnements complexes, les opérateurs humains contribuent à la résilience globale du système par leur capacité à voir et à agir en dehors du champ des capacités de l'automatisation [15]. Pour permettre à l'opérateur humain de rester maître des décisions prises, il est alors nécessaire de s'assurer de sa compréhension aussi bien de la situation que du fonctionnement du système qu'il utilise, et de concevoir des outils d'aide à la décision respectueux de son processus décisionnel. Ces changements appellent une réflexion à la fois sur la manière dont les systèmes intelligents présentent et justifient les connaissances qu'ils apportent, et sur le processus d'interaction humain-IA qui doit être mis en place. Bien que l'humain puisse formellement avoir le droit de décider (d'accepter ou de rejeter les suggestions du système), on peut se demander si ce droit est toujours valable lorsque les implications de ces propositions sont difficiles à comprendre. Pour maintenir les opérateurs dans la boucle de décision, la

prise de décision humaine repose donc de plus en plus sur un prétraitement de l'information qui doit améliorer la compréhension de l'opérateur de la situation autour du dispositif et de son état de fonctionnement. Une étape importante pour préparer l'avenir est l'étude des méthodologies de travail en équipe entre les humains et les systèmes intelligents, qui visent à faire de la machine un collaborateur à part entière du processus décisionnel [6]. La notion d'engagement y occupe une place centrale, l'agent artificiel idéal étant supposé pouvoir prendre des initiatives pour suggérer, contraindre, produire, évaluer et modifier la décision en construction, tout en tenant compte des intentions de l'opérateur humain.

Dans ce cadre, le paradigme « Human-Autonomy Teaming » a récemment émergé pour préparer ces évolutions en cherchant de nouvelles façons de penser les interactions [7]. Une équipe humain-autonomie y est définie comme un collectif interdépendant dans l'activité et le résultat, impliquant un ou plusieurs humains et un ou plusieurs agents artificiels, dans lequel chacun est considéré comme un membre à part entière et occupe un rôle distinct, et dont tous les membres s'efforcent d'atteindre un objectif commun [14]. Pour Sciarra et al., afin qu'un agent artificiel puisse interagir à la manière d'un coéquipier, lui et ses partenaires humains devront être capables d'anticiper et de s'adapter à l'état des autres [17]. Afin d'être véritablement au service du décideur et de lui apporter une valeur ajoutée, les outils de recommandation

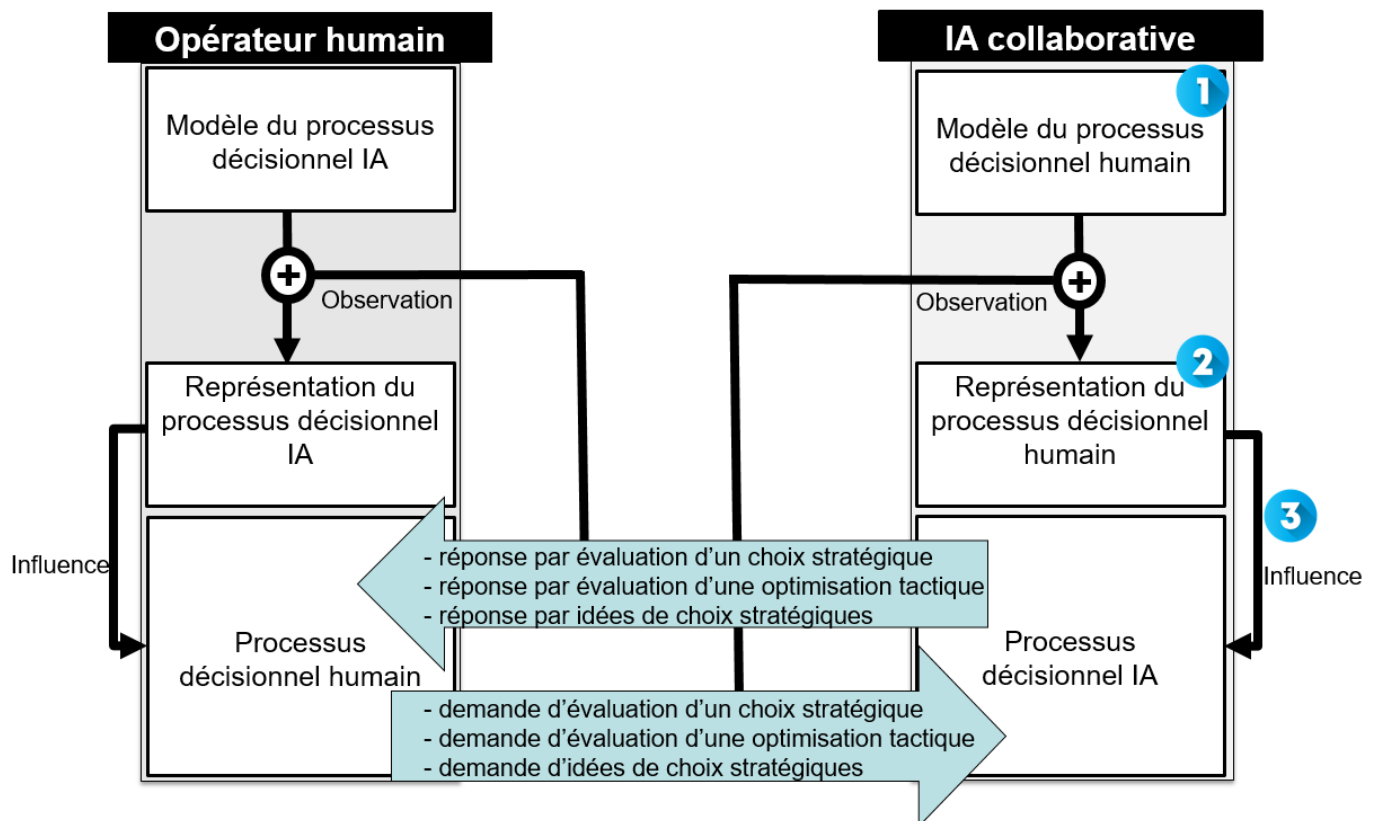


Figure 1 : Modèle de la prise de décision assistée par une IA collaborative qui s'adapte au processus décisionnel de l'opérateur.

ergonomique se doivent d'intégrer les limites cognitives, émotionnelles et sociales qui structurent la prise de décision humaine [8]. Pour ce faire, les systèmes intelligents pourraient chercher à apprendre à connaître leurs coéquipiers par le biais d'une pré-interaction afin de construire un modèle de leurs représentations et qui serait ensuite utilisé pour déterminer la stratégie d'interaction optimale [5]. Une meilleure description du lien entre l'utilisation de l'IA et la manière de comprendre une situation pourrait ainsi faciliter la prédiction du type de solutions recherchées par les opérateurs au cours du processus d'interaction, et conduire à une communication plus contextualisée [1]. Personnaliser l'interaction proposée par un système intelligent à un opérateur en fonction d'une estimation de ses intentions décisionnelles nécessite alors que le système soit capable de bâtir une telle représentation [20].

Nous distinguons à partir de cet examen de la littérature trois étapes nécessaires à la conception d'une IA d'aide à la décision humaine réellement au service du processus décisionnel des opérateurs, que nous synthétisons dans la Figure 1. La première étape est la construction d'un modèle du processus décisionnel des opérateurs pour un cas d'usage, c'est-à-dire à caractériser les grandes étapes d'élaboration d'une solution à ce type de problème donné. La seconde étape consiste à automatiser le suivi d'une instance de ce modèle décisionnel. La caractérisation du lien entre l'utilisation faite de l'IA et les intentions de l'opérateur pourrait aider le système à construire une représentation de ce processus de décision au fur et à mesure de son évolution. La troisième étape vise à identifier des opportunités d'assistance adaptées au processus décisionnel de l'opérateur tel que le système se le représente. Dans la première partie de cet article, nous présentons l'environnement de simulation qui nous permet d'investiguer expérimentalement les modalités de prise de décision en équipe humain-IA, à travers une tâche de replanification de vol de drone. Dans les parties 2, 3 et 4, nous implémentons respectivement les trois étapes de ce modèle afin de concevoir un tel système d'aide à la décision humaine.

1 Environnement de simulation

Ce travail s'appuie sur un environnement de simulation qui met en scène un opérateur aérien militaire chargé de superviser un drone pour effectuer des missions d'observation en zone ennemie. L'objectif des missions est de survoler

plusieurs cibles pour les photographier, puis de quitter la zone ennemie en minimisant les risques pris et le carburant consommé. Les essais correspondent à des missions qui se déroulent sur différents terrains mais avec le même scénario (Figure 2) : 1) le drone se dirige vers la zone ennemie avec un plan de vol initial, 2) de nouvelles menaces sont soudainement détectées, le plan de vol n'est en conséquence plus satisfaisant, 3) l'opérateur interagit avec le système pour définir un nouveau plan de vol, 4) l'opérateur valide un nouveau plan de vol, ce qui termine la tâche de supervision. Les terrains sont représentés dans l'interface par une grille quadrillée aux cases colorées : les cases vertes correspondent aux cibles à aller photographier, les cases grises et noires représentent des zones où le relief est élevé et très élevé, et les cases rouges des menaces dont la dangerosité estimée est symbolisée par une valeur numérique (Figure 3). Les terrains proposent des situations à la fois contrôlées, permettant à différents opérateurs de vivre des situations comparables, mais aussi suffisamment complexes pour que la décision nécessite une assistance du système et soit nécessairement le résultat d'un compromis.

La qualité du plan de vol prévisionnel est représentée selon trois dimensions par trois jauges, et la tâche peut être comprise comme un problème d'optimisation multicritère : l'objectif est de minimiser à la fois la proportion de cibles abandonnées, la quantité de carburant consommé et le risque pris. Cependant, les terrains ont été conçus de telle sorte qu'il n'existe pas de plan concevable qui permette en même temps d'aller photographier toutes les cibles, de consommer peu de carburant et de prendre peu de risque. L'absence de solution satisfaisante au problème fait rentrer en jeu une dimension subjective du jugement humain qui passe alors par l'expression de préférences individuelles sur les critères à dégrader. La conception technique de cet environnement est présentée plus extensivement dans [13]. Pendant la phase d'interaction, L'interface fournit des outils de replanification collaboratifs qui permettent la construction d'un nouveau plan de vol avec différentes répartitions des tâches entre l'opérateur et le système. Un outil de modification manuelle du plan vient limiter l'IA à une assistance calculatoire ; un outil de co-construction repose sur une contribution mixte nécessitant des entrées à la fois de l'opérateur et du système d'IA, notamment pour la sélection des cibles à conserver et leur ordre de passage ; un outil de consultation des suggestions de l'IA

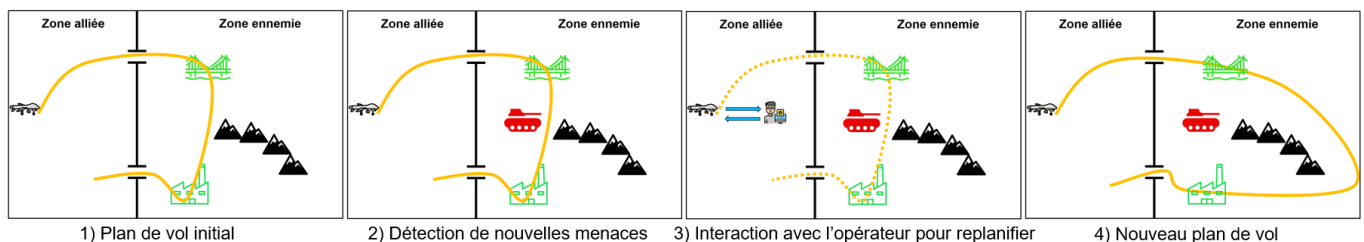


Figure 2 : Scénario de replanification d'un plan de vol de drone par un opérateur aérien.

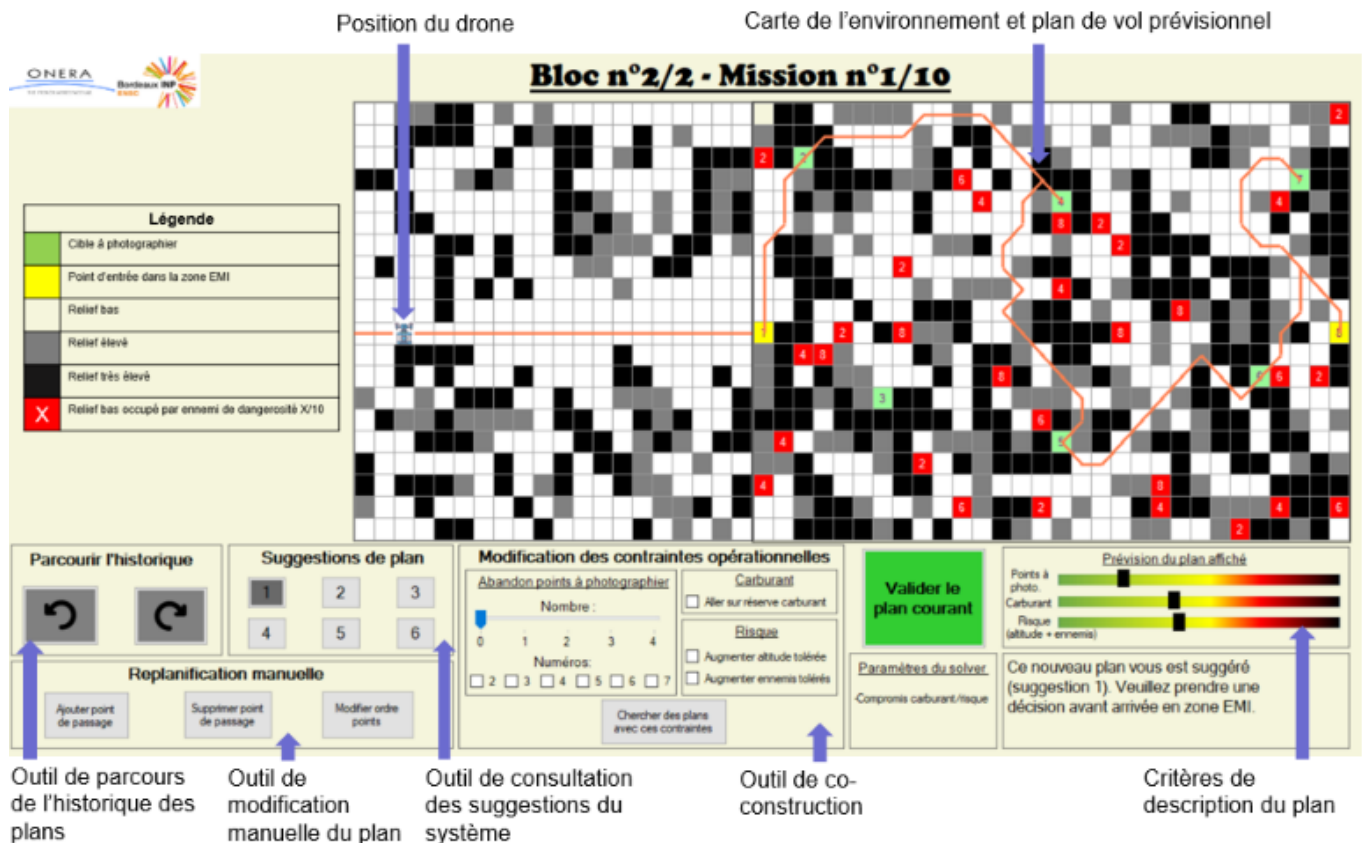


Figure 3 : Interface de l'environnement de simulation pour la tâche de replanification de plan de vol de drone.

place le système en position de leader qui va suggérer des solutions globales ; et un outil de parcours de l'historique des plans construits donne de la flexibilité à l'interaction en permettant de revenir à des trajets que les membres de l'équipe auraient précédemment élaborés.

2 Construction d'un modèle du processus décisionnel des opérateurs

Afin de construire une méthodologie d'assistance à la décision humaine adaptée à son processus décisionnel, la première étape a consisté à déterminer la forme générale de ce processus. Dans ce cadre, nous avons mené une étude expérimentale qualitative dont l'objectif était de déterminer quelles méthodologies de replanification des opérateurs peuvent déployer empiriquement dans l'environnement

présenté de simulation. Le plan expérimental est montré dans la Figure 4. Tout d'abord, au cours d'une phase de pré-test, 24 participants ont réalisé la tâche de replanification sur 20 terrains différents. Cela nous a permis de déterminer les deux terrains présentant le plus de variabilité dans les critères compromis des plans validés. Nous avons conservé ces deux plans pour la suite de l'étude afin de maximiser la complexité de la décision à prendre. Ensuite, nous avons réalisé la phase d'expérimentation proprement dite avec 8 participants (4 femmes, 4 hommes), jeunes chercheurs et ingénieurs de recherche en aéronautique, d'âge moyen 28.1 ans et d'écart-type 5.2 ans. Les participants ont réalisé la tâche de replanification sur les deux terrains sélectionnés, pendant que leur utilisation de l'interface était filmée. Ensuite, un entretien d'autoconfrontation individuel était conduit avec chacun d'entre eux. L'entretien d'autoconfrontation est une méthode

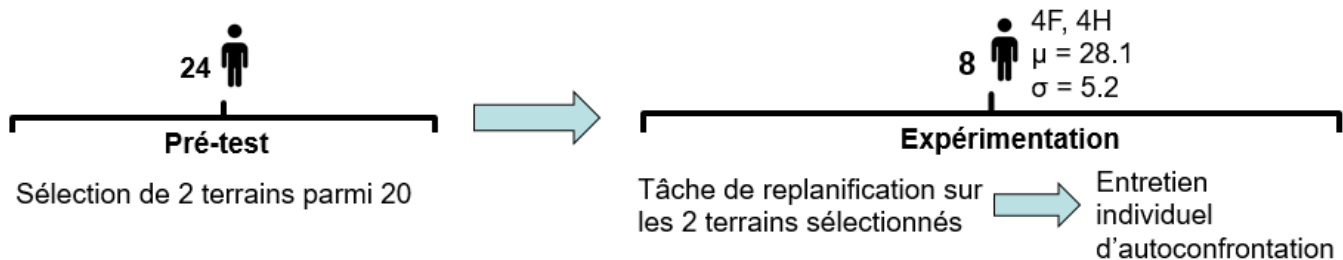


Figure 4 : Plan expérimental de l'étude qualitative.

- a)
 « je vais d'abord regarder **les cibles**, pour savoir **lesquelles virer** »
 « tout d'abord regarder le **positionnement des cibles** vis-à-vis des risques, la **distance entre les cibles** »
 « au début, chercher **quels objectifs** je devrais enlever si je dois en enlever »
 « au début je regardais quelles étaient les **mauvaises cibles**, c'est-à-dire difficile d'accès »
 « il fallait regarder **quel point** coûte cher, c'est-à-dire qui est hyper entouré, difficile d'accès »
- b)
 « ensuite je me demandais si je ne pouvais pas **améliorer** »
 « ensuite je me suis demandé si je pouvais encore **améliorer un peu le chemin** ou pas, **à la marge** »
 « ensuite j'ai cherché si je pouvais pas encore **optimiser des portions, échanger du carburant contre du risque** »
 « j'ai ajouté **une déviation** que j'avais identifiée plus tôt **entre C3 et C4** »
 « puis j'ai ajouté manuellement la **dévi**ation **identifiée avant C7** »
- c)
 « je me suis demandé si finalement je ne pouvais pas **ajouter manuellement C6** »
 « je suis retourné **sur C6** qui correspond à ça »
 « donc je me suis dit qu'il fallait forcément **abandonner C5**, et revenir à l'idée de toute à l'heure »
 « mais finalement, je me suis demandé si on ne pouvait pas **ajouter C7**, mais différemment, en faisant un détour plus long »

Figure 5 : Verbatim de l'explicitation de la méthodologie de replanification fournie par les participants lors des entretiens. Les termes « objectif », « point » ou « Cx » désignent les cibles liées aux terrains. Les mots mis en gras caractérisent le type de modification de plan réalisée ou considérée par l'opérateur, et les mots soulignés l'étape temporelle du raisonnement à laquelle la modification intervient.

d'analyse de l'activité humaine qui consiste à confronter un individu aux traces de son activité et à l'inciter à expliquer sa production par rapport à la réalité de sa pratique [19]. Au cours de ces entretiens, l'expérimentateur a visionné en leur compagnie les enregistrements vidéos de leurs réalisations de la tâche afin de recueillir les intentions et les stratégies décisionnelles liées à chaque utilisation des outils de replanification.

La synthèse des entretiens a révélé une grande variabilité à la fois des outils de replanification utilisés et des compromis de plans validés par les huit participants sur les deux terrains replanifiés. Malgré cette variabilité inter-individuelle, des récurrences dans les stratégies de planification déployées sont observables. La Figure 5 rassemble certains fragments de phrases prononcées par les participants lors des entretiens et qui renseignent sur leur processus décisionnel. Les intentions des opérateurs peuvent être divisées en deux catégories générales qui semblent correspondre à deux type de phases décisionnelles distinctes. D'une part, celles qui se concentrent sur les cibles, à travers des réflexions quant à leurs

localisations relatives et leur proximité aux menaces afin de déterminer lesquelles pourraient être abandonnées (Figure 5a). D'autre part, celles qui portent sur les segments de plan qui relient les cibles entre elles, ce qui consiste à localement prendre plus de risque pour diminuer la consommation de carburant ou l'inverse (Figure 5b).

Ce dilemme entre optimiser l'utilisation des ressources existantes et trouver un moyen d'en créer de nouvelles correspond à un modèle descriptif de la science managériale connu sous le nom de « conflit exploration-exploitation dans le processus cognitif humain » [3, 16]. Ce modèle a également fait l'objet de recherches en neurosciences, où des corrélats neurophysiologiques à ces deux états ont pu être observés, l'exploitation activant les régions du cerveau associées à la recherche de récompense pour évaluer la valeur des choix actuels, et l'exploration faisant appel aux régions associées au contrôle attentionnel pour investiguer des choix alternatifs [10]. Par ailleurs, les marqueurs temporels de raisonnement que nous avons notés tendent à indiquer que ces cycles décisionnels d'exploration et d'exploitation peuvent se

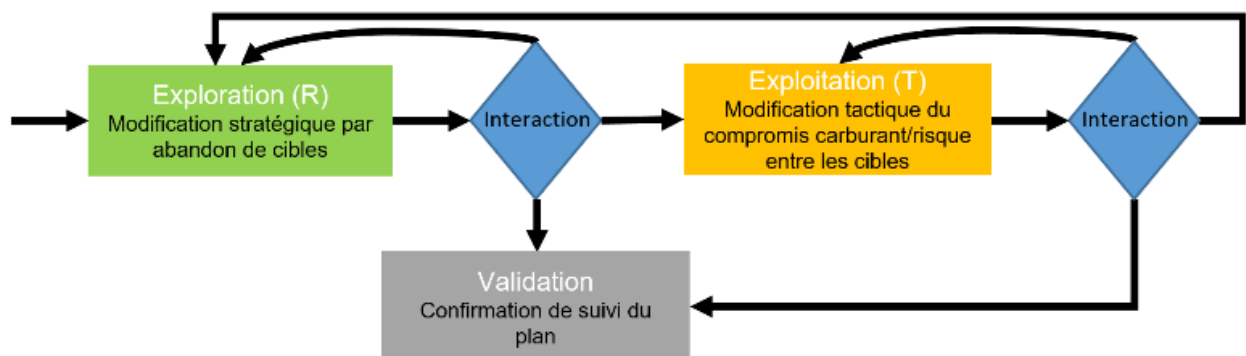


Figure 6 : Processus décisionnel type des opérateurs pour la tâche de replanification de plan de vol.

succéder plusieurs fois au cours de la replanification, et en débutant systématiquement par une phase d'exploration (Figure 5c). Tous les participants ont de fait commencé les interactions par une phase d'exploration, puis se sont engagés dans une alternance de phases d'exploitation et d'exploration jusqu'à la validation du suivi d'un plan (Figure 6). Les deux mots « exploration » et « exploitation » ayant une calligraphie voisine, nous proposons de les distinguer par leur lettre médiane : (R) pour exploration et (T) pour exploitation.

3 Automatisation du suivi de l'évolution du modèle décisionnel pendant son déploiement

Une fois la forme générale du processus décisionnel des opérateurs dans notre environnement de simulation identifiée, la seconde étape était de déterminer une méthodologie pour que le système puisse reconnaître automatiquement ce processus au fur et à mesure de son évolution au cours de l'interaction. Il s'agissait de donner au système la capacité de classer dynamiquement l'opérateur comme étant dans une phase décisionnelle d'exploration ou bien d'exploitation, et de détecter le passage de l'une à l'autre.

Pour cela, nous nous sommes intéressés aux liens possibles entre les types d'outils de replanification utilisés et les phases décisionnelles concomitantes. Nous avons observé que l'outil de parcours de l'historique des plans relève systématiquement de l'exploitation et est utilisé soit pour comparer les propriétés tactiques de plusieurs plans, soit pour retourner à un plan précédent afin qu'il redevienne le plan de travail. A l'inverse, le rôle d'assistance ou d'élaboration de l'IA compris dans les autres outils a pu être utilisé pour de l'exploration ou bien pour de l'exploitation. La consultation d'une suggestion de

plan de l'IA est une action d'exploration lorsque le plan proposé par le système est vu pour la première fois ou lorsque l'opérateur est dans une phase de recherche initiale, mais est une action d'exploitation lorsque l'opérateur y revient après en avoir déjà pris connaissance. Pour les outils de modification manuelle et de co-construction, où l'opérateur participe à l'élaboration de la solution, les critères de recherche de plan qu'il fournit à l'IA contiennent des indices sur son but sous-jacent. Lorsque ces critères de recherche comprennent notamment l'abandon des mêmes cibles que le plan précédent, l'opérateur est dans une phase d'exploitation et emploie l'IA pour de l'optimisation tactique. Lorsque ces critères de recherche portent sur une famille de plans qui n'a jamais été examinée, alors l'IA est employée à des fins d'exploration pour envisager de nouvelles stratégies.

Ainsi, au-delà du type d'outil de replanification utilisé, c'est le ou les critères de compromis du plan que l'opérateur cherche à modifier qui semblent renseigner sur sa volonté d'explorer ou d'exploiter des solutions. Les modifications changeant les cibles à abandonner caractérisent les phases d'exploration, tandis que les modifications conservant les mêmes cibles et portant donc exclusivement sur les quantités de risque pris et de carburant consommé caractérisent les phases d'exploitation. Si une modification du plan implique un changement des cibles abandonnées, alors l'opérateur travaille à un niveau stratégique et se trouve à cette étape dans une phase décisionnelle d'exploration (R). Si la modification du plan n'implique pas de changement des cibles abandonnées, alors cette modification porte exclusivement sur les critères de carburant consommé et de risque pris, ce qui correspond à un travail tactique et une phase décisionnelle d'exploitation (T). La Figure 7 représente l'espace théorique des compromis de plan selon les trois dimensions que nous considérons : cibles

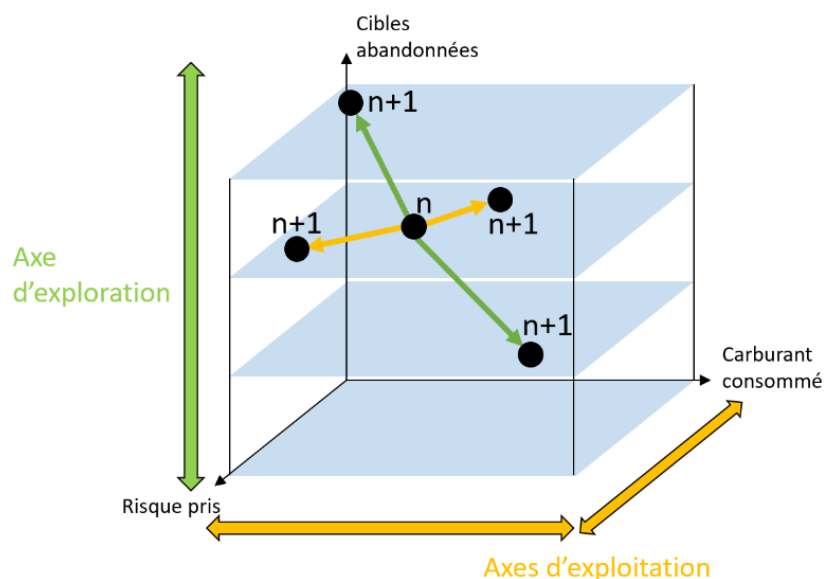


Figure 7 : Exemple théorique de la phase décisionnelle liée à plusieurs possibilités de compromis à l'étape $n+1$ de modification de plan selon une étape n donnée.

Cible(s) abandonnée(s)	3,7	3,4	4	4,5	5	5,7	5	3	3	3,4	3,5	5	5	5	5,7	5	5
Changement de cibles abandonnées ?	oui	non	non	non	non	non	non	non	oui	non	non	non	oui	oui	non	non	oui
Phase de décision inférée	R	R	R	R	R	R	R	R	T	R	R	R	T	T	R	R	T

Figure 8 : Inférence de la phase décisionnelle de l'opérateur à chaque modification du plan de vol pour un essai de l'étude.

abandonnées, carburant consommé et risque pris. Cet exemple théorique montre la phase décisionnelle dans laquelle se trouverait un opérateur selon différentes possibilités de compromis de plan à l'étape de $n+1$ étant donné une étape n : exploration si le compromis se déplace entre autres sur l'axe lié aux cibles abandonnées, exploitation si le compromis se déplace seulement sur les axes de risque pris et de carburant consommé.

À partir de cette règle, la Figure 8 fournit un exemple d'inférence des phases décisionnelles par lesquelles passe un opérateur au cours de la tâche de replanification sur un des essais de l'étude. Chaque colonne du tableau correspond à une utilisation d'un outil de replanification pendant cet essai. La première ligne indique les numéros des cibles qui sont abandonnées par le plan prévisionnel à cette étape. La seconde ligne indique le résultat du critère d'inférence des phases décisionnelles, qui consiste à examiner s'il y a un changement de cibles abandonnées entre la modification de plan précédente et la modification de plan actuelle, c'est-à-dire entre une colonne et la colonne précédente. La troisième ligne en déduit la phase décisionnelle dans laquelle se trouve l'opérateur après chaque interaction.

4 Identification d'opportunités d'assistance adaptées à l'évolution du processus décisionnel suivi

Ces liens entre les modifications des critères de compromis de plan et le processus décisionnel humain nous ont permis de développer une méthodologie de reconnaissance des phases d'exploration et d'exploitation par lesquelles passent les opérateurs au cours de la replanification. Ce processus peut être reconstruit a posteriori une fois la replanification achevée et un nouveau plan de vol validé, mais aussi dynamiquement pendant la réalisation de la tâche. Cette capacité du système à bâtir une représentation décisionnelle de son coéquipier préfigure des modalités d'assistance à la replanification qui s'appuient sur ces mécanismes cognitifs. Cette représentation des intentions humaines permet au système de concevoir l'opérateur comme un agent rationnel et d'éclairer ses croyances, c'est-à-dire l'espace des solutions qu'il a considéré, et ses désirs, c'est-à-dire ses préférences en termes de critères de compromis de plan. Si actuellement dans l'environnement de simulation le système ne participe à l'élaboration de la décision qu'en réponse à des demandes de l'opérateur soit d'idées ou d'évaluation de choix stratégiques, soit d'évaluation d'optimisation tactiques, cela permet d'envisager une interaction plus symétrique à travers la proposition contextualisée d'idées de choix stratégiques ou d'optimisations tactiques. Nous envisageons ces prises d'initiative du système comme des assistances à l'exploration et à l'exploitation, à travers l'ajout sur la carte dans l'interface de deux plans alternatifs que l'opérateur pourrait choisir comme nouveau plan de travail. Afin d'aider l'opérateur à les

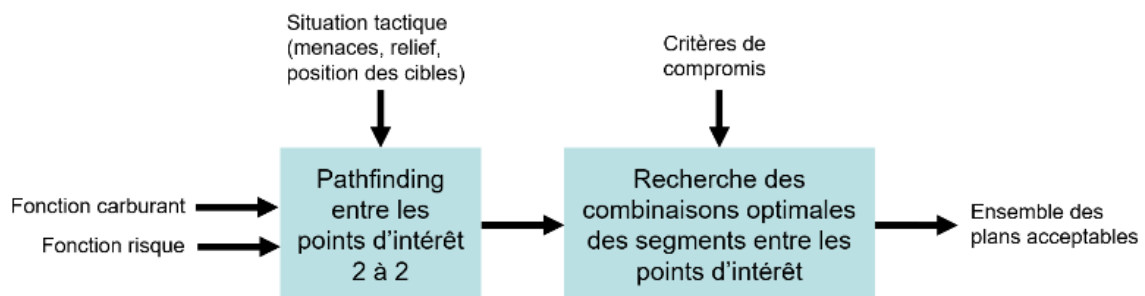


Figure 9 : Algorithme de recherche de plans à suggérer à l'opérateur implémenté dans l'environnement. Premièrement, l'algorithme recherche les segments optimaux reliant les 8 points d'intérêts 2 à 2 (les 6 cibles du terrain, la case d'entrée et la case de sortie de la zone ennemie) grâce à l'algorithme A*. À chaque case du terrain est associé un coût prenant en compte des éléments liés à sa situation tactique et le poids en carburant et en risque correspondant et définis par des fonctions carburant et risque. Deuxièmement, l'algorithme recherche les combinaisons optimales de ces segments entre les points d'intérêts par une méthode de séparation et évaluation. Troisièmement, l'application des critères de compromis aux plans construits permet de délimiter l'ensemble de ceux qui sont acceptables.

juger, leurs critères de compromis respectifs seraient également affichés. Un exemple de l'assistance à l'exploration tel qu'il pourrait être implémenté de cette manière est présenté dans les Annexes 1 et 2.

L'assistance à l'exploration pourrait être mise en œuvre quand deux modifications de plan n et $n-1$ sont classifiées comme correspondant à des phases d'exploration. Afin d'aider l'opérateur dans ses recherches stratégiques, les propositions de plans alternatifs consisteraient à abandonner des cibles différentes de celles qu'il a considérées. L'assistance à l'exploitation, elle, pourrait être mise en œuvre quand une modification n classifiée comme phase d'exploitation succède à une autre $n-1$ classifiée comme phase d'exploration. L'IA participerait à l'optimisation tactique du plan en proposant soit d'autres ordres de passage entre les cibles, soit la modification des segments entre les cibles pour diminuer le risque pris ou le carburant consommé. L'algorithme de recherche de plan actuellement implémenté, schématisé en Figure 9, et qui sert aux différents outils de modification de plan, pourrait être adapté pour la génération de ces propositions de plans alternatifs.

Conclusion

Dans cet article, nous avons présenté une implémentation des trois étapes que nous avons identifiées pour la conception d'une IA d'assistance qui s'appuie sur le processus décisionnel humain. Nous avons établi un modèle du processus décisionnel de l'opérateur pour une tâche de replanification de plan de vol de drone dans un environnement de simulation à travers des entretiens d'autoconfrontation. Ensuite, nous avons analysé comment la manière dont l'opérateur utilise les différents outils permet de révéler dans quelle phase de construction de la solution il se situe, ce qui nous a permis d'automatiser le suivi de ce processus décisionnel. Enfin, nous avons proposé deux opportunités d'assistance adaptées à l'évolution de ce processus : une aide à l'exploration, et une aide à l'exploitation des plans de vol. Ces résultats montrent qu'il est possible de concevoir des outils d'aide à la décision qui laissent une part de créativité à l'opérateur et permettent ainsi d'avoir des indices sur ce qu'il cherche à réaliser et la manière dont il veut atteindre son but. Ceci ouvre la possibilité à l'outil d'assistance d'avoir un comportement plus adaptatif qui prenne mieux en compte les besoins actuels de l'opérateur. Bien que l'étude soit réalisée dans le cadre restreint de la décision liée à la replanification d'un trajet, elle montre l'intérêt d'étendre le paradigme d'interaction humain/système et de ne pas se limiter à une explicitation pour l'humain de ce que fait le système. Une adaptation réciproque des deux agents est envisageable, du moment que les outils mis à disposition permettent d'accéder au moins partiellement à leurs intentions.

Pour autant, on peut identifier la dégradation de l'esprit critique de l'utilisateur comme un risque lié à la mise en place de ces systèmes d'aide plus adaptatifs. Metge et al. ont mis en évidence en ce sens que lorsque le degré de participation d'un système à une prise de décision complexe augmente à travers l'introduction de suggestions, les choix validés par différents utilisateurs deviennent plus homogènes alors que leur sentiment de responsabilité rapporté ne diminue pas significativement [12]. Les utilisateurs peuvent avoir tendance à s'approprier les idées introduites par le système lorsqu'elles sont pertinentes, mais sans avoir nécessairement conscience de cette influence sur leur propre prise de décision. La quantité optimale à fournir de ces nouvelles modalités d'assistance reste une question ouverte à ce stade, et nécessiterait d'investiguer expérimentalement dans quelle mesure elles ne risqueraient pas d'altérer voire de se substituer au processus décisionnel humain en essayant de l'accompagner.

Remerciements

Le projet de recherche dont fait partie cette étude a été subventionné par l'Agence de l'Innovation de Défense (AID) et l'Office National d'Études et Recherches Aérospatiales (ONERA).

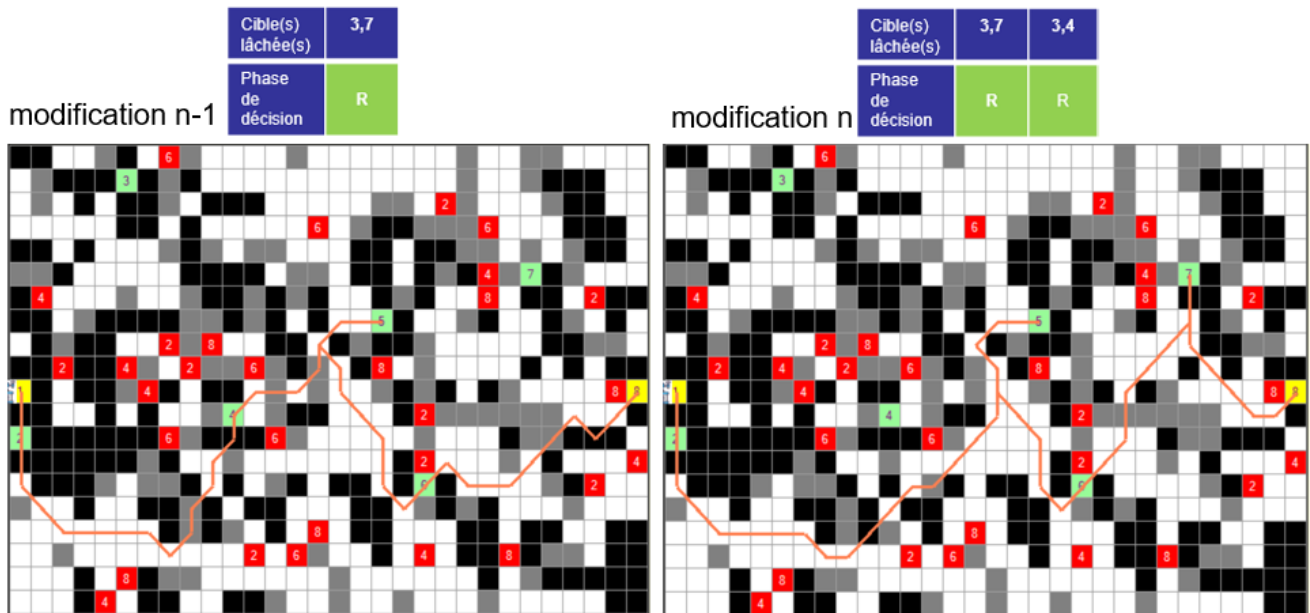
Les figures ont été créées en utilisant des ressources de Flaticon.com.

Références

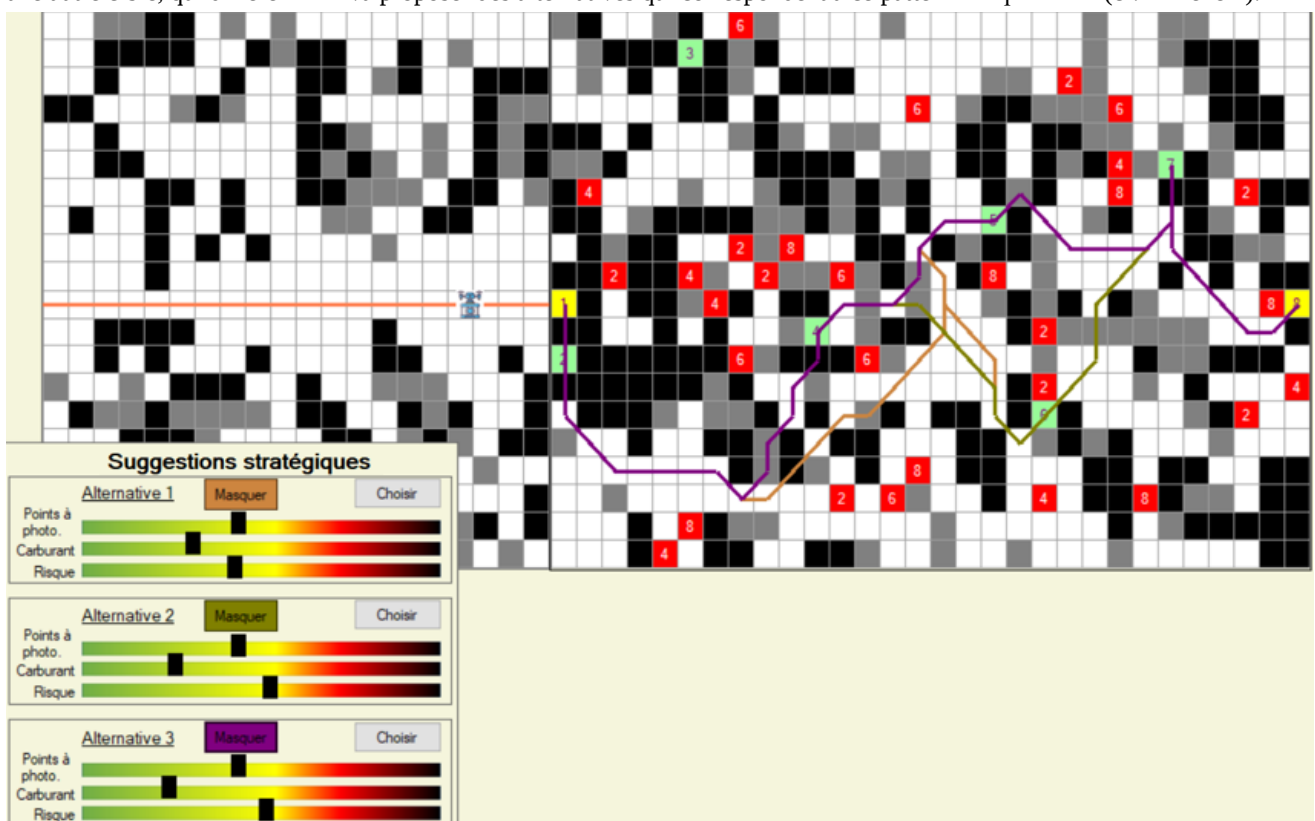
- [1] Alix, C., Lafond, D., Mattioli, J., De Heer, J., Chattington, M., & Robic, P. O. (2021, June). Empowering Adaptive Human Autonomy Collaboration with Artificial Intelligence. In 2021 16th International Conference of System of Systems Engineering (SoSE) (pp. 126-131). IEEE.
- [2] Atyabi, A., MahmoudZadeh, S., & Nefti-Meziani, S. (2018). Current advancements on autonomous mission planning and management systems: An AUV and UAV perspective. *Annual Reviews in Control*, 46, 196-215.
- [3] Berger-Tal, O., Nathan, J., Meron, E., & Saltz, D. (2014). The exploration-exploitation dilemma: a multidisciplinary framework. *PloS one*, 9(4), e95693.
- [4] Briant, R., « La synergie homme-machine et l'avenir des opérations aériennes », *Focus stratégique*, n° 106, Ifri, septembre 2021.
- [5] De Melo, C. M., Files, B. T., Pollard, K. A., & Khooshabeh, P. Social Factors in Human-Agent Teaming.
- [6] Deterding, S., Hook, J., Fiebrink, R., Gillies, M., Gow, J., Akten, M., ... & Compton, K. (2017, May). Mixed-initiative creative interfaces. In *Proceedings of the 2017 CHI Conference Extended Abstracts on Human Factors in Computing Systems* (pp. 628-635).
- [7] Flathmann, C., Schelble, B., Tubre, B., McNeese, N., & Rodeghero, P. (2020, November). Invoking Principles of Groupware to Develop and Evaluate Present and Future Human-Agent Teams. In *Proceedings of the 8th International Conference on Human-Agent Interaction* (pp. 15-24).

- [8] Jiang, J., Karran, A., Coursaris, C. K., Majorique, P., & Léger, J. B. A Situation Awareness Perspective on Human-Agent Collaboration: Tensions and Opportunities.
- [9] Klein, G. (2008). Naturalistic decision-making. *Human factors*, 50(3), 456-460.
- [10] Laureiro-Martínez, D., Brusoni, S., Canessa, N., & Zollo, M. (2015). Understanding the exploration-exploitation dilemma: An fMRI study of attention control and decision-making performance. *Strategic management journal*, 36(3), 319-338.
- [11] Lyons, J. B., Sycara, K., Lewis, M., & Capiola, A. (2021). Human-Autonomy Teaming: Definitions, Debates, and Directions. *Frontiers in Psychology*, 12, 1932.
- [12] Metge, A., Maille, N., & Le Blanc, B. (2021, June). Transition between cooperative and collaborative interaction modes for human-AI teaming. In *CNIA 2021: Conférence Nationale en Intelligence Artificielle* (pp. pp-38).
- [13] Metge, A., Maille, N., & Le Blanc, B. (2021, October). Interface de collaboration humain-IA: application au cas de la replanification de vol de drone. In *ERGO'IA 2021*.
- [14] O'Neill, T., McNeese, N., Barron, A., & Schelble, B. (2020). Human-autonomy teaming: A review and analysis of the empirical literature. *Human Factors*, 0018720820960865.
- [15] Prinzel, L., Ellis, K., Koelling, J., Krois, P., Davies, M., & Mah, R. (2021). Examining the Changing Roles and Responsibilities of Humans in Envisioned Future in- Time Aviation Safety Management Systems. In *78th International Symposium on Aviation Psychology* (p. 346).
- [16] Sinha, S. (2015). The exploration-exploitation dilemma: a review in the context of managing growth of new ventures. *Vikalpa*, 40(3), 313-323.
- [17] Sycara, K., Hughes, D., Li, H., Lewis, M., & Lauharatanahirun, N. (2020, September). Adaptation in human-autonomy teamwork. In *2020 IEEE International Conference on Human-Machine Systems (ICHMS)* (pp. 1-4). IEEE.
- [18] Tessier, C. (2021, June). Éthique et IA: analyse et discussion. In *CNIA 2021: Conférence Nationale en Intelligence Artificielle* (pp. pp-22).
- [19] Theureau, J. (2003). L'entretien d'autoconfrontation comme composante d'un programme de recherche empirique et technologique. *Cahiers de l'INSEP*, 34(1), 81-86.
- [20] Zhang, R., McNeese, N. J., Freeman, G., & Musick, G. (2021). "An Ideal Human" Expectations of AI Teammates in Human-AI Teaming. *Proceedings of the ACM on Human-Computer Interaction*, 4(CSCW3), 1-25.

Annexes



Annexe 1 : Deux modifications de plans successives lors de la tâche de replanification sur un terrain. Classifiées comme deux actions d'exploration, l'assistance à l'exploration est activée. Les deux modifications abandonnent une cible commune, la cible 3, et une autre cible, qui diffère. L'IA va proposer des alternatives qui correspondent à ce pattern d'exploration (cf. Annexe 2).



Annexe 2 : Deux suggestions de plans sont proposées en tant qu'alternative 2 et alternative 3, l'alternative 1 étant le plan à la modification n. Leurs critères de compromis et le masquage des tracés permettent de les comparer. Le plan de l'alternative 1 abandonne les cibles 3 et 4, le plan de l'alternative 2 les cibles 3 et 5, le plan de l'alternative 3 les cibles 3 et 6. On remarque que les alternatives 2 et 3 ont des critères de compromis similaires, avec une dépense de carburant inférieure mais un risque supérieur à l'alternative 1. Ces informations apportées à l'opérateur en réaction à ses modifications de plan précédentes pourraient nourrir sa réflexion stratégique et lui faire gagner du temps dans l'exploration de solutions.

Influence indépendante de l'exploration et de l'exploitation dans les métaheuristiques : application aux systèmes de recommandation

Alexandre Bettinger, Armelle Brun, Anne Boyer

Université de Lorraine, CNRS LORIA

prénom.nom@loria.fr

Résumé

L'exploration et l'exploitation (E&E) d'un espace de recherche sont deux processus fondamentaux en intelligence artificielle. L'influence de l'E&E vise à améliorer ces processus. L'explicabilité de l'E&E est aussi un enjeu important. Cet article introduit un processus indépendant d'influence de l'E&E et des indicateurs pour l'explicabilité de l'E&E. Nos expérimentations sur les algorithmes génétique et par renforcement confirment que (1) le processus d'influence proposé a un impact positif sur l'E&E, (2) les indicateurs proposés représentent l'E&E avec plus de clarté.

Mots-clés

Recommandation, métaheuristique, algorithme évolutionnaire, exploration, exploitation, influence, renforcement

Abstract

The exploration and exploitation (E&E) of a search space are two fundamental processes in artificial intelligence. The influence of E&E aims to improve these processes. The explainability of E&E is also an important issue. This article introduces an E&E independent influence process and indicators for the explainability of E&E. Our experiments on genetic and reinforcement algorithms confirm that (1) the proposed influence process has a positive impact on E&E, (2) the proposed indicators represent E&E more clearly.

Keywords

Recommendation, metaheuristic, evolutionary algorithm, exploration, exploitation, influence, reinforcement

1 Introduction

La recommandation est le fait de filtrer des informations afin de cibler des éléments (*items*, ressources) susceptibles d'intéresser un ou plusieurs utilisateurs. Cette tâche peut être vue comme le traitement d'un vaste espace de recherche qui représente l'ensemble des recommandations possibles. Selon le contexte de la recommandation, une recommandation peut prendre différentes formes telles que des *items* uniques, des ensembles d'*items* (*itemsets*) ou des séquences d'*items*. Dans ce qui suit, nous nous intéressons au cadre de recommandation d'*itemsets* (également appelés vecteurs ou solutions).

L'espace de recherche peut être exploré par un processus d'exploration et d'exploitation (E&E) pour trouver de bonnes recommandations [4]. Explorer signifie s'assurer que de nombreuses régions de l'espace de recherche sont examinées [6], afin de ne pas se retrouver piégé dans un optimum local. Exploiter signifie que si une région prometteuse de l'espace de recherche est trouvée, elle est examinée plus attentivement [6] pour voir si elle contient de bons optima locaux. Dans le cadre de la recommandation, un optimum local est un *itemset* considéré comme optimal par rapport à la zone de l'espace de recherche dans laquelle il se situe, sachant qu'il peut exister de meilleurs *itemsets* dans d'autres zones de l'espace de recherche. La littérature souligne que le domaine des métaheuristiques (MH) aborde la problématique de la proposition de recommandations dans de vastes espaces de recherche [6].

Des études ont été menées pour influencer la dynamique de l'E&E afin d'assurer de meilleurs comportements de recherche et ainsi trouver de meilleures recommandations concernant des critères d'évaluation prédéfinis [6]. Influencer l'E&E implique de savoir quand, pourquoi et comment influencer. En effet, l'E&E doivent être contrôlées régulièrement au travers d'une mesure et les opérations réalisées pour exercer l'influence doivent être justifiées. Notons ici que le contrôle de l'E&E et l'influence de l'E&E sont des mécanismes différents mais font tous deux partie du processus d'influence de l'E&E. Ces éléments contribuent à expliquer les comportements d'E&E afin de comprendre ce qui a conduit aux recommandations proposées. De plus, ils augmentent les capacités d'E&E à trouver de meilleures recommandations.

Dans un cadre de recommandation, un utilisateur peut avoir des préférences explicites ou implicites sur les recommandations attendues, qui peuvent être représentées par des contraintes [27] associées à des attributs d'*items*. Nous représentons les contraintes sous la forme de besoins de maximisation et de minimisation ainsi que de plages de valeurs totales. Par exemple, dans le cas de la recommandation d'un ensemble d'articles pour un week-end, le prix du logement ne doit pas être supérieur à 800\$, de préférence pas supérieur à 650\$, le montant total pour l'ensemble du séjour devant être inférieur à 1500 \$.

La problématique que nous visons à aborder dans ce travail est donc : Comment contrôler, influencer et expliquer

l'E&E dans le domaine des métaheuristiques ? Nos expérimentations sont réalisées dans le cadre de la recommandation par MH.

L'apport de ce travail est triple. Premièrement, nous proposons un processus d'influence indépendant de l'algorithme influencé afin de favoriser la pleine expression de l'influence et de ne pas opérer de modifications structurales et/ou paramétriques à l'intérieur de l'algorithme influencé. Deuxièmement, nous proposons une influence systématique de l'E&E, à chaque itération de l'algorithme influencé et tant que les *itemsets* ne respectent pas un indicateur d'E&E. Enfin, nous introduisons des indicateurs pour mesurer et représenter l'E&E sous de nouveaux angles, ce qui vient servir l'explicabilité de l'E&E.

2 Travaux connexes

Dans de vastes espaces de recherche, filtrer les informations est fastidieux et effectuer une recherche exhaustive est inconcevable. Par conséquent, le ciblage des *items* à recommander nécessite de parcourir partiellement l'espace de recherche à l'aide de mécanismes algorithmiques capables d'explorer et d'exploiter [4]. L'E&E sont les deux comportements clés pour parcourir un grand espace de recherche [6]. L'E&E sont aussi appelées diversification et intensification [19], ou recherche globale et recherche locale [6]. L'E&E sont donc intrinsèquement en conflit [6] car si beaucoup de temps est passé sur l'une, moins de temps est disponible pour l'autre. Eiben et Schippers [11] ont abordé la nécessité de trouver un équilibre entre E&E, considérant ces deux processus comme antagonistes.

Nous proposons une définition des MH, expliquons pourquoi elles sont utilisées pour la recommandation et proposons une synthèse de trois revues [6][29][22] qui traitent des questions liées à l'influence de l'E&E, par compromis ou guidage ou recherche d'équilibre. Bon nombre des points abordés dans ces revues peuvent être généralisés à d'autres domaines de recherche et d'application.

2.1 Métaheuristiques

Les méthodes heuristiques désignent des techniques de résolution de problèmes, d'apprentissage et de découverte où une recherche exhaustive n'est pas possible en raison de problèmes liés à la puissance de calcul, au temps d'exécution mais aussi à la consommation d'énergie. Les MH sont des heuristiques intégrées dans des processus itératifs afin d'augmenter leurs capacités d'E&E [12]. Les MH sont dites stochastiques car elles utilisent de manière itérative des processus aléatoires. La grande majorité de ces méthodes sont utilisées pour résoudre des problèmes d'optimisation bien que le cadre général ne leur soit pas strictement dédié [7]. Les MH se sont avérées bien adaptées aux problèmes d'optimisation multi-objectifs et mono-objectifs car elles sont capables de capturer plusieurs solutions simultanément, tout en tenant compte des contraintes. Notons que traiter un problème d'optimisation mono-objectif consiste à trouver une ou plusieurs bonnes solutions pour un ou plusieurs attributs, en respectant des contraintes et en essayant de satisfaire des métriques de performance et de qualité. Il est

possible de combiner les objectifs d'un problème d'optimisation multi-objectifs pour le convertir en un problème d'optimisation à objectif unique en définissant la fonction de *fitness* comme la somme pondérée des valeurs normalisées des attributs.

Notre contribution relève du domaine des MH car (1) elles sont transparentes et explicables [21] c'est à dire que leur lecture est claire et permet de comprendre comment les données sont collectées et traitées pour aboutir aux résultats, (2) elles sont performantes dans des environnements soumis à de nombreux aléas et (3) elles sont particulièrement efficaces pour optimiser les décisions [32].

2.2 Métaheuristiques et recommandation

Les MH peuvent être utilisées pour résoudre le problème de la recommandation dans d'immenses espaces de recherche. Dans la littérature, l'utilisation des MH n'est pas systématiquement référée à la recommandation mais souvent à l'optimisation qui est, à bien des égards, très proche de la recommandation. En effet, en considérant le problème du sac à dos [20], un des problèmes d'optimisation les plus anciens et les plus connus, nous supposons que les solutions trouvées par les MH pour résoudre ce problème peuvent constituer des recommandations simples et qu'il est possible d'ajouter des contraintes représentant les préférences d'un utilisateur pour fournir des recommandations personnalisées, sans changer l'essence du problème d'optimisation et en apportant le fait que les solutions trouvées constituent des recommandations personnalisées pour l'utilisateur. Par conséquent, un état de l'art sur l'utilisation des MH pour la tâche de recommandation doit prendre en considération toutes les contributions du domaine des MH. Deux revues particulièrement complètes relatives aux MH [9] [10] et des travaux spécifiques référant l'utilisation des MH à la recommandation [18] [26] [28] constituent la base de notre état de l'art sur ce sujet. Dans [26], un outil de recommandation multi-objectifs est proposé. Le système de recommandation génère des parcours d'apprentissage appropriés pour les apprenants à l'aide d'algorithmes MH, notamment l'algorithme génétique (Genetic Algorithm - GA) et l'algorithme d'optimisation des colonies de fourmis (Ant Colony Optimization - ACO). Concrètement, les *items* à recommander (des cours) sont des nœuds dans un graphe, la MH est utilisée pour générer des chemins à travers le graphe et les meilleurs chemins trouvés sont recommandés. Dans [18], un système de recommandation hybride appliqué au jeu de données Movielens est proposé. Concrètement, un algorithme de *clustering* k-means est utilisé pour créer des *clusters* d'utilisateurs, chaque *cluster* ayant un centroïde. Ensuite, un algorithme de colonie d'abeilles artificielles (Artificial Bee Colony - ABC) est utilisé pour optimiser les distances des utilisateurs aux centroïdes et pour les reclasser si nécessaire. Les recommandations données à un utilisateur sont directement liées aux distances aux centroïdes de l'utilisateur, en effet, des contenus d'utilisateurs similaires (films), au regard des distances aux centroïdes, sont recommandés. Dans [28], les auteurs proposent d'améliorer les recommandations d'itinéraires pour les touristes

grâce à des MH. Concrètement, un algorithme de *clustering* k-means est utilisé pour créer des *clusters* de points d'intérêt (Points Of Interest - POI) dans des zones géographiques. Pour chaque *cluster*, un algorithme génétique optimise l'itinéraire pour visiter tous les POI. Ces trois exemples montrent que, comme dit plus haut, la tâche de recommandation peut être considérée comme une sélection optimisée d'*items* dans un espace de recherche.

2.3 Influence de l'exploration et de l'exploitation dans les métaheuristiques

2.3.1 Définitions

L'E&E peuvent être caractérisées comme une recherche globale et locale [1], l'acquisition et l'utilisation de connaissances [30], des processus qui se distinguent par l'intensité du caractère aléatoire [24], ou comme deux types de comportement dans l'acquisition d'informations [3]. Les MH se retrouvent souvent piégées dans un optimum local, la principale raison étant la difficulté à bien arbitrer entre E&E. Une bonne MH doit donc faire de bons compromis entre E&E, mais aucune réponse complète et/ou générale à ce problème n'a été apportée à ce jour.

2.3.2 Contrôle de l'exploration et de l'exploitation

Chercher à contrôler l'E&E revient à poser les questions suivantes : quand, quoi et comment contrôler ? La mesure de la diversité, génotypique ou phénotypique est devenue largement acceptée comme approche de contrôle [23]. La mesure de la diversité génotypique peut être basée sur la différence, la distance, l'entropie, la probabilité ou même les ancêtres/l'histoire. La mesure de la diversité phénotypique peut être basée sur la différence, la distance, l'entropie, la probabilité.

2.3.3 Influence de l'exploration et de l'exploitation

De nombreux facteurs ont un impact sur l'E&E. Une discussion initiale sur l'E&E est menée dans [11] et d'autres travaux ont ensuite poursuivi cette discussion, traitant par exemple de la puissance d'exploration [8] et du redimensionnement de la population [13].

Dans la littérature, l'influence de l'E&E est qualifiée de compromis, d'orientation ou d'équilibre. L'influence est cruciale dans les méthodes d'optimisation car elle améliore l'efficacité de la recherche. La littérature a proposé de nombreuses approches adaptatives [31] et hybrides [19] pour influencer l'E&E. De plus, la plupart des approches reposent sur la diversité par le maintien, le contrôle ou l'apprentissage. Il est important de noter que bien que de nombreux travaux effectués sur la diversité s'attendent à un meilleur équilibre entre E&E, toutes les diversifications ne sont pas utiles. De plus, il faut garder à l'esprit que les phases d'E&E sont toujours imbriquées, ce qui rend leur identification plus complexe, il est également quasiment impossible d'observer clairement leurs contributions respectives dans le processus d'E&E. Finalement, il existe peu de travaux théoriques sur la diversité [8], la plupart étant applicatifs.

2.3.4 Directions de recherche

À partir de cette revue de la littérature, nous pouvons voir qu'il existe encore de nombreuses questions ouvertes concernant l'E&E, telles que définir formellement/mathématiquement les phases d'E&E et l'équation d'équilibre, proposer des métriques (directes) pour mesurer l'E&E, définir quand et comment contrôler l'équilibre entre E&E, mener des analyses sur les différentes approches de maintien, de contrôle et d'apprentissage de la diversité.

2.4 Influence de l'exploration et de l'exploitation en psycho-sociologie

De nombreuses décisions dans la vie des êtres vivants nécessitent un équilibre entre l'exploration de différentes options et l'exploitation de leurs récompenses [22]. Cette dernière ressource explore comment le compromis entre E&E dépend de la conceptualisation de l'E&E, des facteurs environnementaux, sociaux et individuels, de l'échelle à laquelle l'E&E sont considérées, de la relation et des types de transitions entre eux, des objectifs du décideur. En effet, l'E&E sont mieux conceptualisées comme des points sur un continuum et sont omniprésentes à de nombreux niveaux d'abstraction. Elles se produisent à plusieurs échelles dans l'espace et dans le temps. Les décisions sur la manière de se comporter à ces différents niveaux peuvent être basées sur une heuristique connue ou nouvelle. La manière de choisir les heuristiques est donc également importante. De nombreuses approches abordent les problèmes de compromis entre E&E [15] [22].

Trois éléments sont mis en évidence par [22] pour une unification de la recherche en E&E : un continuum E&E [5], différents types de transitions entre E&E [25] et le rôle des objectifs des agents dans ce processus [5]. De nombreuses ouvertures s'ensuivent : (1) proposer une théorie globale pour couvrir les dimensions conceptuelles de l'E&E, (2) saisir les interactions entre les différents niveaux d'abstraction sur lesquels se situe l'E&E, (3) proposer une utilisation hiérarchisée des mécanismes de compromis à différentes échelles et à différents niveaux d'abstraction, (4) prendre en compte les effets sur l'E&E des objectifs des agents pour mieux évaluer les coûts et récompenses associés à l'E&E et ainsi atteindre les meilleurs compromis possibles, (5) proposer une théorie complète sur la variété des transitions qui peuvent se produire entre E&E, et (6) explorer dans quelles mesures les mécanismes de compromis sont sensibles à divers facteurs.

2.5 Conclusion

Pour résumer, de nombreux travaux se sont intéressés à l'arbitrage entre E&E, sans aboutir à une réponse complète ou générale. Le cadre de la recommandation est fortement concerné par ces problèmes. Des contributions proposant une nouvelle façon d'explorer et d'exploiter l'espace de recherche sont attendues. Il en va de même pour les contributions visant à expliquer, mesurer, contrôler et influencer l'E&E.

3 Contributions

Nous proposons un processus pour contrôler et influencer l'E&E d'un algorithme de recherche (appelé algorithme influencé). De notre point de vue, les processus d'influence proposés dans la littérature ne garantissent pas une expression très impactante de l'influence. Une contribution de ce travail vise donc à donner plus de puissance au processus d'influence en concevant un processus indépendant de l'algorithme influencé. Nous présentons nos travaux dans le cadre des algorithmes évolutionnaires (Evolutionary Algorithm - EA) et des algorithmes par renforcement, en raison de leur popularité. Notons que l'indépendance fait que l'influence peut être utilisée sur la plupart des MH, sans adaptation majeure.

Les tableaux de *digits* binaires sont une manière traditionnelle de représenter des *items* ou *itemsets* recommandés [20]. Chaque index d'un tableau correspond à un *item* recommandable. Les *digits* peuvent prendre la valeur 0 (non-recommandation de l'*item*) ou 1 (recommandation de l'*item*). Le nombre de *digits* à 1 correspond au nombre d'*items* recommandés. L'algorithme effectue itérativement des variations : génère, croise et transforme des tableaux, tout en essayant d'optimiser et de garder le meilleur d'entre eux au regard de contraintes prédéfinies et de critères d'évaluation.

De plus, nous proposons de nouveaux indicateurs pour présenter plus clairement le comportement d'E&E.

3.1 Processus d'influence de l'exploration et de l'exploitation, indépendant de l'algorithme influencé

Soulignons que dans la littérature, le processus d'influence a des caractéristiques qui, de notre point de vue, limitent son impact sur l'E&E. Il (1) dépend fortement de l'algorithme influencé, comme dans le cas de l'optimisation des paramètres (Algorithme 1, Figure 1), (2) ne possède pas ses propres opérateurs de variation, (3) ne garantit aucune forme de respect des bornes de l'indicateur d'E&E à chaque itération de l'algorithme influencé. Nous proposons un nouveau processus d'influence qui (1) est dissocié de l'algorithme influencé, (2) a ses propres opérateurs de variation, (3) est appelé à chaque itération de l'algorithme influencé et (4) cherche à garantir la conformité des solutions aux bornes de l'indicateur d'E&E à chaque itération de l'algorithme influencé. Ainsi, seuls les *itemsets* sont éventuellement modifiés et renvoyés à l'algorithme influencé, ce dernier ne présentant aucune modification structurelle et/ou paramétrique. Le processus d'influence proposé est donc générique.

L'indicateur mentionné ci-dessus est la mesure de l'E&E effectuée lors du contrôle. Nous proposons de l'utiliser pour décider de forcer l'exploitation, l'exploration ou les deux avec des variations supplémentaires. Un exemple d'indicateur d'E&E est la diversité des solutions trouvées jusqu'ici par l'algorithme influencé. L'approche que nous proposons est conçue pour utiliser n'importe quel indicateur, dès lors que des limites minimales et/ou maximales associées sont

Algorithm 1 Influence Traditionnelle de l'E&E

- 1: Étape 1 : Population Initiale
 - 2: Étape 2 : Évaluations des Itemsets
 - 3: Étape 3 : Variations des Itemsets (Sélections, Croisements, Mutations)
 - 4: Étape 4 : Évaluation de l'Indicateur d'E&E (ex : avec la diversité)
 - 5: Étape 5 : Ajustement de la Probabilité de Croisement/Mutation. Dépend de l'Évaluation de l'Indicateur. Vise à Forcer l'Exploration ou l'Exploitation
 - 6: Étape 6 : Fin d'Itération
 - 7: Étape 7 : Retour à l'étape 2
-

définies en amont, ainsi que la valeur du paramètre utilisé pour ajuster ces limites (voir plus de détails dans les sections suivantes). Ces valeurs peuvent être déterminées par des tests d'exécution, des valeurs arbitraires ou des informations statistiques de base concernant les données traitées.

Le principe de fonctionnement du processus d'influence repose sur une fonction indépendante appelée systématiquement lors de l'exécution de l'algorithme influencé (Algorithme 2). Dans la littérature, l'influence n'est actionnée qu'une seule fois par itération. En ce qui nous concerne, juste après une décision d'influence, les éléments utilisés pour prendre cette décision pourraient permettre de décider à nouveau d'influencer. Nous proposons donc un contrôle systématique de l'E&E, à chaque itération de l'algorithme influencé mais également après chaque processus de variation d'influence c'est-à-dire tant que les *itemsets* ne respectent pas les limites de l'indicateur d'E&E. Le processus d'influence garantit ainsi la conformité des *itemsets* à l'indicateur d'E&E à chaque itération. Nous considérons que cela permet de renforcer l'impact de l'influence, le problème éventuel étant d'effectuer une influence trop impactante qui défigure ce que fait l'algorithme influencé en matière d'E&E. Nous devons également noter que le processus d'influence augmentera inévitablement le coût de calcul car il apporte des traitements supplémentaires.

Concrètement, toutes les solutions de l'algorithme influencé sont passées en paramètre de la fonction d'influence (Ligne 5), qui va les transformer (Lignes 19, 23) de manière itérative (Lignes 17, 18, 22) selon l'évolution de l'indicateur d'E&E (Ligne 11). Lorsque l'indicateur atteint une valeur prédéfinie ou se repositionne dans un certain intervalle de valeurs (Lignes 14, 15), la fonction d'influence arrête les transformations et renvoie toutes les solutions à l'algorithme influencé, qui continue son exécution jusqu'au prochain appel de la fonction d'influence (Ligne 5). Notons que les variations d'influence se font tant que l'indicateur est hors bornes (Lignes 17, 18, 22). Il est donc nécessaire de replacer les valeurs des bornes en fonction de l'évolution de l'indicateur d'E&E et ainsi converger vers une stabilisation autour des valeurs observées de l'indicateur d'E&E afin de pouvoir stopper les itérations d'influence. Pour cela, après chaque itération de la fonction d'influence, les bornes minimum et maximum de l'indicateur d'E&E sont ajustées

Algorithm 2 EA Appelant la Fonction d'Influence

```

1: FUNCTION ea() # EA Traditionnel
2:   initialisation()
3:   WHILE(condition)
4:     ea_variations()
5:     influence(itemsets)
6:     fin_iteration()
7:   finalisation()
8:
9: FUNCTION influence(itemsets) # Influence de l'E&E
10:  # ind est l'indicateur d'E&E
11:  SET ind TO calculer_ind(itemsets)
12:  SET ind_valeurs TO []
13:  ADD ind TO ind_valeurs
14:  SET min_limit TO calculer_min_limit()
15:  SET max_limit TO calculer_max_limit()
16:  SET ajust_param TO calculer_ajust_param()
17:  WHILE ind <= min_limit or ind >= max_limit :
18:    WHILE ind <= min_limit
19:      exploiter(itemsets) # Opère des variations
20:      ajuster_limites()
21:      maj_sauvegarder_ind()
22:    WHILE ind >= max_limit
23:      explorer(itemsets) # Opère des variations
24:      ajuster_limites()
25:      maj_sauvegarder_ind()

```

(Lignes 20, 24) à l'aide d'un paramètre (Ligne 16) ajouté ou supprimé de la valeur moyenne de l'indicateur d'E&E observée jusqu'à présent.

3.2 Nouveaux indicateurs pour l'influence et l'explicabilité de l'exploration et de l'exploitation

Comme le suggère la littérature, l'E&E mériteraient d'être représentées sous de nouveaux angles [6][29][22]. C'est ce que nous entendons faire en proposant trois types d'indicateurs qui mesurent et représentent les comportements d'E&E et qui peuvent être utilisés au sein de la fonction d'influence introduite précédemment.

3.2.1 Complétion des contraintes (CC)

L'indicateur CC évalue la complétion des contraintes des *itemsets* en traitant les contraintes comme des volumes à remplir. Pour évaluer cet indicateur, pour chaque *item* d'un *itemset* la moyenne des rapports des valeurs de chaque colonne sous contrainte de l'*item* par les contraintes maximales associées est calculée. Cela donne la complétion des contraintes d'un *item* (CC_item). La somme des CC_item d'un *itemset* donnent sa complétion des contraintes ($CC_itemset$). L'indicateur CC est calculé en faisant la moyenne de l'ensemble des $CC_itemset$ (Equation (1)). L'ensemble des $CC_itemset$ est noté CC_itsets .

Plus CC est petit, plus il y a d'espace à combler parmi tous les *itemsets* car il est potentiellement possible d'ajouter plus d'*items* dans certains d'entre eux, tout en respectant

les contraintes. Dans ce cas, il est possible d'affirmer que le processus d'E&E ne permettait pas jusqu'à présent d'insérer ces *items*. A contrario, plus CC est grand, moins il y a d'espace à combler parmi tous les *itemsets* et il est donc plus compliqué d'envisager l'ajout d'autres *items* dans certains d'entre eux, tout en respectant les contraintes.

$$CC = \overline{CC_itsets} \quad (1)$$

3.2.2 E&E power

L'indicateur EEP représente la puissance d'E&E en tant que quantité d'exploration par rapport à l'exploitation. Nous considérons qu'à chaque fois qu'un *digit* est passé à 1 pour la première fois, cela constitue une exploration. Ainsi, une itération découvrant au moins un nouveau *digit* est une itération d'exploration et un vecteur contenant au moins un nouveau *digit* est un vecteur d'exploration. Les autres itérations et vecteurs sont dits d'exploitation.

Nous évaluons le nombre d'itérations d'exploration divisé par le nombre d'itérations d'exploitation ($\#eit$), le nombre de vecteurs d'exploration divisé par le nombre de vecteurs d'exploitation ($\#evect$) et le nombre de *digits* nouvellement découverts divisé par le nombre de *digits* déjà découverts ($\#dnd$). EEP est la moyenne de ces trois éléments (Equation (2)), il représente le ratio d'exploration comme la synthèse de ses valeurs.

$$EEP = (\#eit, \#evect, \#dnd) \quad (2)$$

Un registre des *digits* non découverts peut aussi être utilisé pour définir formellement une phase d'E&E. Soit p une phase d'E&E pouvant prendre les valeurs $\{explore, exploit\}$. Soit ndd le registre des *digits* non encore découverts. Ce registre est mis à jour à chaque itération en enlevant des *digits* une fois qu'ils sont utilisés (passés à la valeur 1) pour la première fois à l'intérieur de vecteurs. Soit ie un marqueur de fin d'itération de l'algorithme influencé. $ndd(ie)$ correspond donc à ndd au marqueur de fin d'itération ie . Nous pouvons donc définir p comme dans Equation (3). Cette équation se lit ainsi : La phase à l'itération $ie + 1$ est une phase d'exploration si de la phase ie à la phase $ie + 1$ la taille du registre des *digits* non encore découverts a diminué. La phase à l'itération $ie + 1$ est une phase d'exploitation si de la phase ie à la phase $ie + 1$ la taille du registre des *digits* non encore découverts est restée inchangée.

$$p(ie + 1) = \begin{cases} explore, & \text{if } \overline{ndd(ie)} > \overline{ndd(ie + 1)} \\ exploit, & \text{if } \overline{ndd(ie)} = \overline{ndd(ie + 1)} \end{cases} \quad (3)$$

3.2.3 Temporalité

Nous proposons de représenter l'alternance des itérations d'E&E comme des écarts entre des valeurs temporelles. Les écarts temporels entre les instants de démarrage des itérations d'exploration sont évalués et font ressortir les écarts minimum (*ming*) et maximum (*maxg*), la différence entre l'écart minimum et l'écart maximum ($DMMG$) (Equation (4)), l'écart moyen (AG) (Equation (5)). Idem pour les itérations d'exploitation. $DMMG$ permet d'évaluer les

écarts temporels les plus extrêmes entre les itérations d'exploration. AG représente implicitement la quantité d'itérations d'exploration qui représente implicitement leur régularité. La régularité attendue dépend du contexte.

$$DMMG = maxg - ming \quad (4)$$

$$AG = \overline{gaps} \quad (5)$$

4 Expérimentations

4.1 Données et mesures d'évaluation

Nos expérimentations ont été réalisées sur le jeu de données très utilisé dans la littérature de la recommandation « Movielens 25M » [14]. Il se compose d'évaluations et de *tags* donnés par des utilisateurs sur des films. Il contient 25 millions d'évaluations et 1 million de *tags* donnés par 162 000 utilisateurs sur 62 000 films. L'ensemble de données comprend également 15 millions de scores de pertinence pour 1 129 *tags*.

Les objectifs et les contraintes des recommandations ont été déterminés arbitrairement en veillant à apporter suffisamment de diversité, de même que les limites des indicateurs d'E&E et les valeurs utilisées pour les mettre à jour.

Une exécution concerne une variation de contrainte, un objectif de recommandation, un cas d'exécution (influencé ou non) et une macro-exécution. Les expérimentations comprennent six variations de contraintes, six objectifs de recommandation, huit cas d'exécution : un non influencé (référence) et sept influencés (qui varient selon leur indicateur d'E&E), et trente boucles de macro-exécution, pour un total de plusieurs milliers d'exécutions, chacune comprenant mille itérations.

Les fonctions d'influence sont appliquées sur un algorithme génétique traditionnel qui sert de référence pour les comparaisons. Les fonctions d'influence sont aussi appliquées sur un algorithme par renforcement traditionnel servant à son tour de référence. L'objectif est de montrer les différences comportementales de l'algorithme de référence lorsqu'il est influencé par le processus proposé et avec différents indicateurs d'E&E.

Pour calculer la fitness d'un *itemset*, les valeurs des attributs des *items* composant cet *itemset* sont sommées dans une fonction de *fitness* à objectif unique afin de simplifier les comparaisons et d'économiser du temps de calcul.

Des tests statistiques traditionnels de type *ranksum* sont effectués afin de s'assurer que les valeurs des critères d'évaluation, qui sont des moyennes sur toutes les macro-itérations, proviennent bien de distributions différentes par rapport à celles de l'algorithme de référence. Il s'agit d'une évaluation traditionnelle en MH.

L'infrastructure informatique Grid'5000 [2] est utilisée pour ces expérimentations.

4.2 Macro-processus

Le macro-processus adopté dans ces expérimentations s'articule autour de différentes étapes. Dans un premier temps, les données d'entrée brutes sont récupérées, correspondant

à l'objectif de recommandation qui est construit à partir des informations de profil. Ces données sont sous la forme de tuples (*item* id, attribut 1, ..., attribut n). Nous choisissons d'utiliser les attributs numériques suivants de Movielens : la note moyenne, la quantité de notes et la pertinence moyenne des *tags*. Les données sont normalisées. Nous choisissons d'associer chaque attribut à une contrainte "min" ou "max". S'il doit être maximisé, nous choisissons de laisser ses valeurs telles quelles et s'il doit être minimisé, ses valeurs prennent le complément à 1 pour représenter le fait que plus la valeur est petite, plus elle aura d'importance lorsque la somme se fera dans la fonction de *fitness*. Les contraintes sont définies sur les attributs sous forme de tuples (valeur minimale totale dans un *itemset*, valeur maximale totale dans un *itemset*). Des poids sont associés aux attributs, un poids représente la contribution moyenne d'un attribut dans la somme totale de toutes les valeurs d'attribut considérées. Un score d'*item* est défini comme la somme pondérée de ses attributs. Un score d'*itemset* est défini comme la somme de ses scores d'*items*. Cette combinaison de choix présentée ci-dessus n'a pas été observée dans la littérature.

L'algorithme de référence (évolutionnaire ou renforcement) est exécuté trente fois (macro-itérations). Lors d'une exécution, les vecteurs générés sont enregistrés dans *all_vectors* et les meilleurs vecteurs sont enregistrés dans *top_vectors*. Enfin, des sélections sont effectuées pour afficher les recommandations finales (à partir de *top_vectors*) et les valeurs moyennes des critères d'évaluation de la performance et de la qualité sont renvoyées pour toutes les exécutions. Divers éléments de visualisation sont enregistrés en conséquence.

4.3 Critères d'évaluation

Nous décidons d'évaluer notre processus d'influence en considérant plusieurs critères liés à la performance, à la qualité et au temps qui sont soit ceux introduits dans la section 3.2, soit proposés dans la littérature [17][16].

4.3.1 Critères de performance

Comme mentionné dans la section 4.2, un score d'*item* est défini comme la somme pondérée de ses valeurs d'attribut normalisées et un score d'*itemset* est défini comme la somme des scores de ses *items*.

Convergence Iterations (CI) : Nombre moyen d'itérations nécessaires pour atteindre les points de convergence.

Max Score Iterations (MSI) : Nombre moyen d'itérations nécessaires pour atteindre les scores maximum.

Convergence Score (CS) : Score maximum moyen atteint sur les points de convergence.

Max Score (MS) : Score maximum moyen atteint.

Execution Time (ET) : Temps d'exécution moyen.

4.3.2 Critères de qualité

All Vectors Coverage (AVC) : Couverture prenant en compte tous les vecteurs générés. Représente le pourcentage de *digits* qui sont au moins une fois à 1 en considérant tous les vecteurs générés. Si chaque *digit* est au moins une fois à 1 dans l'ensemble des vecteurs considérés, la couverture est totale. Plus la couverture est élevée, mieux c'est,

car cela signifie que le processus de recherche a trouvé de nombreux *items*.

Constraints Completion (CC) : Pourcentage moyen de complétion des contraintes. Il représente dans quelle mesure les *items* recommandés remplissent les contraintes.

E&E Power (EEP) : Part de l'exploration par rapport à l'exploitation. Voir la section 3 pour plus de détails.

4.3.3 Critères temporels

Min-Max Gap (DMMG) : Différence des écarts temporels minimum et maximum entre deux itérations d'exploration.

Average Gap (AG) : Écart temporel moyen entre deux itérations d'exploration.

4.4 Résultats - Algorithme génétique

L'algorithme génétique de référence utilise un opérateur de croisement de type "one point" et un opérateur de mutation de type "bit string", ce dernier est effectué tant que l'*itemset* muté respecte les contraintes maximales. Les pourcentages agrégés sont indiqués dans Table 1. Les lignes $\nearrow$, $\rightarrow$ et $\searrow$ font respectivement référence aux pourcentages de résultats des critères d'évaluation augmentés, inchangés et diminués par rapport à l'algorithme de référence, en tenant compte de tous les cas d'exécution. Pour compléter ces informations, nous effectuons des tests statistiques représentés par la ligne *%Pass* indiquant le pourcentage de tests réussis en tenant compte de tous les cas d'exécution. Précisons que l'hypothèse nulle d'un test statistique est que les valeurs des critères d'évaluation, de l'algorithme non influencé et de l'une de ses versions influencées, proviennent de la même distribution. Si l'hypothèse nulle est rejetée, c'est un succès ($p\text{-value} < 0,05$) et sinon c'est un échec ($p\text{-value} \geq 0,05$). Nous considérons qu'un critère d'évaluation permet de tirer des conclusions fiables si son *%Pass* est supérieur à 50%.

TABLE 1 – Résultats - Impact de l'Influence sur l'E&E de l'Algorithme Génétique

	Performance					Qualité			Temporel	
	CI	MSI	CS	MS	ET	AVC	CC	EEP	DMMG	AG
$\nearrow$	62	65	40	63	100	50	44	2	82	77
$\rightarrow$	2	0	56	36	0	50	48	13	0	17
$\searrow$	37	35	4	1	0	0	8	85	18	5
%Pass	17	20	56	68	100	57	52	89	78	96

4.4.1 Critères de performance

CI et *MSI* augmentent majoritairement par l'influence avec respectivement 62% et 65%. Ainsi, l'influence augmente généralement le nombre d'itérations nécessaires pour atteindre la convergence et les scores maximum, ce qui est un inconvénient compte tenu du temps d'exécution mais qui est aussi un avantage vis-à-vis du problème de la convergence prématurée. Dans l'ensemble des cas d'exécution, aucun problème de convergence prématurée n'a été observé. Cette conclusion doit être vue à la lumière de tests statistiques qui ne rejettent pas l'hypothèse nulle. Par conséquent, ces critères d'évaluation ne permettent pas de

tirer des conclusions fiables dans le cadre de ces expérimentations.

CS et *MS* sont soit inchangés soit améliorés par l'influence. *CS* est inchangé avec 56% mais obtient toujours un important ratio d'augmentation avec 40%. *MS* est augmenté avec 63% et obtient un taux de stagnation significatif de 36%. Globalement, les fonctions d'influence conduisent à des scores au moins égaux à ceux de l'algorithme de base et dans de nombreux cas sont capables de les améliorer. Ces résultats sont validés par des tests statistiques. En effet, la plupart d'entre eux rejettent l'hypothèse nulle donc ces critères d'évaluation permettent de tirer des conclusions fiables dans le cadre de ces expérimentations.

Comme attendu, *ET* est raisonnablement augmenté par l'influence, ce qui est validé par les tests statistiques.

4.4.2 Critères de qualité

AVC est inchangé (50%) ou augmente (50%) par l'influence. Une augmentation d'*AVC* signifie que davantage de dimensions (*items*) sont utilisées dans le processus de recherche. Ces résultats sont validés par les tests statistiques, en effet la plupart d'entre eux rejettent l'hypothèse nulle donc ce critère d'évaluation permet de tirer des conclusions fiables dans le cadre de ces expérimentations.

CC reste inchangé (48%) ou augmente (44%) par l'influence, ce qui montre que le processus d'influence proposé contribue à trouver des *itemsets* qui remplissent mieux les contraintes. La qualité des recommandations est ainsi améliorée par l'influence. De plus, ces résultats confirment que l'explicabilité de l'E&E est améliorée car ils soulignent la capacité du processus d'influence à remplir les volumes de contraintes et par conséquent à atteindre des régions spécifiques et des optima de l'espace de recherche. Ces résultats sont validés par les tests statistiques, en effet la plupart d'entre eux rejettent l'hypothèse nulle donc ce critère d'évaluation permet de tirer des conclusions fiables dans le cadre de ces expérimentations.

EEP diminue globalement (85%) par l'influence, ce qui signifie que les fonctions d'influence ont tendance à diminuer la part de l'exploration par rapport à l'exploitation et donc à renforcer l'exploitation. Ces résultats contribuent à l'explicabilité de l'E&E en montrant les grandes orientations du processus de recherche, induites par l'influence. Ces résultats sont validés par les tests statistiques, en effet la plupart d'entre eux rejettent l'hypothèse nulle donc ce critère d'évaluation permet de tirer des conclusions fiables dans le cadre de ces expérimentations.

4.4.3 Critères temporels

Les valeurs des critères temporels sont augmentées par l'influence, avec des ratios de 82% pour *DMMG* et de 77% pour *AG*. Cela signifie que la distribution temporelle des itérations d'exploration devient plus clairsemée et que les différences temporelles extrêmes entre ces itérations ont tendance à être plus prononcées. Cependant, dans certains cas, nous observons l'effet inverse (18% et 5%), qui correspond à des phases d'exploration rapprochées avec des écarts temporels extrêmes réduits. Ces résultats sont validés par les tests statistiques, en effet la plupart d'entre eux

rejetent l'hypothèse nulle, donc ces critères d'évaluation permettent de tirer des conclusions fiables dans le cadre de ces expérimentations. Ces expérimentations montrent que les indicateurs temporels expliquent comment l'influence affecte la régularité de l'E&E.

4.4.4 Conclusion

Les critères étudiés ci-dessus montrent un impact positif du processus d'influence avec des valeurs de ratio significatives pour la plupart des critères : deux critères de performance (*CS*, *MS*), tous les critères de qualité (*AVC*, *CC*, *EEP*) et tous les critères temporels (*DMMG*, *AG*).

En ce qui concerne les critères de performance, le processus d'influence permet à l'algorithme d'atteindre de meilleurs scores au détriment du nombre d'itérations et du temps d'exécution. En ce qui concerne les critères de qualité, le processus d'influence fait que l'algorithme de référence évalue plus de solutions, ce qui correspond à une valeur de *CC* plus élevée dans certains cas. Nous rappelons ici que les préférences des utilisateurs sont encodées dans les contraintes, en conséquence *CC* porte des informations sur la qualité de la recommandation. Comme la valeur de *CC* reste dans presque tous les cas influencés au moins égale à la valeur de *CC* des cas non influencés, nous pouvons conclure que le processus d'influence proposé permet de trouver de meilleures recommandations. Par ailleurs, le processus d'influence a un impact significatif sur *EEP* en favorisant l'exploitation. En ce qui concerne les critères temporels, le processus d'influence conduit à un impact notable sur la régularité de l'E&E.

Pour résumer, le processus d'influence a un impact positif sur le comportement de recherche. Nous notons qu'une analyse détaillée des résultats montre que certains cas d'influence sont plus performants que d'autres. Nous notons également que l'influence proposée n'a pas un impact positif dans tous les contextes, ce qui était attendu car cela dépend de la structure de l'espace de recherche et de l'ensemble des contraintes. Nous soulignons également que les résultats confirment que le processus d'influence n'induit pas de comportements aléatoires car les scores restent au moins égaux à la référence.

De plus, nous confirmons que les indicateurs d'influence proposés aident à fournir des explications concernant l'E&E. En effet, ils représentent les comportements d'E&E sous de nouveaux angles comme le montrent *CC*, *EEP*, *DMMG* et *AG*.

Finalement, nous pouvons conclure que ce travail est un pas en avant pour les processus d'influence de l'E&E ainsi que pour la représentation et l'explicabilité de l'E&E. Ces résultats sont prometteurs pour nos travaux futurs.

4.5 Résultats - Algorithme par renforcement

Pour montrer l'adaptabilité du processus d'influence proposé, nous choisissons de l'appliquer à un algorithme de renforcement traditionnel (Algorithme (3)) dont la structure globale est similaire à l'algorithme génétique. *Q* et *N* représentent respectivement le tableau des scores de renforcement associés à chaque digit et le tableau du nombre de

Algorithm 3 RA Calling The Influence Function

```

1: FUNCTION ra()
2:   initialisation()
3:   WHILE(condition not reached)
4:     ra_variations_with_Q()
5:     influence(itemsets)
6:     update_N_Q()
7:     iteration_end()
8:   finalisation()

```

fois où chaque digit a pris la valeur 1. *N* et *Q* sont mis à jour de manière traditionnelle, chaque digit modifié "*d*" voit sa valeur correspondante incrémentée de un dans *N* (Equation (6)) et par la formule traditionnelle dans *Q* (Equation (7)). Les résultats sont présentés dans Tableau (2).

$$N[d] += 1 \quad (6)$$

$$Q[d] += 1/N[d] * (reward - Q[d]) \quad (7)$$

TABLE 2 – Résultats - Impact de l'Influence sur l'E&E de l'Algorithme par Renforcement

	Performance					Qualité			Temporel	
	CI	MSI	CS	MS	ET	AVC	CC	EEP	DMMG	AG
% ↗	37	30	34	28	100	71	31	19	89	71
% →	0	1	37	47	0	27	40	6	1	28
% ↘	63	69	29	25	0	2	28	76	10	0
%Pass	45	38	54	58	100	67	39	79	93	95

4.5.1 Critères de performance

CI et *MSI* diminuent majoritairement par l'influence avec respectivement 63% et 69%. Ainsi, l'influence diminue généralement le nombre d'itérations nécessaires pour atteindre la convergence et les scores maximum, ce qui est un avantage compte tenu du temps d'exécution mais qui est aussi un inconvénient concernant le problème de convergence prématurée. Dans l'ensemble des cas d'exécution, aucun problème de convergence prématurée n'a été observé.

CS et *MS* sont inchangés (37%, 47%), améliorés (34%, 28%) et diminués (29%, 25%) par l'influence. Dans l'ensemble, un nombre important de cas d'influence conduisent à des scores au moins égaux à ceux de l'algorithme de référence et dans de nombreux cas sont capables de les améliorer.

Comme attendu, *ET* est raisonnablement augmenté par l'influence, dans tous les cas.

4.5.2 Critères de qualité

AVC est inchangé (27%) ou augmente (71%) par l'influence. Une augmentation de *AVC* signifie que davantage de dimensions (*items*) sont utilisés dans le processus de recherche.

CC reste inchangé (40%), augmente (31%) et diminue (28%) par l'influence, ce qui montre que le processus d'influence proposé contribue dans de nombreux cas à trouver

des *itemsets* qui remplissent mieux les contraintes. La qualité des recommandations est améliorée dans ces cas.

EEP diminue globalement (76%) par l'influence, ce qui signifie que les fonctions d'influence ont tendance à diminuer la part de l'exploration par rapport à l'exploitation et donc à renforcer l'exploitation.

4.5.3 Critères temporels

Les valeurs de *DMMG* et *AG* sont augmentées par l'influence (89%, 71%). Cela signifie que la distribution temporelle des itérations d'exploration devient plus clairsemée et que les différences temporelles extrêmes entre ces itérations ont tendance à être plus prononcées.

4.5.4 Conclusion

Ces expérimentations confirment les conclusions tirées pour l'algorithme génétique. Le processus d'influence a un impact positif, ce qui confirme la généralité de l'approche proposée.

5 Discussions et travaux futurs

Nos résultats présentant des conclusions positives pour l'algorithme génétique ainsi que pour l'algorithme par renforcement, nous soulignons ici la généralité du processus d'influence proposé.

Premièrement, l'influence indépendante proposée est adaptable aux MH et plus généralement à tout algorithme d'E&E. En effet, le principe d'une influence itérative indépendante est universel dès lors qu'il est possible de déterminer : (1) le moment auquel appeler la fonction d'influence, (2) la nature de ce qui est passé en paramètre, (3) la manière dont les limites de l'indicateur d'E&E sont ajustées et (4) les variations à effectuer au sein d'une itération d'influence. Pour continuer, les résultats dépendent des objectifs de recommandation, des contraintes et des choix de paramètres d'influence. Modifier ces éléments mènerait à des résultats différents en permettant ou non aux fonctions d'influence d'améliorer plus souvent les performances des algorithmes non influencés et la qualité des recommandations. Dans le cadre des travaux présentés, nous avons utilisé des valeurs diversifiées pour ces éléments mais rien ne nous permet de justifier nos choix par rapport à l'ensemble des combinaisons de choix possibles. Il en va de même pour les volumes des expérimentations, ils sont discutables mais nous les jugeons suffisants au regard des volumes des expérimentations proposés dans l'état de l'art du domaine des MH.

Par ailleurs, le moment où la fonction d'influence est appelée impacte l'E&E et dépend de l'algorithme influencé ainsi que de la façon dont la fonction d'influence est conçue (indicateur d'E&E, opérateurs de variation, bornes, règles d'ajustement). Nous avons choisi de donner le dernier mot à la fonction d'influence en la positionnant avant la fin de chaque itération de l'algorithme influencé, ce choix peut être débattu.

De plus, les variations d'influence peuvent être effectuées de plusieurs façons, selon le contexte de mise en œuvre.

Un autre point important dans l'approche proposée est que les variations d'influence opèrent directement sur les

itemsets de l'itération actuelle de l'algorithme influencé. Ces *itemsets*, éventuellement mutés, sont renvoyés à l'algorithme influencé, qui continue avec eux tels quels. Cela risque potentiellement d'inclure trop d'aléatoire et donc de trop altérer les enrichissements de l'algorithme influencé. Cependant, cela dépend des opérateurs de variation utilisés, de la pertinence des indicateurs d'E&E et d'autres choix de paramètres. Pour déterminer ces éléments, des tests de fonctionnement peuvent être effectués, des valeurs arbitraires peuvent être utilisées, des informations statistiques de base concernant les données traitées peuvent également être exploitées.

Enfin, nos contributions pourraient être analysées par rapport à d'autres critères d'évaluation.

Nos futurs travaux seront marqués par des développements autour de nouveaux processus et indicateurs d'influence de l'E&E tout en continuant à améliorer ceux déjà mis en place. Nous effectuerons également des expérimentations en utilisant d'autres MH pour davantage soutenir les résultats obtenus. De plus, nous avons l'intention de mener des expérimentations plus approfondies et diversifiées dans le cadre de la recommandation pour étudier en profondeur les caractéristiques des recommandations fournies. Enfin, nous poursuivrons nos travaux sur l'explicabilité de l'E&E.

Remerciements

Ce travail est réalisé dans le cadre de la convention PEACE avec le Rectorat français de l'Académie de Nancy-Metz ainsi qu'avec le Ministère de l'Éducation Nationale et de la Jeunesse.

Références

- [1] Anne Auger, Marc Schoenauer, and Olivier Teytaud. Local and global order 3/2 convergence of a surrogate evolutionary algorithm. In *Proceedings of the 7th annual conference on Genetic and evolutionary computation*, pages 857–864, 2005.
- [2] Daniel Balouek and al. Adding virtualization capabilities to the Grid'5000 testbed. In *Cloud Computing and Services Science*, volume 367, pages 3–20. Springer International Publishing, 2013.
- [3] Jie Chen, Bin Xin, Zhihong Peng, Lihua Dou, and Juan Zhang. Optimal contraction theorem for exploration–exploitation tradeoff in search and optimization. *IEEE Transactions on Systems, Man, and Cybernetics-Part A : Systems and Humans*, 39(3) :680–691, 2009.
- [4] Minmin Chen. Exploration in recommender systems. In *Fifteenth ACM Conference on Recommender Systems*, pages 551–553, Amsterdam Netherlands, 2021. ACM.
- [5] Jonathan D Cohen, Samuel M McClure, and Angela J Yu. Should i stay or should i go? how the human brain manages the trade-off between exploitation and exploration. *Philosophical Transactions of the Royal*

- Society B : Biological Sciences*, 362(1481) :933–942, 2007.
- [6] Matej Crepinsek, Shih-hsi Liu, and Marjan Mernik. Exploration and exploitation in evolutionary algorithms : A survey. *ACM Computing Surveys*, 45(35) :1–33, 06 2013.
 - [7] Kenneth A De Jong. Are genetic algorithms function optimizers ? In *PPSN*, volume 2, pages 3–14, 1992.
 - [8] Kenneth A De Jong and William M Spears. A formal analysis of the role of multi-point crossover in genetic algorithms. *Annals of mathematics and Artificial intelligence*, 5(1) :1–26, 1992.
 - [9] Sachin Desale, Akhtar Rasool, Sushil Andhale, and Priti Rane. Heuristic and meta-heuristic algorithms and their relevance to the real world : a survey. *Int. J. Comput. Eng. Res. Trends*, 351(5) :2349–7084, 2015.
 - [10] Tansel Dokeroglu, Ender Sevinc, Tayfun Kucukyilmaz, and Ahmet Cosar. A survey on new generation metaheuristic algorithms. *C&IE*, 2019.
 - [11] Agoston E Eiben and Cornelis A Schippers. On evolutionary exploration and exploitation. *Fundamenta Informaticae*, 35(1-4) :35–50, 1998.
 - [12] Fred W Glover and Gary A Kochenberger. *Handbook of metaheuristics*, volume 57. Springer Science & Business Media, Dordrecht, 2006.
 - [13] Georges R Harik, Fernando G Lobo, et al. A parameter-less genetic algorithm. In *GECCO*, volume 99, pages 258–267, 1999.
 - [14] F Maxwell Harper and Joseph A Konstan. The movie-lens datasets : History and context. *ACM tiis*, 5(4) :1–19, 2015.
 - [15] Thomas T Hills, Peter M Todd, David Lazer, A David Redish, Iain D Couzin, Cognitive Search Research Group, et al. Exploration versus exploitation in space, mind, and society. *Trends in cognitive sciences*, 19(1) :46–54, 2015.
 - [16] Folasade Olubusola Isinkaye, Yetunde O Folajimi, and Bolande Adefowoke Ojokoh. Recommendation systems : Principles, methods and evaluation. *Egyptian informatics journal*, 16(3) :261–273, 2015.
 - [17] Marius Kaminskas and Derek Bridge. Diversity, serendipity, novelty, and coverage : A survey and empirical analysis of beyond-accuracy objectives in recommender systems. *ACM TiiS*, 7, 1 :2, 2017.
 - [18] Rahul Katarya. Movie recommender system with metaheuristic artificial bee. *Neural Computing and Applications*, 30(6) :1983–1990, 2018.
 - [19] Manuel Lozano and Carlos García-Martínez. Hybrid metaheuristics with evolutionary algorithms specializing in intensification and diversification : Overview and progress report. *C&OR*, 37(3) :481–497, 2010.
 - [20] Silvano Martello and Paolo Toth. *Knapsack problems : algorithms and computer implementations*. John Wiley & Sons, Inc., Hoboken, New Jersey, US, 1990.
 - [21] Amy McGovern, Ryan Lagerquist, David John Gagne, G Eli Jergensen, Kimberly L Elmore, Cameron R Homeyer, and Travis Smith. Making the black box more transparent : Understanding the physical implications of machine learning. *Bulletin of the American Meteorological Society*, 100(11) :2175–2199, 2019.
 - [22] K. Mehlhorn, B. R. Newell, P. M. Todd, M. D. Lee, K. Morgan, V. A. Braithwaite, D. Hausmann, K. Fiedler, and C. Gonzalez. Unpacking the exploration-exploitation tradeoff : A synthesis of human and animal literatures. *Decision*, 2(3) :191–215, 2015.
 - [23] Zbigniew Michalewicz and Zbigniew Michalewicz. *Genetic algorithms+ data structures= evolution programs*. Springer Science & Business Media, 1996.
 - [24] Arthur Pchelkin. Efficient exploration in reinforcement learning based on utile suffix memory. *Informatica*, 14(2) :237–250, 2003.
 - [25] Ke Sang, Peter Todd, and Robert Goldstone. Learning near-optimal search in a minimal explore/exploit task. In *Proceedings of the Annual Meeting of the Cognitive Science Society*, volume 33, 2011.
 - [26] Ngo Tung Son, Jafreezal Jaafar, Izzatdin Abdul Aziz, and Bui Ngoc Anh. Meta-heuristic algorithms for learning path recommender at mooc. *IEEE Access*, 9 :59093–59107, 2021.
 - [27] Özge Sürer, Robin Burke, and Edward C Malhouse. Multistakeholder recommendation with provider constraints. In *Proceedings of the 12th ACM Conference on Recommender Systems*, pages 54–62, Vancouver British Columbia Canada, 2018. ACM.
 - [28] Maritzol Tenemaza, Sergio Lujan-Mora, Angelica De Antonio, and Jaime Ramirez. Improving itinerary recommendations for tourists through metaheuristic algorithms : an optimization proposal. *IEEE Access*, 8 :79003–79023, 2020.
 - [29] Junqin Xu and Jihui Zhang. Exploration-exploitation tradeoffs in metaheuristics : Survey and analysis. In *Proceedings of the 33rd Chinese control conference*, pages 8633–8638, Nanjing, China, 2014. IEEE, IEEE.
 - [30] Gary Yen, Fengming Yang, and Travis Hickey. Coordination of exploration and exploitation in a dynamic environment. *International Journal of Smart Engineering System Design*, 4(3) :177–182, 2002.
 - [31] Jun Zhang, Wei-Neng Chen, Zhi-Hui Zhan, Wei-Jie Yu, Yuan-Long Li, Ni Chen, and Qi Zhou. A survey on algorithm adaptation in evolutionary computation. *Frontiers of Electrical and Electronic Engineering*, 7(1) :16–31, 2012.
 - [32] Yong Zheng and Aviana Pu. Utility-based multi-stakeholder recommendations by multi-objective optimization. In *2018 IEEE/WIC/ACM International Conference on Web Intelligence (WI)*, pages 128–135, Santiago, Chile, 2018. IEEE, IEEE.

Vers une éthique processuelle de l'IA

Éric Pardoux^{1,2}, Louis Devillaine³

¹ CNRS, IHRIM (UMR 5317), ENS Lyon

² CNRS, Maison Française d'Oxford (UMIFRE)

³ Univ. Grenoble Alpes, CNRS, Sciences Po Grenoble, Pacte, 38000 Grenoble, France

eric.pardoux[at]ens-lyon.fr / louis.devillaine[at]univ-grenoble-alpes.fr

Résumé

Nous mettons en évidence certaines des dimensions éthiques présentes dans le processus de design des systèmes d'AM. Sur cette base, nous proposons une éthique processuelle qui associe des enjeux éthiques à toutes les étapes de la conception. Cette démarche nous amène à mettre en lumière différents instants du processus de design qui peuvent occasionner des préoccupations d'ordre éthique. Il s'agit de présenter la base de réflexions futures sur la place que l'éthique peut prendre dans la production d'IA en général.

Mots-clés

Éthique de l'IA, Éthique processuelle, Éthique by design.

Abstract

We highlight some of the ethical dimensions present in the design of ML systems. Based on this, we propose to rethink ethics as process undergoing at all stages of the design process. This approach leads us to underline different moments in the design process that may give rise to ethical concerns. Our aim is to present the basis for future reflections on the place that ethics can take in the production of AI.

Keywords

AI ethics, Processual ethics, Ethics by design.

1 Introduction

Il n'est probablement plus réellement nécessaire de démontrer les avancées grandissantes de l'intelligence artificielle (IA) ces dernières années [31]. Ces avancées ne sont pas sans soulever des questions éthiques¹ liées aux dommages que l'introduction de l'IA pourrait provoquer : reproduction de biais sociaux, discrimination, dangers de l'automa-

tion sont quelques-uns des enjeux éthiques que l'on peut citer [35, 41, 51]. Lors de son déploiement, l'IA peut ainsi venir perturber des valeurs morales et sociales comme la justice, l'équité ou la responsabilité². Notre objectif ici est d'introduire une réflexion sur la façon dont la considération des enjeux éthiques de l'IA peut avoir lieu tout au long de son développement. Pour ce faire nous proposerons l'ébauche d'un cadre éthique qui permettra d'appréhender l'éthique de l'IA au travers de son processus de *design*³. Notre ambition est notamment de proposer une approche complémentaire à l'éthique par principes qui nous semble dominer la littérature de l'éthique de l'IA [40]. Nous suggérerons ainsi l'adoption accrue d'une éthique processuelle, accompagnant le *design* des systèmes d'IA (SIA) de leur idée initiale à leurs usages finaux. Comme nous le verrons, cette éthique processuelle de l'IA doit se traduire par une prise de conscience plus marquée des enjeux éthiques souvent implicites et pourtant cruciaux présents tout au long du processus de conception des SIA.

Avant de rentrer dans le détail de ces considérations, il convient de préciser à quoi nous nous référons lorsque nous employons le terme d'IA. Nous ancrons avant tout nos considérations à partir des algorithmes basés sur l'apprentissage machine (AM)⁴. Nous faisons ce choix car il permet d'aborder des applications émergentes de l'IA, qui posent des questions nouvelles à l'éthique — par exemple le problème de l'explicabilité, ou la responsabilité pour les SIA à apprentissage incrémental. De plus, nous considérons que certains des lieux de l'éthique présentés pour les algorithmes d'AM sont communs avec le développement de SIA non basés sur l'AM. En cela, nous espérons fournir

2. Il convient toutefois de souligner que ces perturbations peuvent être bénéfiques si elles promeuvent et améliorent la qualité morale des actions et décisions auxquelles elles contribuent.

3. Nous employons le terme de *design* dans son acception anglaise, afin de recouvrir les différents sens qu'on peut lui donner. Ceux-ci seront explicités dans la suite de l'article.

4. Pour expliciter notre compréhension, nous repartons de la définition donnée par T. Mitchell [36] : « on dit d'un programme informatique qu'il apprend d'une expérience E vis-à-vis d'une classe de tâches T et d'une mesure de performance P si ses performances sur les tâches T, mesurées par P, s'améliorent avec l'expérience E ». En l'occurrence, l'expérience peut consister en des exemples — on parle alors d'apprentissage supervisé, ou non-supervisé — ou bien en un ensemble de règles — apprentissage par renforcement.

1. Par éthique nous entendons ici la façon dont des principes moraux et des valeurs peuvent se traduire dans les actions et les décisions, notamment dans le développement et l'introduction de nouvelles technologies. Il existe une pluralité de perspectives éthiques — par exemple le conséquentialisme, l'approche déontologique ou encore l'éthique des vertus. Au-delà de leurs différences d'implémentation, leur principal point commun réside dans la fixation d'un horizon d'attente par rapport à ce que le réel devrait être (d'un point de vue moral, social ou encore politique). Cela reste le cas, que le jugement éthique des actions et décisions se traduise (i) dans la considération de leurs conséquences, (ii) des principes suivis pour les réaliser ou (iii) au travers des vertus mobilisées par l'agent moral dans leur réalisation.

des réflexions utiles à un cadre plus large qui dépasserait les limites seules de l'AM.

Cet article sera organisé en deux temps principaux : tout d'abord nous formulerons une proposition de reconstruction d'un cadre éthique dans lequel inscrire les pratiques de *design* d'AM. Nous aborderons ensuite plus en détail comment cette posture éthique théorique pourrait s'illustrer dans le cadre du développement de systèmes d'AM (SAM)⁵. Pour ce faire, nous décrirons le processus de *design* de SAM pour faire émerger les lieux des enjeux éthiques en son sein.

2 Où est l'éthique de l'IA ?

2.1 Une exigence éthique sur la technique ?

Les préoccupations éthiques entourant le développement de l'IA engendrent une volonté de régulation qui ne concerne plus seulement les usages. Dans les applications dites à « haut risque » telles que les applications en santé, les cadres éthiques prescrivent tout autant les usages qui devraient être faits des SIA que ce qu'ils devraient être en eux-mêmes. Cela se traduit par exemple par une exigence de certaines caractéristiques : les SIA doivent se révéler être explicables, transparents, justes, équitables, non discriminatoires et tout ceci dans leur conception même⁶. On retrouve ces exigences dans la littérature en éthique de l'IA tout comme dans le cadre de régulations [25].

À l'échelon européen, de nombreuses propositions sont formulées ces dernières années pour parvenir à mettre en place un cadre régulateur commun au développement des SIA à « haut risque »⁷. La dernière proposition en date contient ainsi un article prônant une exigence de transparence⁸.

Au niveau français, l'article 17 de la loi relative à la bioéthique⁹ demande par exemple à ce que les professionnels de santé utilisant un SAM « s'assure[nt] que la personne concernée en a été informée et qu'elle est, le cas échéant, avertie de l'interprétation qui en résulte. ». Dans le même article est également inscrite une exigence pour que les « concepteurs d'un traitement algorithmique [...] s'assurent de l'explicabilité de son fonctionnement pour les utilisateurs ».

Dans un cadre comme dans l'autre, il est intéressant de remarquer l'absence de régulation forte des usages de l'IA, au profit d'un conditionnement technique avant tout. En effet,

5. Par système d'apprentissage machine, nous entendons ici un ensemble hétéroclite d'objets techniques, incluant aussi bien les programmes informatiques qui implémentent les méthodes d'AM que les interfaces humain-machine (IHM) et les systèmes sociotechniques qui les intègrent.

6. Néanmoins, le lien entre ces caractéristiques « éthiques » et des propriétés techniques intrinsèques aux algorithmes n'est pas une évidence. Sur le cas de l'équité voir par exemple l'étude menée par S. Wachter *et al.* [51]. Sur les obstacles techniques à l'élaboration d'un système explicable, se référer à la revue de Dazeley *et al.* [9].

7. Dans cette catégorie large, sont incluses notamment les SIA pour la santé, les systèmes destinés à la justice, à l'éducation, aux prestations sociales, aux moyens de répression, etc. Voir l'annexe III de la *Proposition de règlement du Parlement européen et du Conseil Établissant des règles harmonisées concernant l'IA* (COM/2021/206 final) parue le 21/04/2021.

8. Voir l'article 13 de la proposition citée précédemment.

9. Loi n° 2021-1017 du 2 août 2021 relative à la bioéthique (1), NOR : SSAX1917211L.

certaines applications sont certes proscrites¹⁰, mais aucune mention n'est faite d'un encadrement strict des usages une fois que les SAM sont autorisés et certifiés, ni d'un suivi du système une fois déployé.

La direction prise par la formulation juridique de ces exigences techniques, tout comme la perspective employée dans une large part de la littérature sur l'éthique de l'IA relève selon nous d'une certaine posture éthique. Cette dernière reviendrait à concevoir l'objet et son éthique à partir de son usage — notamment pour déterminer s'il est à « haut risque » ou non — mais agirait avant tout en conditionnant la conception et les caractéristiques techniques des SIA. Une telle pratique pourrait s'interpréter comme renvoyant à une « éthique *by design* »¹¹.

Dans la suite de cet article, nous nous proposons de présenter les enjeux de cette démarche d'éthique *by design*, avant de faire le lien avec le processus de *design* d'IA. Cela nous permettra de mettre en avant les étapes au cours desquelles il peut être bénéfique de mettre en place un questionnement éthique lors du *design* de SAM.

2.2 L'éthique *by design*

L'éthique par *design* est un concept relativement récent. Selon F. Fischer [13], la première occurrence académique du terme d'« *ethics by design* » date de 2010, dans le contexte de la stratégie d'entreprise [39]. Il s'agit alors d'une ouverture des entreprises vers une éthique conséquentialiste. Cette démarche consiste à ne plus seulement constater les changements opérés par les pratiques managériales ou l'introduction de technologies dans un contexte organisationnel, mais à les anticiper pour en prévenir les maux potentiels. Elle dénote ainsi d'une volonté d'opérer une éthique *a priori* par rapport à l'usage, qui soit en anticipation des éventuelles nuisances liées à l'introduction d'une nouvelle technique. Il ne s'agit donc pas seulement de mettre en place des pratiques vertueuses, qui relèveraient d'une déontologie professionnelle au sein des entreprises mais bien de prendre en compte l'intégralité des processus agissant au sein de l'organisation et des conséquences en découlant. Cette première occurrence de l'éthique *by design*, issue de la littérature du management, destine donc la réflexion éthique à toutes les étapes des processus d'organisation et de décision qui peuvent être menés dans les entreprises.

Dans le contexte francophone, le terme de *design* est souvent repris directement de l'anglais : il nous semble important de souligner la diversité des sens dont il peut relever, afin de mieux problématiser les pratiques relevant de l'éthique *by design*. Trois visions principales de l'éthique *by design* peuvent être adoptées par les *designers* en fonction du sens donné au *design* [12].

Une éthique par dessin. Elle peut être entendue en premier lieu dans sa traduction littérale comme une éthique par

10. Pour le règlement européen, il s'agit notamment des SIA ayant recours à des « techniques subliminales » de manipulation, à l'exploitation de vulnérabilités, à la notation sociale ou encore à l'identification biométrique à distance en temps réel.

11. Notre constat semble par ailleurs appuyé par l'existence d'un document de travail de l'UE portant sur des recommandations suivant une certaine acception de l'éthique *by design* pour les chercheurs en IA [8].

dessein, l'intention des *designers* étant supposée se traduire dans le dispositif technique de façon positive. Le dispositif technique est conçu pour être déployé dans un certain cadre où il serait éthique par nature, du fait de l'intention transmise par les *designers*¹². Les futurs usagers seraient empêchés de produire des effets indésirables car ils suivraient cet usage prescrit par intention. Dans une approche d'AM incrémental et non statique, on peut voir les limites d'une telle conception de l'éthique par dessein au travers de l'exemple de Tay, l'agent conversationnel déployé par Microsoft sur Twitter. L'usage détourné du système a dépassé les limites anticipées par les développeurs jusqu'à l'apprentissage d'injures [52]. Le dessein des *designers* s'est ainsi retrouvé débordé par l'usage, montrant les limites de l'absence de contraintes techniques.

Une éthique par définition. L'éthique *by design* s'envisage ici sous le prisme de la conception même de l'objet technique, qui est supposé incarner des valeurs désirables. Son usage dans la société serait maîtrisé par les *designers* au travers de la définition de ses spécifications. C'est l'agencement même des composants techniques qui vient alors conditionner les actes éthiques de ces derniers, pour reprendre la formulation de Fischer, l'éthique devient ainsi « technicisée » [12, p. 64]. On retrouve ici notamment les approches de *Privacy by design* [45], reprises par Fischer. Ceci peut être considéré comme une éthique par définition, car les fonctionnalités de protection de la vie privée sont alors intégrées dès le début du processus de *design* et non pas en seconde intention une fois le cœur du système technique déjà développé. Cette approche est critiquable du fait de la potentielle rigidité du système une fois déployé, qui peut conduire à un rejet de la part des usagers, si les contraintes pour respecter la vie privée deviennent plus fortes que les bénéfices qu'elles fournissent.

Une éthique de la médiation. Une dernière proposition est l'éthique *by design* fondée sur une approche relationnelle entre usager et objet, en incorporant une réflexion sur les conséquences probables de l'interaction du système avec les individus constituant la société. Cette approche intègre en quelque sorte les deux dimensions précédentes : la technique est une médiation entre l'humain et le monde, elle conditionne sa vision de ce dernier et les actions qu'il peut entreprendre à son propos. L'éthique n'est plus centrée sur l'humain mais sur ses relations à son milieu — aussi bien technique que social. En ce sens, il nous paraît bénéfique de favoriser cette éthique de la médiation qui enjoint à systématiquement contextualiser socialement l'usage qui sera fait des techniques développées. Il s'agit donc d'avoir des pratiques du *design* qui ne soient ni trop ouvertes — comme dans une éthique par dessein — ni trop fermée — comme dans des systèmes trop contraints conçus par définition. L'idée est de laisser à l'utilisateur la main sur la paramétrisation des systèmes en prenant en compte les potentiels usages qu'il en aura [28]. Cette approche plus équilibrée

permettrait de diriger les usagers vers un questionnement ou un comportement éthique souhaitable, sans pour autant être trop paternaliste [14]. C'est une façon d'échapper à l'impasse d'une relativisation totale des enjeux éthiques de la technique par la prétendue infinité de ses usages et détournements possibles. L'approche relationnelle promue par la médiation nous paraît transcrire au mieux la réalité des processus de *design*, de réception et d'intégration des systèmes sociotechniques dans le monde réel.

Ces trois approches sont évidemment non exclusives — dans le sens où l'éthique de la médiation peut mobiliser des approches par définition ou par dessein. Elles donnent à voir un panorama étoffé des différentes façons d'aborder la conception d'un système sociotechnique¹³.

2.3 Ce que fait et peut le *design*

Il convient ici à notre sens de nous arrêter un temps sur ce à quoi renvoie l'idée de *design*, dans un cadre général. Comme la typologie de l'éthique *by design* suggérée par Fischer nous le fait entrevoir, ce terme est associé à une diversité de significations et de pratiques. Une première nuance forte est à faire avec la dimension dynamique du *design*, qui est alors analogue à la conception de l'objet technique : le *design* est un processus, un enchaînement d'actions et de décisions qui mène à la création d'un dispositif technique.

Pour autant, le *design* peut également renvoyer au résultat en tant que tel de ce processus. C'est alors le sens d'agencement sociotechnique qui prime avant tout. De par cet agencement particulier à tout dispositif, comme on l'a vu auparavant, certaines actions peuvent être favorisées ou au contraire entravées voire rendues impossibles. Nous retrouvons ici l'hypothèse de L. Winner [53], selon laquelle la dimension politique et éthique des objets découle avant tout de leur organisation et de leurs interrelations plutôt que de leurs caractéristiques intrinsèques. Si la littérature philosophique est divisée sur la façon avec laquelle émergent ces valeurs politiques¹⁴, il reste un constat fort : la technique a notamment pour fonction de modifier l'état du monde et le processus de *design* y participe directement.

En effet, ce dernier implique de chercher à transformer le monde en ce qu'il *devrait être*¹⁵. Le *design* est donc tout autant un processus de composition (d'agencement) d'un nouveau réel qu'un processus de recomposition (réagencement) d'un état passé du monde et du tissu sociotechnique. En cela encore, il possède certaines similitudes avec le processus d'innovation [19]. D'un part, parce que le *design*

13. C'est ici une vision que nous tenons à défendre : les dispositifs techniques ne font sens qu'une fois déployés dans un contexte social donné, co-produisant alors une nouvelle réalité aux dimensions à la fois techniques, sociales, éthiques, juridiques ou encore politiques [24].

14. Citons par exemple, en réponse à Winner, la critique de S. Woolgar et G. Cooper [55] ou encore la reconstruction plus récente de S. Lavelle qui développe cette réflexion à propos de la société numérique [29].

15. « *The engineer, and more generally the designer, is concerned with how things ought to be — how they ought to be in order to attain goals, and to function.* » (« L'ingénieur, et plus généralement le designer, se préoccupe de ce que les choses devraient être — comment elles devraient être dans le but d'atteindre des objectifs, et de fonctionner » [48, pp. 4-5 (notre traduction)]. Kraemer *et al.* proposent une définition similaire dans le cadre de l'éthique des algorithmes [28].

12. Ici et par la suite nous utiliserons ce terme dans un sens large incluant toute personne contribuant de près ou de loin au processus de *design* : cela va des commanditaires aux ingénieures, en passant par les programmeurs ou encore les chercheuses.

possède une dualité : il est un processus tout en étant le résultat de ce même processus. D'autre part, il est également une activité de réorganisation du réel¹⁶ au travers de dispositifs techniques.

Au-delà de cette fonction de transformation, le *design* s'incarne dans une grande pluralité de pratiques. L'identification et la classification de ces dernières peut s'articuler autour de la place laissée à l'usager de l'objet dans la pratique du *design*. Nous insistons sur l'importance de distinguer les usages — qui renvoient aux pratiques réelles et situées d'un système — de l'utilisation — qui est une spécification issue du *design*, indiquant un ensemble restreint d'emplois, le plus souvent acontextuels¹⁷. Penser par les usages possibles, c'est donc anticiper détournements et dévoiements et par là même les soucis éthiques qui pourraient découler entre l'utilisation (projetée) et l'usage (effectif).

Trois grandes approches du *design* coexistent dans cette approche centrée sur l'usage au lieu de l'utilisation¹⁸ : le *design* centré sur l'usager, la démarche de *design* participatif et enfin le *metadesign*. Les deux premières approches sont voisines. Dans le *design* centré sur l'usager, le *designer* projette les usages à venir de son dispositif sur un usager (fictif ou non) dont il essaie d'anticiper les réactions. L'ambition est alors de favoriser certaines actions ou d'en proscrire d'autres de par l'agencement du dispositif. Le *design* participatif (ou *co-design*) suit une démarche similaire mais en intégrant directement au processus de *design* un ou plusieurs usagers du futur dispositif¹⁹.

Le *metadesign*, se détache des deux approches précédentes en prescrivant de façon moins restreinte les usages du dispositif. Il s'agit alors de laisser une plus grande marge de manœuvre à l'usager pour qu'il soit moins contraint par l'usage prescrit, en lui laissant une liberté de création et d'initiative au sein du système sociotechnique déployé. Cette dernière approche est particulièrement intéressante dans le champ du développement logiciel, qui connaît des évolutions fréquentes et une capacité de portage élevée. Pour attrayante que soit cette notion de *metadesign*, il convient de rappeler que le concepteur initial du système conserve tout de même une emprise forte sur les potentialités proposées de même qu'un regard sur les activités qui utilisent son dispositif. Toutefois, on marque là la possibilité d'une nuance forte entre un dispositif technique qui serait implémenté de façon rigide et donnant peu d'emprise à l'usager et une approche laissant les coudées plus franches à ce dernier.

Cette nuance entre les différents types de démarches de *de-*

16. À la suite de V. Flusser, il est intéressant de noter que le *designer* peut être vu comme un « comploteur rusé qui tend ses pièges » (« A designer is a cunning plotter laying his traps. » (notre traduction)) au travers de ce qu'il conçoit [16, p. 17].

17. Autrement dit, l'utilisation c'est l'usage idéalement prescrit par les *designers*, ce qui nous renvoie à un *design* uniquement par dessin ou définition, sans tenir compte des effets de médiation.

18. Nous reprenons cette distinction depuis les travaux de V. Beaubois, présentés lors de la conférence « L'invention de l'usager » le 17/02/2022 et à paraître avant fin 2022.

19. De manière intéressante, cette démarche peut également être étendue en intégrant des professionnels de l'éthique au processus de *design* [50].

sign par l'usage souligne plusieurs éléments importants : l'évaluation éthique dépend des usages possibles, et l'éventail de ces derniers dépend en partie du degré d'ouverture ou de fermeture du système technique. L'intention de transformation du réel qui dirige le processus de *design* peut être partagée entre différents acteurs, qui n'ont pas les mêmes compétences vis-à-vis du processus de conception, ni vis-à-vis de l'usage qui sera fait de son résultat. Il est à noter par ailleurs que ces trois approches ne sont en rien exclusives : des pratiques de *design* participatif où différentes parties prenantes sont à l'œuvre peuvent concevoir un système suivant les principes du *metadesign*.

Si le terme de *design* inclut ainsi en lui-même la volonté de transformer le monde selon une certaine intention, qu'apporte donc la notion d'éthique *by design* ?

Selon nous c'est la posture réflexive qui vient donner toute sa nuance et sa puissance à l'éthique *by design*. L'intention du *designer* porte consciemment ou non des enjeux éthiques qu'il faut interroger tout au long du processus. En s'intéressant à l'éthique par le prisme du *design*, nous suggérons qu'il est nécessaire de clarifier les enjeux éthiques présents tout au long des étapes de conception et de production des dispositifs techniques. Chacun des choix effectués dans le processus de *design* peuvent se trouver répercutés dans l'usage effectif du dispositif finalement produit. Comme nous l'avons déjà souligné, ce processus peut impliquer une pluralité importante d'acteurs : se poser la question de ceux qui peuvent avoir une posture de partie prenante en son sein, c'est aussi pratiquer un choix éthique [2]. Défendre un tel processus de mise en question à la fois technique mais également sociale du *design*, c'est défendre la possibilité de construire un SIA qui puisse se montrer réellement digne de confiance²⁰. Nous allons voir par la suite comment ces enjeux peuvent se traduire plus particulièrement dans le cadre du développement d'un SAM.

3 Une éthique tout au long du *design*

3.1 Pour une éthique processuelle

3.1.1 Procédure ou processus ? Repenser l'éthique

Il s'agit selon nous avant de tout de faire évoluer la façon dont les enjeux éthiques sont considérés durant le *design* de SAM. Par une posture réflexive adoptée tout au long du processus de *design*, nous entendons ainsi plaider en faveur d'un regard critique sur les activités et choix de conception entourant ces systèmes.

Repenser la technique de façon critique implique de décentrer la focale du dispositif technique seul pour l'élargir au tissu social dans lequel il sera intégré. Plutôt qu'une remise en cause frontale des valeurs éthiques à adopter, il s'agit ici de refonder la façon dont nous les mettons en œuvre. Cela revient à ne plus se contenter uniquement de soulever ponctuellement les questions éthiques, lors de la certification ou

20. Ici nous souhaitons mettre l'accent sur le fait qu'il s'agit d'une confiance à construire dans l'intégralité de la chaîne du *design* [27]. L'IHM n'est qu'une incarnation superficielle du *design* global qui traverse le système de ses caractéristiques techniques jusqu'à ses modalités d'intégration au contexte d'usage.

du dépôt d'un projet de recherche par exemple.

Comme Perrin *et al.* le soulignent, une telle pratique ponctuelle de l'éthique peut faire craindre de voir toute réflexion éthique « vidée » de son contenu pour devenir une simple procédure administrative, servant avant tout à protéger les institutions et les chercheur·e·s » [44, p. 237]. Au contraire nous défendons une vision processuelle de l'éthique qui compléterait l'éthique procédurale²¹. Il convient pour cela de distinguer les procédures des processus. Les procédures sont des dispositifs mobilisés pour accomplir des fonctions prédéfinies, elles sont donc de nature fermées, discrètes, statiques et acontextuelles — pour une procédure donnée, sa fonction est accomplie quelle que soit la situation. En contraste, les processus nécessitent une certaine disposition face à une situation donnée : cela implique une capacité réflexive qui les rend ouverts, continus, dynamiques et adaptatifs. C'est sur la base de cette distinction que nous défendons l'intégration plus forte d'une dimension processuelle à l'éthique des *designers*. En percevant l'éthique comme un processus et non comme une simple procédure à accomplir ponctuellement, on évite de réduire l'éthique à une simple condition à remplir, une étape à satisfaire qui serait toujours identique, quel que soit le contexte.

Selon Perrin *et al.*, l'éthique procédurale est emblématique des sciences biomédicales tandis que son pendant processuel s'inspire grandement de l'approche pratiquée en anthropologie et dans les sciences sociales qualitatives [44]. Dans le cadre de la recherche, l'éthique procédurale se réalise avant tout par l'accomplissement de procédures, par exemple dans un dispositif administratif encadrant la recherche ou le développement d'une technique. Elle se traduit par exemple au travers de l'accent mis sur l'obtention du consentement éclairé des usagers ou des sujets de recherche, le passage devant des comités d'éthique ou l'élaboration de dossiers éthiques aux critères stricts et universels. En cela, elle permet tout de même de fixer un cadre favorable à la protection des individus observés. Au contraire, l'éthique processuelle affirme la dimension politique et donc éthique de chaque étape de la recherche [30] : les anthropologues doivent arbitrer leurs décisions en fonction des différentes parties prenantes de leur recherche.

Ce détour par l'anthropologie n'est pas anodin : comme le suggèrent déjà M. C. Elish et d. boyd, l'AM peut se concevoir comme une « ethnographie computationnelle » [11]²². Reconnaître le caractère non neutre de la technique implique une prise de conscience des responsabilités portées par les *designers* vis-à-vis du terrain dans lequel le dispositif devra s'intégrer. Cela implique une compréhension — quasi ethnographique donc — de ce contexte de déploiement. Ce rapport au terrain peut passer par le questionne-

ment constant porté par l'éthique processuelle. Cette dernière se montre ainsi complémentaire à une éthique plus ancrée dans des procédures administratives. L'essentiel de la nuance entre procédure et processus réside dans une approche différente des temporalités de l'éthique²³ : d'une part l'éthique procédurale est discrète, figée dans certaines procédures, d'autre part l'éthique processuelle est continue, tout au long de la recherche ou du développement. L'idée de processus se traduit au travers de dispositifs différents (comme les conférences de citoyens ou les *focus groups* par exemple), plus adaptés aux contextes particuliers, dont l'existence peut venir en soutien des procédures éthiques plus classiques pré-existantes. Cette hybridation entre deux manières d'envisager l'éthique nous paraît être une piste intéressante pour résoudre les tensions entre la visée universelle d'une éthique basée sur des grands principes et le caractère exceptionnel que peuvent revêtir certaines techniques liées à l'AM [2]. La mise en avant de la possibilité d'une éthique processuelle ne doit ainsi pas effacer les effets bénéfiques que les procédures éthiques peuvent avoir pour protéger les futurs usagers d'une technique. L'exercice continu d'un raisonnement éthique tout au long du *design* permet d'intégrer une plus grande variété d'enjeux éthiques²⁴, sans nécessairement négliger la certification finale du SAM [54]. L'éthique processuelle ouvre ainsi la possibilité d'une réflexion sur le choix des procédures éthiques à engager.

Ce retour au présent et aux pratiques telles qu'elles se font n'est en rien original dans la littérature de l'éthique des techniques [20]. Il nous semble toutefois intéressant de spécifier quelques pistes pour imaginer l'intégration de cette éthique au *design* des SAM. Ce projet de recherche est particulièrement large, nous souhaitons donc ici présenter dans un premier temps les principaux lieux qui nous paraissent propices à une réflexion éthique processuelle dans le *design* et la conception de SAM.

3.1.2 Comment intégrer l'éthique processuelle ?

L'éthique processuelle ne se substitue pas à un cadre éthique préexistant. Elle constitue plutôt une manière de décliner une posture éthique particulière²⁵. Il s'agit donc pour nous d'interroger avant tout la façon dont l'éthique se fait plutôt que les principes qu'elle défend. En cela, elle reste compatible avec l'approche dominante d'éthique « par principes ». Cette éthique par principes vient souligner le besoin que les résultats des SAM respectent certaines valeurs fondamentales. Ces principes regroupent tout aussi bien des caractéristiques techniques intrinsèques aux algorithmes que des valeurs liées aux usages qui en sont faits. Parmi ceux qui sont saillants dans le champ de l'IA, on compte no-

21. Par ailleurs, Nurock *et al.* ont également défendu une éthique processuelle à intégrer au *design* [42]. Contrairement à notre perspective, leur compréhension de l'éthique processuelle est marquée par celle du *care*.

22. Il s'agit pour elles d'insister sur le besoin de réflexivité dans les pratiques de l'AM, notamment en portant un regard critique sur les limites éthiques et pratiques des données employées et des modèles construits. Elles soulignent néanmoins la nécessité de construire une méthodologie propre à chaque situation.

23. Il est à souligner que ces catégories d'une part procédurale et d'autre part processuelle recouvrent en elles-mêmes une diversité de pratiques que nous ne détaillons pas ici en détail.

24. À ce propos nous renvoyons notamment à une proposition méthodologique destinée à la méthode agile [56].

25. Ainsi que le suggère explicitement B. Mittelstadt lors qu'il parle d'encourager la vision de « l'éthique comme un processus et non pas comme un solutionnisme technologique » (« *Pursue ethics as a process, not technological solutionism* » (notre traduction)) [37, p. 10].

tamment la transparence, la non-malfaisance, l'équité et la justice, la responsabilité et le respect de la vie privée²⁶. Des propositions émergent pour montrer en quoi ces principes éthiques sont mis en jeu dès le développement des SAM, et non plus seulement dans leurs usages. Morley *et al.* recensent ainsi une centaine d'initiatives pour intégrer l'éthique au sein de la phase de conception [40]. Ils soulèvent cependant le manque d'applicabilité directe des propositions issues de la littérature. Une difficulté réside dans la transcription de principes au sein du *design* des SAM. Des risques surgissent notamment d'une instrumentalisation de ces principes à des fins mercatiques et commerciales par des organisations qui voudraient paraître plus éthiques qu'elles ne le sont réellement [15].

3.2 Intrications entre technique et éthique

Nous nous proposons d'expliciter, à chaque étape, une partie des questionnements éthiques qui, bien que parfois implicites ou non-considérés, sont effectivement compris dans ce processus de conception²⁷. Il est à noter que la succession d'étapes que nous allons proposer est avant tout schématique et peut être rompue au moins de deux manières. D'une part, les frontières qui les délimitent sont poreuses, c'est-à-dire que deux opérations présentées distinctement peuvent en réalité être réalisées simultanément, soit en parallèle, soit au sein d'une activité de travail. Lorsqu'une équipe travaille sur un projet utilisant un algorithme d'AM, il n'est en effet pas rare que le codage de celui-ci débute avant même que les données ne soient récoltées. Des allers-retours entre ces étapes sont d'autre part monnaie courante : il est habituel que des résultats donnés par un modèle prototypique mettent en lumière un échantillon aberrant de la base de données (BdD), qui en sera ensuite évincé. La liste d'étapes proposée n'est donc pas à prendre comme une représentation littérale de la procédure de *design*. Il s'agit plutôt d'une esquisse simplifiée à des fins analytiques. Ce découpage en étapes d'un processus continu permet d'identifier des points de vigilance qui échapperaient à une approche procédurale de l'éthique (qui se niche souvent uniquement dans la certification finale du SAM). L'enjeu est de favoriser une posture réflexive en reconnaissant les situations qui peuvent être dommageables au point de vue éthique.

3.2.1 Formalisation du problème

Le processus de conception débute au moment où il est décidé qu'il serait judicieux de recourir à un SAM pour résoudre un problème identifié. Bien que cette décision ne soit pas en elle-même une étape du processus de conception et qu'elle n'appartienne pas aux développeurs eux-mêmes en général, nous estimons qu'elle en constitue son point d'entrée et qu'à ce titre elle mérite qu'on lui porte une attention particulière. La décision de recours à l'usage de ces systèmes se fonde parfois sur les croyances que les SAM,

capables de traiter un grand nombre d'informations à la fois, pourraient résoudre plus efficacement que des humains la majorité des problèmes [7]. Il appartient dès lors aux développeurs d'interroger la pertinence de l'emploi d'un SAM pour résoudre le problème réel identifié. S'ils sont les mieux placés pour définir ce que peut un système informatique, ils doivent pouvoir tout autant être critiques de ce que ce dernier ne permet pas d'accomplir [11].

Le problème réel est traduit en un problème mathématique que doit résoudre le SAM. Ce problème mathématique consiste, pour le dire simplement, en l'optimisation d'une fonction de coût. Les paramètres de cette fonction sont des variables, quantitatives ou qualitatives, qui correspondent à des attributs formalisés du problème. Pour l'évaluation future du modèle entraîné, des métriques de performance dudit modèle sont également définies. Pourtant, traduire un problème réel en un problème qui soit soluble par un algorithme n'est absolument pas évident. Une difficulté d'ordre éthique apparaît à ce stade, la correspondance entre la réalité du problème et la fonction de coût. Pour l'illustrer, un système chargé de différencier des images de chiens et de loups peut être amené à se fier à la couleur de l'arrière-plan (blanc pour les loups, à cause de la neige) plutôt qu'aux phénotypes de ces animaux [46]. Dans ce cas, le système résout bien mathématiquement le problème, mais cela ne correspond pas aux attentes qui pouvaient être investies en lui. Cela montre l'existence possible d'une divergence entre le problème formulé symboliquement — distinguer des animaux — et computationnellement — distinguer des matrices de pixels représentant pour l'humain des photos d'animaux. Cette problématique est d'autant plus prégnante lorsque l'on touche à des qualités humaines, qu'elles soient physiologiques ou morales. En ce sens, les variables qui servent de critères aux algorithmes de recrutement supposés trouver le meilleur salarié à attribuer à une entreprise donnée paraissent difficiles à définir précisément, et d'autant plus à quantifier [43]. C'est ce que B. Chin-Yee et R. Upshur appellent le problème phénoménologique de l'IA, à savoir l'incapacité à rendre compte de tout un pan de l'expérience humaine de manière quantitative [6]. Le choix, la définition et la mesure de ces variables deviennent dès lors des enjeux éthiques à part entière.

3.2.2 Préparation des données

Une grande partie du travail de conception réside dans la préparation de l'ensemble des données qui serviront à entraîner puis valider le modèle²⁸. Cette étape détermine ce qu'apprend le modèle et les résultats qu'il obtient.

Constitution de la base de données. Dans certains cas, la BdD existe déjà en accès libre et est récupérée, sa constitution ne demande ainsi pas de captation des données. Celle-ci requiert une méthode rigoureuse pour assurer que les données soient collectées dans des conditions similaires.

26. Ce sont en tout cas ceux qui sont les plus mis en avant dans la multitude de guides éthiques proposés par des institutions, qu'elles soient gouvernementales, supra-nationales, académiques ou privées [25].

27. Les phases schématiques de développement d'un SAM proposées ici sont similaires au modèle classique CRISP-DM [4].

28. Il est à noter que la prise en compte des conditions de mise en existence de ces bases de données peut constituer une préoccupation éthique à part entière. Les activités professionnelles de micro-travail, relativement précaires, ne sont pas encore strictement encadrées et consistent en la réalisation de tâches d'annotation ou de vérification fastidieuses, pas toujours claires et qui peuvent être refusées sans motif manifeste par les clients [34].

Quelle que soit la méthode de collecte il reste important d'éviter la perte de sens qui pourrait advenir en isolant les données de leur contexte de production [32].

La connaissance d'un phénomène dépend des données choisies pour le représenter et le mesurer [43]. Une BdD satisfaisante doit être représentative du phénomène qu'elle prétend décrire, ce qui n'est pas un objectif trivial à atteindre lorsque ce phénomène n'est pas entièrement connu. Et même lorsque l'on connaît les propriétés statistiques des variables représentant un phénomène, constituer une base qui les respecte reste une tâche complexe²⁹. Un manque de représentativité de la BdD introduit dès lors un risque de biais qui peuvent se manifester de manière ostensible lors de la mise en pratique du modèle dans l'espace social.

Annotation de la base de données. Dans le cas de l'apprentissage supervisé, il y a une volonté de réemployer des catégorisations théoriques ou statistiques pré-existantes. La BdD est annotée en fonction de ces catégories, le plus souvent manuellement — il y a donc une dépendance aux capacités et à la subjectivité de l'annotateur. Le soin apporté à l'annotation permet de tenter d'éviter la génération de biais, ou devrait du moins rendre compte des limites de la catégorisation effectuée. Une part importante d'arbitraire reste pourtant incluse dans cette étape de classification, sans qu'elle ne soit explicitement décrite en général³⁰.

La finesse de la distribution des classes attribuées aux échantillons de la BdD peut avoir des conséquences directes sur les résultats d'un SAM. Une illustration est donnée par la mort d'E. Herzberg, causée en 2018 par une reconnaissance d'obstacle incorrecte de la part d'une voiture Uber quasi-autonome³¹.

L'annotation de BdD, tout comme sa constitution en général, n'est ainsi pas une activité purement descriptive (épistémique), elle a un pouvoir normatif (éthique) sur l'emploi fait des données. Éthique et épistémologie se rejoignent dans l'exigence d'une correspondance adéquate entre la BdD, son annotation et la réalité. L'évaluation de cette correspondance repose sur l'explicitation de l'origine et de la nature des données et contribue à éclairer au mieux possible les décisions éthiques.

Nettoyage et standardisation de la base de données. Cette procédure consiste à opérer des modifications de la BdD pour s'assurer de la qualité et de la correction de celle-ci. Bien que des systèmes de correction automatisée existent, l'implication d'un humain dans cette tâche est quasiment toujours nécessaire [22]. Grâce à la visualisation et

à l'analyse superficielle des données³², il est possible de trouver des échantillons qui sont considérés comme étant inadéquats. Il peut s'agir d'erreurs de mesure, de retranscription ou de disparités liées à l'utilisation de données provenant de sources différentes. Par ailleurs, les données peuvent subir un pré-traitement pour rendre possible leur analyse statistique par le système³³. Cette calibration est également nécessaire lorsque des données sont issues de sources différentes avec leurs normes spécifiques³⁴.

Le nettoyage des données implique de classer les données entre celles qui seraient correctes et incorrectes, puis d'effectuer des modifications ou des suppressions de données inadéquates. Ce sont les perceptions seules des *designers* sur ce qui correspond ou non au réel qui les amènent à opérer des transformations manuelles de la BdD. Ce qui est préjudiciable ici est d'écarter une donnée sous prétexte qu'elle apparaît anormale alors qu'aucune erreur particulière n'a été commise dans sa collection et qu'elle pourrait légitimement prétendre à être conservée. On observe ainsi à cette étape un risque à faire dériver la normativité de la normalité [5]. Dans le cas d'une BdD de petite taille, il n'est pas inhabituel de devoir effectuer des modifications sur certaines données afin de pouvoir les y inclure³⁵, ce qui introduit une part importante d'arbitraire et fausse la correspondance de cette donnée au réel. Un besoin d'expertise est ainsi nécessaire pour nettoyer avec soin la BdD. Cette expertise ne saurait être attribuée à un concepteur de SAM dont les connaissances sur le phénomène modélisé peuvent être limitées. Des spécialistes du phénomène en question pourraient cependant être associés au projet pour le rendre proprement transdisciplinaire dans une perspective de *design* participatif.

3.2.3 Entraînement du modèle

Durant la phase d'entraînement, les paramètres du modèle sont optimisés par des successions d'opérations mettant en jeu une fraction des données d'entraînement. Ce modèle est mis à l'épreuve sur la base de performances définies par des métriques choisies à l'avance. Un point d'attention concerne le choix de la métrique de performance. La manière dont est mesurée la performance d'un système doit être en accord avec l'usage qui sera fait de celui-ci. Ce problème est très frappant dans le cadre du diagnostic d'une maladie grave par imagerie médicale [28]. Il n'existe pas dans ce cas de critère objectif pour arbitrer entre une minimisation des faux positifs ou des faux négatifs. Le choix doit s'effectuer sur la base du contexte d'usage du SAM (clinique ou recherche?) et de l'éthique associée (déontologique ou utilitariste?). Il y a donc un enjeu d'alignement entre les hypothèses sous-jacentes à l'entraînement et le cadre éthique dans lequel le SAM sera employé. La réflexion qui entoure le choix d'une métrique de performance

29. Le commentaire effectué par Deschamps *et al.* [10] reproche par exemple à Asselborn *et al.* [1] d'avoir construit un dispositif de diagnostic automatisé de la dysgraphie biaisé sur le plan de la représentativité. Dans ce cas de figure, les enfants dysgraphiques de la BdD sont considérés comme étant dysgraphiques à un niveau trop élevé par rapport à l'ensemble des enfants dysgraphiques.

30. Des initiatives sont proposées afin de pallier ce problème [23].

31. La voiture était pilotée par un SAM. Parmi les causes de l'accident, l'enquête a révélé l'incapacité du programme de reconnaissance d'obstacles à attribuer une catégorie adéquate à M^{me} Herzberg, retardant ainsi le déclenchement de la procédure de freinage d'urgence [21]. La classe "piéton à côté de son vélo" n'était pas prise en compte jusqu'ici dans les différentes annotations et l'algorithme n'avait donc pas été entraîné avec des exemples qui représentaient ce cas de figure précisément.

32. Par exemple, par l'examen des distributions d'une certaine variable.

33. Typiquement, les variables peuvent être centrées et réduites pour éviter de devoir comparer des variations sur des plages d'étendues trop différentes.

34. C'est notamment le cas pour des mesures temps-réel effectuées avec des matériels dont la fréquence d'échantillonnage diffère.

35. Ramener arbitrairement à zéro, ou à la moyenne des valeurs.

constitue donc également un véritable enjeu éthique. Pour le résoudre, une démarche de *design* participatif incluant les futurs usagers paraît tout à fait adaptée. Ceux-ci peuvent par eux-mêmes donner conscience aux *designers* des besoins liés à la mise en application concrète du dispositif. Dans une démarche d'éthique processuelle, la présence de différentes parties prenantes permet d'interroger des enjeux éthiques inhérents au projet dont les développeurs n'ont pas nécessairement conscience³⁶.

3.2.4 Test du modèle

Des données non utilisées dans la phase d'entraînement servent à tester le modèle pour vérifier s'il est performant sur de nouvelles données, évaluant ainsi son pouvoir de généralisation. L'étape de test du modèle sert à éviter le surapprentissage, c'est-à-dire que le modèle n'apprenne à résoudre le problème déterminé que sur le jeu de données qui l'a entraîné. Cependant, si l'ensemble de test provient de la même source que l'ensemble d'entraînement, il est possible qu'ils comportent tous deux des biais similaires, même s'ils sont bien disjoints³⁷. Des problèmes de performances ou des biais peuvent alors émerger lors de la transplantation du système vers un autre contexte d'usage. Une approche de *metadesign* où le système serait fait pour laisser le choix de régler l'optimisation du modèle selon la demande du contexte pourrait permettre d'éviter certains de ces écueils. Il convient alors de tester l'adéquation du modèle dans toutes les configurations possibles.

3.2.5 Intégration à un dispositif sociotechnique

Le modèle fonctionnel est inclus dans un dispositif au travers duquel l'utilisateur pourra interagir avec lui. Il est rendu accessible par la construction d'une interface logicielle et/ou matérielle. Cette dernière peut inclure des modules à même de fournir des explications à l'utilisateur³⁸.

De nombreuses questions se posent ici autour de l'interaction humain-machine et de la manière dont elle permet — ou non — le respect de principes éthiques désirables. La préservation de l'autonomie de jugement et d'action des usagers en lien avec l'automatisation de la tâche est notamment un point crucial [47]. L'utilisateur peut également être dupé par des explications qui seraient inexactes ou fallacieuses [26], et être amené à prendre des décisions erronées sur la base de preuves insuffisantes ou trompeuses [38]. Ces enjeux — parmi d'autres — renvoient notamment aux principes d'explicabilité et de transparence des algorithmes. Résoudre ces enjeux doit pouvoir amener à une plus grande maîtrise des dispositifs basés sur l'IA par les usagers³⁹. L'approche du *metadesign* peut permettre d'accomplir un

tel objectif, et ainsi éviter d'autres dérives, comme le paternalisme éthique auquel un processus de *design* trop strict pourrait mener, comme le suggère L. Floridi [14].

3.2.6 Déploiement et usage

Le système est mis à disposition de ses usagers et utilisé dans un contexte particulier. Les préoccupations en matière d'éthique de conception des SAM dépendent en partie des usages qui en découlent. Bien que l'usage qui est fait d'un système peut être indépendant de la volonté de ses *designers*⁴⁰, les éventuelles dérives doivent être anticipées et prises en compte au mieux possible⁴¹. En particulier, il apparaît crucial de se questionner sur la contribution apportée par un SAM dans le cadre d'une prise de décision. Les résultats prédictifs algorithmiques amènent à interrompre et diriger la procédure de prise de décision humaine collective, notamment dans le cadre hospitalier [3]. Cela pose des problèmes de gouvernance et d'attribution de responsabilité. Paradoxalement, ces résultats ne constituent pas en eux-mêmes des directives qui organisent le travail⁴². La prise en compte de la manière dont ces résultats sont présentés et justifiés paraît ainsi favoriser leur bon usage⁴³.

4 Conclusion

Nous avons souhaité contribuer à la discussion sur l'éthique de l'IA à partir de la notion d'éthique *by design*. Après avoir reconnu qu'elle peut s'envisager sous une pluralité de formes (dessin, définition, médiation), nous avons choisi de nous inscrire dans la caractérisation du *design* comme une relation de médiation entre les *designers*, les usagers et leurs milieux. Cette relation implique de considérer les divers degrés avec lesquels les usagers sont pris en compte dans les démarches de *design* : de l'abstraction à l'encapsulation, en passant par l'implication — du centrage sur l'utilisateur au *metadesign* en passant par le *design* participatif. Voir le *design* comme un processus nous a permis de considérer que l'éthique — pour être *by design* — doit symétriquement revêtir une approche processuelle pour considérer l'ensemble des questionnements éthiques inhérents à la conception de systèmes techniques. Elle est en cela complémentaire d'une éthique procédurale fondée sur des évaluations ponctuelles et *a posteriori* des systèmes. Dans le cas de la conception de SAM, une telle éthique processuelle se traduit par une attention continue portée sur de nombreux points de vigilance dont nous avons fourni une liste non exhaustive. La diversité et la complexité des activités de *design* d'un SAM, aussi bien que leurs effets significatifs sur l'espace social, requièrent en effet une vision de l'éthique renouvelée qui s'intègre tout au long du processus de *de-*

36. Sur un tout autre plan, notons que l'AM demande une quantité de calculs très importante et se trouve ainsi être très gourmand en énergie. Le choix du modèle, du nombre de paramètres à optimiser, ou encore le seuil de performance à atteindre influent tous sur le nombre d'opérations en jeu, et donc sur la consommation du modèle [49].

37. En outre, comme l'ensemble de test est fini, il n'est pas évident qu'il contienne l'entièreté des situations réelles d'usage ultérieures du SAM.

38. Il s'agit en général d'explications *a posteriori*, ou *post hoc* [33].

39. C'est quelque part dans la lignée de ce que suggèrent Kraemer *et al.* [28] en plaçant pour une augmentation des potentialités de paramétrage des systèmes par l'utilisateur.

40. On l'a vu avec l'apprentissage incrémental dans le cas de Tay.

41. Nos considérations n'occulent toutefois pas la nécessité d'un cadre éthique pour les usages de l'IA au-delà de son *design*.

42. Même si certains systèmes peuvent en eux-mêmes optimiser l'organisation du travail en automatisant la direction des salariés, comme dans le domaine de la logistique [17].

43. A. Christin remarque que de nombreux salariés s'engagent dans des pratiques de résistance routinière pour contrer et minimiser l'impact des SAM sur leur travail, dans les domaines du journalisme web et de la justice [7], ce qui n'est pas toujours concluant [18].

sign. En cela, nous croyons que notre approche est à même de dépasser le seul cadre des SAM, pour s'ouvrir à celui de l'IA en général. L'approche processuelle de l'éthique *by design* que nous préconisons peut en effet s'appliquer à tout processus de *design*. Néanmoins, notre approche analytique dans le cas des SAM permet de ne pas négliger des enjeux spécifiquement liés à l'AM, comme la possibilité d'un apprentissage continu et dynamique, ou l'inexplicabilité de certains résultats. Nous suggérons qu'une initiation à ces problématiques au sein des formations en ingénierie, ainsi qu'en science des données ou encore en santé publique, ne peut être qu'encouragée. Elle passerait d'abord par une prise de conscience de la portée éthique des pratiques de *design*. Repenser l'éthique de l'IA ce n'est pas seulement un questionnement des principes moraux ou du cadre éthique qu'il s'agit de valoriser mais c'est aussi réfléchir de façon critique à la façon de les mettre en œuvre.

Remerciements

L. D. a travaillé sur ce texte au sein de la chaire « Éthique & IA » soutenue par l'institut pluridisciplinaire en IA MIAI@Grenoble Alpes (ANR-19-P3IA-0003). L. D. a reçu, *via* le projet européen StorAIge, des fonds de l'Entreprise Commune ECSEL (EC) sous le numéro d'agrément n°101007321⁴⁴. É. P. a reçu le soutien financier du CNRS à travers les programmes interdisciplinaires de la MITI. Nous tenons à remercier T. GUYET, T. REVERDY et T. MÉNISIER, ainsi que les relecteurs anonymes pour leurs retours qui nous ont permis d'améliorer ce texte.

Références

- [1] Thibault ASSELBORN et al. « Automated human-level diagnosis of dysgraphia using a consumer tablet ». In : *npj Digital Medicine* 1.42 (2018), p. 1-9.
- [2] Kristine BÆRØE, Maarten JANSEN et Angeliki KERASIDOU. « Machine Learning in Healthcare : Exceptional Technologies Require Exceptional Ethics ». In : *The American Journal of Bioethics* 20.11 (2020), p. 48-51.
- [3] Simon BAILEY et al. « Dismembering organisation ». In : *Current Sociology* 68.4 (2020), p. 546-571.
- [4] Peter CHAPMAN et al. *CRISP-DM 1.0 : Step-by-step data mining guide*. SPSS, 2000.
- [5] Arthur CHARPENTIER. « L'éthique de la modélisation dans un monde où la normalité n'existe plus ». In : *Risques* 112 (2017), p. 117-121.
- [6] Benjamin CHIN-YEE et Ross UPSHUR. « Three Problems with Big Data and Artificial Intelligence in Medicine ». In : *Perspectives in Biology and Medicine* 62.2 (2019), p. 237-256.
- [7] Angèle CHRISTIN. « Algorithms in practice ». In : *Big Data & Society* 4.2 (2017).
- [8] Brandt DAINOW et Philip BREY. *Ethics By Design and Ethics of Use Approaches for Artificial Intelligence*. Sous la dir. d'EUROPEAN COMMISSION DG RESEARCH AND INNOVATION. 2021.
- [9] Richard DAZELEY et al. « Levels of explainable artificial intelligence for human-aligned conversational explanations ». In : *Artificial Intelligence* 299 (2021), p. 103525.
- [10] Louis DESCHAMPS et al. « Methodological issues in the creation of a diagnosis tool for dysgraphia ». In : *npj Digital Medicine* 2.36 (2019), p. 1-3.
- [11] M. C. ELISH et danah BOYD. « Situating methods in the magic of Big Data and AI ». In : *Communication Monographs* 85.1 (2017), p. 57-80.
- [12] Flora FISCHER. « L'éthique *by design* du numérique ». In : *Sciences du Design* n°10.2 (2019), p. 61.
- [13] Flora FISCHER. « Les normativités des technologies numériques : approche d'une éthique « *by design* » ». Thèse de doct. UTC Compiègne, 2020.
- [14] Luciano FLORIDI. « Tolerant Paternalism : Pro-ethical Design as a Resolution of the Dilemma of Toleration ». In : *Science and Engineering Ethics* 22.6 (2016), p. 1669-1688.
- [15] Luciano FLORIDI. « Translating Principles into Practices of Digital Ethics : Five Risks of Being Unethical ». In : *Philosophy & Technology* 32.2 (2019), p. 185-193.
- [16] Vilem FLUSSER. *The Shape of Things*. Reaktion Books, 1999.
- [17] David GABORIEAU. « « Le nez dans le micro ». Répercussions du travail sous commande vocale dans les entrepôts de la grande distribution alimentaire ». In : *La Nouvelle Revue du Travail* 1 (2012).
- [18] Ari Brendan GALPER. « Accommodation-through-Bypassing : Overcoming Professionals' Resistance to the Implementation of Algorithmic Technology ». Sociologie. MIT, 2020.
- [19] Benoît GODIN. « Making sense of innovation : from weapon to instrument to buzzword ». In : *Quaderni* 90 (2016), p. 21-40.
- [20] Xavier GUCHET. « L'éthique des techniques, entre réflexivité et instrumentalisation ». In : *Revue française d'éthique appliquée* 2.2 (2016), p. 8-10.
- [21] Andrew J. HAWKINS. « Serious safety lapses led to Uber's fatal self-driving crash, new documents suggest ». In : *The Verge* (2019).
- [22] Joseph HELLERSTEIN. « Quantitative Data Cleaning for Large Databases ». In : *United Nations Economic Commission for Europe*. T. 25. 2008, p. 1-42.
- [23] Ben HUTCHINSON et al. « Towards Accountability for Machine Learning Datasets ». In : *Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency*. ACM, 2021, p. 560-575.

⁴⁴ L'EC est soutenue par le programme de recherche et d'innovation Horizon 2020 de l'Union européenne et par la France, la Belgique, la République tchèque, l'Allemagne, l'Italie, la Suède, la Suisse et la Turquie.

- [24] Sheila JASANOFF. *The Ethics of Invention : Technology and the Human Future*. W.W.NORTON, 2016.
- [25] Anna JOBIN, Marcello IENCA et Effy VAYENA. « The global landscape of AI ethics guidelines ». In : *Nature Machine Intelligence* 1.9 (2019), p. 389-399.
- [26] Harmanpreet KAUR et al. « Interpreting Interpretability ». In : *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*. ACM, 2020, p. 1-14.
- [27] Charalampia (Xaroula) KERASIDOU et al. « Before and beyond trust : reliance in medical AI ». In : *Journal of Medical Ethics* (2021).
- [28] Felicitas KRAEMER, Kees van OVERVELD et Martin PETERSON. « Is there an ethics of algorithms ? » In : *Ethics and Information Technology* 13.3 (2011), p. 251-260.
- [29] Sylvain LAVELLE. « Politiques des artefacts. » In : *Cités* 39.3 (2009), p. 39.
- [30] Rena LEDERMAN. « The Ethical Is Political ». In : *American Ethnologist* 33.4 (2006), p. 545-548.
- [31] Raymond S. T. LEE. *Artificial Intelligence in Daily Life*. Springer Singapore, 2020.
- [32] Sabina LEONELLI. *La recherche scientifique à l'ère des Big Data*. Sesto San Giovanni : Mimésis, 2019.
- [33] Zachary C. LIPTON. « The Mythos of Model Interpretability ». In : *2016 ICML Workshop on Human Interpretability in Machine Learning*. 2016.
- [34] Clément Le LUDEC et al. « Quel statut pour les petits doigts de l'intelligence artificielle ? » In : *Les Mondes du travail* (2020), p. 99-110.
- [35] David LYELL et Enrico COIERA. « Automation bias and verification complexity : a systematic review ». In : *Journal of the American Medical Informatics Association* 24.2 (2016), p. 423-431.
- [36] Tom M. MITCHELL. *Machine Learning*. McGraw-Hill series in computer science. McGraw-Hill, 1997.
- [37] Brent MITTELSTADT. « Principles alone cannot guarantee ethical AI ». In : *Nature Machine Intelligence* 1.11 (2019), p. 501-507.
- [38] Brent MITTELSTADT et al. « The ethics of algorithms : Mapping the debate ». In : *Big Data & Society* 3.2 (2016), p. 205395171667967.
- [39] Stephanie L. MOORE. *Ethics by design*. HRD Press Inc, 2010.
- [40] Jessica MORLEY et al. « From What to How : An Initial Review of Publicly Available AI Ethics Tools, Methods and Research to Translate Principles into Practices ». In : *Science and Engineering Ethics* 26.4 (2019), p. 2141-2168.
- [41] Eirini NTOUTSI et al. « Bias in data-driven artificial intelligence systems—An introductory survey ». In : *WIREs Data Mining and Knowledge Discovery* 10.3 (2020).
- [42] Vanessa NUROCK, Raja CHATILA et Marie-Hélène PARIZEAU. « What Does “Ethical by Design” Mean ? » In : *Reflections on Artificial Intelligence for Humanity*. Springer International Publishing, 2021, p. 171-190.
- [43] Samir PASSI et Solon BAROCAS. « Problem Formulation and Fairness ». In : *Proceedings of the Conference on Fairness, Accountability, and Transparency*. ACM, 2019, p. 39-48.
- [44] Julie PERRIN et al. « En quête d'éthique ». In : *TSANTSA – Journal of the Swiss Anthropological Association* 25 (2020), p. 225-267.
- [45] Philippe PUCHERAL et al. « La Privacy by design : une fausse bonne solution aux problèmes de protection des données personnelles soulevés par l'Open data et les objets connectés ? » In : *LEGICOM* 56.1 (2016), p. 89-99.
- [46] Marco Tulio RIBEIRO, Sameer SINGH et Carlos GUESTRIN. « "Why Should I Trust You?" » In : *KDD '16 : Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining*. 2016, p. 1135-1144.
- [47] Monika SIMMLER et Ruth FRISCHKNECHT. « A taxonomy of human-machine collaboration ». In : *AI & Society* 1 (2021), p. 239-250.
- [48] Herbert SIMON. *The sciences of the artificial*. Cambridge, Mass : MIT Press, 1996.
- [49] Denis TRYSTRAM, Romain COUILLET et Thierry MÉNISSIER. « Apprentissage profond et consommation énergétique : la partie immergée de l'IA-ceberg ». In : *The Conversation* (2021).
- [50] Peter-Paul VERBEEK. « Accompanying technology : Philosophy of Technology after the Ethical Turn ». In : *Techné* 14 (2010), p. 49-55.
- [51] Sandra WACHTER, Brent MITTELSTADT et Chris RUSSELL. « Why fairness cannot be automated : Bridging the gap between EU non-discrimination law and AI ». In : *Computer Law & Security Review* 41 (2021), p. 105567.
- [52] Jane WAKEFIELD. « Microsoft chatbot is taught to swear on Twitter ». In : *BBC News* (2016).
- [53] Langdon WINNER. « Do Artifacts Have Politics ? » In : *Daedalus* (1980), p. 121-136.
- [54] Philip Matthias WINTER et al. *Trusted Artificial Intelligence : Towards Certification of Machine Learning Applications*. 2021. DOI : 10 . 48550 / ARXIV.2103.16910.
- [55] Steve WOOLGAR et Geoff COOPER. « Do Artefacts Have Ambivalence ». In : *Social Studies of Science* 29.3 (1999), p. 433-449.
- [56] Niina ZUBER et al. *Empowered and Embedded : Ethics and Agile Processes*. 2021.

L'intelligence artificielle peut-elle être saisie par le droit de l'Union européenne ?

Anaëlle Martin

Docteur en droit public à l'Université de Strasbourg, CEIE

anaelle.martin@outlook.com

Résumé

La présente étude a pour but d'exposer les rapports complexes entre l'Union européenne et l'intelligence artificielle, sous l'angle normatif, à l'heure où de nombreuses autorités ont publié des textes disparates visant à donner un cadre réglementaire — éthique et juridique — à l'utilisation des systèmes d'intelligence artificielle. Plus d'un an après la publication de la proposition de la Commission sur des règles harmonisées (législation sur l'intelligence artificielle) les institutions européennes éprouvent toujours des difficultés à appréhender en termes juridiques cette nouvelle technologie.

Mots-clés

Intelligence artificielle, Union européenne, Réglementation, Droit, Éthique

Abstract

The purpose of this study is to expose the complex relationship between the European Union and the artificial intelligence, from a normative perspective, at a time when many authorities have published a disparate set of texts aimed at providing a regulatory framework — ethical and legal — for the use of AI systems. About a year after the Commission's proposal for harmonised rules on artificial intelligence (Artificial Intelligence Act) European institutions are still facing difficulties in the legal understanding of this new technology.

Keywords

Artificial intelligence, European Union, regulation, Law, Ethics

1. Introduction

Dans une proposition de résolution de 2016¹, le Parlement européen n'hésitait pas à se référer aux célèbres *lois d'Asimov* pour attirer l'attention de la Commission européenne sur la nécessité d'adopter un cadre réglementaire au développement de la robotique et de l'intelligence artificielle. Au regard de leur caractère fictif, ces prétendues lois ne sauraient offrir une source d'inspiration pertinente pour l'élaboration d'un cadre légal susceptible d'atteindre le double objectif de promouvoir l'adoption de l'intelligence artificielle tout en tenant compte des risques associés à l'utilisation des nouvelles technologies. Si, en ce domaine, le droit a bien été saisi par la science-fiction, comme l'atteste la connotation juridique des *lois d'Asimov*, l'intelligence artificielle peut-elle être saisie par le

droit, en particulier le droit de l'Union, sans que les institutions européennes empruntent pour cela à la fiction ? Il est vrai que dans la sphère médiatique, l'intelligence artificielle ou son acronyme « l'IA » (qui sera retenu dans les développements qui suivent) appelle généralement un traitement plus folklorique et fantaisiste que scientifique. Ce constat vaut aussi bien pour les promesses — souvent excessives — que les risques — parfois fantasmagoriques — qui lui sont attachés. La question de savoir si l'Union européenne parviendra à adopter une réglementation convaincante de l'IA est d'actualité depuis que la Commission a publié, le 21 avril 2021, une première version de l'*AI Act*. La proposition répond aux appels du Parlement et du Conseil européens visant à l'adoption de mesures harmonisées pour assurer le bon fonctionnement du marché des systèmes d'IA, en mettant en balance leurs risques et bénéfices. De façon ambitieuse, le futur règlement vise à faire de l'Union un acteur mondial dans le développement d'une IA sûre, fiable et éthique². Dans sa proposition, la Commission affirme que les règles en matière d'IA devraient être « axées sur le facteur humain, de manière à ce que les personnes puissent avoir confiance dans le fait que la technologie est utilisée d'une façon sûre et conforme à la loi, notamment en ce qui concerne le respect des droits fondamentaux ». Ce projet législatif qui promeut une approche coordonnée des implications humaines de l'IA s'appuie sur les craintes suscitées par l'imprévisibilité, l'opacité et l'autonomie de certains systèmes d'IA, ainsi que les risques d'erreurs et de biais. Une intervention à l'échelle de l'Union s'impose pour éviter que des approches nationales hétérogènes créent une insécurité juridique et ralentissent la pénétration de l'IA sur le marché européen. En outre, seul un cadre réglementaire commun de l'IA peut garantir des conditions de concurrence équitables des entreprises, renforcer la compétitivité des États membres et protéger la souveraineté numérique de l'Europe. L'*AI Act* vise aussi à éviter qu'une disparité de normes nationales fragilise les valeurs de l'Union. Dans la section 2, nous rappelons l'état du droit positif, plus d'un an après la publication de l'*AI Act*. La section 3 énumère les principaux obstacles à l'adoption d'une réglementation commune adaptée aux défis posés par l'IA dans l'Union. La section 4 analyse l'approche singulière de la Commission pour appréhender l'IA en interrogeant les prémisses normatives de sa proposition. Dans la section 5, nous mettons en lumière les collisions entre la règle de droit et les nouvelles technologies du numérique en montrant que les concepts juridiques sont mis à mal et menacés par la logique et le langage de l'IA.

¹ Projet de rapport contenant des recommandations à la Commission concernant des règles de droit civil sur la robotique, Commission des affaires juridiques (2015/2103(INL)).

² Proposition de règlement du Parlement européen et du Conseil établissant des règles harmonisées concernant l'intelligence artificielle (législation sur l'intelligence artificielle).

2. L'état du droit positif à l'échelle européenne en matière d'IA

Par droit positif européen, nous entendons le droit de l'Union européenne (l'UE³) tel qu'il est (*de lege lata*) et non tel qu'on souhaiterait qu'il soit (*de lege ferenda*). Il s'agit d'adopter ici une approche descriptive — et non prescriptive — du droit de l'IA, tel qu'il s'élabore actuellement dans les instances européennes. Aussi nous bornerons-nous à présenter l'état des lieux du droit de l'UE en matière du numérique, ce qui implique de situer le projet législatif sur l'IA dans un corpus normatif plus vaste. Les orientations politiques de l'UE, pour la présente décennie, visent une « Europe adaptée à l'ère du numérique », pour reprendre l'expression employée par la Présidente de la Commission européenne Ursula von der Leyen⁴. L'objectif de la Commission est de renforcer la souveraineté technologique de l'Europe en mettant l'accent sur les technologies, les infrastructures et les données. Cette ambition a poussé les institutions européennes à adopter des textes majeurs, à la suite du célèbre RGPD⁵, au premier rang desquels la proposition de règlement sur l'intelligence artificielle du 21 avril 2021 (*AI Act*⁶) — dont il sera principalement question dans cette étude — ainsi que les règlements sur les données (*Data Act*⁷), la gouvernance des données (*Data Governance Act*⁸), les services numériques (*Digital Services Act*⁹), les marchés numériques (*Digital Markets Act*¹⁰) et la cybersécurité (*Cybersecurity Act*¹¹).

À ce stade de la présentation, trois précisions s'imposent. Tout d'abord, l'*AI Act* ne doit pas être analysé isolément car ses dispositions renvoient à l'application d'autres actes de l'UE, comme le RGPD dont il est un complément. La proposition de règlement sur l'IA s'inscrit ainsi dans le cadre plus large d'un ensemble de mesures visant à résoudre les problèmes liés à l'utilisation de l'IA. L'*AI Act* vise à assurer la cohérence avec d'autres initiatives de l'UE — en cours ou prévues — comme la révision de la législation sur les produits¹² (la directive *Machines*¹³ et la directive sur la sécurité des produits¹⁴). Ces actions normatives s'efforcent de favoriser la mise en place

d'un écosystème de confiance pour l'IA au sein de l'UE.

Ensuite, eu égard à la complexité du processus législatif dans l'UE, il importe de souligner que les divers textes adoptés par les institutions en matière du numérique ne sont pas au même stade normatif car si certains sont applicables depuis quelques années déjà, comme le RGPD (depuis le 25 mai 2018), ou définitivement adoptés, comme le *Cybersecurity Act* (en vigueur depuis juin 2019), la plupart sont sous une forme embryonnaire ou en cours de négociations. C'est le cas de l'*AI Act* dont les dispositions font, à l'heure actuelle, l'objet d'après débats. S'agissant des autres législations en cours d'adoption, en lien avec le règlement sur l'IA, relevons que le Parlement européen et le Conseil de l'Union ont conclu un accord, le 24 mars 2022, sur le *DMA*, et le 23 avril 2022, sur le *DSA*. Rappelons que le *DSA* vise à encadrer les plateformes numériques pour lutter contre la désinformation et la haine sur les réseaux sociaux¹⁵, tandis que le *DMA* tend à prévenir les abus de position dominante sur les marchés numériques¹⁶. En raison de l'ambition de ces textes de s'ériger comme une muraille *anti-GAFAM*, c'est fort à propos que la Présidente de la Commission a pu qualifier leur adoption d'« historique ». En effet, ces derniers sont appelés, à l'instar du RGPD, à servir de modèles au niveau mondial en matière de réglementation numérique, et à aligner les concurrents de l'UE sur ses règles de marché. Le volet *données* de la stratégie européenne du numérique a également conduit la Commission à dévoiler, le 25 novembre 2021, le *Data Governance Act*, dans le but de définir des règles de gouvernance en créant des structures pour faciliter les données, et, le 23 février 2022, le *Data Act*, afin de fixer les règles d'accès aux données. Si le premier projet a été adopté le 16 mai 2022 par le Conseil¹⁷, le second est en cours de discussions¹⁸. De façon plus sectorielle, mais tout aussi ambitieuse, la Commission cherche à déployer des espaces européens des données dans des secteurs stratégiques comme la santé. À ce titre, l'espace européen des données de santé (*European Health Data Space*) qui aura pour objet d'améliorer l'utilisation des données de santé dans l'Union, à des fins de recherche et d'innovation, constitue le troisième temps du plan de la Commission sur les données.

³ Le droit de l'UE, anciennement « droit communautaire », composé du droit primaire (traités de l'UE), du droit dérivé (directives et règlements) ainsi que de la jurisprudence de la Cour de justice de l'UE (CJUE) est à distinguer du droit européen tel qu'il découle de la jurisprudence de la Cour européenne des droits de l'Homme (CEDH) établie par le Conseil de l'Europe. Pour des non juristes, des confusions sont possibles dans la mesure où les institutions de l'UE et celles du Conseil de l'Europe, qui sont deux organisations dites « régionales », ont produit une littérature conséquente sur l'IA.

⁴ Une Union plus ambitieuse. Mon programme pour l'Europe. Orientations politiques pour la prochaine Commission européenne 2019-2024.

⁵ Règlement (UE) 2016/679 relatif à la protection des personnes physiques à l'égard du traitement des données à caractère personnel et à la libre circulation de ces données.

⁶ Proposition de règlement du Parlement européen et du Conseil établissant des règles harmonisées concernant l'intelligence artificielle (législation sur l'intelligence artificielle).

⁷ Proposition de règlement du Parlement européen et du Conseil fixant des règles harmonisées pour l'équité de l'accès aux données et de l'utilisation des données (règlement sur les données).

⁸ Proposition de règlement du Parlement européen et du Conseil sur la gouvernance européenne des données (acte sur la gouvernance des données).

⁹ Proposition de règlement du Parlement européen et du Conseil relatif à un marché intérieur des services numériques (DSA).

¹⁰ Proposition de règlement du Parlement européen et du Conseil relatif aux marchés contestables et équitables dans le secteur numérique (DMA).

¹¹ Règlement (UE) 2019/881 du Parlement européen et du Conseil du 17 avril 2019 relatif à l'ENISA et à la certification de cybersécurité des technologies de l'information et des communications.

¹² Proposal for a regulation of the European Parliament and of the Council on machinery products.

¹³ Directive 2006/42/CE du Parlement européen et du Conseil du 17 mai 2006 relative aux machines et modifiant la directive 95/16/CE.

¹⁴ Directive 2001/95/CE du Parlement européen et du Conseil du 3 décembre 2001 relative à la sécurité générale des produits.

¹⁵ Le DSA vise à instaurer un cadre juridique efficace et proportionné à la modération des contenus par les plateformes et à renforcer les obligations des places de marché en ligne.

¹⁶ Le DMA a pour objet la régulation économique des grandes plateformes numériques, afin de doter l'UE d'un instrument de régulation *ex ante* dédié aux contrôleurs d'accès, et qui vise à assurer la contestabilité et l'équité des marchés numériques.

¹⁷ Le 6 avril 2022 par le Parlement européen.

¹⁸ Au-delà des aspects économiques relatifs aux données générées par les objets connectés, ce texte devra s'inscrire dans le paradigme de protection des données tel qu'il résulte du RGPD, ce qui pose, comme souvent, un problème d'articulation entre les différentes normes.

Enfin, une dernière précision mérite d'être apportée sur la notion même de « droit de l'IA » (car c'est bien de ça qu'il s'agit). L'*AI Act* n'étant pas de génération spontanée, ce projet de règlement prend appui sur une vaste littérature consacrée à l'IA. Mais contrairement à la plupart des textes générés jusqu'à présent par les institutions européennes, ainsi que les groupes d'experts mandatés par l'UE, le projet de règlement de la Commission revêt une valeur juridique forte dans la mesure où il s'agit d'établir un cadre normatif contraignant pour le développement et l'utilisation de l'IA en Europe. L'*AI Act* entend se démarquer, ontologiquement, des résolutions, recommandations, avis, lignes directrices et livres verts ou blancs, lesquels sont dépourvus de portée autre que politique ou symbolique. Ainsi, si la proposition législative sur l'IA s'inscrit dans un paquet comprenant une communication¹⁹ (droit souple), ledit paquet comporte une autre proposition de réglementation sur les machines²⁰ (droit dur).

Pour avoir une idée claire de l'état du droit positif de l'UE en matière d'IA, il est nécessaire de garder à l'esprit les trois remarques précédentes, à savoir la question de l'articulation des différentes dispositions de la réglementation européenne du numérique (l'*AI Act* s'insère dans un vaste réseau normatif en constante évolution) ; la nécessité de distinguer le droit en vigueur du droit en cours d'élaboration (l'*AI Act* n'étant pas un texte définitif, des modifications sont possibles) ; la nécessité de distinguer le *soft law* (recommandation, lignes directrices) du *hard law* (l'*AI Act* appartenant à cette dernière catégorie).

Ces précisions étant faites, il est loisible de dresser un rapide panorama de la stratégie européenne dans laquelle s'inscrit la proposition de règlement sur l'IA. Parallèlement aux stratégies nationales sur le numérique²¹, l'UE poursuit, à son échelle, une orientation que la Commission a révélée en avril 2018²². Dans ce document d'une vingtaine de pages, la Commission affirme que tous les ingrédients sont réunis pour que l'UE joue un rôle prédominant dans la révolution de l'IA « sur la base de ses propres valeurs ». Estimant que les économies les plus développées ont adopté des approches de l'IA qui reflètent leurs systèmes politique, économique, culturel et social, la Commission souligne la nécessité pour les Européens d'unir leur force afin que l'Europe fasse partie de la transition numérique. Le contexte international compétitif oblige les États membres à agir au niveau de l'Union (plutôt qu'au niveau national) dans la mesure où les concurrents, américain

et asiatique, ont adopté des stratégies visant à en faire des leaders mondiaux. L'urgence à agir est d'autant plus grande que le Vieux Continent accuse un retard en matière d'investissements privés, à l'heure où les entreprises extra-européennes investissent massivement et exploitent de grandes quantités de données. Le 19 février 2020, la Commission publiait, en même temps qu'elle dévoilait sa stratégie sur les données²³, son *Livre blanc* sur l'IA²⁴. Il y est affirmé que l'UE souhaite élaborer un cadre de normes et de principes pour une économie numérique éthique. La transition numérique exige un cadre politique et des infrastructures appropriées pour permettre à l'ambition européenne de s'exprimer, en créant un marché unique européen de la donnée, tout en tenant compte des risques liés au caractère potentiellement discriminatoire des algorithmes. À en croire la Commission, la spécificité de la stratégie européenne, par rapport à celle des autres leaders nationaux, réside dans l'importance accordée au respect des droits de l'homme²⁵. L'UE entend, par conséquent, élaborer un ensemble de principes numériques pour renforcer les droits de ses citoyens. Cette ambition d'une « IA made in Europe », pour reprendre l'expression consacrée par le Conseil de l'Union²⁶, peut s'appuyer sur des règles de droit positif efficaces relatives à la protection des données. Le RGPD en est le pilier central dans la mesure où ce règlement garantit la libre circulation des données à caractère personnel à l'intérieur de l'UE, tout en assurant un niveau élevé de protection²⁷.

À cet égard, l'actuelle présidence du Conseil de l'UE (janvier-juin 2022²⁸) est l'occasion, pour la France, de peser dans les négociations sur l'*AI Act* et de réaffirmer sa stratégie pour une souveraineté numérique européenne²⁹. À l'instar du *DSA* et du *DMA*, qui furent au menu de la présidence française du Conseil, l'*AI Act* est au centre des préoccupations de la France qui plaide en faveur d'une adoption rapide d'une version finale de ce texte. Si le développement d'une IA *digne de confiance* est l'une de ses priorités, la France soutient l'émergence, à l'échelle européenne, d'un écosystème propice à l'émergence d'entreprises innovantes, susceptibles d'être les champions technologiques de demain. En outre, il ressort des débats que les négociateurs français souhaitent rendre le règlement sur l'IA plus flexible, notamment en ce qui concerne son recours par les forces de l'ordre. La position française sur l'aspect répressif de l'*AI Act* n'étant pas partagée par l'ensemble des États membres, les discussions s'annoncent difficiles.

¹⁹ Communication de la Commission au Parlement européen, au Conseil, au Comité économique et social européen et au Comité des régions favoriser une approche européenne en matière d'intelligence artificielle, Bruxelles, le 21.4.2021 COM(2021) 205 final.

²⁰ Proposal for a regulation of the European Parliament and of the Council on machinery products.

²¹ Van Roy, V. (2020). AI Watch - National strategies on Artificial Intelligence: A European perspective in 2019, JRC Technical Report.

²² Communication de la Commission au Parlement européen, au Conseil européen, au Conseil, au Comité économique et social européen et au Comité des régions. L'intelligence artificielle pour l'Europe SWD(2018) 137 final.

²³ Communication de la Commission au Parlement européen, au Conseil, au Comité économique et social européen et au Comité des régions : Une stratégie européenne pour les données. COM(2020) 66 final.

²⁴ LIVRE BLANC Intelligence artificielle. Une approche européenne axée sur l'excellence et la confiance. COM(2020) 65 final.

²⁵ Cette affirmation est à nuancer dans la mesure où la plupart des stratégies nationales sur l'IA insistent sur la nécessité de respecter les droits fondamentaux.

²⁶ Conseil de l'Union européenne, Intelligence artificielle? Conclusions sur le plan coordonné dans le domaine de l'intelligence artificielle – 6177/19, 2019.

²⁷ Les dispositions sur la prise de décision fondée sur le traitement automatisé, dont le profilage, prévoient que les personnes ont le droit d'obtenir des informations utiles sur la logique sous-jacente. Le RGPD confère aux personnes concernées le droit de ne pas faire l'objet d'une décision fondée exclusivement sur un traitement automatisé, sauf dans certains cas. En conférant aux citoyens des droits supplémentaires et en exigeant plus de transparence, le règlement garantit une plus grande responsabilité des acteurs du traitement des données. Sur une base harmonisée, il dote les autorités indépendantes de solides pouvoirs d'exécution et instaure un nouveau système de gouvernance. Enfin, le RGPD garantit une égalité des conditions de concurrence pour les sociétés exerçant leurs activités sur le marché européen, indépendamment du lieu d'établissement.

²⁸ Conférence de presse du président de la République Présentation de la présidence française du Conseil de l'Union européenne : <https://www.elysee.fr/admin/upload/default/0001/12/26ab8ecfe7127e8fd3100f18dc4a38a16d47e69.pdf>

²⁹ <https://www.intelligence-artificielle.gouv.fr/fr/strategie-nationale/la-strategie-nationale-pour-l-ia>.

3. Les obstacles à l'adoption d'un cadre normatif harmonisé et efficace de l'IA

L'adoption définitive d'un règlement visant à établir des règles harmonisées en matière d'IA afin de préserver les droits fondamentaux et les valeurs de l'UE ne se fera pas sans d'intenses controverses, eu égard à l'extrême sensibilité du sujet. À l'heure où ces lignes sont écrites, des milliers d'amendements ont été déposés par les acteurs institutionnels de l'UE. Si la Commission s'est inspirée, dans sa proposition législative, des résolutions du Parlement européen³⁰, il faudra compter avec les positions des États membres, et leur stratégie nationale respective, pour espérer doter l'IA d'un cadre légal à l'échelle de l'Europe. Il est vrai que l'objectif poursuivi par l'UE relève de la gageure puisqu'il s'agit de donner aux citoyens européens une confiance dans les systèmes d'IA, en leur apportant des garanties solides en termes de droits et libertés, tout en encourageant les opérateurs à développer l'IA afin de faire de l'UE un leader mondial dans ce secteur hautement stratégique. Nous mentionnerons, dans les sous-sections suivantes, deux types de limites ou vulnérabilités susceptibles de faire obstacle à une adoption rapide de l'*AI Act*.

3.1 Les obstacles technico-juridiques

Sans qu'il s'agisse, à proprement parler, d'un obstacle à l'adoption d'un futur règlement, il importe de relever l'abondance de textes produits en matière d'IA ces dernières années. Les multiples chartes sur l'IA, dépourvus de caractère contraignant³¹, tendent à brouiller les frontières entre le droit et l'éthique, notamment lorsque les droits fondamentaux y sont mentionnés. Sans évoquer la doctrine dont le rôle est précisément de produire de la littérature, les institutions de l'UE, au premier chef desquelles la Commission et le Parlement, se sont elles aussi attachées à multiplier les communications, tantôt pour partager des craintes sur les risques liés à l'exploitation des algorithmes, tantôt pour souligner les bénéfices à recourir à certains systèmes d'IA. L'inflation de cette littérature grise sur l'IA, accentuée par les recommandations des groupes d'experts, ne concerne pas que l'UE puisque l'on retrouve au sein d'instances nationales et internationales, publiques et privées, un phénomène similaire d'accroissement de textes programmatiques ou incitatifs. Si la volonté de contrôler ou réguler l'IA est louable, la production normative qui tend, pour des raisons évidentes, à privilégier la souplesse du *soft law*, au détriment de la rigidité d'un cadre réglementaire classique, contribue à saper la qualité, l'autorité et la légitimité d'un droit perçu, de plus en plus, comme *floou*.

Il s'agit peut-être là d'un premier indice de l'impuissance du juridique à saisir l'IA dans toute sa complexité technique. Et c'est ce que la Commission semble implicitement admettre lorsqu'elle met en avant l'importance des codes de conduite. Ces codes qui peuvent inclure des engagements volontaires liés à la durabilité environnementale, l'accessibilité pour les personnes handicapées et la diversité des équipes sont, dans certaines conditions, créés et mis en oeuvre par les fournisseurs eux-mêmes. Nous aurons l'occasion d'y revenir. Bornons-nous, à ce stade, à souligner que le caractère abstrait et l'indétermination relative de la norme de droit s'avèrent problématiques pour appréhender les outils numériques.

Le principal obstacle « technico-juridique », à l'adoption d'un règlement européen sur l'IA, réside assurément dans le choix des définitions retenues par la Commission, à commencer par celle de l'IA. La première version de l'*AI Act* opte ainsi pour l'expression « système d'IA » plutôt qu'« IA », cette dernière renvoyant davantage à une discipline théorique. Les systèmes d'IA sont définis, à l'article 3, comme des logiciels développés au moyen de techniques énumérées à l'annexe I³² qui « peuvent, pour un ensemble donné d'objectifs définis par l'homme, générer des résultats tels que des contenus, des prédictions, des recommandations ou des décisions influençant les environnements avec lesquels il interagit »³³. Il est permis de s'interroger sur l'opportunité d'ajouter la précision selon laquelle les objectifs doivent être « définis par l'homme ». Cette précaution semble faire écho aux fictives lois de la robotique d'Asimov auxquelles le Parlement européen s'était, dans une de ses résolutions, montré particulièrement sensible. Cette définition provisoire pourrait être allégée, ou précisée, notamment dans ses aspects les plus superficiels comme le fait que le résultat influence l'environnement. Il en est de même de la liste indicative des résultats attendus (recommandation, prédiction, décision) qui semble trahir un tâtonnement conceptuel de la part des rédacteurs. Définir l'IA est central pour l'Union car les définitions déterminent le champ d'application du texte et visent à permettre à certains produits et services de circuler librement sur le marché européen, à condition d'être conformes aux prescriptions contenues dans le futur règlement. L'enjeu est d'autant plus grand qu'une définition large risquerait d'engendrer de l'insécurité juridique et de potentiels conflits d'interprétations. Une telle approche serait préjudiciable pour l'innovation et le commerce des produits basés sur l'IA. Une définition souple permettrait, quant à elle, de tenir compte des progrès techniques et de la dynamique inhérente à la matière technologique, mais un minimum de précision est nécessaire si l'on souhaite garantir la sécurité juridique.

³⁰ Les 3 résolutions du 20 octobre 2020 adoptées par le Parlement européen sont les suivantes : Résolution contenant des recommandations à la Commission sur un régime de responsabilité civile pour l'intelligence artificielle (2020/2014(INL)) ; Résolution contenant des recommandations à la Commission concernant un cadre pour les aspects éthiques de l'intelligence artificielle, de la robotique et des technologies connexes (2020/2012(INL)) ; Résolution sur les droits de propriété intellectuelle pour le développement des technologies liées à l'intelligence artificielle (2020/2015(INI)).

³¹ Contrairement à la Charte des droits fondamentaux de l'Union qui lie les institutions européennes et les États membres lorsqu'ils appliquent le droit de l'UE.

³² L'annexe I liste 3 catégories de systèmes : les systèmes auto-apprenants (*machine learning*), les systèmes logiques et les systèmes statistiques.

³³ On relèvera que la présente définition diffère sensiblement des précédentes, la Commission ayant proposé, par le passé, de définir l'IA comme des systèmes soit purement logiciels, agissant dans le monde virtuel, soit intégrés à des dispositifs matériels. Cette première définition a été retravaillée par le groupe d'experts « de haut niveau », lequel a proposé de définir l'IA comme des « systèmes logiciels (et éventuellement matériels) conçus par des êtres humains et qui, ayant reçu un objectif complexe, agissent dans le monde réel ou numérique en percevant leur environnement par l'acquisition de données, en interprétant les données structurées ou non structurées collectées, en appliquant un raisonnement aux connaissances, ou en traitant les informations, dérivées de ces données et en décidant des meilleures actions à prendre pour atteindre l'objectif donné. Les systèmes d'IA peuvent soit utiliser des règles symboliques, soit apprendre un modèle numérique. Ils peuvent adapter leur comportement en analysant la manière dont l'environnement est affecté par leurs actions ».

Il reste que par rapport aux droits nationaux, le droit de l'UE se caractérise par son pragmatisme et une vision technico-fonctionnel des domaines qu'il appréhende. Droit du marché intérieur, le droit européen reste marqué du sceau de l'économie et par sa finalité intégrative. Ainsi, si la définition de l'IA, posée à l'article 3 de l'*AI Act* pêche par le caractère vague de la formulation, le règlement comporte une annexe énumérant les différentes « techniques et approches d'IA », comme l'apprentissage automatique (a), les approches fondées sur la logique et les connaissances (b) et les approches statistiques (c). La technique de l'énumération est une approche se voulant à la fois plus concrète et plus souple qu'une définition *in abstracto*. Plus concrète en ce qu'elle permet aux opérateurs de savoir, sans ambiguïté, si leur système relève du champ d'application du règlement, puisqu'il leur suffit de se référer à la liste figurant dans l'annexe I. Plus souple dans la mesure où cette liste est modifiable à tout moment et selon une procédure simplifiée. Une telle définition de l'IA présente également l'avantage, pour l'Union, de garder la maîtrise (relative) de son texte. Encore doit-il passer l'étape de son adoption. Or, la définition retenue soulève d'ores et déjà des objections et suscite un débat parmi les États membres, et au sein de ces derniers. La précédente présidence du Conseil³⁴, à savoir celle de la Slovence (juillet-décembre 2021) avait ainsi cherché à retravailler l'aspect définitionnel de l'*AI Act*. L'État en question avait, dans un premier temps, proposé de restreindre la définition de l'IA à l'apprentissage automatique. Les textes de compromis arrêtés par la présidence slovène ont, par la suite, mis l'accent sur la définition des systèmes d'IA pour mieux les distinguer des programmes logiciels classiques. Les systèmes d'IA y étant définis comme ayant la capacité de traiter des données pour déduire la manière d'atteindre un ensemble d'objectifs définis par l'homme par apprentissage, raisonnement ou modélisation.

En tout état de cause, le choix d'une définition relativement large de l'IA, comme celle proposée par la Commission, ne signifie pas que l'UE entend réguler tous les systèmes d'IA. L'article 2 du projet indique que le règlement ne s'applique pas aux systèmes utilisés à des fins militaires. Dans la mesure où l'IA a connu un développement fulgurant dans ce domaine, et face aux craintes de l'emploi de systèmes d'armes létales « autonomes », d'aucuns jugent nécessaire d'élaborer un cadre juridique au niveau de l'UE. Mais pour espérer y parvenir, les États membres doivent d'abord se mettre d'accord sur les chantiers législatifs en cours. Or, la Commission européenne peine à faire émerger son projet en raison de son degré élevé de technicité, accentué par le problème juridique des rapports du futur règlement avec le droit pré-existant, rendant difficile le consensus. Le manque de précision de l'*AI Act* sur la question de l'articulation de ses dispositions avec le RGPD, mais également la directive « police-justice »³⁵, peut également être rangé parmi les limites technico-juridiques.

Les défis définitionnels concernent également d'autres notions comme les systèmes d'IA dits à « haut risque », lesquels appellent un régime juridique plus exigeant que les systèmes d'IA perçus comme présentant un risque faible (obligations de transparence³⁶) ou minimum (codes de conduite³⁷). Si les premiers sont soumis à des obligations contraignantes (évaluation de conformité³⁸), les seconds bénéficient d'un traitement allégé, notamment lorsque les fournisseurs sont seulement « encouragés » à élaborer des codes de conduite destinés à favoriser l'application « volontaire » de certaines exigences aux systèmes autres que les systèmes d'IA à haut risque. Eu égard aux enjeux, le travail sur la définition et la classification/catégorisation revêt une grande importance pour les opérateurs. À ce titre, l'article 6 renvoie, à l'instar de l'article 3, à des annexes (annexe II et III) pour énumérer les systèmes d'IA considérés à « haut risque ». L'*AI Act* s'appuie sur une technique légistique hybride qui combine une stratégie se réclamant de la « neutralité technologique »³⁹ à une approche « par les risques » afin de réguler l'intensité de la réglementation (en fonction du degré du risque identifié). Il ressort, là encore, des documents publiés par les présidences successives du Conseil (slovène, française) que les définitions sont jugées vagues et appellent des orientations pratiques sur l'identification des critères des systèmes dits « à haut risque ».

Le choix pragmatique de l'énumération, tant pour ce qui concerne la définition des systèmes d'IA, en général, que les systèmes d'IA à « haut risque », facilite l'appréhension (et la compréhension) de cette technologie par l'UE. Sans doute permet-elle au législateur européen de saisir les systèmes d'IA, en en cernant les contours, pour mieux les contrôler et les réguler. Cette technique n'est toutefois pas sans risque. Outre qu'elle conduit « à une forme de réification de l'IA, de liste et de stock de produits estimés à risques dont les critères d'arbitrages seront fluctuants », elle permet également de légitimer, par ce biais, des systèmes d'IA au régime juridique incertain⁴⁰. Il en va ainsi des polygraphes pour analyser l'état émotionnel d'une personne physique, pour la gestion de la migration et le contrôle aux frontières⁴¹, lesquels ne sont pas reconnus dans tous les ordres juridiques des États membres. En plus de poser des difficultés sur le plan technique, juridique et scientifique (en raison de l'absence de consensus) cette méthode est discutable sous l'angle éthique. Ajoutons que la difficulté de qualifier *en droit* l'IA, afin de lui conférer un régime juridique adapté, est aggravée par le fait qu'il s'agit d'adopter une définition commune à l'ensemble des États membres. L'ambition d'une approche uniforme, pour l'UE tout entière, est nécessairement entravée par des obstacles politiques qui ne cessent de se dresser à mesure que les négociations avancent. En effet, plus qu'un objet technique, une technologie ou un domaine scientifique, le terme d'IA désigne un être de fiction et de fantasme.

³⁴ La présidence du Conseil de l'Union européenne est une présidence tournante. Actuellement, et jusqu'à la fin du mois de juin, la France assure cette fonction.

³⁵ Directive (UE) 2016/680 du Parlement européen et du Conseil du 27 avril 2016 relative à la protection des personnes physiques à l'égard du traitement des données à caractère personnel par les autorités compétentes à des fins de prévention et de détection des infractions pénales.

³⁶ Article 52 de l'*AI Act*.

³⁷ Article 69 de l'*AI Act*.

³⁸ L'article 16 du projet de règlement détaille les obligations incombant aux fournisseurs de systèmes d'IA à haut risque, notamment l'évaluation de la conformité.

³⁹ Vincent Gautrais, *Neutralité technologique: rédaction et interprétation des lois face aux technologies*, Thémis, Montréal, mai 2012.

⁴⁰ Benbouzid, B., Meneceur, Y. & Smuha, N. quatre nuances de régulation de l'IA: Une cartographie des conflits de définition. *Réseaux*, 232-233, 29-64 (2022)

⁴¹ Voir l'Annexe III, 7 a) et b).

3.2 Les obstacles éthico-politiques

Lors des deux dernières présidences tournantes du Conseil, la Slovénie et la France ont partagé des textes de compromis dans lesquels les États membres réaffirmaient leur compétence exclusive en matière de sécurité nationale, insistant pour que les systèmes développés pour la recherche scientifique soient exclus du champ d'application de l'*AI Act*, tout comme ceux développés à des fins militaires. La sensibilité politique d'un domaine comme l'IA s'exprime également dans l'approche d'équilibriste de la Commission consistant à rechercher un compromis entre des exigences minimales pour encadrer les risques liés à l'IA (sur la sécurité, la santé et les droits fondamentaux⁴²), et l'établissement d'un cadre suffisamment souple pour ne pas entraver son développement. La poursuite d'objectifs aussi diamétralement opposés s'apparente à la recherche de la quadrature du cercle, comme tendent à le montrer les réactions contradictoires suscitées, dans la société, par la publication du projet de règlement. Si les associations de défense des libertés jugent le texte laxiste sur la garantie des droits de l'homme⁴³, entrepreneurs et investisseurs du numérique s'inquiètent de l'impact négatif sur les *startups*. Les deux camps s'accordent, par ailleurs, à dénoncer le manque de clarté du texte de la Commission. L'UE étant, avant tout, un marché, les critiques formulées par les opérateurs économiques ne sont pas à prendre à la légère car une réglementation trop contraignante pourrait décourager les entreprises d'investir dans l'IA. Des reproches similaires ont été émis contre le RGPD, dont les dispositions, axées sur la protection des droits des personnes, ont été accusées d'entraver le développement de certaines technologies⁴⁴.

Cette tension entre protection des principes éthiques et juridique, d'une part, et innovation, d'autre part, parcourt les positions, et oppositions, des États membres. Des divergences de stratégies sur l'IA risquent d'envenimer les discussions et freiner l'adoption de l'*AI Act*. En tant que technologie hautement politique, l'IA présente un enjeu de souveraineté nationale et internationale. La création des conditions favorables à son développement est toutefois éclipsée, dans la sphère publique et médiatique, par celle des libertés individuelles et collectives. Les « experts » s'alarment, régulièrement, des effets négatifs du numérique sur les individus et la société. Ainsi, le fait que l'UE autorise, même dans des conditions strictes, le recours aux systèmes d'identification biométrique à distance suscite la controverse. À cet égard, l'utilisation de cette technologie à des fins répressives divise les États. Si l'Allemagne et la Finlande souhaitent traiter cette question séparément, la France, qui siège actuellement à la présidence du Conseil, fait pression pour assouplir son usage, au titre de l'*AI Act*, par les forces de

l'ordre. Les autorités françaises demandent à élargir les cas dans lesquels la reconnaissance biométrique peut être utilisée, en supprimant la référence au caractère imminent de la menace et en élargissant les possibilités d'utiliser l'IA pour localiser un suspect dans le cadre d'une enquête policière (sans que l'infraction pénale relève nécessairement du champ du mandat d'arrêt européen). De profondes divergences sont également apparues entre la France et l'UE, sur la question de la recherche d'un équilibre entre le respect de la vie privée et la sécurité publique. Dans un arrêt du 6 octobre 2020⁴⁵, le juge européen, en réponse à une question du Conseil d'État, appelait à poser un cadre protecteur, en vertu de la directive *ePrivacy* et du RGPD. L'enjeu était de taille car il s'agissait de mettre en cause, sur la base du droit de l'UE, le droit français qui prévoyait l'obligation de conserver de façon indifférenciée les données de connexion pour la poursuite des infractions pénales. Était aussi en cause, l'utilisation par les services de renseignement, de certains algorithmes. En retenant une approche protectrice des données, le droit de l'UE conduisait à condamner l'obligation de conservation généralisée de celles-ci. Cette lecture de la CJUE heurtait le droit constitutionnel et la politique sécuritaire de la France. Le 21 avril 2021⁴⁶, soit le jour de la publication de l'*AI Act*, le Conseil d'État, sous la pression du gouvernement, décidait de faire primer les exigences de la Constitution tenant au respect de la sauvegarde de la Nation, la recherche des auteurs d'infractions et la lutte contre le terrorisme pour s'opposer à la vision respectueuse des droits fondamentaux, prônée par la Cour de Luxembourg.

Éminemment politiques, ces questions touchent aussi, et surtout, à l'éthique. Si leur développement n'est pas contrôlé, certains systèmes d'IA pourraient causer d'irréversibles préjudices à nos sociétés, en refaçonnant le modèle de la démocratie et de l'État de droit. Ces craintes ont sous-tendu certaines critiques émises à l'encontre du projet de règlement. Certains observateurs ont fait remarquer qu'au-delà de la prise en considération des droits de l'homme, l'*AI Act* ne tenait pas suffisamment compte des préjudices sociétaux⁴⁷. D'autres craintes moins fondées⁴⁸, et d'une légitimité toute relative, bien que fortement relayées, pointent une prétendue lacune du projet de règlement : l'IA générale. Dans la lignée des spéculations de Nick Bostrom⁴⁹, dont les thèses ont pénétré jusqu'aux groupes d'experts mandatés par l'UE⁵⁰, Max Tegmark n'a pas hésité à alerter le Parlement européen sur le danger existentiel d'une IA devenue folle. Mais en concentrant autant d'énergie sur une menace, somme toute chimérique, l'UE ne risque-t-elle pas, pour reprendre une formule appropriée, « un hiver politique de l'IA »⁵¹ ? Il semble utile d'interroger, à cet égard, les prémisses normatives sur lesquelles s'appuie l'*AI Act*.

⁴² L'IA présente des menaces pour la vie privée et autorise des dérives comme la discrimination, la notation sociale ou l'usage de techniques de manipulation.

⁴³ Le manque de fermeté sur l'interdiction de l'identification biométrique et la multiplication de dérogations pour les systèmes d'IA à haut risque ont été critiqués.

⁴⁴ Dans le secteur de la santé, par exemple, les règles de consentement entravent le traitement des données médicales utilisées à des fins d'études.

⁴⁵ CJUE, (grande chambre) 6 octobre 2020, *La Quadrature du Net e.a.*, aff. C-511/18.

⁴⁶ CE, 21 avril 2021, French Data Network et autres n° 393099.

⁴⁷ SMUHAN A. (2021), « Beyond the individual: governing AI's societal harm », *Internet Policy Review*, vol. 10, n° 3.

⁴⁸ Nous estimons que l'argument de la singularité technologique, sur lequel s'appuient certaines critiques, a été magistralement réfuté. J.-G. Ganascia, *Le mythe de la singularité. Faut-il craindre l'Intelligence Artificielle?*, Le Seuil, 2017.

⁴⁹ BOSTROM N. (2014), *Superintelligence: Paths, Dangers, Strategies*, OUP, OXFORD.

⁵⁰ AO governance as a new European Union external policy tool. Study Requested by the AIDA committee <http://www.europarl.europa.eu/supporting-analyses>.

⁵¹ Calo, R. (2018). Artificial Intelligence Policy: A Primer and Roadmap. *University of Bologna Law Review*, 3(2), 180–218. <https://doi.org/10.6092/issn.2531-6133/8670>.

4. Les prémisses normatives de l'AI Act

La lecture du projet de règlement sur l'IA ne laisse pas de place au doute quant au postulat posé par la Commission. L'édifice normatif repose sur un raisonnement partant de l'existence d'un risque technologique qu'il s'agit d'identifier et d'évaluer : selon la finalité de leur utilisation, le texte classe les systèmes d'IA en différentes catégories allant du risque inacceptable, élevé, faible (ou limité) jusqu'au risque minimal. Le cadre normatif proposé par la Commission, dans sa proposition législative, se situe assurément au confluent de l'éthique et du droit. En appréhendant les systèmes d'IA de façon sectorielle et hiérarchisée, en fonction de leur degré de dangerosité, la Commission entend exclure ou encadrer les plus risqués. Si, dans son principe, l'approche est difficilement critiquable⁵², il est permis d'interroger le choix des critères qui ont guidé cette classification (présence de biais, risque de discrimination, problème de transparence). La notion de *haut risque* est suffisamment vague pour englober les atteintes à l'intégrité physique et la santé des personnes, comme les menaces sur les droits fondamentaux (respect de la vie privée, droit à la dignité humaine, non-discrimination). Mais si la nature du danger varie selon ces hypothèses, la Commission ne semble pas leur appliquer un traitement différencié.

La question de savoir ce que révèle, ou dissimule, la « pyramide des risques » et ce qu'elle implique concrètement mérite d'être approfondie. Tout d'abord, on relèvera que bien qu'elle ait cherché à préserver une certaine neutralité technologique — comme l'atteste sa définition évolutive des systèmes logiciels — la classification *a priori* témoigne d'une prise de position certaine de la Commission⁵³. Il importe de rappeler que la manière dont celle-ci a détecté et présenté les risques n'était pas la seule envisageable⁵⁴. L'inévitable part d'arbitraire de l'énumération est amenée à évoluer au gré des débats, et des prises de positions, tout aussi subjectives, des uns et des autres⁵⁵. Au-delà des contingences attachées au processus législatif, il demeure que l'approche de la Commission est centrée sur la personne humaine de sorte que l'on peut raisonnablement qualifier sa vision de l'IA d'« humano-centrée ». En attestent les nombreuses références aux « implications humaines », à la « dignité humaine », au « facteur humain », ou encore au « contrôle humain » de l'IA. Cette approche ressortait déjà, quoique plus explicitement, des lignes directrices établies par le groupe d'experts de haut niveau. Ces derniers ne s'étaient-ils pas prononcés en faveur de systèmes d'IA « centrés sur l'humain », « au service de l'humanité et du bien commun, avec pour objectif d'améliorer le bien-être et la liberté des êtres humains »?⁵⁶ On relèvera toutefois que contrairement aux experts qui renaient une

acception généreuse du bien commun en considérant « la société au sens large » et « les autres êtres sensibles et l'environnement comme des parties prenantes tout au long du cycle de vie de l'IA », la Commission retient une vision étroite de la société en envisageant les seuls êtres humains, à l'exclusion des animaux, et sans réellement prendre en compte les risques pour l'environnement⁵⁷.

Tant l'accent mis sur le respect des droits fondamentaux (inhérents à la personne humaine) que l'approche par les risques semblent découler de cet « anthropocentrisme ». Une autre conséquence, quelque peu paradoxale, semble être la tendance à humaniser l'IA, cette dernière se voyant affublée de qualités proprement *humaines* : ainsi, si l'*AI Act* se contente de souhaiter une IA digne de confiance, certaines communications vont jusqu'à évoquer sa bienveillance. En plus de comporter une forte dose de construction idéologique, un tel angélisme dissimule mal la stratégie de la Commission consistant à faire accepter la transition numérique dans l'Union. En effet, pour que l'UE puisse profiter des avantages économiques de l'IA, et devenir un leader mondial, cette technologie disruptive doit être acceptée par les citoyens. Rappelons que l'anthropocentrisme confine souvent à un anthropomorphisme de l'IA elle-même, comme il ressort de certaines communications de l'UE se référant au *Frankenstein* de Mary Shelley, au mythe antique de *Pygmalion*, au *Golem* de Prague ou encore au *robot* de Karel Čapek⁵⁸. S'il est vrai, comme l'affirme le Parlement européen, que les humains ont toujours « rêvé de construire des machines intelligentes, le plus souvent des androïdes à figure humaine » et que « l'humanité se trouve à l'aube d'une ère où les robots, les algorithmes intelligents, les androïdes et les autres formes d'intelligence artificielle, de plus en plus sophistiqués, semblent être sur le point de déclencher une nouvelle révolution industrielle »⁵⁹, la société ne doit pas négliger d'étudier les effets de cette transformation sur elle-même.

5. Des concepts juridiques mis à mal

Si l'*AI Act* constitue une avancée pour l'UE, dans sa volonté de régir l'IA, il reste que le règlement illustre également les limites auxquelles se heurte le droit lorsqu'il vise à se saisir d'une technologie aussi rebelle aux catégories normatives classiques. Le recours aux standards habituels, notamment ceux issus du droit dérivé et de la Charte, est peu compatible avec l'écriture numérique, le langage algorithmique et, plus généralement, l'environnement informatique dans lequel s'insèrent et se déploient les systèmes d'IA. Les sous-sections suivantes traitent de la quasi-faillite des concepts au cœur du futur droit de l'IA, préfigurant peut-être une reconfiguration souterraine des rapports dialectiques entre l'UE et l'IA.

⁵² D'aucuns pourraient discuter de l'opportunité d'une approche aussi dogmatique en ce que cette position pourrait priver l'UE de certaines pratiques avantageuses.

⁵³ Annexe III du Règlement.

⁵⁴ Il eût été possible, par exemple, de les distinguer en risques « à court terme » et « à long terme », avec un régime juridique différencié. Les possibles préjudices sociétaux, mentionnés précédemment, appartenant à la seconde catégorie.

⁵⁵ Le choix d'inclure certains types d'infractions dans la liste des exceptions et d'en exclure d'autres, sans explications, est discutable. De nombreuses autorités, comme le Comité européen de la protection des données ou le Contrôleur européen de la protection des données ont appelé à un régime plus conséquent et des interdictions plus fermes.

⁵⁶ Lignes directrices en matière d'éthique. Pour une IA digne de confiance. Groupe d'experts indépendants de haut niveau sur l'IA, juin 2028.

⁵⁷ Ce qui explique la place réduite de l'éthique environnementale et l'occultation totale de l'éthique animale, les animaux n'étant jamais mentionnés dans l'*AI Act*.

⁵⁸ Résolution du Parlement européen du 16 février 2017 contenant des recommandations à la Commission concernant des règles de droit civil sur la robotique (2015/2103(INL)) (2018/C 252/25)

⁵⁹ *Ibid.*

5.1 La protection des données personnelles

À l'instar du RGPD, le projet de règlement sur l'IA vise à renforcer la protection des données à caractère personnel. L'utilisation des systèmes d'IA reste donc subordonnée aux exigences de la Charte et du droit dérivé de l'UE. Au regard toutefois des problèmes d'opacité, de complexité et de dépendance vis-à-vis des données, l'emploi de certains systèmes d'IA peut porter atteinte à certains principes comme le droit au respect de la vie privée et, son corolaire, la protection des données personnelles⁶⁰. D'après le RGPD, la protection des données personnelles s'impose pour toute information concernant une personne physique identifiée ou identifiable. Il n'y a donc pas lieu d'appliquer de garanties particulières aux informations anonymes. Pour autant, les progrès technologiques ont modifié les conditions dans lesquelles une personne physique peut être identifiable. L'anonymisation s'avère illusoire dès lors qu'une combinaison d'informations anonymes permet de (ré-)identifier une personne. En outre, les données faisant directement référence aux personnes ne sont pas les seules susceptibles de causer un préjudice. La collecte de données à des fins d'influence sur le comportement a, elle aussi, un impact significatif sur les personnes physiques (objets connectés). À la lumière de ces nouveaux défis, c'est bien le *concept juridique* de données personnelles qui tend à être remis en question. Des auteurs ont préconisé, en conséquence, de modifier le cadre légal pour prendre en compte les « inférences » opérées à partir des données personnelles⁶¹, plutôt que de se focaliser sur les données et leur accès, comme le fait le droit de l'UE. Comme il est possible d'inférer des informations sur les personnes à partir de données anonymes, la logique juridique consistant à encadrer l'accès aux données personnelles se révèle largement inefficace. Si des voix se sont élevées pour inviter les régulateurs à changer de paradigme — pour concentrer la réglementation sur l'usage des données et des connaissances inférées sur les personnes — telle n'est pas la voie suivie par la Commission dans son projet de règlement sur l'IA.

5.2 L'interdiction des discriminations

Le projet de règlement sur l'IA prétend compléter le droit existant en matière de non-discrimination en prévoyant des exigences spécifiques visant à « réduire le risque de discrimination algorithmique, en particulier s'agissant de la conception et de la qualité des jeux de données ». Afin de lutter contre diverses sources de risques, la proposition vise à renforcer la non-discrimination et l'égalité entre les femmes et les hommes⁶². La Commission propose de classer comme étant à haut risque les systèmes utilisés pour des questions liées à l'emploi car ils ont une incidence sur les perspectives de carrière et les moyens de subsistance des personnes. Certains systèmes d'IA peuvent, en effet, perpétuer des schémas historiques de discrimination, par exemple à l'égard des femmes, de certains groupes d'âge, des personnes

handicapées, ou de certaines personnes en raison de leur origine ethnique ou de leur orientation sexuelle. Les systèmes d'IA utilisés pour évaluer la note de crédit ou la solvabilité des personnes devraient également être classés à haut risque car ils déterminent l'accès à des ressources financières ou à des services essentiels. Les systèmes d'IA peuvent, là aussi, conduire à discriminer certaines personnes ou groupes et perpétuer des schémas historiques d'inégalités sociales. La définition juridique de la discrimination et son régime d'interdiction, au cœur du droit dérivé de l'UE, présentent des limites conceptuelles lorsque les données en cause sont de simples « traces comportementales ». Force est d'admettre que le langage algorithmique entre en conflit avec le langage juridique lorsque la variable *sensible* (genre, origine) n'est pas utilisée par les concepteurs mais qu'en pratique, l'algorithme conduit, d'une certaine façon, à discriminer une personne en raison de son appartenance à un groupe social. Comment le droit peut-il appréhender cette supposée discrimination ? Est-il envisageable de prohiber l'usage de variables corrélées indirectement au genre ou à l'origine ethnique ? Est-il permis, par ailleurs, de tolérer des discriminations, à un niveau individuel, si l'*effet disparate* se révèle faible à un niveau agrégé ? Ces questions mettent en lumière les limites du droit et des régimes juridiques classiques pour appréhender ces formes inédites de discriminations. Les atteintes au droit au respect de la vie privée ou les pratiques discriminatoires devraient permettre d'engager la responsabilité des auteurs. Or, en matière d'IA, il est extrêmement complexe d'identifier les causalités et, partant, de chercher la responsabilité des acteurs impliqués dans la chaîne décisionnelle.

5.3 L'engagement de la responsabilité

Comme toute technologie transformatrice, l'IA soulève des questions juridiques, au premier chef desquelles la responsabilité, en particulier pour les systèmes à haut risque. S'agissant du régime légal, l'UE peut s'appuyer sur les règles en matière de responsabilité du fait des produits défectueux⁶³ et de protection des données personnelles⁶⁴. Le projet de règlement sur l'IA rappelle, en outre, qu'en matière de responsabilité civile pour l'IA, le Parlement a adopté une résolution spécifique⁶⁵. La Commission est consciente que certains systèmes d'IA, en particulier ceux qui permettent une prise de décisions autonomes, exigent de revoir le régime de sécurité et le droit civil relatif à la responsabilité. Les robots évolués et les produits de l'internet des objets, par exemple, peuvent agir d'une façon non envisagée lors de leur mise en service. Cette problématique pointe les limites du régime de responsabilité. Le droit positif semble, en effet, peu armé pour offrir un traitement satisfaisant de l'exigence de réparation lorsque le fait générateur du préjudice est un système d'IA. La directive sur la responsabilité du fait des produits, par exemple, ne permet pas de répondre adéquatement aux défis posés par les technologies émergentes. De même, la directive *Machines* ne traite pas de certains aspects des technologies

⁶⁰ Articles 7 et 8 du règlement précité.

⁶¹ S. Wachter et B. Mittelstadt, « A right to reasonable inferences: re-thinking data protection law in the age of big data and AI », *Colum. Bus. L. Rev.*, 2019, 494.

⁶² Articles 21 et 23 de l'*AI Act*.

⁶³ Directive 85/374/CEE du Conseil du 25 juillet 1985 relative au rapprochement des dispositions législatives, réglementaires et administratives des États membres en matière de responsabilité du fait des produits défectueux.

⁶⁴ S'agissant de la sécurité des réseaux et des systèmes d'information, l'UE dispose des règles les plus strictes en matière de protection des données personnelles.

⁶⁵ Résolution sur un régime de responsabilité civile pour l'intelligence artificielle, 2020/2014(INL).

numériques. Ce qui rend difficile l'application d'un régime de responsabilité classique est la pluralité des acteurs impliqués, de la mise au point de l'IA (programmation, fabrication, fourniture des données) à sa mise en oeuvre (exploitation, possession, utilisation). L'intrication rend ardue la recherche des responsabilités, c'est-à-dire des schémas de causalité. La Commission recommande ainsi que le fournisseur — personne physique ou morale — assume la responsabilité de la mise sur le marché/en service d'un système d'IA à haut risque. Compte tenu des risques pour la sécurité, la Commission propose de définir des responsabilités spécifiques pour les utilisateurs. Ces derniers devraient être tenus d'utiliser ces systèmes conformément à la notice. Concernant l'évaluation de la conformité des systèmes d'IA, celle-ci devrait être réalisée par le fournisseur sous sa responsabilité, à l'exception des systèmes d'identification biométrique, pour lesquels l'intervention d'un organisme notifié est requise. S'agissant des fabricants, la Commission prévoit qu'ils ont la responsabilité de la conformité de certains systèmes avec le règlement. Les importateurs et distributeurs devraient aussi être soumis à des obligations puisqu'il est prévu qu'ils s'assurent, lorsqu'un système à haut risque est sous leur responsabilité, que les conditions de stockage/transport ne compromettent pas sa conformité au règlement. Il reste que le concept de responsabilité est mis en difficulté par les systèmes d'IA dans la mesure où un régime légal ne peut s'appliquer que dans les hypothèses où il est possible de définir le produit et ses relations. S'agissant de la responsabilité du concepteur de l'algorithme, que faire lorsque le produit est un vaste réseau socio-technique au sein duquel les causalités sont imprécises ?⁶⁶ De même, peut-on engager la responsabilité des constructeurs d'un algorithme à l'origine de discriminations lorsqu'elles ne sont pas délibérées et qu'elles sont difficilement compréhensibles ? La question est d'autant plus complexe que la vérification n'est pas aisée. Les nouvelles technologies ne sauraient se développer au détriment des individus et de leur sécurité. C'est pourquoi, certains préconisent de s'orienter vers un régime de responsabilité propre à l'IA, de responsabilité objective aggravée. Dans la mesure où ce dernier serait axé sur la réalisation d'un risque caractérisé résultant d'une activité d'un système d'IA, un tel régime implique l'identification préalable des risques prévisibles. La question demeure pour les risques imprévisibles. Sur qui doit-on faire peser ces risques ? Sur la population ou ceux qui en sont à l'origine et en tirent profit ? Il importerait de procéder à un partage équitable des responsabilités en s'attachant à identifier, via la traçabilité, le degré d'implication de chaque acteur dans la survenue du dommage. Cela impliquerait, en vertu de l'*AI Act*, un minimum de transparence afin de veiller au respect des obligations incombant à l'utilisateur et au fournisseur⁶⁷.

5.4 L'exigence de transparence et d'explicabilité

Ainsi que l'admet la Commission, dans les considérants de l'*AI Act*, l'exercice des droits fondamentaux procéduraux, comme le droit à un recours effectif, pourrait être entravé par l'usage d'un système d'IA insuffisamment transparent, explicable et documenté. Sont ainsi considérés à haut risque, les systèmes destinés à être utilisés dans un contexte répressif où la transparence revêt une grande importance, notamment pour garantir que des comptes soient rendus et que des recours efficaces puissent être exercés. Afin de faciliter la vérification, la transparence, au même titre que les exigences de documentation et traçabilité, est souvent invoquée. Bien que commode⁶⁸, la notion n'est pas sans poser difficulté, ne serait-ce que parce que le principe est susceptible de porter atteinte au droit à la protection de la propriété intellectuelle. En outre, si elle est mobilisée pour tenter de remédier à l'opacité qui rend l'IA trop complexe ou incompréhensible, l'exigence de transparence n'est pas la panacée. On a pu souligner qu'elle ne suffisait pas à comprendre les systèmes, ou les expliquer, dans la mesure où la signification ne leur est pas interne (mais en relation/interaction avec les données). Rendre les systèmes d'IA transparents, n'équivaut donc pas à les rendre responsables. Pour que les utilisateurs puissent « interpréter » les résultats, les systèmes d'IA devraient être accompagnés d'une documentation pertinente et inclure des informations sur les risques. Si le déficit de transparence des méthodes d'apprentissage machine, perçues comme des boîtes noires, constitue un défi scientifique pour la communauté, celui d'explicabilité pose des « problèmes opérationnels, juridiques et éthiques »⁶⁹. Dans la mesure où ces méthodes ne représentent pas une causalité entre les paramètres d'entrées et de sortie, la validation des boîtes noires diffère, sous l'angle épistémologique, de celle mise en place pour la modélisation d'un phénomène physique. La nécessité de clarifier les notions d'interprétabilité et d'explicabilité qui, malgré l'absence de consensus, font l'objet d'une vaste littérature, y compris au sein de l'UE, a été soulignée. De même que le risque que des « explications soient plus persuasives qu'informatives »⁷⁰. Loin d'être théoriques, ces questions intéressent le monde juridique, notamment depuis que la loi bioéthique a consacré, au sein du code de la santé publique⁷¹, une information des patients dans l'hypothèse du recours à un dispositif médical basé sur un système de *machine learning*. L'explicabilité a été préférée à la transparence dans la mesure où la publication du code source d'un algorithme ne permet pas, à elle seule, aux d'en comprendre « la logique générale de fonctionnement »⁷². Comprendre les mécanismes et logiques algorithmiques n'est, faut-il le souligner, pas une tâche aisée, ni pour les professionnels de la santé, ni pour ceux du droit, et son traitement doctrinal dépasse l'objet de cet article. Nous nous bornerons à préciser que l'*AI Act* fait de l'utilisateur du système d'IA, le destinataire des obligations d'information⁷³.

⁶⁶ J.-M. John-Mathews, *AI Ethics in Practice, Challenges and Limitations*, Thèse, 2021.

⁶⁷ Article 13 de l'*AI Act*.

⁶⁸ La transparence est souvent présentée comme un moyen de garantir le respect des droits fondamentaux, notamment la non-discrimination, la protection de la vie privée et des données à caractère personnel, ainsi qu'une bonne administration.

⁶⁹ C. Denis; F. Varenne. Interprétabilité et explicabilité de phénomènes prédits par de l'apprentissage machine. *ROIA*, Volume 3 (2022) no. 3-4, pp. 287-310.

⁷⁰ *Ibid.*

⁷¹ Article L.4001-3 du code de la santé publique.

⁷² Conseil d'État, Révision de la loi de bioéthique : quelles options pour demain ?, 28 juin 2018.

⁷³ Articles 3 et 13 de l'*AI Act*.

5.5 De l'AI Act au rêve de Leibniz : intelligence artificielle et jurisprudence naturelle

Dans cette dernière sous-section, et en guise de conclusion, nous esquissons une approche visant à situer l'*AI Act* dans un paradigme tendant à rapprocher le langage juridique — celui de l'UE — et la logique — celle de Leibniz. Cette hypothèse heuristique s'inspire des travaux d'auteurs ayant proposé l'approche « droit et mathématiques » pour souligner le fait que le droit tend, de plus en plus, à préserver sa légitimité en cherchant à épouser le modèle des sciences informatiques⁷⁴. Notre thèse consiste à montrer les affinités électives entre le droit de l'UE — caractérisé par sa technicité et sa teneur économique — et le projet du célèbre juriste-logicien de réconcilier logique et rhétorique en fondant une théorie de la *jurisprudence naturelle*. Rappelons que le rêve leibnizien reposait sur l'idée que le droit était la première science à appliquer la logique aux actions humaines, permettant d'assimiler logiques probabiliste et juridique. Cette logique du probable (axiomatique du contingent) se distingue de la logique du nécessaire qui régit les mathématiques. Ce faisant, nous voudrions faire ressortir le paradoxe suivant : en cherchant à *saisir l'IA*, pour reprendre le titre de notre article, le droit de l'UE pourrait bien se voir *saisi* par celle-ci. Les éléments susceptibles d'étayer cette hypothèse sont au nombre de cinq : en premier lieu, l'approche par les risques, prônée par la Commission, repose sur la probabilité d'un préjudice éventuel. Les systèmes d'IA sont appréhendés en étudiant un risque de préjudice eu égard « à sa probabilité d'occurrence ». Cette approche singulière tend à rapprocher le droit de l'UE de la logique probabiliste, chère à Leibniz, au cœur de certains systèmes d'IA. En second lieu, le fait que Leibniz ait souhaité résoudre les problèmes juridiques par le calcul trouve un écho dans la validation du recours aux métriques et aux seuils probabilistes pour la gestion des risques. La délégation aux machines de questions éthiques, via les outils de débiaisement, aurait sans doute plu au théoricien du meilleur des mondes possibles. En troisième lieu, la façon dont l'*AI Act* définit l'IA, sous forme de listes exposées en annexes, n'est pas sans rappeler la méthode leibnizienne consistant à rationaliser les « *définitions ou explications des termes juridiques* », lesquelles « *doivent être traitées dans un ouvrage spécial, sans aucun mélange avec des préceptes ou des règles ; cela peut-être appelé: classification du droit* ». Leibniz précise que « *les tableaux sont très commodes dans ce domaine et il faut que d'un seul regard toute la connaissance soit d'abord disposée dans un tableau général* »⁷⁵. En quatrième lieu, la technique de mise en balance des risques et bénéfices, au cœur du projet de la Commission, répond précisément au souhait leibnizien de réduire à un calcul (probabilités) toute question de droit. Leibniz préconisait de fonder le raisonnement sur l'analyse dont les mathématiciens donnent des échantillons (définitions, axiomes). Estimant toutefois que dans la *jurisprudence naturelle*, « les raisons ne doivent pas être comptées mais pesées », il s'étonnait que personne n'ait, à son époque, fourni « cette balance qui doit servir à peser la force des raisons »⁷⁶.

Cette remarque revêt, à l'ère de l'IA, et sous l'angle du droit de l'UE, une dimension toute particulière. En dernier lieu, il est permis de relever, à l'heure où des voix s'élèvent contre le règne des algorithmes⁷⁷, que la transition numérique que l'UE encourage s'inscrit dans un projet politique et social fantasmé, il y a plus de trois siècles, par le philosophe de Leipzig. En effet, la rationalisation (mathématisation) des sciences morales semble permise dès lors que pour former un jugement infaillible, il suffit, pour Leibniz, de « déterminer ce qui est le plus probable ex datis ». Il ajoutait que cette approche serait « d'un secours admirable même en politique et en médecine, pour raisonner sur les symptômes et circonstances données d'une manière constante et parfaite ». La transformation numérique à laquelle l'Europe se prépare, sous l'égide de la Commission, paraît bien prendre pour modèle la vision leibnizienne de la société, de son droit et de sa gouvernance. Que l'on pense, par exemple, au futur espace européen des données de santé.

6. Références

- [1] Andrade, S. Intelligence artificielle: réflexion sur la responsabilité du fait des logiciels d'aide à la décision médicale, mémoire (2021)
- [2] Benbouzid, B., Meneceur, Y. & Smuha, N. Quatre nuances de régulation de l'intelligence artificielle: Une cartographie des conflits de définition. *Réseaux*, 232-233, 29-64 (2022)
- [3] Bostrom, N. Superintelligence: Paths, Dangers, Stratégies, OUP, Oxford (2014)
- [4] Calo, R. Artificial Intelligence Policy: A Primer and Roadmap. *University of Bologna Law Review*, 3(2), 180-218. <https://doi.org/10.6092/issn.2531-6133/8670> (2018)
- [5] Denis, C., Varenne, F. Interprétabilité et explicabilité de phénomènes prédits par de l'apprentissage machine. *ROIA*, 3 (3-4), pp. 287-310. 10.5802/roia.32. hal- 03640181 (2022)
- [6] Ganascia, J.-G. Le mythe de la singularité. Faut-il craindre l'Intelligence Artificielle? Le Seuil (2017)
- [7] Garapon, A., Lassègue, J. Justice digitale, PUF (2018)
- [8] Gautrais, V. Neutralité technologique: rédaction et interprétation des lois face aux technologies, Thémis (2012)
- [9] Giovanna Palermo, A. Logique juridique et logique probabiliste à l'époque moderne. Thèse (2013)
- [10] John-Mathews, J.-M. AI Ethics in Practice, Challenges and Limitations. Thèse (2021)
- [11] Meneceur, Y., Barbaro, C. Artificial intelligence and the judicial memory: the great misunderstanding. *AI Ethics* 2, 269-275 (2022)
- [12] Pégny, M. Pour un développement des IAs respectueux de la vie privée dès la conception. hal-03104692 (2021)
- [13] Wachter, S., Mittelstadt, B. « A right to reasonable inferences: re-thinking data protection law in the age of big data and AI », *Colum. Bus. L. Rev.*, 494 (2019)

⁷⁴ A. Garapon, J. Lassègue, Justice digitale, PUF 2018.

⁷⁵ Leibniz, Nova Methodus docendaeque jurisprudentiae, 1667.

⁷⁶ Leibniz, Lettres à Thomas Burnet, 11 février 1697, Hanovre. Journal Electronique des Probabilités et des Statistiques, vol. 2, n°1, juin 2006.

⁷⁷ Meneceur, Y., Barbaro, C. Artificial intelligence and the judicial memory: the great misunderstanding. *AI Ethics* 2, 269-275 (2022).

Le problème du décor revisité : un modèle logique pour le diagnostic d'erreurs humaines

Valentin Fouillard^{1,2}, Nicolas Sabouret¹, Safouan Taha², Frédéric Boulanger²

¹ Université Paris-Saclay, CNRS, LISN, 91405, Orsay, France

² Université Paris-Saclay, CNRS, ENS Paris-Saclay, CentraleSupélec, LMF, 91190, Gif-sur-Yvette, France

prenom.nom@universite-paris-saclay.fr

Résumé

Cet article présente un modèle logique pour étudier et diagnostiquer les erreurs de décision commises par des humains dans des systèmes critiques comme les avions. Notre approche s'appuie sur la révision de croyance pour calculer les états mentaux possibles de l'opérateur à partir des observations et actions enregistrées. Nous montrons que notre modèle capture différentes formes d'incohérences qui correspondent à des erreurs de raisonnement possibles de l'opérateur. Certaines de ces erreurs sont liées à un problème de distorsion du décor, c'est-à-dire à des déviations du comportement attendu selon le principe d'inertie connu sous le nom de « problème du décor ».

Mots-clés

Révision de croyance, diagnostic, erreur humaine

Abstract

This paper presents a logic-based framework to study and put forward explanations for erroneous decision-making by human operators. Our approach is based on belief revision to compute the operator's possible mental states from recorded actions and observations. We demonstrate how our model captures different forms of inconsistencies in the logic model, which correspond to possible reasoning errors by the operator. Some of these errors can be related to a Frame Distortion Problem, namely a deviation from the expected behaviour under the inertia principle known as the Frame Problem.

Keywords

Belief revision, diagnosis, human error

1 Introduction

L'erreur est humaine, c'est pourquoi nous la retrouvons dans de nombreuses situations : erreur d'interprétation, distraction, violation des règles, etc [11]. Ces erreurs peuvent conduire à des accidents graves dans des domaines critiques comme la sûreté nucléaire, le transport aérien, la médecine ou l'ingénierie spatiale. Comprendre ce qui s'est passé dans le cas d'un incident ou d'un accident est donc essentiel pour éviter qu'une telle situation se répète et pour concevoir de

nouveaux systèmes qui prennent en compte ces facteurs humains [17]. La littérature du domaine parle d'« analyse de l'erreur humaine » (*human error analysis* en anglais).

Pour mener à bien cette analyse de l'erreur humaine, les experts doivent comprendre précisément la situation et poser des hypothèses sur les états mentaux des opérateurs humains pour expliquer leurs erreurs. Prenons l'exemple du transport aérien. Dans ce domaine, les enregistreurs de vol (les « boîtes noires ») permettent aux experts d'avoir accès à l'ensemble des informations des instruments de vol, à l'état de l'appareil et aux conversations dans le cockpit. Ils utilisent ces informations pour comprendre ce que les pilotes ont compris de la situation et pour expliquer pourquoi ils ont effectué des actions erronées.

La thèse que nous défendons dans cet article est que la modélisation formelle en logique constitue un fondement solide pour développer des outils semi-automatiques d'analyse de l'erreur humaine. En effet, le problème de l'identification des erreurs de décision dans une séquence d'actions est similaire à un problème de diagnostic, un domaine de l'IA très actif depuis les années 80 et pour lequel de nombreux modèles logiques ont été proposés [19, 16, 18]. Cependant l'utilisation de ces modèles logiques pour l'analyse de l'erreur humaine est difficile car ces modèles implémentent des raisonnements complètement rationnels, comme le principe d'inertie des croyances lié au problème du décor de McCarthy et Hayes [15]. Au contraire, les facteurs humains qui rentrent en jeu dans les erreurs de prise de décision ont tendance à casser ces principes. Ainsi le diagnostic logique des erreurs humaines nécessite de revoir ces principes rationnels pour permettre des déviations des comportements attendues. C'est ce que nous nommons le *problème de la distorsion du décor*.

Cet article propose un modèle logique qui supporte la distorsion du décor pour construire une série d'états mentaux menant à un raisonnement erroné. Il est organisé de la manière suivante. La section 2 introduit la problématique de recherche. La section 3 décrit notre modèle qui permet de représenter la situation d'accident et de calculer les états mentaux successifs en prenant en compte l'inertie des croyances. La section 4 montre, à travers différentes incohérences, dont la « distorsion du décor » que nous avons men-

tionnée ci-dessus, comment nous pouvons diagnostiquer les erreurs humaines dans notre modèle. Enfin, les deux dernières sections abordent les travaux connexes (section 5) et des perspectives de notre modèle (section 6)

2 Définition du problème

Notre objectif est d'utiliser le diagnostic pour retrouver les séquences d'états mentaux d'un opérateur, représenté formellement comme un agent, qui permettent d'expliquer un accident. Pour cela, nous utilisons le *consistency based diagnosis* présenté dans la section suivante.

2.1 Consistency-based diagnosis

Le diagnostic consiste à déterminer quels composants d'un système ont causé un comportement anormal sachant un ensemble d'observations. Dans ce contexte, le *consistency based diagnosis* proposé par Reiter [19] est un modèle logique pour le diagnostic automatique qui consiste à retrouver la cohérence entre la description du comportement du système et les observations. Ce modèle comporte trois ensembles logiques :

- *SD* un ensemble de formules logiques qui décrit le système,
- *ASS* un ensemble de prédicats qui décrit les *hypothèses* de la forme $\neg ab(c)$, i.e le composant c est supposé se comporter normalement,
- *OBS* une conjonction de prédicats qui décrit une observation du système.

Quand $SD \cup ASS \cup OBS$ est incohérent, un diagnostic Δ est un ensemble minimal d'*hypothèses* tel que $SD \cup (ASS \setminus \Delta) \cup OBS$ est cohérent. En d'autres termes, un diagnostic est un ensemble minimal d'éléments dont on doit supposer qu'ils ont un comportement anormal pour retrouver la cohérence avec les observations.

Notre modèle d'analyse d'erreur humaine est basé sur cette approche de diagnostic par recherche de la cohérence (*consistency-based diagnosis*) avec quelques adaptations que nous présentons dans les sections suivantes. La principale modification provient de ce que nous voulons étudier les états mentaux, qui peuvent différer de la vérité du monde. C'est pourquoi nous considérons qu'un état mental est composé de différents éléments (croyances, règles de raisonnement, observations...) qui sont par défaut ceux qui correspondent à la réalité du monde physique. Toutefois, l'état mental d'un agent humain peut être altéré par différents facteurs (erreurs de raisonnement, occultation d'informations), et nous considérons qu'un diagnostic d'erreur humaine est un ensemble minimal d'éléments qu'il faut retirer de l'état mental de l'agent pour retrouver la cohérence avec les actions qu'il a effectuées. Nous présentons dans la section 3 ce modèle et l'algorithme de diagnostic qui va avec.

2.2 Le problème du décor

Le problème du décor (*frame problem* en anglais) a été introduit par John McCarthy et Patrick Hayes [15] en 1969. C'est un problème très connu en modélisation logique qui peut être décrit comme la difficulté à représenter les effets d'une action en logique sans avoir à représenter explicite-

ment tous les non-effets évidents [22]. Pour illustrer ce propos, considérons une formalisation simple, en logique du premier ordre, du *Yale Shooting Problem* (YSP) proposé par Hanks et McDermott [10] :

<i>init</i>	=	$\text{alive}(0), \neg \text{loaded}(0)$
<i>actions</i>	=	$\text{Load}(0), \text{Wait}(1), \text{Shoot}(2)$
<i>rules</i>	=	$\text{Load}(0) \rightarrow \text{loaded}(1),$ $(\text{Shoot}(2) \wedge \text{loaded}(2)) \rightarrow \neg \text{alive}(3)$

0,1,2,3 correspondent à des instants successifs. Les prédicats dans *init* décrivent la situation initiale : la dinde est vivante et le pistolet est chargé. Les prédicats dans *actions* modélisent le déroulement des actions (l'agent charge le pistolet, attend et tire) et les prédicats dans *rules* modélisent la physique du monde : tirer sur la dinde avec un pistolet chargé la tue.

Avec cette modélisation, le bon sens nous dit que $\neg \text{alive}(3)$ est vrai. Toutefois, les effets de l'action *Wait* ne sont pas définis dans le modèle et ne peuvent être déduit sans des règles logiques qui indiquent que $\text{loaded}(t+1) = \text{loaded}(t)$ quand rien d'autre ne dit le contraire. Écrire ces règles manuellement pour chaque fait initial et toutes les actions possibles est une charge trop lourde dans le cas général. C'est ce qu'on appelle le problème du décor.

Une première solution à ce problème a été proposée par McCarthy [14] en utilisant la *circumscription* et en sélectionnant les modèles (c'est-à-dire les ensembles de propositions) qui minimisent le nombre de changements dans le monde. Toutefois, cela ne marche pas sur le YSP où deux solutions minimales peuvent être trouvées : la solution conforme à notre intuition, dans laquelle $\text{loaded}(2) \wedge \neg \text{alive}(3)$ est vrai, mais aussi une solution contre-intuitive $\neg \text{loaded}(2) \wedge \text{alive}(3)$ dans laquelle le pistolet se décharge magiquement pendant l'action *Wait*.

D'autres solutions ont été proposées au fil du temps pour résoudre ce problème, comme les axiomes de Reiter sur les états successeurs [20] ou l'occlusion de Sandewall [21]. Toutes les solutions consistent à ajouter des éléments pour distinguer les prédicats ou propositions qui doivent sortir de l'inertie des clauses du décor. En d'autres termes, ils permettent une description de comment le monde change ou ne change pas sachant les actions occurrentes.

Comme tout autre modèle logique d'actions et de changements, notre modèle d'analyse d'erreur humaine fait face au problème du décor. Mais en plus de cela, nous devons faire face à un autre problème : celui de la *distorstion du décor* que nous illustrons dans la prochaine section en utilisant une version remaniée du YSP.

2.3 Le problème de la distorsion du décor

Supposons, dans le YSP, que l'agent ne voulait pas tuer la dinde. Nous sommes face à une situation d'erreur de prise de décision et nous voulons comprendre pourquoi l'agent a finalement appuyé sur la détente. Pour modéliser cette situation, nous ajoutons un désir dans la liste des prédicats :

$$\text{desires} = \text{desire}(\text{alive}(3))$$

Cela indique que l'agent veut que $\text{alive}(3)$ soit vrai. Du point de vue d'un observateur extérieur, en faisant l'hypothèse que le problème du décor est résolu dans le modèle (i.e que l'agent devrait croire à $t = 2$, que le pistolet reste chargé), il n'y a aucune raison que l'agent choisisse de tirer ($\text{Shoot}(2)$). Mais il l'a fait. La question est alors : qu'est ce qui peut expliquer ce comportement apparemment irrationnel ?

Une explication possible serait que l'agent a oublié que le pistolet était chargé. Cela implique que l'agent croyait que le pistolet était chargé au pas de temps 1, mais plus au pas de temps 2 : $\text{loaded}(1) \wedge \neg \text{loaded}(2)$. Pourtant cela n'est pas attendu avec l'inertie des croyances de l'agent, bien que ce soit une explication plausible du comportement irrationnel de l'agent.

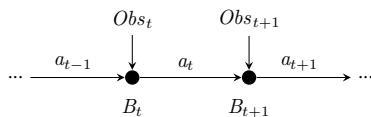
Cet exemple montre que le diagnostic des erreurs humaines selon un modèle logique nécessite de capturer des déviations à tous les niveaux, y compris dans le principe d'inertie des croyances.

Cette situation n'est pas une simple construction de l'esprit. Considérons par exemple le crash du Mont-Saint-Odile [3]. Les pilotes ont effectué une erreur de programmation du pilote automatique avec une vitesse verticale trop grande parce qu'ils avaient oublié qu'ils avaient précédemment configuré le système en mode Vertical Speed (vitesse de descente) au lieu de Flight Path Angle (angle de descente). Les humains peuvent réellement oublier ce qu'ils avaient fait précédemment et ainsi aller à l'encontre du principe d'inertie des croyances. Concrètement, la base de croyance de l'agent a été modifiée entre les deux pas de temps, sans aucune raison !

Ce que nous appelons le *problème de la distorsion du décor* consiste à trouver un diagnostic par recherche de la cohérence (*consistency based diagnosis*) à une situation qui considère des déviations possibles de l'inertie du décor. Dans la prochaine section, nous présentons notre modèle logique qui permet de réaliser ce type de diagnostic.

3 Un modèle logique pour l'analyse des erreurs humaines

Notre approche repose sur la modélisation des actions comme des transitions entre les états mentaux de l'agent, de manière similaire au *calcul des situations* [20], à ceci près que nous modélisons les changements de l'état mental et non du monde lui-même. Nous ajoutons également les observations à chaque pas de temps, qui correspondent aux nouvelles informations reçues par l'agent depuis le monde.



3.1 Modélisation de situation

Nous décrivons les croyances et actions de l'agent par un ensemble de propositions $p_t \in \mathcal{P}$. Chaque proposition p est

indexée par un pas de temps t . Par exemple :

$$\begin{aligned} \neg \text{alive}_3 &\rightarrow \text{L'agent croit que la dinde est morte à } t = 3 \\ \text{Load}_2 \in \text{Act} &\rightarrow \text{l'agent charge le pistolet à } t = 2 \end{aligned}$$

L'utilisation d'une valeur d'indice temporel t' égale, plus grande ou plus petite que le pas de temps courant t nous permet ainsi de représenter des croyances à propos du présent, du passé ou du futur.

Nous définissons l'ensemble des croyances initiales $\mathcal{Init}$ qui représente les croyances des agents au pas de temps $t = 0$. Par exemple $\text{alive}_0 \in \mathcal{Init}$ indique que l'agent croit initialement que la dinde est vivante. Nous considérons aussi les ensembles de prédicats $\text{Obs}_0, \dots, \text{Obs}_n$, chaque Obs_t correspond à des observations possibles au pas de temps t . Par exemple, si $\text{loaded}_1 \in \text{Obs}_1$, alors il est possible que l'agent observe que le pistolet est chargé au pas de temps 1. Enfin, nous considérons la liste $\mathcal{T} = [a_0, \dots, a_n]$ avec $a_i \in \text{Act}$, qui représente la liste des actions effectuées par l'agent (une action par pas de temps). Par exemple, $\text{Load}_0 \in \mathcal{T}$ indique le fait que l'agent a chargé le pistolet au pas de temps 0.

Nous utilisons des formules logiques pour modéliser le raisonnement et le comportement de l'agent. Nous notons $\mathcal{R}$ l'ensemble de toutes les règles que l'agent peut utiliser. Chaque règle $R_i \in \mathcal{R}$ obéit à la grammaire suivante :

$$R_i ::= \varphi \mid [\varphi] a \mid a :: \varphi \quad \varphi \in \mathfrak{F}(\mathcal{P}, \wedge, \neg), \quad a \in \text{Act}$$

où φ est une formule logique propositionnelle utilisant des propositions atomiques de $\mathcal{P}$ ainsi que tous les connecteurs logiques classiques ($\wedge, \vee, \rightarrow$, etc). $[\varphi] a$ indique que la pré-condition φ doit être vraie pour effectuer l'action a . Enfin, $a :: \varphi$ indique que la post-condition φ est vraie après que a a été effectué. L'indice temporel t dans une règle est toujours une variable libre. Par exemple :

$$R_1 \equiv \text{Load}_t :: \text{loaded}_{t+1}$$

modélise l'effet de l'action « charger le pistolet » à n'importe quel pas de temps t .

Pour notre algorithme d'analyse de l'erreur humaine qui sera présenté dans la prochaine section, nous définissons deux fonctions :

- *prec* : $\text{Act} \rightarrow \mathcal{R}$ qui retourne l'ensemble des règles R_i qui définissent les pré-conditions d'une action donnée.
- *effect* : $\text{Act} \rightarrow 2^{\mathcal{P}}$ qui retourne l'ensemble des propositions qui apparaissent dans les post-conditions (i.e après le séparateur $::$ dans une règle) d'une action donnée.

Nous modélisons les buts de l'opérateur (les désirs qui doivent être satisfaits à chaque pas de temps) par un ensemble $\mathcal{D}$ de littéraux négatifs ou positifs avec un indice temporel libre. Par exemple, $\text{alive}_t \in \mathcal{D}$ indique que l'agent veut que la dinde soit en vie à chaque pas de temps. Ainsi, chaque situation dans laquelle $\neg \text{alive}_t$ peut être inféré sera considérée comme incohérente du point de vue de l'agent.

3.2 États mentaux

À partir de la modélisation des situations présentée ci-dessus, notre objectif est de construire un diagnostic sous la forme d'une suite d'états mentaux $B_0 \dots B_n$ dans laquelle chaque état mental B_t est un ensemble de propositions dans $\mathcal{P}$ et de règles dans $\mathcal{R}$ qui décrivent les croyances, observations, désirs et règles de raisonnement de l'agent pour le pas de temps t .

L'état mental initial est :

$$B_0 = \text{Init} \cup \mathcal{R} \cup \mathcal{D}$$

Chaque état successif B_t est construit à partir de l'état précédent B_{t-1} en ajoutant les observations Obs_t et l'action effectuée a_t :

$$B_t = B_{t-1} \cup Obs_t \cup \{a_t\}$$

Nous notons $B_t \vdash \varphi$ pour indiquer que φ est une conséquence de B_t dans la logique propositionnelle classique. Par exemple :

$$\frac{B_t \vdash \psi \rightarrow \varphi, \quad B_t \vdash \psi}{B_t \vdash \varphi} \quad \frac{a_t \in B_t, \quad [\varphi] a_t \in B_t}{B_t \vdash \varphi}$$

Comme nous nous reposons sur la logique propositionnelle, nous utilisons un solveur SAT pour calculer la valeur de vérité des propositions. Concrètement, $B_t \vdash \varphi$ si et seulement si $B_t \wedge \neg\varphi$ est insatisfaisable.

Nous disons que B_t est incohérent quand celui-ci est insatisfaisable, ce qui est équivalent à $B_t \vdash \perp$. L'objectif de notre modèle de diagnostic est de « réparer » ces états de croyances incohérents.

3.2.1 Les prédicats *known*

Une première difficulté que nous devons résoudre est la gestion des propositions qui ne sont pas connues par l'agent. En effet, puisque que nous gérons des croyances et non des faits, une proposition φ est *connue* par un agent au pas de temps t si et seulement si $B_t \vdash \varphi$ ou $B_t \vdash \neg\varphi$. Dans le cas contraire (φ n'apparaît pas dans l'état mental), nous disons qu'elle est *inconnue*. Le raisonnement de l'agent peut ne pas être le même selon qu'une proposition est connue ou non.

Considérons par exemple l'état mental B_t comprenant les règles et observations suivantes :

$$\mathcal{R} \equiv \left\{ \begin{array}{l} R_1 \equiv \varphi \rightarrow \neg\psi \\ R_2 \equiv \neg\varphi \rightarrow \neg\gamma \end{array} \right\} \quad Obs_t \equiv \{\gamma, \psi\}$$

Le comportement que nous souhaitons dans notre modèle est le suivant :

- Si φ est connu dans B_t , φ doit être considéré comme vrai ou faux par l'agent et, par conséquent, $\neg\psi$ ou $\neg\gamma$ doit être inféré.

Cela mène à une incohérence avec Obs_t (et donc à la nécessité d'une correction, qui est l'objet de notre travail de diagnostic).

- À l'inverse, si φ est inconnu, le comportement attendu est qu'il n'y ait pas d'incohérence : l'agent ne sait rien sur φ , donc il ne peut rien en déduire et donc son état mental n'est pas contradictoire.

Malheureusement, le solveur SAT peut inférer φ et $\neg\varphi$ par les formules ci-dessus (il cherche en effet les valeurs possible pour nier la formule passée en paramètre) et il renvoie donc que B_t est incohérent.

Pour implémenter le comportement souhaité (pas d'incohérence trouvée quand φ est inconnu), nous introduisons les prédicats *known* dans notre modèle. Notre objectif est de forcer le modèle à appliquer les règles de raisonnement uniquement sur les propositions connues. Concrètement, pour chaque proposition φ , $known_\varphi$ indique que φ est connu par l'agent. Chaque proposition φ ou $\neg\varphi$ qui apparaît dans une règle $R_k \in \mathcal{R}$ est transformée respectivement en $\varphi \wedge known_\varphi$ ou $\neg\varphi \wedge known_{\neg\varphi}$. Ainsi l'exemple précédent devient :

$$\begin{array}{lcl} R_1 & \equiv & (\varphi \wedge known_\varphi) \rightarrow (\neg\psi \wedge known_{\neg\psi}) \\ R_2 & \equiv & (\neg\varphi \wedge known_{\neg\varphi}) \rightarrow (\neg\gamma \wedge known_{\neg\gamma}) \end{array}$$

En conséquence, une proposition est forcée à connu (*known*) seulement si c'est une conséquence d'une règle avec une prémisse qui est *known* et vraie.

Les croyances initiales dans *Init*, les observations et les actions à chaque pas de temps t sont forcées à être *known* dans B_t afin de déclencher les règles de raisonnement. Par inférence, toutes les croyances dérivées par de telles règles de raisonnement sont mises à *known* aussi. Les autres prédicats restent inconnus de l'agent tout au long du processus, sauf s'ils peuvent être déduits.

3.3 Inertie du modèle : le problème du décor

Le Problème du Décor nous dit que tout ce que l'agent connaît (ce qui est défini par les prédicats *known* à $t-1$) doit être aussi connu à t , sauf si un changement est induit par les observations ou les effets d'une action. Par exemple, si $alive_0 \in B_0$, nous attendons que $alive_1 \in B_1$ sauf indication contraire.

Pour résoudre ce problème, nous introduisons deux prédicats qui forcent une proposition à garder la même valeur et le même état de connaissance :

- $keep^{(v)}$ pour maintenir la valeur de vérité d'une proposition φ :

$$keep_{\varphi_t}^{(v)} \equiv (\varphi_t \longleftrightarrow \varphi_{t-1})$$

- $keep^{(k)}$ pour maintenir l'état de connaissance d'une proposition φ :

$$keep_{\varphi_t}^{(k)} \equiv (known_{\varphi_t} \longleftrightarrow known_{\varphi_{t-1}})$$

Nous définissons les propositions $keep_{\varphi_t}^{(k)}$ et $keep_{\varphi_t}^{(v)}$ pour chaque prédicat qui apparaît dans les croyances initiales, les observations et les effets des actions. Nous les appelons *prédicats primitifs* et notons $\mathfrak{P}_t$ l'ensemble des nouvelles propositions primitives qui doivent être maintenues pour le

pas de temps t :

$$\begin{aligned}\mathfrak{P}_1 &= \mathcal{I}nit \\ \mathfrak{P}_{t>1} &= Obs_{t-1} \cup effect(a_{t-1})\end{aligned}$$

Les propositions *keep* qui doivent être maintenues au pas de temps t accumulent les propositions *keep* précédentes du pas de temps $t-1$ ainsi que les nouvelles. Nous définissons les ensembles $\mathfrak{K}_t^{(v)}$ et $\mathfrak{K}_t^{(k)}$ qui contiennent respectivement toutes les propositions *keep*^(v) et *keep*^(k) qui doivent être dans la base de croyance pour le pas de temps t :

$$\begin{aligned}\mathfrak{K}_t^{(v)} &= \mathfrak{K}_{t-1}^{(v)} \cup \bigcup_{\varphi \in \mathfrak{P}_t} \{keep_\varphi^{(v)}\} \\ \mathfrak{K}_t^{(k)} &= \mathfrak{K}_{t-1}^{(k)} \cup \bigcup_{\varphi \in \mathfrak{P}_t} \{keep_\varphi^{(k)}\}\end{aligned}$$

Ajouter ces propositions *keep* dans l'état mental nous permet de traiter le problème du décor. À partir des prédicats primitifs et des règles de l'agent, nous pouvons retrouver toutes les autres propositions par inférence. La prochaine section montre comment les actions et changements peuvent modifier ces règles *keep*.

Pour faciliter la lecture, les prédicats *known* et les propositions *keep* sont omises dans les prochains exemples.

3.4 La révision de croyance dans un agent rationnel

Notre premier objectif est de modéliser un agent rationnel qui reçoit des observations et effectue des actions. En d'autres termes, chaque B_t successif doit être cohérent afin de représenter le fait que l'agent *avait une bonne raison d'effectuer l'action* sachant ce qu'il pouvait observer et croire. Toutefois, deux cas d'incohérence peuvent apparaître dans le modèle :

- (1) Une observation dans Obs_t n'est pas cohérente avec l'état mental précédent B_{t-1}
- (2) Une action a_t n'est pas cohérente avec les observations, désirs et l'état mental précédent

Ce sont deux problèmes distincts. Résoudre (1) revient à faire une *révision de croyance* [8]. Le but est de déterminer quelle croyance ignorer (celle à retirer du système) face à des informations contradictoires. Ce problème a été étudié par [1] et a résulté en la définition d'AGM, un ensemble d'axiomes qui caractérise un opérateur de révision de croyance *minimal*. L'idée derrière cette minimalité est que les agents rationnels ont tendance à garder leurs anciennes croyances.

Résoudre (2) correspond à trouver un *consistency based-diagnosis* que nous avons présenté dans la sous-section 2.1 (*i.e* le comportement de l'agent n'est pas celui attendu). Toutefois, Wassermann montre qu'un opérateur de révision de croyance peut être utilisé pour effectuer une *consistency based-diagnosis* et vice versa [27] car ces deux méthodes recherchent une solution minimale.

Pour cette raison, nous proposons d'utiliser le *Minimal Correction Set* (MCS) pour implémenter l'opérateur de révision de croyance qui permet de résoudre les deux problèmes.

Pour un système $\Phi = \{\phi_1, \phi_2 \dots \phi_n\}$ donné, $M \subseteq \Phi$ est un MCS de Φ ssi :

- $\Phi \setminus M$ est cohérent
- $\forall \phi_i \in M, (\Phi \setminus M) \cup \{\phi_i\}$ est incohérent

Pour calculer le MCS nous utilisons l'algorithme de [13] que nous avons implémenté avec le SMT-solver z3 [5]. Nous notons $\mathfrak{M}(\Phi, shielded)$ la sortie de cet algorithme. Φ est l'ensemble qui doit être corrigé, et $shielded \subset \Phi$ est un ensemble de propositions et règles qui ne peuvent être retirées du système par l'algorithme de MCS (c.-à-d. $\mathfrak{M}(\Phi, shielded) \cap shielded = \emptyset$). La prochaine section montre comment nous utilisons $\mathfrak{M}$ pour l'analyse des erreurs humaines dans notre algorithme.

Il faut souligner que nous n'obtenons pas nécessairement un MCS unique pour un Φ donné. En d'autres termes, il existe de multiples solutions pour retrouver la cohérence dans l'état mental B_t :

$$B_t = (B_{t-1} \cup Obs_t \cup \{a_t\}) \setminus M$$

où $M \in \mathfrak{M}(\Phi, shielded)$. Par conséquent, à partir d'un état mental B_{t-1} , nous obtenons un arbre d'états où chaque branche correspond à un « choix de révision possible » et donc un état mental possible. Chaque chemin dans l'arbre est alors un *scénario*, c'est à dire, une suite d'états mentaux qui est cohérent avec les actions constatées de l'agent.

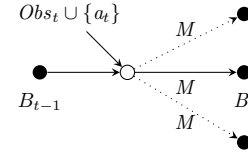


FIGURE 1 – Calcul de B_t (un cercle noir représente un état cohérent, un blanc un état potentiellement incohérent)

Dans la prochaine section, nous nous appuyons sur ces scénarios et les MCS successifs pour diagnostiquer les erreurs humaines dans la situation d'accident que nous modélisons.

4 Analyse de l'erreur humaine

La difficulté de ce travail d'analyse de l'erreur humaine est de distinguer ce qui correspond à la révision de croyance (cas 1 dans la section précédente) de ce qui relève d'une prise de décision erronée (cas 2) dans les séries des MCS d'un scénario.

4.1 Révision de croyance vs erreur humaine

Certains choix de révision dans un scénario correspondent en effet à une révision de croyance rationnelle dans le contexte de l'inertie du décor. Cela se produit dans deux cas :

- (1) La base de croyance doit être révisée à cause des effets attendus d'une action (p. ex. si l'agent charge le pistolet, la valeur de $loaded_t$ doit changer, ce qui est en contradiction avec $keep_{loaded_t}^{(v)}$).
- (2) La base de croyance doit être révisée à cause d'une nouvelle information (p. ex. si l'agent croit initia-

lement que le pistolet est chargé et observe ensuite qu'il ne l'est plus).

Toutes les autres incohérences qui apparaissent dans les états mentaux correspondent à des erreurs dans le raisonnement de l'agent. Toutefois, nous pouvons à nouveau distinguer deux types d'incohérences. Le premier type vient directement du scénario (p. ex. les actions qui n'auraient pas dû être effectuées connaissant les observations). Nous parlons alors d'*incohérences locales*. Le deuxième type correspond au problème de la distorsion du décor. Dans ce qui suit, nous présentons notre algorithme pour séparer les trois types d'incohérences : les incohérences locales, les distorsions du décor, et les révisions de croyance « normales » qu'il faut accepter comme ne relevant pas de l'erreur humaine.

4.2 Un algorithme pour l'analyse de l'erreur humaine

Notre algorithme fonctionne en trois phases. Nous commençons par détecter les MCS qui correspondent à des incohérences locales. Nous détectons ensuite les changements dans la base de croyances liés à l'inertie du décor (qui peuvent créer des incohérences et des MCS, mais qui ne sont pas des erreurs). Enfin nous détectons les MCS qui correspondent à des distorsions du décor.

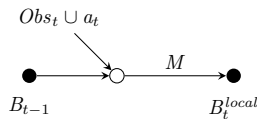
4.2.1 Incohérence locale

En commençant par B_{t-1} , avant d'ajouter $\mathcal{R}_t^{(v)}$ et $\mathcal{R}_t^{(k)}$ dans la base de croyance, nous vérifions les incohérences locales en calculant :

$$B_t^{local} = (B_{t-1} \cup Obs_t \cup \{a_t\}) \setminus M \quad (1)$$

avec $M \in \mathfrak{M}(B_{t-1} \cup Obs_t, \{a_t\})$

Ceci nous permet de capturer les situations où un agent effectue des actions qui ne sont pas cohérentes avec l'état courant des observations et des croyances. L'état B_t^{local} qui en résulte est cohérent avec l'action mais peut contenir des corrections de la base de croyance qui correspondent à une prise de décision erronée.



Un exemple classique d'une telle situation est le crash du vol Rio-Paris en 2009 [4]. Les pilotes recevaient des informations contradictoires et prirent la mauvaise décision (l'avion était en décrochage et ils ont tiré le manche, ce qui a maintenu l'appareil en décrochage). Voici une représentation simplifiée de cette situation dans notre modèle :

$$\begin{aligned} \mathcal{Init} &= \emptyset \\ \mathcal{Obs}_1 &= \{\text{alarm}_1, \text{acceleration}_1\} \\ \mathcal{R} &= \left\{ \begin{array}{l} R_1 \equiv \text{alarm}_t \rightarrow \text{stall}_t \\ R_2 \equiv \text{acceleration}_t \rightarrow \text{overspeed}_t \\ R_3 \equiv \text{overspeed}_t \rightarrow \text{Pull}_t \\ R_4 \equiv \text{stall}_t \rightarrow \text{Push}_t \\ R_5 \equiv \text{Pull}_t \wedge \text{Push}_t \rightarrow \perp \end{array} \right\} \\ \mathcal{T} &= \{\text{Pull}_1\} \end{aligned}$$

Dans cet exemple, les observations de l'alarme et de l'accélération sont incohérentes entre elles (elles mènent à violer R_5 si on suit $R_1 \dots R_4$). Le MCS qui est cohérent avec l'action Pull_1 consiste à garder les observations et règles liées à la survitesse au lieu du décrochage. Cela correspond à ce qui s'est vraiment passé : les enregistreurs du cockpit montrent que les pilotes croyaient qu'ils étaient en survitesse et essayaient d'en sortir en tirant le manche.

4.2.2 Incohérence du décor

Pour chaque B_t^{local} possible, nous introduisons les propositions $keep^{(v)}$ qui sont responsables de l'inertie des croyances. Nous calculons ensuite :

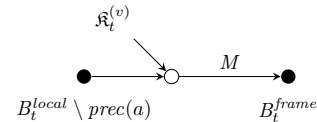
$$B_t^{frame} = ((B_t^{local} \setminus prec(a)) \cup \mathcal{R}_t^{(v)}) \setminus M \quad (2)$$

avec $M \in \mathfrak{M}((B_t^{local} \setminus prec(a)) \cup \mathcal{R}_t^{(v)}, B_t^{local} \setminus prec(a))$

Cela nous permet de capturer les corrections dans la base de croyance qui sont liées au problème du décor. En effet, en calculant les MCS seulement sur les propositions *keep* en protégeant tout ce qui est dans B_t^{local} , nous détectons les incohérences entre les *keep* et les croyances dérivées des observations et des effets d'action. Plus précisément, l'ajout de $\mathcal{R}_t^{(v)}$ maintient la valeur de vérité des croyances précédentes, mais laisse leur état de connaissance (les prédicats *known*) libre. Si une incohérence est détectée, nous savons qu'elle vient de ces croyances dérivées.

Toutefois, comme nous considérons que l'action a est faite à l'instant t , ses pré-conditions et les *known* associés doivent être vrais, ce qui au final force à vrai toutes les propositions qui peuvent être inférées par ces pré-conditions, ce que nous ne voulons pas. En effet, de telles propositions peuvent être incohérentes avec les déductions des règles *keep* que nous voulons appliquer à ce stade pour traiter seulement le problème du décor. C'est pourquoi nous retirons temporairement $prec(a)$, l'ensemble des règles de pré-condition de l'action a , du calcul de l'état mental (il sera réintroduit dans la phase suivante).

Toutes les incohérences capturées à ce stade correspondent à des révisions de croyances rationnelles dans le contexte de l'inertie du décor. Ce ne sont pas des erreurs humaines. Les états B_t^{frame} qui en résultent sont tous des états mentaux cohérents où les incohérences locales et les révisions de croyances « normales » ont été résolues.



La situation suivante illustre cette phase :

$$\begin{aligned} \mathcal{Init} &= \{\text{cloud}_0, \text{sun}_0\} \\ \mathcal{Obs}_1 &= \{\neg \text{cloud}_1\} \\ \mathcal{R} &= \{R_1 \equiv [\neg \text{cloud}_t \wedge \text{sun}_t] \text{GoOut}_t\} \\ \mathcal{T} &= \{\text{GoOut}_1\} \end{aligned}$$

Sortir (action GoOut_1) est un comportement rationnel dans cette situation, bien que cela crée une incohérence avec les

prédicats *known* et les propositions *keep* que nous avons introduites pour le Problème du Décor.

En effet, les pré-conditions de *GoOut* sont $\neg \text{cloud}_1$ et sun_1 qui ne sont pas en contradiction avec les croyances et les observations, le MCS est donc vide pour l'incohérence locale et nous avons :

$$B_1^{\text{local}} = \{\text{cloud}_0, \text{sun}_0, R_1, \neg \text{cloud}_1, \text{GoOut}_1\}$$

Pour calculer B_1^{frame} , nous retirons la règle de précondition R_1 et ajoutons les $\text{keep}^{(v)}$:

$$\{\text{cloud}_0, \text{sun}_0, \neg \text{cloud}_1, \text{GoOut}_1, \text{keep}_{\text{sun}}^{(v)}, \text{keep}_{\text{cloud}}^{(v)}\}$$

Le MCS sur les *keep* contient $\text{keep}_{\text{cloud}}^{(v)}$ afin de rendre $\neg \text{cloud}_1$ possible, et nous avons finalement :

$$B_1^{\text{local}} = \{\text{cloud}_0, \text{sun}_0, \neg \text{cloud}_1, \text{GoOut}_1, \text{keep}_{\text{sun}}^{(v)}\}$$

On a bien traité ici uniquement le problème du décor, en supprimant la règle d'inertie de la proposition *cloud* afin de lui permettre de changer conformément aux observations.

4.2.3 Incohérence de la distorsion du décor

Pour chacun des états mentaux B_t^{frame} possibles, nous introduisons maintenant les $\text{keep}^{(k)}$, qui maintiennent l'état de connaissance des croyances (cela ne s'applique qu'aux croyances, pas aux propositions que l'on peut en dériver). L'analyse de cet ensemble de propositions nous permet de détecter les incohérences suivantes :

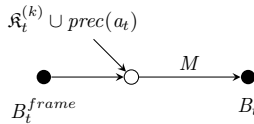
- Un changement de croyance inattendu, que le problème du décor aurait dû éviter.
- Un changement de croyance imposé par les pré-conditions de l'action effectuée ou par des désirs de l'agent.

Concrètement, nous calculons :

$$B_t = (B_t^{\text{frame}} \cup \text{prec}(a) \cup \mathfrak{R}_t^{(k)}) \setminus M \quad (3)$$

avec $M \in \mathfrak{M}(B_t^{\text{frame}} \cup \mathfrak{R}_t^{(k)} \cup \text{prec}(a), \{a_t\})$

Les états B_t résultants sont des états de croyances cohérents qui reflètent les éventuelles erreurs de décision capturées par les MCS.



L'exemple type est le crash du Mont-Saint-Odile présenté en Section 2.3. Les pilotes peuvent régler la vitesse verticale soit en pieds par minute (Vertical Speed) ou en degré d'angle de descente (Flight Path Angle). Quand l'indicateur est en FPA, un affichage de 33 veut dire que l'avion descend avec un angle de 3.3 degrés, ce qui correspond à environ 800 pieds par minute en termes de vitesse verticale. Quand l'indicateur est en VS, le même affichage veut dire que l'avion descend à la vitesse de 3300 pieds par minute (soit quatre fois plus vite) !

En voici une représentation simplifiée dans notre modèle :

$$\begin{aligned} \text{Init} &= \{\text{onVS}_0, \neg \text{onFPA}_0\} \\ \text{Obs}_1 &= \emptyset \\ \mathcal{R} &= \left\{ \begin{array}{l} R_1 = \text{display}(33)_t \wedge \text{onVS}_t \rightarrow \neg \text{goodAngle}_t \\ R_2 = \text{display}(33)_t \wedge \text{onFPA}_t \rightarrow \text{goodAngle}_t \\ R_3 = \text{SetValue}(x)_t :: \text{display}(x)_t \\ R_4 = \text{onVS}_t \leftrightarrow \neg \text{onFPA}_t \end{array} \right\} \\ \mathcal{D} &= \{\text{goodAngle}_t\} \\ \mathcal{T} &= \{\text{SetValue}(33)_1\} \end{aligned}$$

L'enquête a montré que les pilotes savent initialement que l'indicateur est en VS (Vertical Speed) mais « oublie » cette information et règle la valeur sur 33, ce qui aurait été la bonne valeur s'ils avaient été sur le mode FPA (Flight Path Angle). En appliquant notre algorithme sur ce modèle, nous avons un MCS vide pour l'incohérence locale car l'agent ne peut pas déduire $\neg \text{goodAngle}_1$ vu qu'il n'observe pas onVS_1 au temps 1. Nous avons alors :

$$B_1^{\text{local}} = \left\{ \begin{array}{l} \text{onVS}_0, \neg \text{onFPA}_0, \text{goodAngle}_0, \\ \text{goodAngle}_1, \text{SetValue}(33)_1, \\ R_1, R_2, R_3, R_4 \end{array} \right\}$$

Pour calculer B_1^{frame} nous ajoutons les $\text{keep}^{(v)}$, c'est-à-dire les clauses qui gardent la valeur de vérité des propositions entre deux pas de temps. Aucune incohérence n'est détectée du fait que $\text{known}_{\text{onVS}_1}$ est libre et peut être donc faux pour que le système soit toujours cohérent. Nous avons alors :

$$B_1^{\text{frame}} = \left\{ \begin{array}{l} \text{onVS}_0, \neg \text{onFPA}_0, \text{goodAngle}_0, \\ \text{goodAngle}_1, \text{SetValue}(33)_1, \\ R_1, R_2, R_3, R_4, \\ \text{keep}_{\text{onVS}_1}^{(v)}, \text{keep}_{\text{onFPA}_1}^{(v)} \end{array} \right\}$$

Pour calculer B_1 en prenant en compte la distorsion du décor, nous ajoutons les $\text{keep}^{(k)}$, c'est-à-dire les clauses qui gardent la valeur de connaissance (*known*) entre deux pas de temps : $\{\text{keep}_{\text{onVS}_1}^{(k)}, \text{keep}_{\text{onFPA}_1}^{(k)}\}$ Nous avons alors une incohérence avec la proposition goodAngle_1 car $\neg \text{goodAngle}_1$ peut être maintenant déduite par R_1 . Un MCS possible est alors de retirer $\{\text{keep}_{\text{onVS}_1}^{(k)}, \text{keep}_{\text{onFPA}_1}^{(v)}\}$ et ainsi avoir :

$$B_1 = \left\{ \begin{array}{l} \text{onVS}_0, \neg \text{onFPA}_0, \text{goodAngle}_0, \\ \text{goodAngle}_1, \text{SetValue}(33)_1, \\ R_1, R_2, R_3, R_4, \\ \text{keep}_{\text{onVS}_1}^{(v)}, \text{keep}_{\text{onFPA}_1}^{(k)} \end{array} \right\}$$

Nous avons par cette dernière étape traité le problème de la distorsion du décor : l'inertie du décor n'est pas respectée par l'agent afin d'être cohérent avec son désir d'avoir un bon angle d'approche, en oubliant notamment qu'il était en mode Vertical Speed.

Dans cette section, nous avons illustré notre algorithme avec 3 exemples distincts mais sur une seule étape temporelle. Dans le cas général, les différents MCS obtenus à chacun des pas de temps créent des branches, ce qui produit pour des cas d'études réels des arbres d'états mentaux de grande taille, qui décrivent de nombreux scénarios possibles. Nous discutons de ce point dans les perspectives.

5 Travaux connexes

Le diagnostic en Intelligence Artificielle, c'est-à-dire la recherche d'explications à une situation en utilisant la modélisation logique, est étudié depuis les années 80 [19]. Toutefois à notre connaissance aucun des modèles proposés jusqu'ici n'est appliqué dans le contexte de l'analyse d'erreurs de prise de décision humaine. Bien que plusieurs travaux proposent des solutions pour diagnostiquer des systèmes dynamiques (*i.e* en prenant en compte des actions et des changements) [16, 24] tous supposent que les solutions doivent respecter l'inertie du décor. Le problème de la distorsion du décor qui apparaît lors de l'analyse des erreurs humaines n'est pas pris en compte. L'objectif de notre modèle est de dépasser cette limitation.

Pourtant, depuis quelques années, la recherche en IA s'est intéressée à la modélisation des erreurs de raisonnement humain ou, plus généralement, aux limites du raisonnement humain, à des fins de prédiction en simulation. Par exemple, [26] utilise un automate fini pour simuler la dynamique d'opinion sur la vaccination. Il permet d'expliquer une décision de non-vaccination alors que l'ensemble des informations rationnelles à la disposition de l'agent devrait l'amener à prendre la décision d'accepter la vaccination. Dans un contexte différent, [12] propose les *Synthetic Cognitive Models* pour simuler la prise de décision dans un contexte militaire en prenant en compte la rationalité limitée de l'humain. Enfin, les auteurs de [2] utilisent le paradigme BDI pour implémenter des fonctions probabilistes qui mènent à des croyances erronées dans une situation de feux de forêt. Tous ces modèles proposent une solution viable pour simuler des erreurs de prise de décision par des humains mais ils ne peuvent pas être utilisés dans un objectif de diagnostic dans un cas général.

Une autre approche pour capturer les croyances et les raisonnements erronés consiste à s'affranchir de l'hypothèse d'*omniscience logique* telle que définie dans [9], c'est-à-dire la capacité à inférer toutes les conséquences d'une croyance φ . Par exemple, dans [23], les auteurs proposent un modèle basé sur les *mondes impossibles* (*i.e* les mondes qui ne sont pas fermés sous conséquences logiques) pour simuler des erreurs de raisonnement. Pour cela ils associent à chaque règle de raisonnement une *consommation de ressources*, ce qui limite l'application de longs raisonnements. Toutefois, le calcul de tous les *mondes impossibles* pour sélectionner le plus plausible nécessite une puissance de calcul exponentielle. De plus leur modèle ne considère pas les actions et les changements.

Toutes ces approches donnent des solutions intéressantes bien que partielles à notre problème : elles ne prennent pas en compte le problème de la distorsion du décor et ne sont pas adaptées pour un objectif de diagnostic. Notre modèle reprend les idées proposées par ces différents auteurs pour modéliser des erreurs de prise de décision et calculer des diagnostics qui prennent en compte les erreurs humaines. De plus, l'utilisation d'un solveur SMT nous permet de traiter à la fois des variables continues et des variables discrètes, et d'augmenter ainsi l'expressivité du mo-

dèle, comme l'a montré [7].

6 Conclusion et perspectives

Notre modèle utilise un diagnostic basé sur la révision de croyance pour calculer les suites d'états mentaux possibles qui peuvent expliquer des erreurs de prise de décision humaine. Ces états mentaux sont cohérents avec les observations et les actions effectuées par l'agent et prennent en compte le *problème de la distorsion du décor*, à savoir le fait que les croyances humaines peuvent être modifiées sans cause externe. Alors que le problème du décor stipule que les propositions doivent être maintenues quand elles ne sont pas modifiées par une action, l'analyse des erreurs humaines doit prendre en compte des changements spontanés dans le décor tout en maintenant un comportement rationnel de la part de l'agent.

Nous avons implémenté ce modèle à l'aide d'un solveur SMT. Notre algorithme construit un arbre où chaque chemin correspond à un scénario possible pour les actions observées.

Pour le moment, notre modèle calcule l'ensemble des scénarios possibles mais n'identifie pas le plus plausible. Par exemple, dans le modèle complet du crash Rio-Paris, nous trouvons plus de 6000 scénarios, ce qui est beaucoup pour une analyse par un expert humain. Pour pallier cette limitation, nous envisageons d'étendre notre algorithme de manière à filtrer l'ensemble des scénarios et extraire les erreurs humaines « classiques », identifiées dans la littérature en sciences humaines sous le terme de « biais cognitifs » [25]. Une première proposition a été faite dans [6] pour définir des motifs logiques permettant d'identifier quelques biais cognitifs dans les accidents. Notre proposition à terme est d'inclure ces motifs dans notre modèle et de les étendre pour capturer d'autres biais afin de donner un ordre de priorité aux scénarios présentés aux experts.

À plus long terme, notre objectif est de fournir des outils pour aider les experts à comprendre et anticiper les situations d'accident. Nous pensons qu'en utilisant la modélisation logique, nous pouvons générer des scénarios non anticipés par les experts humains. Nous voulons aussi que notre modèle permette de faire la distinction entre une conception ergonomique médiocre et susceptible de conduire à des erreurs (comme l'accident du Mont-Saint-Odile) et les erreurs humaines qui peuvent être identifiées et évitées par la formation. Ces erreurs peuvent provenir de la charge cognitive et peuvent être modélisées par la rationalité limitée. Nous avons l'intention d'étendre notre cadre logique pour capturer ce phénomène cognitif majeur.

Références

- [1] Carlos E Alchourrón, Peter Gärdenfors, and David Makinson. On the logic of theory change : Partial meet contraction and revision functions. *The journal of symbolic logic*, 50(2) :510–530, 1985.
- [2] Maël Arnaud, Carole Adam, and Julie Dugdale. The role of cognitive biases in reactions to bushfires. In *ISCRAM*, Albi, France, May 2017.

- [3] BEA. Bea f-ed920120. Technical report, Bureau d'Enquêtes et d'Analyses pour la sécurité de l'aviation civile, 1993.
- [4] BEA. Bea f-cp090601. Technical report, Bureau d'Enquêtes et d'Analyses pour la sécurité de l'aviation civile, 2012.
- [5] Leonardo De Moura and Nikolaj Bjørner. Z3 : An efficient smt solver. In *International conference on Tools and Algorithms for the Construction and Analysis of Systems*, pages 337–340. Springer, 2008.
- [6] Valentin Fouillard, Nicolas Sabouret, Safouan Taha, and Frédéric Boulanger. Catching cognitive biases in an erroneous decision making process. *IEEE International Conference on Systems, Man and Cybernetics (SMC)*, 2021.
- [7] Alban Grastien. Diagnosis of hybrid systems with SMT : opportunities and challenges. *ECAI 2014*, pages 405–410, 2014.
- [8] Peter Gärdenfors and Hans Rott. *Belief Revision*. Cambridge Tracts in Theoretical Computer Science. Cambridge University Press, 1992.
- [9] Joseph Y. Halpern and Riccardo Pucella. Dealing with logical omniscience : Expressiveness and pragmatics. *Artificial Intelligence*, 175(1) :220–235, 2011. John McCarthy's Legacy.
- [10] Steve Hanks and Drew McDermott. Nonmonotonic logic and temporal projection. *Artificial intelligence*, 33(3) :379–412, 1987.
- [11] Erik Hollnagel. *Cognitive reliability and error analysis method (CREAM)*. Elsevier, 1998.
- [12] Jonathan Kulick and Paul K Davis. Modeling adversaries and related cognitive biases. *Modeling Adversaries and Related Cognitive Biases*, 2003.
- [13] Mark H Liffiton and Kareem A Sakallah. Algorithms for computing minimal unsatisfiable subsets of constraints. *Journal of Automated Reasoning*, 40(1) :1–33, 2008.
- [14] John McCarthy. Applications of circumscription to formalizing common-sense knowledge. *Artificial intelligence*, 28(1) :89–116, 1986.
- [15] John McCarthy and Patrick J Hayes. Some philosophical problems from the standpoint of artificial intelligence. *Machine Intelligence*, page 463–502, 1969.
- [16] Sheila A McIlraith. Explanatory diagnosis : Conjecturing actions to explain observations. In *Logical Foundations for Cognitive Agents*, pages 155–172. Springer, 1999.
- [17] Atsuo Murata, Tomoko Nakamura, and Waldemar Karwowski. Influence of cognitive biases in distorting decision making and leading to critical unfavorable incidents. *Safety*, 1(1) :44–58, 2015.
- [18] Gabriele Paul. Approaches to abductive reasoning : an overview. *Artificial intelligence review*, 7(2) :109–152, 1993.
- [19] Raymond Reiter. A theory of diagnosis from first principles. *Artificial Intelligence*, 32(1) :57 – 95, 1987.
- [20] Raymond Reiter. The frame problem in the situation calculus : A simple solution (sometimes) and a completeness result for goal regression. In *Artificial and Mathematical Theory of Computation*, pages 359–380. Citeseer, 1991.
- [21] Erik Sandewall. Cognitive robotics logic and its metatheory : Features and fluents revisited. *Electron. Trans. Artif. Intell.*, 2 :307–329, 1998.
- [22] Murray Shanahan. The Frame Problem. In Edward N. Zalta, editor, *The Stanford Encyclopedia of Philosophy*. Metaphysics Research Lab, Stanford University, Spring 2016 edition, 2016.
- [23] Anthia Solaki, Francesco Berto, and Sonja Smets. The logic of fast and slow thinking. *Erkenntnis*, 86(3) :733–762, 2021.
- [24] Michael Thielscher. A theory of dynamic diagnosis. *Electronic Transactions on Artificial Intelligence*, 1(4) :73–104, 1997.
- [25] Amos Tversky and Daniel Kahneman. Judgment under uncertainty : Heuristics and biases. *Science*, 185(4157) :1124–1131, 1974.
- [26] Marina Voinson, Sylvain Billiard, and Alexandra Alvergne. Beyond rational decision-making : modelling the influence of cognitive biases on the dynamics of vaccination coverage. *PloS one*, 10(11), 2015.
- [27] Renata Wassermann. An algorithm for belief revision. In *Proceedings of the Seventh International Conference on Principles of Knowledge Representation and Reasoning*, pages 345–352, 2000.

Optimisation du Positionnement de Voitures en Autopartage basée sur la Prédiction de leur Utilité

G. Martin^{1,2}, M. Donain¹, E. Fromont^{2,3}, T. Guns⁴, L. Roze⁵, A. Termier²

¹ Stellantis

² Univ Rennes, Inria, CNRS, IRISA

³ Institut Universitaire de France (IUF)

⁴ Vrije Universiteit Brussel / KULeuven

⁵ Univ Rennes, INSA, Inria, CNRS, IRISA

{gregory.martin, elisa.fromont, laurence.roze, alexandre.termier}@irisa.fr

Résumé

Le succès d'un service d'autopartage en free-floating dépend d'une bonne répartition des véhicules dans la ville, c'est-à-dire où et quand les utilisateurs en ont besoin. Cela nécessite de prédire la demande des utilisateurs, ce qui représente un défi en raison du peu de données disponibles et de la variabilité de ces demandes. L'objectif de ces prédictions est d'aider à calculer la meilleure relocalisation possible de voitures pour le jour suivant. Ainsi il est nécessaire de modéliser à la fois une tâche de prédiction de la demande et une tâche d'optimisation du placement des véhicules. Comme l'optimisation du placement des voitures implique un raisonnement sur le nombre de voitures à affecter dans des zones précises de la ville, nous proposons de convertir le problème de prédiction de la demande en prédiction de l'utilisation attendue d'une voiture lorsqu'elle est placée dans une zone de la ville. Nous abordons les défis liés à la modélisation de cette prédiction et du problème de la relocalisation des véhicules de la flotte du service d'autopartage. Ainsi, une méthode d'optimisation linéaire en nombres entiers est proposée pour résoudre ce problème de relocalisation tout en tenant compte de la prédiction de l'utilisation des véhicules selon leur placement et des distances de relocalisation. Les expériences se basent sur deux ensembles de données provenant de deux services d'autopartage en free-floating urbains et nous montrons comment notre méthode peut améliorer les stratégies de relocalisation et donc la rentabilité des services.

Mots-clés

Autopartage, Optimisation Linéaire en Nombres Entiers, Régression.

Abstract

The success of a free-floating car-sharing service depends on a good allocation of the vehicles across the city, i.e., where and when they are needed by citizens. This requires predicting the demand across the geographical regions and across time, which is challenging due to the sparsity and variability of the data. Furthermore, the purpose of these pre-

dictions is to help compute the best possible car relocation for the next day, hence the need to model both the prediction task and the optimization task in a compatible way. As the allocation optimization involves reasoning about the number of cars to assign to geographical regions, we propose to convert the prediction problem into predicting the expected utilization of a car when added to a region. We discuss the challenges in modeling both the machine learning and the relocation problem, and we propose a mixed-integer linear programming method that solves the relocation problem while taking into account the model predictions and relocation distances. We experiment with two datasets from citywide car sharing companies and show how our method can increase the allocation strategies and hence profitability of the services.

Keywords

Carsharing, Integer Programming, Regression

1 Introduction

L'autopartage en *free-floating* a émergé comme une nouvelle méthode de transport public dans les zones urbaines. Il contribue à réduire l'empreinte carbone des villes en permettant à des voitures d'être partagées par de multiples utilisateurs en ville. Contrairement aux services d'autopartage *one-way*, il n'y a pas de stations prédéfinies dans un service en *free-floating* : les voitures peuvent être prises et garées sur n'importe quelle place de parking, généralement en ville et dans ses environs.

Quand les utilisateurs d'un service d'autopartage ont besoin d'une voiture, il est indispensable que des véhicules assez proches soient disponibles. À cette fin, les opérateurs emploient du personnel, appelé *jockeys*, pour repositionner les véhicules dans la ville. Cette opération est généralement effectuée de manière empirique, principalement la nuit et à petite échelle, avec une connaissance approximative des zones à forte demande où les véhicules doivent être déplacés. Plusieurs articles ont constaté l'absence de stratégies de relocalisation précise, et ont conçus des méthodes

qu'un opérateur doit appliquer [16] ou des incitations afin que les utilisateurs envisagent de terminer leurs trajets dans des zones sélectionnées par l'opérateur [2].

Dans cet article, nous cherchons à maximiser l'utilisation de la flotte de voitures d'un service afin de rendre le service durable et rentable pour l'opérateur. Nous considérons l'utilisation de jockeys pour relocaliser les voitures durant la nuit, comme seule action quotidienne concrète prise par le service d'autopartage pour augmenter l'utilisation de la flotte de véhicules. Par conséquent, l'objectif est de proposer un emplacement pour toutes les voitures de la flotte au début de la journée, de manière à maximiser l'utilisation durant la journée prédite des voitures.

Il existe plusieurs stratégies pour résoudre ce problème. La plus précise consisterait à prévoir les déplacements exacts des clients pour un jour donné, c'est-à-dire le nombre de déplacements, et pour chaque déplacement, l'origine, la destination et l'heure de départ. En sachant cela, nous pourrions envisager les différentes manières de relocaliser les voitures au début de la journée afin de maximiser la satisfaction du client. Ainsi, cela nous permettrait de modéliser la dynamique de la flotte afin de simuler différentes incitations au stationnement et stratégies de relocalisation des voitures. Cependant cela est en pratique extrêmement difficile à réaliser. À la place, nous proposons une stratégie alternative moins fine, mais efficace qui est décomposée en deux parties. Tout d'abord, nous prédisons l'utilisation prévue des voitures lorsqu'elles sont placées dans des zones bien définies en ville au début de la journée, la prédiction est faite en minutes d'utilisation. Ensuite en fonction de ces prédictions, une stratégie de relocalisation des voitures est conçue afin de maximiser l'utilisation totale attendue du service en une journée. Le problème de relocalisation est présenté sous la forme d'un problème d'optimisation linéaire en nombres entiers sur les prédictions, assujéti à des contraintes pratiques sur la relocation des véhicules.

La Section 2 expose les travaux de l'état de l'art. Une modélisation générale de la ville, une stratégie de prédiction de l'utilité d'une voiture et une méthodologie d'affectation des voitures sont exposées dans la Section 3. La prédiction de l'utilité des voitures et la relocalisation des véhicules sont évaluées dans la Section 4 avec deux ensembles de données propriétaires provenant de services réels. Nous concluons dans la Section 5.

2 État de l'Art

Prédiction des trajets pour les services de partage de véhicules. Dans un contexte de vélopartage, [7] proposent une méthode pour une stratégie de relocalisation par l'opérateur où les vélos peuvent être relocalisés à tout moment de la journée. Leur solution n'est pas directement applicable à notre contexte d'autopartage, car la relocalisation des voitures n'est possible que la nuit dans notre cas. De plus, les données sont beaucoup plus rares dans un contexte d'autopartage. Une autre stratégie pour localiser les voitures tous les matins serait de pouvoir prédire la demande des clients et tous les déplacements effectués au cours d'une

journée. Cette stratégie est liée à la prédiction des matrices d'origine-destination. [8] ont proposé une telle méthode pour un service de vélopartage *one-way*, où les vélos peuvent être pris dans une station et laissés dans une autre. Dans notre contexte d'autopartage, le nombre de points d'origine et de destination à considérer est beaucoup plus élevé que pour les services de vélopartage et le nombre de trajets est beaucoup plus faible : il en résulte des tenseurs binaires origine/destination très creux, en considérant le temps comme troisième dimension. Les propriétés de nos tenseurs empêchent également d'utiliser des méthodes statistiques standard telles que (S)ARIMA [15] pour des prédictions (saisonnnières) de couple origine-destination d'intérêt. Par rapport au vélopartage, une autre difficulté vient de la dépendance entre les déplacements en voiture et la forte influence des conditions de circulation non saisonnières (travaux de construction, événements spéciaux, conditions météorologiques extrêmes) et des particularités de la ville (présence de bâtiments culturels, d'écoles, de supermarchés, etc.).

Dans [13], il est montré que la location des voitures dépend de facteurs comme la géographie ou la condition sociale. Ces facteurs sont difficiles à prendre en compte sans une grande quantité d'informations exogènes. Nous discutons de l'utilité de certains d'entre eux dans la Section 4. Le problème de la prédiction de déplacements binaires dans un tenseur peut être lié à la prédiction de liens dans un graphe/réseau dynamique [1]. Cependant, la définition d'un score qui refléterait l'affinité entre deux zones impliquant l'existence d'un voyage à un moment donné, de manière similaire à ce qui peut être fait pour les personnes et les liens dans un réseau social, est difficile et source d'erreurs dans notre contexte. En particulier parce que dans un contexte de *free-floating*, chaque paire possible de noeuds doit être considérée, ce qui empêche l'utilisation de scores de similarité basés sur la topologie du réseau comme ceux présentés dans [10].

Stratégies de relocalisation pour les services de partage de véhicules. Plusieurs travaux dans la littérature sur le vélopartage [4, 5, 11] abordent le problème de la relocalisation des vélos. Par exemple, dans [11], les auteurs optimisent conjointement le nombre de trajets par jour et certains coûts de relocalisation. Ainsi est prise en compte la distance totale parcourue par les camions relocalisant les vélos, sachant qu'ils peuvent déplacer plusieurs vélos par trajet. Dans [4, 5] les auteurs améliorent cette solution en proposant des relocalisations continues dans la journée tout en prenant en compte la demande des usagers au cours de la journée.

De nombreuses stratégies de relocalisation ont également été proposées pour les voitures : dans [17], les auteurs proposent une optimisation linéaire en nombres entiers (*ILP*) pour optimiser l'utilisation des voitures dans un service d'autopartage *one-way* où les relocalisations se font en continu pendant la journée. Leur modèle d'*ILP* vise à réduire le nombre de demandes non satisfaites durant la journée, tout en minimisant le nombre de trajets de relocalisa-

tions effectués par le personnel. Les stratégies proposées dans [3] et [16] ont été conçues pour un service d'autopartage en *free-floating*, mais elles ne considèrent que le problème de relocalisation en soi : étant donné en entrée une distribution initiale des voitures dans la ville et une distribution finale attendue, elles minimisent les coûts de relocalisation. Cependant il n'y a aucune prise en compte des bénéfices apportés par les voitures bien localisées.

Comme nous l'avons expliqué dans l'introduction, nous prenons ces solutions comme inspiration sans pouvoir les appliquer directement à notre problème d'autopartage en *free-floating*. En effet dans notre problème il n'y a pas de stations fixes, la flotte contient peu de voitures, par rapport à des vélos d'un service de vélopartage, un jockey ne peut déplacer qu'une voiture à la fois, les relocalisations ont lieu pendant la nuit et nous sommes intéressés par l'optimisation du nombre de minutes conduites et non du nombre de trajets.

3 Méthodologie

Notre objectif est de proposer une affectation pour le lendemain matin de toutes les voitures d'un service d'autopartage en *free-floating* de telle sorte à ce qu'elle maximise leur utilisation prédite ce jour-là. Nous détaillons d'abord la modélisation géographique de la ville, puis la manière dont le problème de prédiction de l'utilité d'une voiture est réalisé et enfin nous définissons le modèle de relocalisation des voitures selon la prédiction de leur usage.

3.1 Modélisation de la Ville et des Trajets

Pour des raisons d'efficacité et de généralité, il n'est pas possible de modéliser une ville comme étant avec l'ensemble de toutes les places de stationnement. D'après les résultats de l'étude [14], nous avons discrétisé la surface de la ville en une grille hexagonale où chaque hexagone (aussi appelé *cellule* par la suite) a 500 mètres pour rayon. Cette distance représente ce qu'est prêt à parcourir un client trouver une voiture disponible [14]. L'ensemble des hexagones de cette grille est noté K . Toutes les positions situées dans la surface d'une cellule sont étiquetées avec l'indice k de l'hexagone correspondant. Notez que les hexagones inutiles sont supprimés, ce sont ceux dans lesquels aucun déplacement n'a commencé ou ne s'est terminé. Dans la suite, nous pouvons également regrouper, par clustering, les cellules (hexagones) qui se comportent de manière similaire selon certains critères donnés. Notez que cela revient simplement à avoir un petit nombre K de cellules différentes.

3.2 Modélisation de l'Utilité

Il est supposé que les données historiques sont un ensemble de *trajets* utilisateur, à savoir un ensemble de tuples tel que : $\text{trajets} = \{(id_{vehicule}, ts_{depart}, ts_{arrivee}, gps_{depart}, gps_{arrivee}, distance, duree)\}$. Les trajets de jockeys, c'est-à-dire les déplacements des employés dans le but de repositionner les véhicules pendant la nuit, sont exclus des données. À partir de l'horodatage (ts), des caractéristiques supplémentaires sont déduites telles que des informations météorologiques sur la température

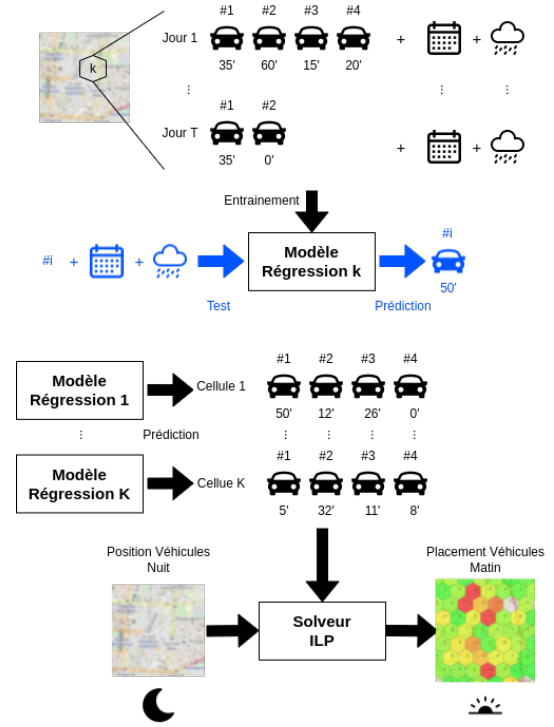


FIGURE 1 – Vue d'ensemble de notre approche : un modèle de régression est entraîné pour chaque cellule afin de prédire, à partir de données historiques, le temps d'utilisation de chaque voiture (par ordre de départ) sachant des attributs temporels et météorologiques. Le solveur d'optimisation utilise à la fois les régresseurs appris pour prédire le temps d'utilisation de chaque voiture et le positionnement de la flotte avant la phase de relocalisation pour obtenir le meilleur positionnement des voitures pour le lendemain matin.

et les précipitations. Des informations de lieux de départ et d'arrivée (*gps*) sont fournies sous forme de coordonnées GPS, qui sont mises en correspondance avec les cellules géographiques précédemment définies. De plus, des informations supplémentaires sur le trajet sont disponibles, comme les kilomètres parcourus (*distance*) ou la durée du trajet en minutes (*duree*). Un résumé de ces informations est présenté en haut de la Figure 1.

Utilité Réelle. À partir des données historiques d'une seule journée, nous voulons calculer le temps d'utilisation total, c'est-à-dire l'utilité totale, $U_k^*(n)$ dans chaque cellule k étant donné la présence de n voitures disponibles dans cette cellule au début de la journée. Il n'est pas suffisant de compter uniquement l'utilité des déplacements de cette cellule vers d'autres cellules, car les voitures peuvent avoir été utilisées pour des trajets ultérieurs au cours de la journée. À partir de ce constat, l'utilité totale est calculée pour chaque véhicule :

$$U^*(id_{vehicule}) = \sum_{t \in \text{trajets}(id_{vehicule})} utilite(t)$$

où $\text{trajets}(id_{vehicule})$ est l'ensemble des trajets pour une

journée dont l'identifiant du véhicule est $id_{vehicule}$ et l'utilité $utilite(t)$ d'un trajet t correspond à la durée du trajet. Soit $matin_vehicules(k)$ l'ensemble des identifiants des véhicules qui se trouvaient dans la cellule k au début de la journée, avec $|matin_vehicules(k)| = n$. L'utilité réelle d'une cellule k qui a eu n véhicules initiaux est mesurée tel que :

$$U_k^*(n) = \sum_{id_{vehicule} \in matin_vehicules(k)} U^*(id_{vehicule})$$

Notez que cette valeur ne peut être calculée avec les données historiques que pour $n = |matin_vehicules(k)|$.

Utilité Estimée. Notre objectif est maintenant de prédire $U_k(n)$ pour toute cellule k et toute valeur n , sans savoir quels trajets seront effectués et quels véhicules seront utilisés pour ceux-ci. Ainsi, il s'agit d'un problème de régression où les données d'apprentissage sont les utilités réelles définies ci-dessus. Cependant, comme seuls les $U_k^*(n)$ sont observés pour un n par cellule par jour, cela donne un ensemble d'apprentissage creux et peu susceptible de conduire à un modèle de régression correct pour des valeurs arbitraires de n .

Ainsi le problème de régression de $U_k^*(n)$ est décomposé, afin d'avoir plus de données d'entraînement et des prédictions plus fines. Les voitures sont d'abord ordonnées par ordre de première utilisation (première voiture, deuxième voiture, etc.). Le but est de modéliser la relation temporelle entre les véhicules par rapport au moment de leur première utilisation dans la journée. Il est maintenant possible de définir la valeur ajoutée de la i^{eme} voiture comme $U_k^i = utilite(id_{vehicule}^i)$ avec $id_{vehicule}^i$ l'identifiant du véhicule initialement dans la cellule k qui a été ordonné (indice i) par horodatage de première utilisation. Nous avons alors $U_k^*(n) = \sum_{i \in 1..n} U_k^i$. À partir de cette reformulation, nous identifions l'objectif de prédiction du *gain d'utilité* U_k^i apporté par la i -ème voiture dans une cellule k . La décomposition de la valeur d'utilité et ce problème de régression plus fin permettent d'avoir plus de données d'entraînement, et donc de mieux estimer $U_k^*(n)$ pour différentes valeurs de n (notamment pour $n < |matin_vehicules(k)|$).

Prédiction. L'objectif est de prédire U_k^i pour chaque cellule k et rang i en utilisant les données d'apprentissage comme définies précédemment, cette prédiction est appelée $\hat{U}_k^i$. Comme les cellules peuvent avoir un comportement différent, un modèle de régression f_k sera appris par cellule. Comme attributs, nous utilisons le rang i et un ensemble fs d'attributs exogènes comprenant des attributs météorologiques, des attributs temporels et les observations d'utilité des jours précédents (car les données sont séquentielles). Par conséquent, un régresseur f_k est appris par cellule tel que

$$\hat{U}_k^i = f_k(i, fs).$$

Nous discutons des modèles de régression à utiliser dans la Section 4.3.

Clustering des Cellules. Au lieu d'apprendre un prédicteur f_k par cellule, il est possible d'utiliser les similarités dans le comportement de l'utilité pour différentes cellules.

En regroupant les cellules, davantage de données sont disponibles pour l'entraînement de chaque modèle et, idéalement, les prédictions sont plus précises. Dans la Section 4 est étudié l'effet du *clustering* des cellules dont le comportement est similaire en termes d'utilité, avec un modèle f_c appris par *cluster* c de cellule.

3.3 Optimisation de la Relocalisation des Véhicules

Relocalisation Naïve. Étant donné $\hat{U}_k^i$, l'utilisation prédite de la i ème voiture située dans la cellule k , et une limite Γ sur le nombre de voitures disponibles, le problème de relocalisation des voitures peut se résumer à déterminer la meilleure assignation de ces Γ voitures dans les différentes cellules. Il faut noter que si $\hat{U}_k^i$ était linéaire en i , par exemple $\hat{U}_k^i = i \cdot w$ pour une certaine valeur w , alors le problème de relocalisation serait un problème de sac à dos non borné de capacité Γ , où chaque cellule représente un élément. Cependant cette simple hypothèse de linéarité ne tient pas, une formulation différente est utilisée pour modéliser le problème de relocalisation. Plus précisément, nous utilisons l'optimisation linéaire en nombres entiers (*ILP*) pour le modéliser comme un problème d'affectation des véhicules aux cellules, où la capacité totale Γ et l'ordre temporel i des véhicules dans chaque cellule doivent être respectés. Des variables booléennes V_{ki} sont utilisées pour indiquer s'il y a un i -ème véhicule dans la cellule k . Le nombre total de véhicules dans une cellule k est donc $\sum_i V_{ki}$. Avec K l'ensemble de tous les identificateurs de cellule définis précédemment, m étant une limite supérieure du nombre de véhicules autorisés dans une cellule et $I = \{1..m\}$. Soit $V = (V_{ki})_{k \in K, i \in I}$, alors la formulation *ILP* du problème est :

$$\underset{V}{argmax} \quad \sum_{k \in K} \sum_{i \in I} \hat{U}_k^i V_{ki} \quad (1)$$

$$s.t. \quad \sum_{k \in K} \sum_{i \in I} V_{ki} \leq \Gamma \quad (2)$$

$$V_{ki} \geq V_{ki+1} \quad \forall k \in K, i \in I \setminus \{m\} \quad (3)$$

$$V_{ki} \in \{0, 1\} \quad \forall k \in K, i \in I \quad (4)$$

L'équation (2) est une contrainte de capacité, tandis que l'équation (3) garantit qu'il n'y a pas d'« écart » dans le rang des variables indicatrices, c'est-à-dire que si $V_{ki} = 0$, alors tout $j > i$ doit être 0 également. Par exemple, s'il n'y a pas de troisième véhicule dans une cellule, il n'y a pas non plus de quatrième véhicule, cinquième véhicule, etc.

Relocalisation basée sur les Jockeys. La formulation précédente calcule le placement optimal, mais ne tient pas compte du fait que 1) à la fin de la journée, les véhicules sont déjà situés dans certaines cellules, et 2) les voitures doivent être relocalisées par des jockeys, d'où la nécessité de prendre en compte leur capacité limitée et les coûts associés à chaque relocalisation.

Pour répondre à la remarque 1), nous désignons par $s_{ki} \in \{0, 1\}$ la présence ou absence d'une i -ème voiture dans la

cellule k à la fin du jour précédent. Pour répondre à la remarque 2), nous incluons une limite γ sur le nombre total de relocalisations possibles de véhicules que les jockeys peuvent effectuer. De plus les jockeys ne pouvant pas se téléporter, ils doivent donc être déplacés par une voiture dite « balai » vers la voiture suivante après avoir relocalisé une voiture précédente. D'un point de vue opérationnel, ces trajets en voiture-balai doivent être minimisés, c'est pourquoi est introduit q_{kl} le coût du déplacement d'un jockey par une voiture-balai de la cellule k à la cellule l ; et p le prix par minute payé par les utilisateurs qui louent une voiture afin que les profits/coûts de la relocalisation puissent être équilibrés. L'objectif est maintenant de trouver la relocalisation qui conduit au plus grand revenu prédit possible moins les coûts de dépose des jockeys, tout en devant effectuer au plus γ relocalisations. Pour modéliser ce problème sous la forme d'un programme linéaire en nombres entiers, nous utilisons deux ensembles de variables de décision booléennes : F_{ki} indique que la i -ème voiture doit être retirée cellule k , tandis que T_{ki} indique qu'une voiture doit être relocalisée dans la cellule k pour être sa i -ème voiture. En utilisant ces variables binaires, la présence ou non d'une voiture à la i -ème position dans la cellule k est calculée par $s_{ki} - F_{ki} + T_{ki}$, qui remplace V_{ki} dans la formulation précédente. Enfin, E_{kl} représente le nombre de trajets de voiture-balai nécessaires entre les cellules k et l , avec q_{kl} le coût d'un de ces trajets de staff.

En intégrant ce résultat dans le problème d'optimisation précédent et en ajoutant les contraintes supplémentaires appropriées, on obtient la formulation suivante :

$$\begin{aligned} \underset{F, T, E}{\operatorname{argmax}} \quad & \sum_{k \in K} \sum_{i \in I} p \cdot \hat{U}_{ki} \cdot (s_{ki} - F_{ki} + T_{ki}) \\ & - \sum_{k \in K} \sum_{l \in K} q_{kl} \cdot E_{kl} \end{aligned} \quad (5)$$

such that :

$$s_{ki} - F_{ki} \geq 0 \quad \forall k \in K, i \in I \quad (6)$$

$$s_{ki} + T_{ki} \leq 1 \quad \forall k \in K, i \in I \quad (7)$$

$$\sum_{k \in K} \sum_{i \in I} T_{ki} = \sum_{k \in K} \sum_{i \in I} F_{ki} \quad (8)$$

$$\sum_{k \in K} \sum_{i \in I} T_{ki} \leq \gamma, \sum_{k \in K} \sum_{i \in I} F_{ki} \leq \gamma \quad (9)$$

$$(s_{ki} - F_{ki} + T_{ki}) \geq (s_{ki+1} - F_{ki+1} + T_{ki+1}) \quad \forall k \in K, i \in I \setminus \{m\} \quad (10)$$

$$\sum_{l \in K} E_{kl} = \sum_{j \in I} T_{kj} \quad \forall k \in K \quad (11)$$

$$\sum_{k \in K} E_{kl} = \sum_{i \in I} F_{li} \quad \forall l \in K \quad (12)$$

$$F_{ki} \in \{0, 1\}, T_{ki} \in \{0, 1\} \quad \forall k \in K, i \in I \quad (13)$$

$$E_{kl} \in \{0, \dots, |I|\} \quad \forall k, l \in K \quad (14)$$

La fonction objectif à l'équation (5) maximise le revenu obtenu grâce à l'utilité prédite par rapport au placement des véhicules tout en prenant en compte les coûts des trajets de voiture-balai pour emmener les jockeys entre les voitures à déplacer. L'équation (6) indique que seuls les véhicules présents à la fin de la journée peuvent être déplacés pendant la nuit, tandis que l'équation (7) indique que si un véhicule est déjà présent à cet emplacement, un autre véhicule ne peut pas être ajouté. Ces deux contraintes indiquent effectivement que, selon s_{ki} , soit F_{ki} , soit T_{ki} sera égal à 0. L'équation (8) indique que le nombre de prises et de déposes de jockeys doivent être les mêmes, tandis que l'équation (9) garantit que le nombre de voitures relocalisées est inférieur au budget de relocalisation γ . L'équation (10) garantit qu'il n'y a pas d'« écart » dans le positionnement final, cette contrainte est obtenue en substituant $s_{ki} - F_{ki} + T_{ki}$ dans l'équation (3) de la première formulation. Enfin, l'équation (11) s'assure qu'à chaque voiture relocalisée en k , alors une voiture-balai doit récupérer le jockey. De la même manière, l'équation (12) assure qu'il y a une voiture-balai qui dépose un jockey pour chaque voiture retirée de l .

4 Expérimentations

4.1 Données

Les expériences sont réalisées sur deux jeux de données provenant de deux services réels d'autopartage en *free-floating*. Ces jeux de données étant sensibles pour l'entreprise dont ils sont issus, nous les appelons « ServiceA » et « ServiceB ».

Le premier jeu de données représente 1138246 trajets effectués entre mi-2018 et début-2019. Les données couvrent 243 jours pour une moyenne de 4684 trajets quotidiens. Les données ont été séparées en trois sous-ensembles : ensemble d'entraînement, de validation et de test. L'ensemble d'entraînement regroupe les trajets des premiers 145 jours, l'ensemble de validation sont les 48 jours suivants et l'ensemble de test sont les derniers 50 jours.

Le second jeu de données est constitué de 130219 trajets effectués entre début-2019 et début-2020. Les données couvrent 306 jours pour une moyenne de 425 trajets quotidiens. Comme pour le premier jeu de données, trois sous-ensembles ont été créés : l'ensemble d'apprentissage correspond aux 204 premiers jours, l'ensemble de validation aux 51 jours suivants et enfin l'ensemble de test aux 51 derniers jours.

Chaque trajet dans les deux jeux de données possède l'identifiant du véhicule, les positions GPS de départ et d'arrivée et les horodatages de départ et d'arrivée. Déduit des trajets, des attributs temporels sont rajoutés, tels que le fait qu'un jour soit ouvré ou non et le jour de la semaine.

Deux ensembles de données météorologiques sont utilisés comme données exogènes pour entraîner les régresseurs. Ils sont récupérés à partir des bulletins météorologiques horaires SYNOP de la station météorologique la plus proche pour chaque service. À partir des bulletins SYNOP, nous extrayons la température (°C), l'humidité relative (%), la

pression (hPa), la vitesse du vent (km/h), la couverture nuageuse (%) et la quantité de pluie (mm). Certains rapports horaires étant manquants, nous avons calculé les valeurs manquantes en considérant un changement linéaire entre les données connues avant et après l'entrée manquante.

4.2 Algorithme pour le Clustering des Cellules

Nous comparons une stratégie sans *clustering* (*No Clustering*) où chaque fonction f_k est entraînée pour toutes les cellules k de la grille de la ville. Une comparaison est faite à la stratégie basée sur un *K-Medoids clustering*. Ainsi, chaque cellule est représentée par un vecteur tel que la i -ième coordonnée du vecteur est le temps d'utilisation moyen en minutes de la voiture de rang i . Comme la taille du vecteur décrivant chaque cellule dépend du rang le plus élevé de la voiture rencontré dans cette cellule particulière, tous les vecteurs de cellule n'ont pas la même taille. Ainsi, la distance entre des vecteurs de tailles différentes, i et $i + h$ (avec $h \geq 0$), n'est calculée que sur les i premiers rangs de voiture communs entre les deux vecteurs. La distance euclidienne entre ces deux parties communes divisée par sa longueur commune est donnée à l'algorithme de clustering. Le nombre optimal de clusters est trouvé en utilisant l'ensemble de validation pour chaque modèle et ensemble d'attributs. L'utilité d'une cellule k devient $U_k^i = f_c(i, fs)$ où f_c est l'utilité prédite du cluster c qui contient la cellule k (avec fs les attributs).

La zone desservie par « ServiceA » est couverte par 151 cellules et la zone desservie par « ServiceB » est couverte par 199 cellules. Ce plus grand nombre de cellules dans le second cas (avec moins de voitures dans la flotte) aura des conséquences dans les tâches de prédiction et d'optimisation suivants.

4.3 Algorithme de Prédiction d'Utilité

Les régresseurs testés sont 1) la moyenne des valeurs historiques (référence), 2) un régresseur par *gradient boosting* (GBR) et 3) un régresseur par machine à vecteur de support (SVR). Différents ensembles d'attributs (FS_1 , FS_2 , FS_3 , FS_4) ont été utilisés pour l'entraîner des modèles.

Attributs. Le premier ensemble d'attributs FS_1 ne contient aucun attribut exogène, de sorte que seul le rang i de chaque voiture est utilisé pour entraîner les régresseurs. Le deuxième ensemble d'attributs FS_2 contient des informations exogènes : le *jour de la semaine*, le fait que ce jour soit un *jour ouvrable ou non* et les *informations météorologiques*. Les informations météorologiques sont la moyenne des bulletins horaires pour chaque attribut météorologique et jour. Ainsi FS_2 contient donc 13 attributs. Le troisième ensemble d'attributs FS_3 en contient deux concernant l'historique en plus de FS_2 pour une cellule donnée k avec un rang donné i : les deux dernières (dans le temps) valeurs d'utilité $f_k(i)$ connues dans les données. Enfin, le quatrième ensemble FS_4 contient l'état de charge des véhicules électriques au début de la journée, en plus de FS_3 . Cet ensemble n'est disponible que pour le jeu de données de « ServiceB ».

Modèles. Notre modèle de référence (*Moy*) est la moyenne des valeurs d'utilité historique pour chaque cellule, pour un rang donné i pour tous les jours de l'ensemble d'apprentissage. Ce modèle de référence utilise uniquement FS_1 .

Les implémentations de Scikit-learn¹ sont utilisés pour le *Gradient Boosting Tree* (GBR) et la machine à vecteur de support (SVR) avec des paramètres choisis par validation. Chaque modèle prédit directement l'utilité d'une voiture dans chaque cellule pour les quatre ensembles d'attributs mentionnés précédemment et avec ou sans clustering préalable des cellules. Chaque attribut a été normalisé en soustrayant la moyenne et en divisant par l'écart-type.

À titre d'information, d'autres approches et stratégies d'apprentissage automatique ont été testées, mais ont donné des résultats peu convaincants. Il s'agit notamment d'une approche alternative en deux étapes : un modèle de classification est d'abord utilisé pour prédire si une voiture quitte ou non la cellule et, dans le cas échéant, un deuxième modèle est utilisé pour prédire l'utilité de cette voiture. Comme le premier modèle de classification n'est pas parfait (environ 70% de précision), l'erreur est propagée à la deuxième étape et l'ensemble du processus en deux étapes donne des résultats moins bons que les stratégies de prédiction actuelles. Une deuxième approche testée n'utilise pas i comme attribut et prédit directement l'ensemble du vecteur d'utilité pour chaque cellule, c'est-à-dire qu'elle utilise des modèles de régression multiple. Cette approche a obtenu de moins bons résultats que le modèle de référence *Mean*, et n'est donc pas retenu.

4.4 Métrique d'Évaluation

Nous utilisons le *Mean Absolute Error* (MAE) pour mesurer le nombre moyen de minutes surestimées ou sous-estimées par chaque modèle de régression entraîné pour chaque cellule pendant un jour, défini comme : $MAE = \frac{1}{|K|} \sum_{k \in K} \sum_{i \in I} |obs_{ki} - pred_{ki}|$, avec obs_{ki} l'utilité réelle de la $i^{ème}$ voiture dans la cellule k et $pred_{ki}$ l'utilité prédite de la $i^{ème}$ voiture dans la cellule k .

Nous utilisons aussi un ratio pour chaque jour entre la MAE et le temps d'utilisation réel moyen par cellule, appelé ici le *Ratio Mean Absolute Error* (RMAE) :

$$RMAE = \frac{MAE}{(1/|K|) \cdot \sum_{k \in K} \sum_{i \in I} obs_{ki}}$$

4.5 Stratégies de Relocalisation

Stratégies. Afin d'évaluer notre approche *ILP* pour relocaliser les véhicules en ville pour le lendemain matin, comme décrit dans la Section 3.3, nous allons comparer le revenu attendu de quatre stratégies différentes. Nous désignons par *Historique* la stratégie basée sur la position historique de la flotte chaque matin, et son utilisation réelle au cours de la journée. Le revenu correspondant est le nombre de minutes d'utilisation total multiplié par le prix payé par les utilisateurs pour la location d'une voiture. *Optim MU* désigne une stratégie de relocalisation qui maximise (avec notre modèle

1. <https://scikit-learn.org/>

Clustering	Attributs	ServiceA		
		Moy	GBR	SVR
Aucun Clustering	FS_1	288 ± 27 (45%)	288 ± 27 (45%)	289 ± 28 (45%)
	FS_2	N/A	288 ± 25 (44%)	283 ± 21 (44%)
	FS_3		290 ± 24 (45%)	284 ± 21 (44%)
KMedoids	FS_1	289 ± 27 (45%)	289 ± 28 (45%)	288 ± 28 (45%)
	FS_2	N/A	278 ± 25 (43%)	278 ± 25 (43%)
	FS_3		278 ± 24 (43%)	277 ± 24 (43%)
Clustering	Atributs	ServiceB		
		Moy	GBR	SVR
Aucun Clustering	FS_1	90 ± 25 (92%)	90 ± 25 (91%)	92 ± 31 (90%)
	FS_2	N/A	97 ± 26 (99%)	101 ± 32 (99%)
	FS_3		97 ± 25 (98%)	100 ± 32 (98%)
	FS_4		94 ± 25 (95%)	100 ± 32 (98%)
KMedoids	FS_1	90 ± 25 (92%)	91 ± 26 (92%)	91 ± 31 (88%)
	FS_2	N/A	90 ± 26 (91%)	98 ± 31 (95%)
	FS_3		90 ± 25 (90%)	95 ± 31 (93%)
	FS_4		85 ± 26 (84%)	95 ± 31 (92%)

TABLE 1 – Performance *MAE* (la plus faible est la meilleure) avec l'écart-type et *RMAE* de la prédiction d'utilité des voitures par jour et par (groupe de) cellule(s), en utilisant une base historique (*Moy*) et deux modèles de régression (*GBR* et *SVR*) avec différents ensembles d'attributs (FS_1 , FS_2 , FS_3 et FS_4) dans le cas du « ServiceA » puis du « ServiceB ».

d'ILP et notre prédiction d'utilité) les utilités des voitures sans tenir compte des coûts. *Optim SO* désigne la stratégie, décrite dans la Section 3.3, qui utilise un objectif unique. Cet objectif maximise le revenu de l'utilisation du service moins les coûts de trajets des voitures-balais. *Optim DO* désigne une stratégie d'optimisation en deux étapes : la première maximise l'utilité de la flotte, comme pour *Optim MU*, tandis que la seconde, après avoir fixé l'emplacement optimal des véhicules, minimise les coûts des trajets de voiture-balai.

Prédiction de l'Utilité. L'utilité utilisée pour évaluer le revenu gagné par le service avec notre approche est un mélange entre l'utilité historique et l'utilité prédite : lorsque le solveur ILP propose une solution de relocalisation différente de l'historique, l'utilité prédite est utilisée pour les voitures pour lesquelles cette valeur est inconnue.

Conversion d'Utilité et Estimation des Coûts. Pour toutes les stratégies, nous supposons que le prix p payé par l'utilisateur par minutes de conduite est constant. Ainsi le revenu attendu de la flotte est : l'utilité totale multipliée par p . Le coût d'un trajet de la voiture-balai q_{kl} , décrit dans la Section 3.3, est calculé en fonction du salaire horaire brut d'un jockey c_j qui doit être doublé (lors du repositionnement d'un jockey de k à l , il y a deux jockeys dans la voiture-balai : celui qui doit être emmené et le conducteur de la voiture-balai), avec le coût c_v de fonctionnement d'une voiture-balai par km, s la vitesse moyenne (en km/h) d'une voiture à l'intérieur de la ville et la distance d_{kl} (en km) entre la cellule k où un jockey est pris en charge par la voiture-balai et la cellule l où il est déposé :

$$q_{kl} = \frac{d_{kl} \cdot (2 \cdot c_j)}{s} + d_{kl} \cdot c_v$$

De plus après une observation du fonctionnement du Servi-

ceA, une hypothèse est faite sur un maximum de 70 voitures relocalisées chaque nuit, donc $\gamma = 70$, ce qui correspond à 7 jockeys qui peuvent relocaliser 10 voitures chacun.

4.6 Étude de Cas

Nous voulons évaluer à la fois la partie prédiction et la partie relocalisation de notre approche pour les deux services. Dans le premier cas, nous évaluons la qualité de la prédiction de $f_k(i, fs)$ pour une cellule k , un rang de voiture i et un ensemble de caractéristiques fs avec les différents régresseurs décrits dans la Section 4.3. Dans le deuxième cas, nous calculons le revenu quotidien attendu de chaque système d'autopartage lorsque l'utilité est prédite et utilisée pour résoudre le problème de relocalisation.

4.6.1 Performance des Modèles de Régression

Les colonnes du Tableau 1 présentent la performance des régresseurs utilisés, c'est-à-dire la moyenne historique (*Moy*), le *Gradient Boosting Tree* (*GBR*) et la machine à vecteur de support (*SVR*). Les lignes décrivent les résultats avec ou sans clustering en utilisant les trois (ou quatre si disponibles) ensembles d'attributs présentés précédemment. La *MAE* (en min) et *RMAE* (en %) sont calculés en fonction de l'utilité réelle dans chaque cellule et chaque jour et au nombre prédit d'utilité sur l'ensemble de test. Comme indiqué précédemment, la performance de la moyenne historique n'a pas été calculée en utilisant les ensembles d'attributs FS_2 , FS_3 et FS_4 car sa validité statistique serait discutable.

ServiceA. On peut remarquer qu'avec FS_1 , la moyenne historique donne l'une des pires *MAE* (~288 mins), ce qui montre le potentiel des modèles d'apprentissage automatique et des attributs que nous avons choisis lorsque les données sont suffisantes.

L'ajout d'attributs temporels (le jour de la semaine et le

fait qu'il s'agisse d'un jour ouvrable ou non) et météorologiques augmente de manière significative, dans la plupart des cas, les performances des régresseurs (les lignes avec FS_2 présentent des MAE et $RMAE$ plus faibles que les lignes avec FS_1) et en particulier lors de l'utilisation d'un algorithme de clustering *K-Medoids* pour regrouper les cellules. L'ajout d'informations sur les valeurs passées dans FS_3 profite davantage au modèle de régression *SVR* (le modèle de régression *GBR* conserve des résultats stables lors de l'utilisation d'un *K-Medoids*). Dans tous les cas sauf le plus simple (*Moy* et/ou seulement FS_1 comme attributs), l'utilisation d'un clustering sur les cellules similaires permet de réduire, la plupart du temps de manière significative, l'erreur tout en gardant la variance des erreurs stable. Le regroupement de cellules similaires permet aux régresseurs de s'entraîner sur un plus grand nombre de données avec les mêmes caractéristiques, ce qui permet de mieux généraliser l'ensemble d'entraînement. Ces résultats permettent de choisir pour notre approche de régression le régresseur *SVR* avec un clustering *K-Medoids* des cellules similaires tout en utilisant l'ensemble d'attributs FS_3 : cette combinaison offre une réduction de 2% du $RMAE$ par rapport à la référence *Moy*, ainsi qu'une réduction de 10 minutes (en moyenne par cellule et par jour) de la MAE . Cette combinaison sera utilisée pour les expériences suivantes.

ServiceB. Cet ensemble de données pose un défi beaucoup plus important à l'algorithme de prédiction que le premier. La première raison est que le nombre de véhicules est plus faible dans ce service (500 par rapport à 600 dans le ServiceA) alors que, comme nous l'avons déjà dit, le nombre de cellules couvertes par le service est plus grand (199), donc le nombre de cellules où il y a très peu d'activité est beaucoup plus grand que dans le ServiceA ce qui nuit à toutes les performances moyennes par cellule. En outre, et pour les mêmes raisons, le service est également moins actif, car il peut être gênant pour l'utilisateur de parcourir une longue distance à pied pour louer une voiture disponible. De ce fait, le nombre de voitures dans chaque cellule (correspondant à la taille du vecteur de rang à prédire) est beaucoup plus faible et l'utilité de chaque voiture est souvent (à un taux de 55%) nulle dans le jeu de données.

Cela peut expliquer les performances relativement bonnes de la référence (*Moy*) avec une erreur de ~ 90 mins qui correspond à près de 92% de la moyenne par jour et cellule du total des minutes de conduite, par rapport à la prédiction des modèles d'apprentissage automatique du Tableau 1. Pour ce jeu de données, l'algorithme *GBR* est plus performant que le *SVR*. Comme pour l'autre ensemble de données, les résultats sont meilleurs lorsque les cellules sont regroupées et le fait de disposer d'informations sur l'état de charge de la voiture (FS_4) permet à l'algorithme *GBR* d'obtenir les meilleures performances. Les meilleurs résultats de prédiction fournissent une MAE de ~ 85 minutes, ce qui correspond à 84% de la moyenne par jour et à la cellule du total des minutes conduites. Comme nous le verrons dans la section suivante, ces résultats de prédiction loin d'être parfaits, nous permettent tout de même d'améliorer le service global ce qui nous laisse espérer une marge d'amélioration beau-

coup plus importante.

4.6.2 Performance de la Relocalisation

À partir de maintenant, nous appelons « placement optimal » la solution optimale à notre problème d'optimisation trouvée par le solveur propriétaire Gurobi [6]. L'utilité (nombre de minutes d'utilisation) est estimée en utilisant le meilleur modèle de régression trouvé dans l'expérience précédente. Ainsi, dans le cas du ServiceA, un algorithme de clustering *K-medoids* des cellules de la ville est suivi d'un régresseur *SVR*, utilisant l'ensemble d'attributs FS_3 . Dans le cas du ServiceB, un algorithme de clustering *K-medoids* est appliqué suivi d'un régresseur *GBR*, utilisant l'ensemble d'attributs FS_3 car la formulation de la relocalisation ne prend pas en compte l'état de charge.

Pour toutes les stratégies, les revenus quotidiens totaux estimés par les deux services sont calculés. Ils sont représentés par jour dans la Figure 2 et les moyennes journalières sont indiquées pour les deux ensembles de données dans le Tableau 2. Nous nous attendons à ce que le revenu soit plus élevé pour *Optim MU* puisqu'il ne prend pas en compte les coûts de relocalisation, mais cette stratégie est également moins réaliste. Nous nous attendons également à ce que *Optim SO* donne un revenu plus important que *Optim DO* puisque la stratégie en une étape peut prendre en compte les coûts de relocalisation de la voiture, ce qui devrait aider à détecter les relocalisations de véhicules non rentables. Dans ces deux derniers cas, le coût des jockeys est intégré à la valeur de revenu affiché.

ServiceA. Le Tableau 2 montre que pour ce service, toutes les méthodes d'optimisation contribuent (en moyenne) à augmenter les revenus de l'entreprise. Des détails supplémentaires dans la Figure 2 (en haut) montrent qu'il n'y a que trois jours où ce n'est pas le cas. Néanmoins le gain est de presque 10% par rapport au service actuel (*Historique*) pour le reste de l'ensemble de test. Pour ce service, étonnamment, il n'y a presque aucune différence entre les stratégies de relocalisation. Cela est probablement dû au fait que le service est très utilisé et a donc une utilité qui est disproportionnée par rapport à de potentiels coûts liés aux relocalisations. Cela explique également pourquoi il n'y a pas de différence entre les stratégies d'optimisation en deux étapes (*Optim DO*) et en une seule étape (*Optim SO*) puisque le coût n'influence pas significativement les relocalisations.

ServiceB. Les résultats pour le ServiceB, présentés dans le Tableau 2 et dans la Figure 2 (en bas), sont en grande partie cohérents avec ce que nous attendions. De manière plus surprenante, nous pouvons constater que la stratégie *Optim DO* donne des résultats inférieurs à la référence *Historique*, tandis que les autres stratégies d'optimisation sont meilleures (+9% pour *Optim MU* et +5% pour *Optim SO*). Cependant les résultats sont moins impressionnants que pour le ServiceA. La raison en est l'inverse de celle de la première ville : comme le service est moins utilisé de manière absolue, la proportion des coûts de relocalisation est plus grande et la variation dans le plan de ramassage des voitures-balais compte plus sur la fonction objectif de la stratégie. Cela montre également l'importance de la straté-

Stratégie	ServiceA		ServiceB	
	Revenu (Base 100)	Nb Relocalisation	Revenu (Base 100)	Nb Relocalisation
Historique	100	N/A	100	N/A
Optim MU	109	53	109	42
Optim SO	107	49	105	42
Optim DO	107	53	99	42

TABLE 2 – Revenu quotidien moyen et nombre de relocalisations pour quatre stratégies de relocalisation dans le cas du ServiceA et du ServiceB.



FIGURE 2 – Revenus quotidiens estimés par les services fonctionnant dans le a) ServiceA (en haut) et le b) ServiceB (en bas) (une graduation d’abscisse correspond à un dimanche), avec la stratégie Historique (ligne bleue pointillée horizontale à 100) comme indice de référence (base 100). La ligne pleine rouge, les lignes pointillées jaunes et vertes sont respectivement les revenus normalisés (sur la référence) fournis par les stratégies *Optim SO*, *Optim DO* et *Optim MU*. Le coût des jockeys est pris en compte dans les valeurs de revenu pour *Optim SO* et *Optim DO*.

gie d’optimisation en une étape (*Optim SO*) par rapport à celle en deux étapes (*Optim DO*) lorsque les coûts sont élevés. Pour ServiceB, notre stratégie la plus réaliste (*Optim SO*) donne donc une augmentation de 5% du bénéfice de l’entreprise.

5 Conclusion

Nous avons abordé le problème de l’optimisation de la relocalisation des véhicules dans un service de partage de voitures en *free-floating* afin d’augmenter son utilisation quotidienne et sa rentabilité. Nous avons proposé de prédire l’utilisation des voitures et montré comment ces prédictions pouvaient être utilisées avec un modèle d’optimisation li-

néaire en nombres entiers pour optimiser l’utilisation des véhicules tout en prenant en compte les coûts des trajets de voiture-balai pour les jockeys. Bien que la phase de prédiction puisse être encore améliorée en introduisant davantage de données exogènes, notre approche peut déjà augmenter l’efficacité et la rentabilité du ServiceA et du ServiceB jusqu’à 7% et 5% respectivement.

Notre solution de relocalisation pourrait bénéficier d’une formulation plus poussée pour optimiser l’itinéraire de la voiture-balai et de tous les jockeys (ces deux problèmes sont liés aux problèmes de tournées de véhicules [12]). Cela serait utile, mais non trivial : l’optimisation conjointe des tournées des véhicules et de notre problème de relocalisation nécessite l’ajout de contraintes supplémentaires et les problèmes sont d’une ampleur non triviale pour des solveurs ILP génériques ; et la combinaison avec le routage complet de la voiture-balai et des jockeys serait plus difficile. Une autre direction serait de s’inspirer de [9], et d’utiliser le résultat de l’optimisation comme fonction de perte pour la formation des prédicteurs, plutôt que de former d’abord les prédicteurs et d’optimiser ensuite les prédictions, afin d’obtenir des prédictions plus axées sur la décision au prix de temps de calcul plus élevé.

Remerciements

Ce travail a été réalisé en collaboration entre le Groupe Stelantis et Inria à travers l’*OpenLab Artificial Intelligence* et a été soutenu par le programme CIFRE de l’ANRT.

Références

- [1] Charu C. Aggarwal and Karthik Subbian. Evolutionary network analysis : A survey. *ACM Comput. Surv.*, 47(1):10 :1–10 :36, 2014.
- [2] Alfred Benedikt Brendel, Julian Brennecke, Patryk Zapadka, and Lutz Kolbe. A Decision Support System for Computation of Carsharing Pricing Areas and its Influence on Vehicle Distribution. *ICIS 2017 Proceedings*, December 2017.
- [3] Carl Axel Folkestad, Nora Hansen, Kjetil Fagerholt, Henrik Andersson, and Giovanni Pantuso. Optimal charging and repositioning of electric vehicles in a free-floating carsharing system. *Computers & Operations Research*, 113 :104771, January 2020.
- [4] Supriyo Ghosh, Jing Yu Koh, and Patrick Jaillet. Improving Customer Satisfaction in Bike Sharing Sys-

- tems through Dynamic Repositioning. In *IJCAI*, pages 5864–5870, 2019.
- [5] Supriyo Ghosh, Pradeep Varakantham, Yossiri Adulyasak, and Patrick Jaillet. Dynamic repositioning to reduce lost demand in bike sharing systems. *Journal of Artificial Intelligence Research*, 58 :387–430, 2017.
- [6] LLC Gurobi Optimization. Gurobi optimizer reference manual, 2020.
- [7] Pierre Hulot, Daniel Aloise, and Sanjay Dominik Jena. Towards Station-Level Demand Prediction for Effective Rebalancing in Bike-Sharing Systems. In *Proc. of the 24th ACM SIGKDD Int. Conf. on Knowledge Discovery & Data Mining*, pages 378–386, 2018.
- [8] Yexin Li and Yu Zheng. Citywide bike usage prediction in a bike-sharing system. *IEEE Transactions on Knowledge and Data Engineering*, 32(6) :1079–1091, 2020.
- [9] Jayanta Mandi, Emir Demirovic, Peter J. Stuckey, and Tias Guns. Smart predict-and-optimize for hard combinatorial optimization problems. In *The Thirty-Fourth AAAI Conference on Artificial Intelligence*, pages 1603–1610, 2020.
- [10] Víctor Martínez, Fernando Berzal, and Juan-Carlos Cubero. A Survey of Link Prediction in Complex Networks. *ACM Comput. Surv.*, 49(4) :69 :1–69 :33, December 2016.
- [11] Tal Raviv, Michal Tzur, and Iris A. Forma. Static repositioning in a bike-sharing system : models and solution approaches. *EURO Journal on Transportation and Logistics*, 2(3) :187–229, 2013.
- [12] M. W. P. Savelsbergh and M. Sol. The General Pickup and Delivery Problem. *Transportation Science*, 29(1) :17–29, February 1995.
- [13] Stefan Schmöller, Simone Weikl, Johannes Müller, and Klaus Bogenberger. Empirical analysis of free-floating carsharing usage : The Munich and Berlin case. *Transportation Research Part C : Emerging Technologies*, 56 :34–51, 2015.
- [14] Rene Seign and Klaus Bogenberger. Prescriptions for the successful diffusion of carsharing with electric vehicles. In *Conference on Future Automotive Technology*, 2013.
- [15] W.W.S. Wei. *Time Series Analysis : Univariate and Multivariate Methods*. Pearson Addison Wesley, 2006.
- [16] Simone Weikl and Klaus Bogenberger. A practice-ready relocation model for free-floating carsharing systems with electric vehicles – Mesoscopic approach and field trial results. *Transportation Research Part C : Emerging Technologies*, 57 :206–223, 2015.
- [17] Rabih Zakaria, Mohammad Dib, Laurent Moalic, and Alexandre Caminada. Car relocation for carsharing service : Comparison of CPLEX and greedy search. In *2014 IEEE Symp. on Comp. Intel. in Vehicles and Transportation Systems*, pages 51–58, December 2014.

Intelligence Artificielle & Génie Civil : enjeux et cas d'usages

R. Leclercq¹, G. Magnaval¹

¹ Socotec Monitoring, 9 rue Léon Blum 91120 Palaiseau

Résumé

Le secteur de la construction est une des industries les moins digitalisées. Pourtant, il fait face à des défis majeurs pour nos sociétés sur la sécurité des travailleurs, la gestion des plannings et la productivité. Les industries manufacturières et de la logistique ont utilisé l'Intelligence Artificielle (IA) comme levier d'innovation avec succès. Malgré ses avantages, de nombreux défis subsistent avant sa généralisation dans la construction.

Cet article vise à présenter des applications de l'IA, examiner les techniques utilisées et identifier les opportunités et les défis dans le secteur de la construction. L'étude dresse une revue critique de la littérature disponible sur les applications de l'IA dans l'industrie de la construction sur l'ensemble du cycle de vie des ouvrages.

Mots-clés

Intelligence Artificielle, Apprentissage Automatique, Construction, Génie Civil, Opportunités IA.

Abstract

Construction industry is one of the least digitized industries in the world. Yet it faces complex challenges such as cost and time overruns, health and safety, productivity, and labor shortages. Artificial Intelligence (AI) is currently revolutionizing industries such as manufacturing, retail, and telecommunications. Despite its benefits, many challenges must be addressed before it can be implemented in the industry.

This paper aims to present applications of AI, review current state of the art, and identify opportunities and challenges in the construction industry. The study provides a critical review of the available literature on the applications of AI in the construction industry across the entire life cycle of structures.

Keywords

Artificial intelligence, Machine learning, Construction Industry, Civil Engineering, AI challenges.

1 Introduction

1.1 Contexte

L'Intelligence Artificielle (IA) est aujourd'hui omniprésente dans nos vies personnelles et professionnelles : chatbots, reconnaissance faciale, diagnostic médical, ... Bien qu'initiée au milieu du XX^e siècle, son déploiement massif est récent. Il est le fruit des avancées en Machine Learning et Deep

Learning, facilitées par l'explosion des volumes de données ainsi que des capacités de calcul.

D'autre part, les enjeux auxquels fait face le Génie Civil sont nombreux : sécurité des ouvriers, vieillissement des infrastructures, dérèglement climatique ou encore faible productivité. Ainsi, depuis l'après-guerre, le secteur voit sa productivité horaire stagner, voire baisser. Dans le même temps, d'autres industries comme l'agriculture ou la logistique ont vu leur productivité horaire être respectivement multipliée par 16 et 8 (cf. Figure 1).

Bien que représentant 7% du PIB mondial [1], la maturité digitale du secteur est encore aujourd'hui limitée [2]. Pourtant, les industriels ainsi que de nouveaux acteurs commencent à se positionner et explorer les potentiels à des fins exploratoires (Vinci avec son incubateur Léonard, Bouygues et le TunnelLab, Kattera, Google, The Boring Company).

Même si les masses de données structurées sont moindres comparativement aux secteurs du retail ou de la publicité, le secteur du Génie Civil n'échappe pas à l'accroissement des systèmes d'informations. La quantité d'informations collectées exigent le développement d'outils pour en extraire des connaissances utilisables. L'IA apporte une contribution significative pour exploiter ces données en offrant un avantage compétitif significatif par rapport aux méthodes traditionnelles [3]. Ainsi, la quatrième révolution industrielle initiée dans le secteur manufacturier se base sur le concept d'Industrie 4.0 avec pour sous-jacent technologique l'automatisation, la robotique et l'IA. Dans l'aéronautique par exemple, les compagnies aériennes ont recours à l'IA pour optimiser la consommation de kérosène avec des réductions de coûts jusqu'à 7% [4].

Pour autant, l'industrie de la construction n'a pas encore adopté l'IA comme un levier d'innovation malgré les potentiels annoncés. En effet, depuis plusieurs décennies, des recherches et ouvrages ont été publiés pour montrer la valeur de l'IA sur les cas d'usage du secteur. On peut citer le Structural Health Monitoring (SHM), la maintenance prédictive des ouvrages, la recherche opérationnelle de sols pollués [5] ou encore la planification des ressources.

Cet article a pour objet de montrer le potentiel de l'IA pour le Génie Civil avec l'identification de cas d'usage sur l'ensemble du cycle de vie des ouvrages. Il vise à montrer les opportunités futures et les limites connues à l'adoption de l'IA.

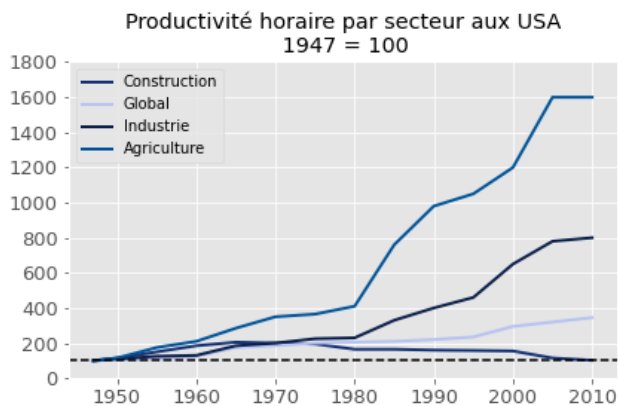


Figure 1: Evolution de la productivité dans la construction (McKinsey Global Institute [6])

1.2 Méthodologie de recherche

Pour établir l'état de l'art passé et présent de l'IA dans la construction, une revue de littérature a été faite en se basant sur des canaux multiples :

- Des bases de données d'articles scientifiques comme Institute of Electrical and Electronics Engineers (IEEE), Association for Computing Machinery (ACM) ou ScienceDirect en utilisant le moteur de recherche Dimensions.ai.
- Des études prospectives issus des grands cabinets de conseils comme McKinsey, Deloitte, PwC.
- Des publications issues de journaux spécialisés dans le domaine comme Le Moniteur ou Construction Cayola.

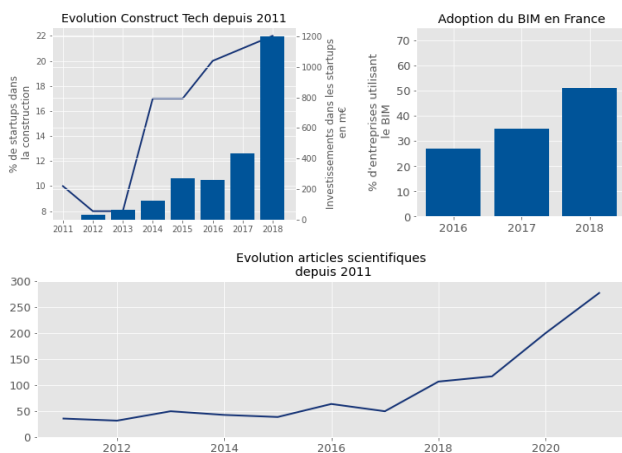


Figure 2: Ecosystème IA dans la construction depuis 2011

L'analyse du nombre de publications scientifiques est issue d'une recherche par mots-clés (*Computer vision OR Machine learning OR Expert System OR Natural Language Processing OR Artificial Intelligence OR Clustering OR Deep Learning OR Convolutional Neural Network OR Recurrent Neural Network AND Construction Industry*) sur l'agrégateur de base de données Dimensions.ai. Seuls les articles de journaux dans le domaine du Génie Civil sont comptabilisés. Les résultats indiquent une augmentation significative du nombre de publications depuis 2011 (x7).

Comme l'illustre la Figure 2, cette augmentation est corrélée aux investissements en hausse dans l'écosystème des start-ups [7 – 8]. Ceci témoigne d'une adoption progressive des innovations dans l'industrie.

2 Approches et enjeux de l'IA en génie civil

2.1 Périmètre du Génie civil

Pour bien définir le périmètre de cet article, il est important de comprendre le domaine du Génie Civil. Ce domaine d'application est très vaste et intègre les travaux publics et le bâtiment. Il mobilise des compétences dans les domaines des structures, de la géotechnique, de l'hydraulique, du transport et de l'environnement. Cette diversité du champ de compétences a pour conséquence la spécialisation des profils et donc la complexification du management des projets. En effet, la filière s'est structurée avec des acteurs multiples (cf. Figure 3) :

- Promoteurs
- Organismes de contrôle
- Bureaux d'études
- Entreprise générales (constructeurs)
- Maitrise d'ouvrage (coordinateur)
- Fournisseurs
- Gestionnaires

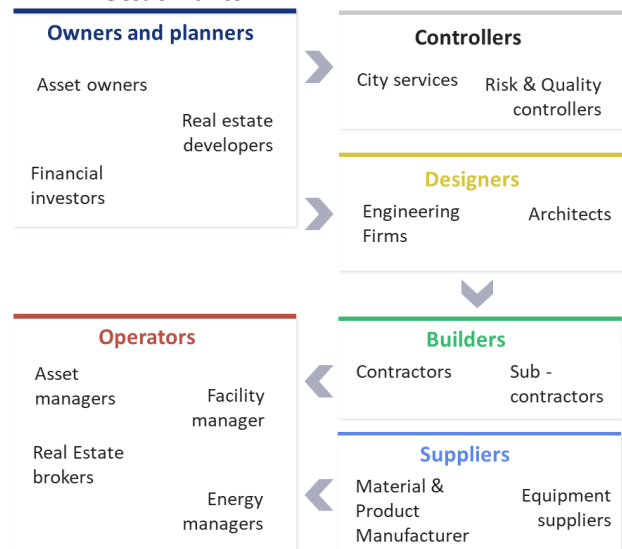


Figure 3: Rôles et fonctions des acteurs de la filière du génie civil [9]

En plus d'une multiplicité de rôles, la filière compte un tissu de PME et TPE avec 472 000 entreprises pour 293 milliards d'euros de chiffres d'affaires (8% du PIB). Ses entreprises emploient plus de 1,3 million de salariés en équivalent temps plein [10].

Les caractéristiques de ce tissu industriel ne jouent pas en la faveur du développement des technologies, dont l'IA. La diversité des acteurs, des entreprises et des compétences génère un silotage des systèmes d'informations. Or, la centralisation des données sur des référentiels communs est un des prérequis du développement de solutions d'IA. L'émergence du BIM et de l'IoT décrits en partie 2.3 et 2.4 sont des opportunités pour lever ces barrières.

2.2 Approche traditionnelle et approche IA

Dans le domaine de la construction, la gestion des projets est établie selon un cycle séquentiel : initialisation, planning, exécution, suivi et clôture. L'approche traditionnelle de la gestion de projet met l'accent sur les processus linéaires, la documentation, la planification initiale et la hiérarchisation des priorités. Selon cette méthodologie, les exigences et les méthodes sont fixées à l'avance et il n'y a pas de place pour l'itération. Cela conduit à des retards de livraison et des problèmes de budget qui concernent près d'un tiers des projets selon KPMG [11].

En plus de suivre une méthodologie linéaire, les modes d'actions sont issus de processus normés, qualifiés et éprouvés. Pour la conception, les règles de calcul sont fixées en Europe par les Eurocodes pour garantir les mêmes standards de sécurité sur l'ensemble des projets. Les procédés constructifs et les actions de maintenance sont basés sur un mélange de lois physiques et empiriques. En effet, si les lois de résistance des matériaux ou de thermique du bâtiment sont bien connues, les hypothèses qu'elles recouvrent ne sont pas toujours vérifiées en pratique. Les calculs associés sont parfois trop complexes compte-tenu des réalités des budgets et planning. Dès lors, les retours d'expériences formalisés dans des facteurs de sécurité entrent en jeu pour simplifier la résolution des problèmes. Ceci conduit à un surdimensionnement des structures et à un processus non optimisé.

Dans le domaine de l'IA, la méthodologie d'action est agile. L'itération est au cœur de l'apprentissage avec une amélioration continue des modèles. Ce principe assure un produit final centré sur les besoins réels du terrain tout en facilitant l'adoption avec un déploiement plus rapide des innovations. Contrairement aux codes de calculs, l'IA se base sur les données observées et des modèles mathématiques.

Pour favoriser l'adoption de l'IA dans le secteur, il est nécessaire de prendre conscience des différences d'approche. Pour autant, ces dernières sont complémentaires et peuvent être mises en compétition.

2.3 Lien avec le BIM et l'IoT

Depuis quelques années, le secteur a largement adopté le Building Information Modeling (BIM) (cf. Figure 2). Le BIM représente un tournant technologique majeur pour le Génie Civil en désilotant les sources de données. Selon le comité américain BILS [12], le BIM se définit comme « une représentation numérique des caractéristiques physiques et fonctionnelles d'un ouvrage ». Le BIM est une source de données partagées pour les informations relatives à une installation, qui constitue une base fiable pour les décisions à prendre au cours de son cycle de vie, défini comme allant de la première conception à la démolition. L'un des prérequis au BIM est la collaboration de différentes parties prenantes à différentes phases du cycle de vie d'une installation pour insérer, extraire, mettre à jour ou modifier des informations dans la BIM afin de soutenir et de refléter les rôles de cette partie prenante.

Cette technologie est aujourd'hui ouverte et interopérable

avec des standards communs comme le standard IFC (Industry Foundation Classes). Un ouvrage peut ainsi être décrit selon les mêmes attributs indépendamment des entreprises intervenant sur le projet.

La technologie BIM est aujourd'hui en phase d'adoption massive. Elle représente un investissement important pour les entreprises du secteur, à la fois en termes de R&D mais aussi de gestion du changement (formation et acculturation). L'IA doit capitaliser sur le BIM pour en exploiter la pleine valeur.

Une autre source de données nouvelles pour le secteur est l'internet des objets connectés (IoT). La baisse des coûts et le développement des réseaux basses fréquences (LORA, Sigfox) ont permis la démocratisation des capteurs avec des usages variés : compteurs d'eau pour mesurer la consommation, balises GNSS pour localiser les matériaux sur chantier, accéléromètres sur des machines pour suivre leur utilisation, capteurs de déformation pour contrôler la santé structurelle d'un ouvrage.

2.4 Gestion des données en génie civil

Les sources de données sont aujourd'hui en pleine mutation avec le BIM, l'IoT mais aussi les évolutions des Systèmes d'Informations des entreprises. Cependant, de nombreux processus sont encore aujourd'hui sous format papier ou non structurés.

La filière avec un écosystème de start-up et de grands groupes travaille pour digitaliser les processus en structurant des données exploitables. Les logiciels de gestion documentaires ou de planning assurent une gouvernance plus facile des données. Par ailleurs, des initiatives de mises à disposition d'Open Data émergent : le BRGM avec des données géologiques et environnementales ou l'IGN avec des plans et données LIDAR.

3 L'IA présente sur l'ensemble du cycle de vie des actifs

La vie d'un ouvrage commence bien avant sa construction. La vie d'un actif (infrastructure ou bâtiment) suit 4 phases :

- Conception
- Construction
- Exploitation
- Démolition et revalorisation

Dans cette partie, les opportunités d'adoption de l'IA avec des applications est explicitée pour chaque étape du cycle de vie. Les cas d'étude qui ne sont pas spécifiques au secteur de la construction (pilotage financier, ressources humaines, satisfaction clients) ne sont pas mentionnés, même s'ils peuvent être pertinents pour les entreprises du domaine.

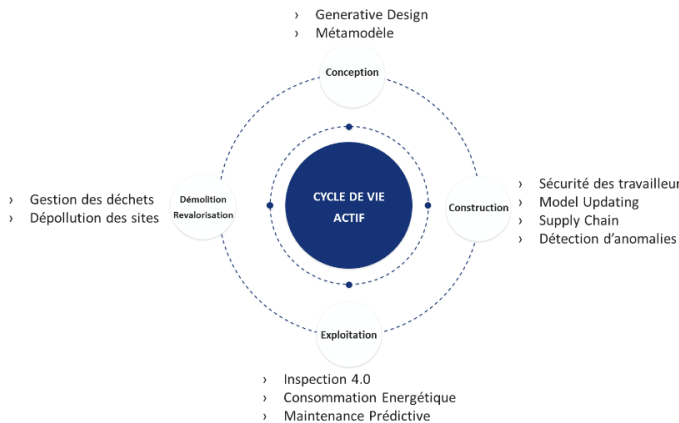


Figure 4: Le cycle de vie des ouvrages en génie civil

3.1 Conception des ouvrages

La conception est la première étape du projet. Cette étape inclut :

- La planification du projet (étude financière, planning, étude de faisabilité)
- La conception avec des études architecturales et structurales ainsi que d'impact sur l'environnement (études environnementales, urbanistiques, etc.).

La responsabilité (y compris pénale) de cette phase est confiée au bureau d'étude qui a la charge de conduire les études dans les règles de l'art. Pour mener les analyses, les ingénieurs ont souvent recours aux modèles numériques pour simuler le comportement de l'ouvrage dans des cas de charges définies par les codes de calculs.

En pratique, la combinaison d'une formulation complexe du problème et de simulations chronophages signifie que l'optimisation est rarement utilisée dans la conception des bâtiments et des structures civiles. Les modèles numériques utilisés lors des études sont fortement non linéaires avec un grand nombre de variables.

L'utilisation de l'IA à travers des métamodèles est une opportunité d'optimisation pour les bureaux d'études.

En effet, les métamodèles établissent un lien statistique entre les données d'entrée et les données de sortie, qui sont recueillies en exécutant la simulation d'un système complexe. Dès lors, le recours au métamodèle diminue significativement le temps d'exécution. L'exploration plus approfondie de l'espace de conception est ainsi possible pour estimer la meilleure combinaison de variables.

Tseranidis et al. [13] ont exploré l'application de l'approximation pilotée par les données à la conception des bâtiments.

Un exemple d'application des métamodèles est présenté par Zhang et al. [14] avec l'utilisation d'une régression multivariée par spline adaptative (MARS) sur un problème de liquéfaction des sols soumis à un séisme.

3.2 Construction

Une fois que la phase de conception est suffisamment avancée, la phase de construction est initiée. Les équipes sont alors réunies sur chantier, les matériaux et les machines sont acheminés et les travaux peuvent débuter. Certains éléments

sont fabriqués en usine avant d'être assemblés sur chantier.

La responsabilité de la phase de construction est portée par les entreprises générales qui assurent que les processus constructifs employés sont conformes aux études et réglementations.

La première exigence des entreprises de construction est la sécurité des travailleurs. Selon l'assurance maladie, le secteur du BTP présente un des plus forts taux d'accidents professionnels avec 56 accidents du travail pour 1000 salariés (contre 34 en moyenne). Des manquements en matière de sécurité et de protection de la santé continuent d'être constatés sur les chantiers, pouvant mener à des accidents graves. La réduction des risques est un enjeu fort pour les entreprises. Les acteurs commencent à mettre en place des systèmes complets avec le monitoring, la visualisation et la notification d'action. Ces systèmes sont améliorés par l'IA qui est une solution pour identifier les risques en amont et réduire le temps d'action. Par exemple, Ding et al. [15] ont développé un modèle d'apprentissage profond hybride qui intègre un réseau de neurones convolutifs (CNN) et une mémoire à long terme (LSTM) pour identifier automatiquement les comportements dangereux des employés sur les chantiers de construction. Le modèle CNN est appliqué à chaque image pour capturer les caractéristiques spatiales obtenues à partir de la vidéo et le réseau LSTM est utilisé pour comprendre les informations temporelles à partir des images continues qui sont générées.

Un autre enjeu fort pendant la construction est d'assurer la sécurité et la durabilité de l'ouvrage. Pour s'assurer des hypothèses prises en phase de conception, les données mesurées sur site par des capteurs permettent une confrontation de la mesure avec ce qui est attendu. L'exploitation de l'ensemble de ces données de capteurs peut se faire en utilisant l'IA comme système d'alerte. Sur les séries temporelles, l'identification de motifs inhabituels suivant une approche supervisée a été employé par Garcia et al. [16] dans le cas du suivi de moisissure pour la construction en bois. Les auteurs ont proposé plusieurs modèles (LSTM, OneClassSVM) pour identifier les anomalies sur les données de 16 capteurs de température et d'humidité à différentes localisations de l'ouvrage.

L'IA a changé la donne pour la chaîne d'approvisionnement du commerce de détail en réduisant les temps d'arrêt de la fabrication, en réduisant l'offre excédentaire et en augmentant la prévisibilité des expéditions. Les exemples d'apprentissage supervisé [17] deviendront directement applicables au Génie Civil à mesure que la modularisation et la préfabrication se répandront. La généralisation de ces procédés constructifs va augmenter les quantités de modules produits hors-sites, ce qui nécessitera une meilleure coordination de la chaîne d'approvisionnement pour contrôler les coûts et les flux de trésorerie globaux.

3.3 Exploitation

Lorsque le chantier est terminé, y compris la mise en œuvre d'éventuels aménagements et l'installation d'éventuels équipements, l'ouvrage est mis en service. À partir de ce

moment, le propriétaire (ou son représentant) devient responsable de l'exploitation et de l'entretien de son bien. Cette étape constitue la phase la plus longue du cycle de vie et présente des opportunités pour l'IA. En effet, la profondeur des données collectées est supérieure avec des conditions d'acquisition moins sujets aux aléas de la construction.

Pour s'assurer de la sécurité des ouvrages, des inspections périodiques sont réalisées. Ces inspections sont normées et ont pour objet de qualifier l'état de l'infrastructure ainsi que d'évaluer les actions de maintenance à réaliser. Sur site, une des missions de l'inspecteur est de relever l'ensemble des défauts. Cette tâche est fastidieuse et source d'erreurs. Des algorithmes de computer vision assistent l'inspecteur pour segmenter et quantifier des défauts comme les fissures ou la présence de végétation. Des chercheurs du laboratoire de l'EPFL ont notamment développé un algorithme CNN avec une architecture encodeur-décodeur pour l'identification des fissures sur du béton [18]. Dans un futur plus lointain, ces algorithmes de détection seront couplés avec des drones pour automatiser l'inspection. Cette automatisation permettra de réduire l'incertitude liée à l'inspecteur et d'augmenter la fréquence des inspections.

L'IA est également utile pour optimiser la maintenance des ouvrages à travers la maintenance prédictive ou prescriptive. Ces développements peuvent se faire à l'échelle d'un ouvrage ou celle d'un parc. Hamida et al. [19] ont développé une méthode pour définir une politique de maintenance à l'échelle de 7000 ponts au Québec. Dans cette étude, des modèles d'espace-état (SSM) sont implémentés pour modéliser la détérioration du parc de ponts à partir des notes d'inspections visuelles. Les estimations de la détérioration des ponts peuvent aider à examiner l'efficacité des interventions précédentes et la tendance à long terme de l'état du réseau ainsi que de jeter les bases de la planification des futures actions d'entretien.

Avec près de 45% des consommations énergétiques en France [20], l'optimisation de l'énergie en exploitation est un enjeu majeur pour le secteur du bâtiment. Bien que les méthodes physiques visant à calculer le comportement de consommation d'énergie d'un bâtiment soient précises, elles ne sont pas pratiques dans certaines applications en raison de la nécessité d'inspecter et de recueillir en permanence des données pour tous les paramètres d'entrée. Chokor et al. [21] ont proposé un nouveau système d'évaluation du rendement du système Leadership in Energy and Environmental Design (LEED) basé sur l'étude d'un modèle de prédiction personnalisé. Ce modèle est piloté par l'énergie réellement consommée des bâtiments. Les résultats montrent que le modèle de régression par gradient boosting est supérieur aux autres modèles de régression, ce qui aide les énergéticiens à faire de meilleurs choix tout au long du cycle de vie du projet.

3.4 Démolition et revalorisation de l'actif

Lorsque l'ouvrage ne répond plus aux critères suffisants d'opérabilité et/ou de sécurité, les deux choix évidents qui

s'offrent au propriétaire sont la réparation ou la démolition et la reconstruction.

Les deux enjeux principaux de la déconstruction sont la gestion des déchets et la remise en état du site. Compte-tenu des volumes de déchets générés, la revalorisation des matériaux est un enjeu économique et écologique. Kuritcyn et al. [22] ont proposé plusieurs algorithmes supervisés (Random Forest et C-LibSVM) pour classifier les déchets sur la base d'analyse spectrale dans le domaine du visible et de l'infrarouge.

Pour pérenniser les sites, la dépollution est souvent obligatoire. Il s'agit souvent d'identifier les zones de sols pollués et donc les volumes de terres à excaver. Quacha et al. [23] ont montré que des algorithmes de recherche opérationnelle peuvent être employés pour optimiser le volume des terrains à retraiter. Les auteurs ont proposé un modèle de processus gaussien en prenant en compte l'agrégation des données expérimentales et les discontinuités physiques en l'appliquant sur un site industriel avec 118 échantillons de sols. L'exploitation des données avec cette approche par IA contribue à la fois à l'optimisation des volumes à retraiter mais aussi à la réduction des incertitudes.

4 Défis à relever dans le futur pour une adoption massive de l'IA

Si les perspectives d'innovation dans le domaine des technologies de la construction sont prometteuses, il reste des défis à relever. Un des problèmes majeurs évoqué au paragraphe 2.4 révèle que la représentation de l'information est encore inadéquate. D'importants défis fondamentaux doivent encore être résolus avant que les innovations technologiques puissent commencer à incorporer largement les techniques d'IA dans les outils logiciels en production. Parmi les principaux défis, on citera le développement des jeux de données, l'explicabilité et l'acculturation.

4.1 Développement des jeux de données

Pour s'améliorer, l'IA nécessite des jeux de données fiables et qualifiés. Or, le Génie Civil pêche dans la gouvernance des données. Cela s'explique par plusieurs éléments, à la fois culturels et techniques. Premièrement, l'utilisation du papier est encore d'actualité aujourd'hui. Il n'est pas rare qu'un ouvrier ou un inspecteur transmettent des informations sous ce format plutôt que via un formulaire numérisé. Un travail de structuration et de développement d'outils logiciels est en cours dans la filière.

Deuxièmement, les structurations de données sont généralement spécifiques à une discipline, représentant les bâtiments avec la sémantique d'une seule vision professionnelle (comme l'architecture, l'ingénierie structurelle ou les systèmes de fluides). En tant que telle, la collaboration pluridisciplinaire à l'aide de modèles est difficile. La plupart des équipes utilisent des types de modèles distincts. Troisièmement, la sélection des données suivies n'est pas faite selon le paradigme d'une exploitation par l'IA par la suite. Dès lors, de nombreuses relations et propriétés d'objets sont encore implicites dans les modèles de données, laissées à l'interprétation intelligente de leurs

utilisateurs humains. De même, les spécifications de conception et les codes de construction définissent généralement des paramètres qui sont des compilations complexes de contraintes géométriques qui sont très difficiles à exprimer à l'aide d'ensembles de règles "si-alors".

Un travail de long terme est à conduire pour que des jeux de données massifs soient exploitables pour l'IA.

En parallèle des efforts à consentir pour développer des bases de données à grande échelle, l'IA pour le Génie Civil s'inscrit aujourd'hui dans la mouvance du *small data*. Les données à dispositions sont faibles en nombre mais résultent d'un processus intellectuel. C'est le cas notamment des mesures de capteurs structurels. Il n'est pas possible d'équiper densément les ouvrages : l'emplacement des capteurs est choisi aux points critiques. Ce choix résulte souvent d'une analyse complexe du type modèle aux éléments finis. Un autre élément explicatif du *small data* en génie civil est l'unicité des chantiers et des structures. En effet, contrairement aux usines, chaque environnement est différent. Cela pose les questions de représentativité et de transfert de connaissance entre les projets.

En phase de construction, un verrou supplémentaire est la connectivité Internet. En effet, l'ensemble des zones ne sont pas bien couvertes et sont sujettes à des coupures de courant [24]. Par exemple, les capteurs et les actionneurs communiquent des informations qui doivent être calculées en temps réel pendant la construction. Il est pertinent de chercher des moyens de résoudre ce problème de manière efficace et effective. L'utilisation des technologies de communication 4G (LTE/max) a permis de résoudre ce problème dans une large mesure. L'émergence de la 5G offre une fiabilité encore plus grande pour les chantiers de construction grâce à son débit de données élevé, la réduction de la latence, les économies d'énergie, la réduction des coûts, la plus grande capacité du système et la connectivité massive des appareils.

4.2 Interprétabilité et Explicabilité

Le domaine du Génie Civil vise à construire les infrastructures les plus pérennes en garantissant la sécurité des utilisateurs. Cela soulève des enjeux de compréhension des décisions et d'arbitrages pris lors de chaque étape du cycle de vie. Ce point renvoie aux notions d'interprétabilité et d'explicabilité qu'il convient de définir.

Selon Doshi-Velez et Kim [25], l'interprétabilité désigne « la capacité d'expliquer ou de présenter en termes compréhensibles pour un humain ». Quant à elle, l'explicabilité est associée à la logique et à la mécanique internes d'un système d'apprentissage automatique. Plus un modèle est explicable, plus la compréhension que les humains en ont est profonde.

Si, en pratique, les ingénieurs Génie Civil peuvent percevoir les modèles numériques aux éléments finis comme des boîtes-noires, des experts sont capables d'expliquer les lois sous-jacentes des simulations, même les plus complexes. Pour instaurer la confiance dans les systèmes d'IA et éviter

les biais potentiels liés aux données, il est essentiel que les praticiens de la construction comprennent comment le système prend ses décisions. Il est donc nécessaire d'utiliser l'IA explicable (XAI) pour produire des modèles explicables et permettre aux humains de comprendre, de faire confiance et de gérer les systèmes. Cette explicabilité est d'autant plus importante que la conception des bâtiments et infrastructures peut mettre en péril des vies humaines.

4.3 Culture et gestion des talents

Il est connu que le secteur de la construction est l'un des secteurs les moins numérisés et qu'il est lent à adopter les nouvelles technologies. Cela s'explique par la nature risquée et coûteuse de la plupart des processus de construction, où les petites erreurs peuvent entraîner des conséquences énormes. Contrairement à des secteurs comme l'industrie manufacturière, les sites de construction sont pour la plupart uniques et différents, ce qui nécessite une IA capable d'apprendre et de s'adapter rapidement dans des environnements changeants. Dès lors, les technologies d'IA à déployer dans le secteur de la construction doivent être utilisables dans différents projets ou sites de construction et testées de manière approfondie pour convaincre les entrepreneurs et les entreprises de construction de les utiliser. Il peut s'agir de tirer parti d'autres technologies numériques comme la blockchain pour améliorer la confiance et la transparence.

Sur le marché mondial, il y a une pénurie de talents formés à l'IA. Il est encore plus difficile de trouver les profils AI+X ayant une expérience dans la construction pour construire des solutions personnalisées visant à résoudre les nombreux problèmes du secteur. Or, contrairement à d'autres domaines, la connaissance du secteur est un atout majeur, à la fois pour la gestion du changement mais aussi pour la compréhension du problème. En effet, poser le problème d'optimisation sur une problématique de maintenance prédictive d'un pont du XIXe siècle est plus complexe que sur l'optimisation de publicités en ligne.

Pour favoriser l'adoption de l'IA dans ce secteur au cœur des enjeux climatiques et d'une importance économique majeure, les investissements publics et privés pour acculturer et former les professionnels de la construction est nécessaire.

5 Conclusion

Cet article présente les enjeux et cas d'usages de l'Intelligence Artificielle (IA) tout au long du cycle de vie des bâtiments et infrastructures.

L'Intelligence Artificielle n'en est qu'à ses débuts dans le domaine du Génie Civil. Pour autant, les cabinets de conseils comme Deloitte, McKinsey et PwC prédisent tous une accélération forte des investissements dans ces technologies. Depuis une dizaine d'années, des chercheurs ont ouverts la voie en mettant en application des solutions d'IA. Ces cas d'usages sont des preuves de concept qui démontrent que les retours sur investissement d'autres domaines sont possibles dans le secteur de la construction.

Pour autant, l'adoption massive se fera en considérant les défis du secteur : la structuration de jeux de données interoperables et fiables, l'explicabilité des modèles et la

gestion du changement.

Par ailleurs, la pertinence de l'IA est encore renforcée par d'autres tendances émergentes telles que le BIM et l'IoT. Avec l'augmentation des données générées tout au long du cycle de vie des bâtiments et l'émergence d'autres technologies numériques, l'IA a la capacité d'exploiter ces données et de tirer parti des capacités des autres technologies pour améliorer les processus de conception, de construction, d'exploitation et de démolition.

6 Références

- [1] AGENDA, Industry. Shaping the future of construction a breakthrough in mindset and technology. WEF, Geneva, 2016.
- [2] YOUNG, Damali, PANTHI, Kamalesh, et NOOR, Omar. Challenges Involved in Adopting BIM on the Construction Jobsite. EPiC Series in Built Environment, 2021, vol. 2, p. 302-310.
- [3] ABIOYE, Sofiat O., OYEDELE, Lukumon O., AKANBI, Lukman, et al. Artificial intelligence in the construction industry: A review of present status, opportunities and future challenges. Journal of Building Engineering, 2021, vol. 44, p. 103299.
- [4] RAVAL, Chintan, How Artificial Intelligence is transforming the Aerospace Industry, eInfochips, <https://www.einfochips.com/blog/how-artificial-intelligence-is-transforming-the-aerospace-industry/>, 2021.
- [5] GOULET, James-A. Probabilistic machine learning for civil engineers. MIT Press, 2020.
- [6] SVEIKAUSKAS, Leo, ROWE, Samuel, MILDENBERGER, James, et al. Productivity growth in construction. Journal of Construction Engineering and Management, 2016, vol. 142, no 10, p. 04016045.
- [7] RICAUD, Yan, Innovation et BTP : la transformation du secteur en marche, PwC, 2018.
- [8] ALVAREZ, Laureano, DE REYNA, Galo, Construction Predictions, Deloitte, 2020.
- [9] STOWE, Ken, LÉPINOY, Olivier, et KHANZODE, Atul. Innovation in the construction project delivery networks in Construction 4.0. In : Construction 4.0. Routledge, 2020. p. 62-88.
- [10] MORENILLAS, Noémie, SKENARD, Gabriel, Les Entreprises en France, INSEE, 2019.
- [11] ARMSTRONG, Geno, GILGE, Clay, Global Construction Survey 2019, KPMG, 2019.
- [12] COMMITTE, NIMBS. National Building Information Modeling Standard. 2010.
- [13] TSERANIDIS, Stavros, BROWN, Nathan C., et MUELLER, Caitlin T. Data-driven approximation algorithms for rapid performance evaluation and optimization of civil structures. Automation in Construction, 2016, vol. 72, p. 279-293.
- [14] ZHANG, Wengang, GOH, Anthony TC, et ZHANG, Yanmei. Multivariate adaptive regression splines application for multivariate geotechnical problems with big data. Geotechnical and Geological Engineering, 2016, vol. 34, no 1, p. 193-204.
- [15] DING, Lieyun, FANG, Weili, LUO, Hanbin, et al. A deep hybrid learning model to detect unsafe behavior: Integrating convolution neural networks and long short-term memory. Automation in construction, 2018, vol. 86, p. 118-124.
- [16] GARCÍA FAURA, Álvaro, ŠTEPEC, Dejan, CANKAR, Matija, et al. Application of Unsupervised Anomaly Detection Techniques to Moisture Content Data from Wood Constructions. Forests, 2021, vol. 12, no 2, p. 194.
- [17] WANG, Mudan, WANG, Cynthia Changxin, SEPASGOZAR, Samad, et al. A systematic review of digital technology adoption in off-site construction: Current status and future direction towards industry 4.0. Buildings, 2020, vol. 10, no 11, p. 204.
- [18] REZAIE, Amir, ACHANTA, Radhakrishna, GODIO, Michele, et al. Comparison of crack segmentation using digital image correlation measurements and deep learning. Construction and Building Materials, 2020, vol. 261, p. 120474.
- [19] HAMIDA, Zachary et GOULET, James-A. A stochastic model for estimating the network-scale deterioration and effect of interventions on bridges. Structural Control and Health Monitoring, 2021, p. e2916.
- [20] Ministère de la transition écologique, Energie dans les bâtiments, <https://www.ecologie.gouv.fr/energie-dans-les-batiments#:~:text=Le%20secteur%20du%20b%C3%A2timent%20repr%C3%A9sente,climatique%20et%20la%20transition%20%C3%A9nerg%C3%A9tique>, 2021.
- [21] CHOKOR, Abbas et EL ASMAR, Mounir. Data-driven approach to investigate the energy consumption of LEED-certified research buildings in climate zone 2B. Journal of Energy Engineering, 2017, vol. 143, no 2, p. 05016006.
- [22] KURITCYN, Petr, ANDING, Katharina, LINß, Elske, et al. Increasing the safety in recycling of construction and demolition waste by using supervised machine learning. In : Journal of Physics: Conference Series. IOP Publishing, 2015. p. 012035.
- [23] QUACH, Alyssa Ngu-Oanh, TABOR, Lucie, DUMONT, Dany, et al. A machine learning approach for characterizing soil contamination in the presence of physical site discontinuities and aggregated samples. Advanced Engineering Informatics, 2017, vol. 33, p. 60-67.
- [24] LOUIS, Joseph et DUNSTON, Phillip S. Integrating IoT into operational workflows for real-time and automated decision-making in repetitive construction operations. Automation in Construction, 2018, vol. 94, p. 317-327.
- [25] DOSHI-VELEZ, Finale et KIM, Been. Towards a rigorous science of interpretable machine learning. arXiv preprint arXiv:1702.08608, 2017.

Évaluation et Production de Plongements de Mots à Partir de Contenus Web Français à Grande Échelle

H. Abdine¹, C. Xypolopoulos¹, M. Kamal Eddine¹, M. Vazirgiannis^{1,2}

¹ École Polytechnique, LIX

² Athens University of Economics and Business, AUEB

{hadi.abdine, christos.xypolopoulos, moussa.kamal-eddine, michalis.vazirgiannis}@polytechnique.edu

Résumé

Les représentations distribuées de mots sont couramment utilisées dans de nombreuses tâches de traitement du langage naturel, en ajoutant que les vecteurs de mots pré-entraînés sur d'énormes corpus de textes ont atteint une performance importante dans de nombreuses tâches NLP. Cet article présente plusieurs vecteurs de mots de qualité supérieure pour la langue française, deux d'entre eux étant entraînés sur des données françaises massives recueillies au cours de cette étude et les autres sur un corpus français déjà existant utilisé pour pré-entraîner le model FlauBERT. Nous évaluons également la qualité de nos vecteurs de mots proposés et des vecteurs de mots français existants sur la tâche d'analogie de mots français. De plus, nous effectuons l'évaluation sur de multiples tâches NLU qui montrent l'amélioration importante des performances des vecteurs de mots pré-entraînés durant cette étude par rapport aux plongements existants et aléatoires. Enfin, nous avons créé une application web pour tester et visualiser les plongements de mots obtenus¹. Ceux derniers sont disponibles au public, ainsi que le code d'affiliation sur les tâches NLU², et le code de la démonstration³.

Mots-clés

Plongements de mots, Word2Vec, pré-entraînement, apprentissage profond, NLP, FLUE, analogie de mots.

Abstract

Distributed word representations are popularly used in many tasks in natural language processing, adding that pre-trained word vectors on huge text corpus achieved high performance in many different NLP tasks. This paper introduces multiple high-quality word vectors for the French language where two of them are trained on massive crawled French data and the others are trained on an already existing French corpus used to pretrain FlauBERT model. We also evaluate the quality of our proposed word vectors and the existing French word vectors on the French word analogy task. In addition, we do the evaluation on multiple NLU tasks that shows the important performance enhancement of the pre-trained word vectors during this study compa-

red to the existing and random ones. Finally, we created a demo web application to test and visualize the obtained word embeddings¹. The produced French word embeddings are available to the public, along with the fine-tuning code on the NLU tasks², and the demo code³.

Keywords

Word embeddings, Word2Vec, pretraining, deep learning, NLP, FLUE, word analogy.

1 Introduction

Les plongements de mots sont un type de représentation des mots qui est devenu très important et largement utilisé dans différentes applications du traitement du langage naturel. Par exemple, les plongements de mots pré-entraînés ont joué un rôle essentiel dans l'obtention de performances remarquables avec des modèles d'apprentissage profond sur des problèmes difficiles de compréhension du langage naturel. Ces vecteurs de mots sont produits à l'aide des modèles de réseaux de neurones non supervisés basés sur l'idée que la signification d'un mot est liée au contexte dans lequel il apparaît. Ainsi, des mots similaires en signification ont une représentation similaire. Plus précisément, ils ont une représentation proche dans l'espace. En conséquence, la qualité des vecteurs de mots est directement liée à la quantité et à la clarté du corpus sur lequel ils ont été entraînés. De nombreuses techniques et algorithmes ont été proposés depuis 2013 pour apprendre ces représentations distribuées de mots. Nous nous concentrons principalement sur Word2Vec [14], GloVe [18], et FastText [3] qui utilisent deux approches : CBoW et SkipGram. Ces méthodes ont été utilisées pour apprendre des vecteurs de mots à partir d'énormes corpus. La plupart des vecteurs de mots disponibles publiquement sont pré-entraînés sur des textes en anglais. Cependant, les vecteurs de mots pour les autres langues sont peu nombreux, inexistant ou formés à l'aide d'un très petit corpus qui ne peut produire des vecteurs de mots de bonne qualité. Dans cet article, nous nous intéressons à la création des plongements de mots français statiques avec un benchmark qui les compare avec d'autres plongements de mots. Ainsi, nous n'incluons pas les plongements de mots contextuels tels que FlauBERT [12], CamemBERT [13] et BARThez [7] dans la comparaison. Ces modèles pré-entraînés basés sur

1. nlp.polytechnique.fr/

2. github.com/hadi-abdine/WordplongementsEvalFLUE/

3. github.com/hadi-abdine/FrenchWordplongementsDemo/

l'attention, même s'ils ont de meilleures performances, utilisent plus de mémoire et de puissance de calcul. Ce travail présente des plongements de mots français entraînés sur un large corpus français collecté du web avec plus d'un million de noms de domaine. De plus, nous entraînons les plongements Word2Vec sur d'autres corpus et ressources françaises afin de les comparer avec ceux entraînés sur le corpus français collecté. Ensuite, nous évaluons ces vecteurs de mots sur l'ensemble de questions créées dans [9] pour la tâche d'analogie de mots. Enfin, nous évaluons les vecteurs de mots français sur certaines tâches du benchmark FLUE [12] et nous comparons les résultats avec ceux des vecteurs français FastText [9] et Word2Vec [8] existants produits à partir des corpus Common Crawl, Wikipedia et FrWac. Nous discutons par la suite la signification de la tâche d'analogie des mots sur la qualité de leurs plongements.

2 Travaux antérieurs

Depuis que les représentations distribuées des mots ont été introduites, de nombreux vecteurs de mots pré-entraînés pour de nombreuses langues sont produits. Par exemple, des vecteurs de mots anglais ont été appris à l'aide d'une partie de l'ensemble de données de Google News et publiés avec Word2Vec [14]. Ensuite, des représentations de mots ont été entraînées pour 157 langues, dont la française, en utilisant FastText [9] pour apprendre des vecteurs de mots sur le corpus Common Crawl et Wikipedia. En 2015, les auteurs de [8] ont pré-entraîné plusieurs modèles Word2Vec pour la langue française en utilisant le corpus FrWac [1] et le corpus FrWiki Dump (données françaises issues de Wikipedia). Depuis l'apparition de Word2Vec, l'évaluation des vecteurs de mots était basée sur la tâche d'analogie des mots introduite dans [14]. Cette tâche évalue la relation linéaire entre les vecteurs de mots pour vérifier la qualité de leurs plongements. Suivant cette idée, une liste de questions d'analogie en français pour évaluer les représentations de mots est présentée dans l'article [9]. Enfin, dans [12], les auteurs ont présenté un benchmark général pour évaluer les systèmes de NLP français nommé FLUE contenant de diverses tâches pour évaluer les modèles de compréhension naturelle français (NLU).

Notre objectif principal est d'évaluer la qualité des plongements de mots statiques pour la langue française et de fournir de nouveaux plongements entraînés sur un grand ensemble de données diversifiées qui sont performants dans les tâches d'analogie de mots, ainsi que dans les tâches de compréhension de la langue. De nombreux articles montrent que les modèles basés sur l'attention pré-entraînée comme BERT [6] et ELMo [19] surpassent les plongements de mots statiques. Mais cette amélioration des performances par rapport aux plongements statiques s'est faite au prix d'une utilisation moins efficace des ressources informatiques et de temps d'apprentissage et d'inférence plus longs, ainsi que d'une moindre interprétabilité et d'une plus grande dégradation de l'environnement. Même s'ils ne sont pas aussi expressifs ou puissants que les modèles de plongements contextuels, ils restent utiles dans la recherche sur le traitement du langage

naturel [10].

3 Crawling du web français

Pour entraîner les plongements de mots, nous avons décidé de collecter un énorme corpus brut français de multiples domaines sur le web en utilisant un crawler qui parcourt en permanence, de façon autonome et automatique, les différents sites et pages et sauvegarde leur contenu.

Graines. Nos graines initiales ont été sélectionnées en trouvant les pages Web les plus populaires sous le domaine ".fr". Pour ce faire, nous avons utilisé une liste publique⁴ qui contient les noms de domaine de sites web de différents genres tels que les sites d'information, Wikipedia et les médias sociaux. La frontière a ensuite été mise à jour avec les nouveaux liens découverts pendant l'exploration.

Crawler. Pour le crawling, nous avons choisi Heritrix⁵, un crawler web open-source supporté par Internet Archive. Notre configuration est constituée d'un seul nœud avec 25 threads crawlant pendant une période de 1,5 mois.

Sortie. La sortie générée par Heritrix suit le format de fichier WARC, avec des fichiers fractionnés à 1 Go chacun. Ce format a été traditionnellement utilisé pour enregistrer les "web crawls" sous forme de séquences de blocs de contenu recueillis sur le World Wide Web. Les blocs de contenu d'un fichier WARC peuvent contenir des ressources dans n'importe quel format, par exemple des images binaires ou des fichiers audiovisuels qui peuvent être intégrés ou liés à des pages HTML.

Extraction. Pour extraire le texte intégré dans la sortie du Heritrix, nous avons utilisé un outil warc-extractor⁶. Cet outil est utilisé pour analyser les enregistrements avec le WARC et ensuite analyser les pages HTML. Au cours de cette étape, nous avons également intégré le module de détection de langue FastText⁷ qui nous permet de filtrer le texte HTML par langue.

Dédoublonnage. Pour éliminer les données redondantes dans le corpus, nous avons utilisé l'outil de déduplication d'Isaac Whitfield⁸, le même que celui utilisé sur le corpus Common Crawl [16] basé sur un algorithme de hachage très rapide et résistant aux collisions. Le corpus final après déduplication a une taille totale de 33 Go, (171,701,319) lignes, (5,073,407,023) mots et (12,464,568) unigrammes uniques.

Ethique. Heritrix est conçu pour respecter le fichier robots.txt (fichier écrit par les propriétaires de sites Web pour donner des instructions sur leur site aux robots d'exploration), les directives d'exclusion et les balises de META nofollow.

4. Sites web les plus populaires avec le domaine .fr

5. <https://github.com/internetarchive/heritrix3>

6. <https://github.com/alexeygrigorev/warc-extractor>

7. <https://github.com/vinhkhuc/JFastText>

8. <https://github.com/whitfin/runiq>

Plongements	# Vocab	Outil	Méthode	Corpus	Dimension
fr_w2v_web_w5	0.8M	Word2Vec	cbow	fr_web	300
fr_w2v_web_w20	4.4M	Word2Vec	cbow	fr_web	300
fr_w2v_fl_w5	1M	Word2Vec	cbow	flaubert_corpus	300
fr_w2v_fl_w20	6M	Word2Vec	cbow	flaubert_corpus	300
cc.fr.300	2M	FastText	skip-gram	CC+wikipedia	300
frWac_200_cbow	3.6M	Word2Vec	cbow	frWac	200
frWac_500_cbow	1M	Word2Vec	cbow	frWac	500
frWac_700_sg	184K	Word2Vec	skip-gram	frWac	700
frWiki_1000_cbow	66K	Word2Vec	cbow	wikipedia dump	1000

TABLE 1 – Résumé des modèles utilisés dans nos expériences

4 Apprentissage des plongements de mots

Afin d'apprendre les plongements de mots français, nous avons utilisé Word2Vec de Gensim pour produire quatre modèles CBoW (Continuous Bag of Words) de vecteurs de mots. Deux de ces modèles ont été entraînés sur une portion de 33 Go du corpus français utilisé pour pré-entraîner FlauBERT[12]. Ce corpus est composé de 24 sous-corpus collectés à partir de différentes sources, avec des sujets et des styles d'écriture variés. Le premier modèle Word2Vec est noté fr_w2v_fl_w5 entraîné en utilisant Word2Vec CBoW avec une taille de fenêtre de 5 et une fréquence de coupure de 60. En d'autres termes, nous ne considérons la formation d'un plongement pour un mot, que s'il apparaît au moins 60 fois dans le corpus. Le second est fr_w2v_fl_w20 formé en utilisant Word2Vec CBoW avec une taille de fenêtre de 20 et une fréquence de coupure de 5. Les deux autres modèles ont été entraînés sur le corpus français déduplicé de 33 Go collecté sur le Web (section 3). Le premier est noté fr_w2v_web_w5 entraîné sur Word2Vec CBoW encore avec une taille de fenêtre de 5 et une fréquence de coupure de 60. Le second est fr_w2v_web_w20 entraîné sur Word2Vec CBoW avec une taille de fenêtre de 20 et une fréquence de coupure de 5. Tous les modèles examinés dans cet article ont un dimension de plongement de 300. Le tableau 1 contient les détails de chaque plongement de mots utilisée dans cette étude : cc.fr.300 représente les vecteurs de mots Fasttext français [9] de dimension 300 formés sur le corpus Common Crawl [16], les modèles frWac_200_cbow, frWac_500_cbow et frWac_700_sg [8] représentent les vecteurs Word2Vec français de dimensions 200, 500 et 700 respectivement qui sont entraînés sur le corpus frWac [2] (un corpus de 1,6 milliards de mots construit à partir du Web en limitant le crawl au domaine .fr) et finalement le modèle frWiki_1000_cbow [8] qui représente les vecteurs Word2Vec de dimension 1000 entraînés sur la partie française du jeu de données frWiki dump (une copie complète de tous les wikis de Wikimedia écrits en français).

5 Évaluation avec l'analogie des mots

Dans ce travail, la première méthode d'évaluation des vecteurs de mots est la tâche d'analogie [14]. Dans cette tâche, étant donné trois mots, A, B et C, nous pouvons prédire un

mot D tel que la relation entre A et B est la même entre C et D, en supposant qu'une relation linéaire entre les vecteurs de paires de mots indique la qualité de leurs plongements. Par exemple, si la relation entre les représentations des mots **king** et **man** est similaire à la relation entre les représentations de **queen** et **woman**, cela impliquera de bonnes plongements de mots. Par exemple, dans cette évaluation, nous utilisons les vecteurs de mots x_A , x_B et x_C de trois mots A, B et C pour calculer le vecteur $x_B - x_A + x_C$, et le vecteur le plus proche dans le dictionnaire du vecteur résultant sera considéré comme le vecteur du mot D. La performance des vecteurs de mots est finalement calculée en utilisant la précision moyenne sur le jeu de données de test. Pour évaluer nos plongements de mots français présentées dans la section 3, nous utilisons le jeu de données d'analogie française créé dans [9]. Ce dernier contient (31,688) questions d'analogie. Nous comparons également les résultats de nos vecteurs de mots avec ceux du FastText français (cc.fr.300), également produit dans [9] et les vecteurs de mots français entraînés sur frWacky et frWiki.

Plongements	Accuracy
fr_w2v_web_w5	41.95%
fr_w2v_web_w20	52.50%
fr_w2v_fl_w5	43.02%
fr_w2v_fl_w20	45.88%
cc.fr.300	63.91%
frWac_500_cbow	67.98%
frWac_200_cbow	54.45%
frWac_700_sg	55.52%
frWiki_1000_cbow	0.87%

TABLE 2 – Précision sur la tâche d'analogie

Résultats. Le tableau 2 montre la précision de la tâche d'analogie pour tous les modèles. Les résultats indiquent que les vecteurs de mots FastText et les plongements de mots entraînés sur frWacky français sont meilleurs que les vecteurs de mots CBoW Word2Vec produits au cours de cette étude dans la tâche d'analogie. Cependant, les auteurs du [11] ont prouvé que la propriété de relation géométrique par rapport aux analogies ne tient pas en général et qu'elle est probablement fortuite plutôt que systématique. En conclusion, nous avons décidé d'évaluer les vecteurs de mots sur des tâches de compréhension du langage naturel (NLU) telles que la

classification de textes, la paraphrase et l'inférence linguistique qui sont présentées dans [12] au lieu de s'appuyer uniquement sur l'analogie de mots comme d'autres études de plongements de mots statiques.

Application web. Pour visualiser et examiner nos plongements de mots, nous avons créé une application web qui contient les outils suivants :

1. Examen de l'analogie des mots : il prend en entrée trois mots, A, B et C, et il calcule les dix vecteurs les plus proches de $x_B - x_A + x_C$, puis affiche les mots correspondants comme nous pouvons le voir sur la figure 1.
2. Un calculateur de similarité en cosinus : il prend deux mots et calcule la similarité en cosinus entre les vecteurs correspondants.
3. Un outil de mots similaires : il trouve les dix mots les similaires à un mot d'entrée.
4. Un outil de visualisation : en utilisant T-SNE et k-means, cet outil affiche dans un espace vectoriel 2-D les n mots les plus proches du mot w distribués en k clusters, où n , w et k sont choisis par l'utilisateur. Par exemple, dans la figure 2, nous voyons une partie du graphe pour $n=200$, $w=paris$ et $k=8$.

6 Évaluation sur les tâches de FLUE

FLUE est un benchmark d'évaluation de la compréhension de la langue française [12], (French Language Understanding Evaluation), créé pour évaluer les performances des modèles NLP français. Il contient de nombreux jeux de données qui varient en termes de sujets, de niveau de difficulté, de taille et de degré de formalité. Cette section présente les résultats d'affinage de nos différents plongements Word2Vec et les compare avec les résultats d'affinage des autres plongements de mots mentionnés dans le tableau 1 sur certaines tâches de FLUE. Nous incluons également un plongement sans pré-entraînement (initialisation aléatoire de la couche de plongements des mots) pour la comparaison.

6.1 Classification des textes

Description des données. Dans cette tâche, le jeu de données utilisé est la partie française de CLS (Cross Lingual Sentiment) [20] qui se compose d'avis Amazon divisés en trois sous-ensembles : livres, DVD et musique.

Suivant [12], chaque sous-ensemble est binarisé en considérant tous les avis supérieurs à trois étoiles comme positifs, ceux inférieurs à trois étoiles comme négatifs, et en éliminant ceux qui ont trois étoiles. En outre, chaque sous-ensemble contient un ensemble équilibré de formation et de test contenant environ 1000 critiques positives et 1000 critiques négatives chacun. Enfin, un jeu de données de validation est formé pour chaque sous-ensemble en prenant une division aléatoire de 20 % des données d'entraînement.

Description du modèle. Le modèle utilise une couche de plongements de mots dont les poids initiaux et la taille des plongements varient en fonction de plongements de mots évalués. Cette couche est suivie d'un Bi-LSTM à deux

couches de 1500D (par direction). Enfin, nous utilisons la tête de classification utilisée dans [6] qui est composée des couches suivantes : une couche de dropout avec un taux de dropout de 0.1, une couche linéaire, une couche d'activation à tangente hyperbolique, une seconde couche de dropout avec le même taux de dropout et une seconde couche linéaire avec une dimension de sortie de deux (identique au nombre de classes).

Nous entraînons le modèle pendant 30 époques pour les différents plongements de mots tout en effectuant une recherche par quadrillage sur trois taux d'apprentissage différents : 5e-5, 2e-5 et 5e-6. Le modèle le plus performant sur le jeu de données de validation est choisi pour l'évaluation sur le jeu de données de test.

Plongements	Livres	Musique	DVD
w2v_0	67.20%	67.20%	62.15%
fr_w2v_web_w5	79.38%	79.49%	80.29%
fr_w2v_web_w20	81.55%	79.75%	80.75%
fr_w2v_fl_w5	82.16%	81.00%	79.64%
fr_w2v_fl_w20	82.38%	82.58%	82.43%
cc.fr.300	75.30%	74.35%	72.43%
frWac_500_cbow	81.30%	80.80%	79.59%
frWac_200_cbow	79.65%	75.40%	80.75%
frWac_700_sg	77.55%	75.20%	78.43%
frWiki_1000_cbow	66.30%	70.05%	60.94%

TABLE 3 – Précision sur le CLS français

Résultats. Le tableau 3 présente la précision sur le jeu de données de test pour chaque plongement de mots : w2v_0 représente les plongements des mots aléatoires. Les résultats démontrent l'importance des modèles pré-entraînés. Comme nous le voyons, tous les plongements de mots pré-entraînés surpassent de loin les plongements non pré-entraînés (aléatoires) avec une différence qui peut aller jusqu'à plus de 20% en termes de précision. En outre, nous voyons clairement une contradiction dans les résultats entre la tâche d'analogie et l'analyse des sentiments. Autrement dit, une meilleure performance sur la tâche d'analogie ne signifie pas une meilleure performance sur la tâche NLU.

6.2 Paraphrase

Description des données. Cette tâche utilise la partie française de PAWS-X [22] qui étend le jeu de données adversariales multilingues pour l'identification de paraphrases. Cette tâche vise à identifier si une paire de phrases avec des variations dans le choix des mots et la grammaire ont le même sens ou non. Le jeu de données est obtenu à partir de la traduction de paires de phrases en anglais provenant de Wikipédia et de Quora, avec un taux de chevauchement lexical élevé et jugé par des humains. Le jeu de données utilisé contient (49,401) échantillons d'entraînement, (1,992) échantillons de validation et (1,985) échantillons de test.

Description du modèle. Pour affiner les vecteurs de mots sur PAWS-X, nous utilisons le modèle d'inférence séquentielle amélioré (ESIM) [4]. Ce modèle est formé de trois composants essentiels : l'encodage des entrées, la modélisa-

père - homme + femme →

mère, 0.699
 fille, 0.633
 grand-mère, 0.613
 petite-fille, 0.611

FIGURE 1 – Exemple d’analogie de mots en utilisant l’outil d’analogie de notre application web (avec `fr_w2v_web_w20`).

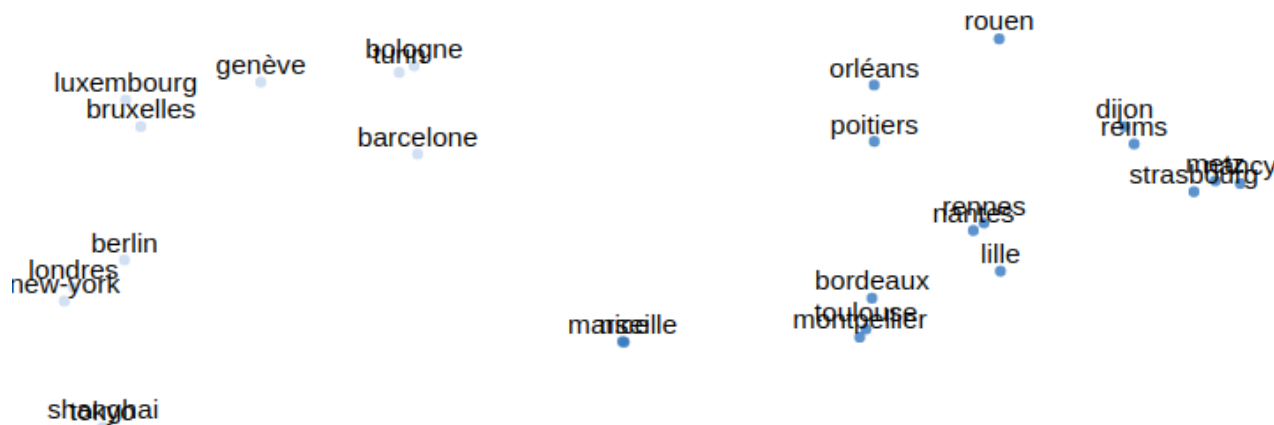


FIGURE 2 – 2 des 8 clusters de la requête : "Les 200 mots les plus proches de **paris**" en utilisant l'outil de visualisation de notre application web (avec fr_w2v_web_w20)

tion des inférences locales et la composition des inférences. Après la couche de plongements des mots, la couche d'encodage d'entrée est composée d'un BiLSTM 1500D à une couche qui encode les informations d'inférence locale des phrases A et B d'une longueur respective de l_a et l_b pour donner la sortie $\overline{A}$ et $\overline{B}$. Le composant suivant est le modèle d'inférence locale où :

1. Nous calculons la matrice de similarité E entre les deux phrases représentant les poids d'attention, telle que $e_{ij} = \bar{a}_i^T \bar{b}_j$.
2. Nous calculons la pertinence locale entre les deux phrases où chaque mot sera représenté comme une somme pondérée des informations pertinentes. De manière plus détaillée, $\tilde{a}_i = \sum_{j=1}^{l_b} \text{softmax}(e_{ij}) \bar{b}_j$ et $\tilde{b}_j = \sum_{i=1}^{l_a} \text{softmax}(e_{ij}) \bar{a}_i$.
3. Nous nous retrouvons l'amélioration de l'inférence d'inférence locale par une concaténation présentée par les auteurs qui prévoient d'améliorer l'inférence locale. Par exemple, les vecteurs résultants seront $m_a = [\bar{a}; \tilde{a}; \bar{a} - \tilde{a}; \bar{a} \odot \tilde{a}]$ et $m_b = [\bar{b}; \tilde{b}; \bar{b} - \tilde{b}; \bar{b} \odot \tilde{b}]$.

La quatrième couche est une couche de projection qui est composée d'une couche linéaire avec une dimension cachée de 1500D, d'une activation ReLU et d'une couche de dropout avec un taux de 0.1

Enfin, la couche de composition d'inférence est composée d'une couche BiLSTM 1500D à 1 couche qui donne les vecteurs V_a et V_b en sortie. Nous avons utilisé la même tête

de classification que dans la tâche de classification de texte dont l'entrée est le vecteur concaténé de «average pooling» et de «max pooling» de V_a et V_b .

Plongements	Accuracy
w2v_0	73.49%
fr_w2v_web_w5	77.12%
fr_w2v_web_w20	78.02%
fr_w2v_fl_w5	75.2%
fr_w2v_fl_w20	74.24%
cc.fr.300	62.80%
frWac_500_cbow	71.47%
frWac_200_cbow	74.80%
frWac_700_sg	67.39%
frWiki_1000_cbow	69.76%

TABLE 4 – Précision sur le PAWS-X français en utilisant l’ESIM

Résultats. La précision finale de chaque vecteur de mots est indiquée dans le tableau 4. Une fois encore, on peut constater que non seulement les résultats sont opposés entre les tâches NLU et la tâche d’analogie, mais aussi que les vecteurs de mots aléatoires surpassent les vecteurs de mots qui ont obtenu la deuxième meilleure précision dans la tâche d’analogie.

6.3 L'interprétation en langage naturel

Données et description du modèle. La tâche d'interprétation en langue naturelle (ou Natural Language Inference)

française de FLUE utilise la partie française du jeu de données XNLI [5]. Il contient des données NLI pour 15 langues. Chaque paire de phrases dans les jeux de données de test et de validation est annotée manuellement par des humains avec trois classes d'inférence : entaillement, neutre, et sans entaillement (contradiction). La partie française de jeu de données d'apprentissage de XNLI est obtenue par traduction automatique. Dans cette tâche, l'objectif est de déterminer s'il existe une relation d'implication, de contradiction ou de neutralité entre une phrase p appelée prémisses et une autre phrase h appelée hypothèse. Notez que le même échantillon peut être utilisé deux fois avec un ordre inverse entre les deux phrases avec une classe différente. Le jeu de données se compose de (92,702) échantillons d'entraînement, (2,491) échantillons de validation et (5,010) échantillons de test. Le modèle utilisé pour affiner les vecteurs de mots dans cette tâche est ESIM avec les mêmes configurations et paramètres que ceux décrits dans la section 6.2.

Plongements	Accuracy
w2v_0	61.37%
fr_w2v_web_w5	67.71%
fr_w2v_web_w20	68.27%
fr_w2v_fl_w5	69.41%
fr_w2v_fl_w20	69.57%
cc.fr.300	64.70%
frWac_500_cbow	63.82%
frWac_200_cbow	63.74%
frWac_700_sg	60.78%
frWiki_1000_cbow	61.34%

TABLE 5 – Précision sur le XNLI français en utilisant l'ESIM

Résultats. Nous rapportons les précisions finales des différents vecteurs de mots sur le jeu de données français de XNLI dans le tableau 5. Les résultats continuent de montrer les avantages des poids pré-entraînés et la surperformance des plongements de mots CBoW Word2Vec produits dans cette étude par rapport aux vecteurs de mots déjà existants.

6.4 Désambiguïsation du sens du nom

Description de données Cette tâche (NSD) est proposée par [12] pour la désambiguïsation du sens des mots (WSD) en français qui cible uniquement les noms. Elle est basée sur la partie française de la tâche WSD dans [15] pour créer le jeu de données de test composé de 306 phrases et (1,445) noms français annotés avec les clés de sens de WordNet et vérifiés manuellement. Le jeu de données d'apprentissage est obtenu en utilisant le meilleur système de traduction automatique anglais-français de l'outil fairseq [17] pour traduire en français le corpus SemCor et WordNet Glosses.

Description du modèle Nous avons utilisé les mêmes classificateurs que ceux présentés par [21] qui transmettent la sortie de nos vecteurs de mots dans une pile de 6 couches d'encodeurs de transformer et prédisent le sens du mot par une couche softmax à la fin. Nous avons utilisé les mêmes paramètres que dans [12].

Plongements	Single		
	Mean	Std	Ensemble
w2v_0	47.85%	± 1.17	52.87%
fr_w2v_web_w5	50.76%	± 1.4	53.77%
fr_w2v_web_w20	50.28%	± 0.92	53.2%
fr_w2v_fl_w5	50.16%	± 1.41	53.36%
fr_w2v_fl_w20	50.65%	± 1.62	52.46%
cc.fr.300	49.28%	± 1.5	52.39%

TABLE 6 – F1 score (%) sur la tâche de NSD

Résultats. Pour chaque modèle de plongements de mots, nous reportons dans le tableau 6 les valeurs de la moyenne et de l'écart-type des scores F1 (%) des 8 modèles individuellement et le score F1 (%) de l'ensemble des modèles en faisant la moyenne de la sortie de la couche Softmax de tous les modèles. Nous observons dans cette tâche que, même si fr_w2v_web_w5 surpasse légèrement les autres vecteurs de mots, nous avons un score F1 (%) très similaire pour tous les modèles. Nous pouvons dire que, pour cette tâche, les vecteurs de mots pré-entraînés n'améliorent pas le score final. Nous pensons que ces vecteurs ont la même nature quelque soit le contexte du mot alors que le but de la tâche est d'étudier les différentes significations d'un mot, est la raison de la similitude des scores entre ces modèles.

7 Conclusion

Dans ce travail, nous contribuons des plongements de mots entraînés sur des données extraites du web français et d'autres vecteurs de mots entraînés sur un ensemble de données mixtes formées à partir de sources et de sujets divers et variés en français, y compris Common Crawl et Wikipedia. Ces vecteurs de mots peuvent être utilisés comme poids initiaux pour divers modèles d'apprentissage profond, ce qui peut améliorer considérablement les performances par rapport à l'utilisation de poids aléatoires dans la couche de plongements des mots. En outre, ces plongements de mots dans des domaines généraux pourraient être pré-entraînés en plus sur des données de domaines spécifiques tels que la santé et le domaine juridique afin d'adapter les poids à un contexte approprié. De plus, nous avons utilisé un benchmark de quatre tâches NLP pour comparer la qualité de quatre plongements de mots français produites dans cette étude et de cinq autres plongements lexicaux existantes. Toutes les ressources présentées sont disponibles pour la communauté des chercheurs via notre application web. Enfin, nous avons montré que la tâche d'analogie de mots n'était pas fiable pour juger de la qualité des plongements de mots et de leur capacité à être performantes dans les tâches NLU. La raison des résultats sur la tâche d'analogie de mots pourrait faire l'objet d'une étude future.

Remerciements

Cette recherche a été soutenue par la chaire ANR AML/HELAS (ANR-CHIA-0020-01).

Références

- [1] Marco Baroni, Silvia Bernardini, Adriano Ferraresi, and Eros Zanchetta. The wacky wide web : A collection of very large linguistically processed web-crawled corpora. *Language Resources and Evaluation*, 43 :209–226, 09 2009.
- [2] Marco Baroni, Silvia Bernardini, Adriano Ferraresi, and Eros Zanchetta. The wacky wide web : A collection of very large linguistically processed web-crawled corpora. *Language Resources and Evaluation*, 43 :209–226, 09 2009.
- [3] Piotr Bojanowski, Edouard Grave, Armand Joulin, and Tomas Mikolov. Enriching word vectors with subword information. *Transactions of the Association for Computational Linguistics*, 5 :135–146, Dec 2017.
- [4] Qian Chen, Xiaodan Zhu, Zhen-Hua Ling, Si Wei, Hui Jiang, and Diana Inkpen. Enhanced LSTM for natural language inference. In *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 1 : Long Papers)*, pages 1657–1668, Vancouver, Canada, July 2017. Association for Computational Linguistics.
- [5] Alexis Conneau, Ruty Rinott, Guillaume Lample, Adina Williams, Samuel Bowman, Holger Schwenk, and Veselin Stoyanov. Xnli : Evaluating cross-lingual sentence representations. *Proceedings of the 2018 Conference on Empirical Methods in Natural Language Processing*, 2018.
- [6] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. BERT : Pre-training of deep bidirectional transformers for language understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics : Human Language Technologies, Volume 1 (Long and Short Papers)*, pages 4171–4186, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics.
- [7] Moussa Kamal Eddine, Antoine J. P. Tixier, and Michalis Vazirgiannis. Barthez : a skilled pretrained french sequence-to-sequence model, 2021.
- [8] Jean-Philippe Fauconnier. French word embeddings, 2015.
- [9] Edouard Grave, Piotr Bojanowski, Prakhar Gupta, Armand Joulin, and Tomas Mikolov. Learning word vectors for 157 languages. In *Proceedings of the Eleventh International Conference on Language Resources and Evaluation (LREC 2018)*, Miyazaki, Japan, May 2018. European Language Resources Association (ELRA).
- [10] Prakhar Gupta. Learning computationally efficient static word and sentence representations. Technical report, EPFL, 2021.
- [11] Sammy Khalife, Leo Liberti, and Michalis Vazirgiannis. Geometry and analogies : A study and propagation method for word representations. In Carlos Martín-Vide, Matthew Purver, and Senja Pollak, editors, *Statistical Language and Speech Processing*, pages 100–111, Cham, 2019. Springer International Publishing.
- [12] Hang Le, Loïc Vial, Jibril Frej, Vincent Segonne, Maximin Coavoux, Benjamin Lecouteux, Alexandre Al-lauzen, Benoit Crabbe, Laurent Besacier, and Didier Schwab. FlauBERT : Unsupervised Language Model Pre-training for French. In *LREC*, Marseille, France, 2020.
- [13] Louis Martin, Benjamin Muller, Pedro Javier Ortiz Suárez, Yoann Dupont, Laurent Romary, Éric de la Clergerie, Djamé Seddah, and Benoît Sagot. CamemBERT : a tasty French language model. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 7203–7219, Online, July 2020. Association for Computational Linguistics.
- [14] Tomáš Mikolov, Kai Chen, Greg Corrado, and Jeffrey Dean. Efficient estimation of word representations in vector space. In Yoshua Bengio and Yann LeCun, editors, *1st International Conference on Learning Representations, ICLR 2013, Scottsdale, Arizona, USA, May 2-4, 2013, Workshop Track Proceedings*, 2013.
- [15] Roberto Navigli, David Jurgens, and Daniele Vannella. SemEval-2013 task 12 : Multilingual word sense disambiguation. In *Second Joint Conference on Lexical and Computational Semantics (*SEM), Volume 2 : Proceedings of the Seventh International Workshop on Semantic Evaluation (SemEval 2013)*, pages 222–231, Atlanta, Georgia, USA, June 2013. Association for Computational Linguistics.
- [16] Pedro Javier Ortiz Suarez, Benoit Sagot, and Laurent Romary. Asynchronous pipelines for processing huge corpora on medium to low resource infrastructures. In Piotr Bański, Adrien Barbaresi, Hanno Biber, Evelyn Breiteneder, Simon Clematide, Marc Kupietz, Harald Lungen, and Caroline Iliadi, editors, *Proceedings of the Workshop on Challenges in the Management of Large Corpora (CMLC-7) 2019. Cardiff, 22nd July 2019*, pages 9 – 16, Mannheim, 2019. Leibniz-Institut für Deutsche Sprache.
- [17] Myle Ott, Sergey Edunov, Alexei Baevski, Angela Fan, Sam Gross, Nathan Ng, David Grangier, and Michael Auli. fairseq : A fast, extensible toolkit for sequence modeling. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics (Demonstrations)*, pages 48–53, Minneapolis, Minnesota, June 2019. Association for Computational Linguistics.
- [18] Jeffrey Pennington, Richard Socher, and Christopher Manning. GloVe : Global vectors for word representation. In *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pages 1532–1543, Doha, Qatar, October 2014. Association for Computational Linguistics.

- [19] Matthew Peters, Mark Neumann, Mohit Iyyer, Matt Gardner, Christopher Clark, Kenton Lee, and Luke Zettlemoyer. Deep contextualized word representations. In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics : Human Language Technologies, Volume 1 (Long Papers)*, pages 2227–2237, New Orleans, Louisiana, June 2018. Association for Computational Linguistics.
- [20] Peter Prettenhofer and Benno Stein. Cross-language text classification using structural correspondence learning. In *Proceedings of the 48th Annual Meeting of the Association for Computational Linguistics*, pages 1118–1127, Uppsala, Sweden, July 2010. Association for Computational Linguistics.
- [21] Loïc Vial, Benjamin Lecouteux, and Didier Schwab. Sense Vocabulary Compression through the Semantic Knowledge of WordNet for Neural Word Sense Disambiguation. In *Global Wordnet Conference*, Wroclaw, Poland, 2019.
- [22] Yinfei Yang, Yuan Zhang, Chris Tar, and Jason Baldridge. Paws-x : A cross-lingual adversarial dataset for paraphrase identification. *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, 2019.

Des clusters de tweets aux tags de descriptions : présentation d'un évènement par la caractérisation de ses manifestations

Olivier Gracianne^{1,2}, Anaïs Halftermeyer¹, Thi-Bich-Hanh Dao¹

¹ Univ. Orléans, INSA Centre Val de Loire, LIFO EA 4022

² Atos France

olivier.gracianne@etu.univ-orleans.fr, [anaïs.halftermeyer, thi-bich-hanh.dao]@univ-orleans.fr

Résumé

Notre travail aborde le problème de caractérisation de sous-événements dans des tweets par la description de leur groupement. Nous nous appuyons sur la représentation vectorielle de tweets dans un espace de plongement et les clusterisons. Nous proposons deux méthodes pour sélectionner des candidats parmi les mots des tweets pour construire un ensemble de tags de description. À partir de ces tags, nous proposons de construire une description par cluster en utilisant un modèle déclaratif en programmation linéaire en nombres entiers. Les expérimentations sur un jeu de données réelles montrent l'intérêt de notre approche.

Mots-clés

Twitter, plongement de mots/documents, clustering, description de clusters, programmation linéaire en nombre entier

Abstract

Our work deals with the sub-event characterization in tweets problem through the description of their grouping. We leverage vector representation of tweets in embedding space and cluster them. We propose two methods to select descriptor candidates among tweets' words to build a description tag set. From these tags, we propose to build a description per cluster using an integer linear programming declarative model. Experiments on real data show the interest of our approach.

Keywords

Twitter, word/document embedding, clustering, cluster description, integer linear programming

1 Introduction

L'utilisation massive des réseaux sociaux a produit d'énormes jeux de données pouvant être explorés et exploités avec une grande variété d'outils et de techniques. Chacun de ces médias a ses spécificités qui le rend plus ou moins adapté à exploiter pour une tâche donnée. Twitter offre des données qui ont des propriétés intéressantes, notamment pour la caractérisation d'événements. En particulier, les messages y sont courts et postés peu de temps après l'évènement en lui-même, par une grande variété de

sources. Cette plate-forme fournit une quantité conséquente de données disponibles publiquement à travers une API facile d'utilisation¹. Depuis août 2020, elle est devenue encore plus populaire avec une hausse significative des limites de requête pour les académiques.

Cela a permis aux chercheurs de constituer facilement des jeux de données avec de grands ensembles de tweets. Naturellement, le nombre d'approches s'appuyant sur ce type de données a augmenté dans le même temps.

La détection et la description d'événements sont devenues plus actives alors même qu'il s'agissait de tâches déjà exploitées.

La plate-forme Twitter permet la diffusion de contenus très diversifiés sur des sujets tout autant variés, parmi lesquels les événements du monde réel. Ainsi ces événements sont révélés par le prisme du tweet. Toutefois caractériser ces événements dans le volume et l'hétérogénéité de la masse informationnelle disponible sur ce média implique de devoir détecter des composants de grain plus fin. Ces manifestations mesurables de l'évènement cible se reflètent dans des sous-ensembles de tweets que nous appelons *sous-événements*.

Dans ces travaux, nous considérons des tweets émergeant autour d'un évènement, annoncé à l'avance ou non, étant publiés dans la fenêtre temporelle de celui-ci, et qu'il est possible de regrouper en plusieurs sous-événements structurés. Si nous prenons l'exemple d'une tempête, ses sous-événements constitutifs et a priori détectables pourraient être des vents violents et des crues menant à des voies de communication coupées. Nous émettons l'hypothèse que le clustering du contenu des tweets ayant un lien avec l'évènement permet de mettre en évidence certains de ses sous-événements. Nous représentons finalement ces derniers avec leur description, qui est un ensemble de tags, attribuée à un cluster de tweets. Ces groupements sont construits avec des messages contenant un mot ou un hashtag les liant à l'évènement.

Dans cet article, nous présentons notre approche bout-en-bout qui détecte et décrit des sous-événements en tirant avantage d'outils de TAL bien connus. Nos contributions

1. <https://api.twitter.com/2/tweets/search/all> est par exemple le point d'accès à cette API permettant de faire une recherche sur tous les tweets indexés par l'API.

sont :

- Le problème de description formulé avec une approche déclarative, en utilisant les plongements des mots des tweets comme descripteurs ;
- Le pipeline de traitement bout-en-bout utilisé ;
- Nous expérimentons notre méthode sur un jeu de données en français que nous avons collecté en ciblant la tempête Alex ayant frappé la France en octobre 2020. Nos résultats montrent la validité de notre méthode pour sélectionner des mots sémantiquement importants par rapport à leur cluster.

L'organisation du papier est la suivante. Nous faisons un tour d'horizon des travaux connexes à notre approche en section 2. Notre approche est présentée dans la section 3 et les expérimentations dans la section 4. Nous présentons nos conclusions en section 5 et discutons des perspectives de notre travail.

2 Travaux connexes

La détection d'évènements sur Twitter et la description de clusters sont deux domaines de recherches actifs et plusieurs approches existent déjà pour ce média. Nous présentons ici les axes de recherches principaux pour les deux.

2.1 Détection d'évènements sur Twitter

Toutes les approches portées à notre connaissance comportent une détection de brusques changements du comportement des utilisateurs. Un évènement détectable sur ce média, virtuel ou du monde réel, découle toujours de ce type de variations. Leur nature peut concerner le contenu des tweets, l'utilisation des fonctions "j'aime" et retweet ou encore la dynamique des communautés.

Ces dernières peuvent être représentées par des graphes d'utilisateurs, comme dans l'approche proposée dans [1] qui s'appuie sur des représentations en graphes des communautés d'utilisateurs et permet ainsi de détecter des mouvements de protestation. [15] propose de suivre l'évolution des communautés pour détecter les évènements significatifs pour leurs utilisateurs.

Nous avons choisi de nous orienter vers des méthodes s'appuyant sur le contenu des tweets, dans le but de pouvoir tirer parti de leur représentation en espace sémantique. Nous ne tenons pas compte des caractéristiques sociales de ce réseau.

La détection d'évènements basée sur le contenu des tweets peut être de différentes natures. Quantifier et/ou qualifier numériquement l'intérêt porté aux mots permet d'identifier ceux qui sont utiles à la détection d'évènements. Dans [17], les auteurs proposent un score permettant de traduire un tweet en signal et utilisent des méthodes de traitement du signal pour repérer les mots indiquant des évènements. Les travaux dans [11] se basent sur le traçage des n-grammes. Ceux qui apparaissent souvent obtiennent un score élevé et les tweets dans lesquels ils sont présents sont clusterisés. Les clusters d'évènements sont identifiés grâce à la connaissance extérieure (base de donnée de tierce partie).

Nous nous inscrivons dans la lignée de ces travaux, en tirant en plus parti des représentations de tweet par plongement.

Les approches à base de calcul de distribution de probabilité d'apparition d'un tweet pour un sujet sont nombreuses. Qu'elles se concentrent sur les sujets liés à des évènements [19] ou qu'elles en intègrent de plus généraux [18], la détection de l'évènement consiste à mesurer la brusque hausse de la probabilité d'affecter un tweet à un sujet. Toutefois, ces approches ne prennent pas en compte la sémantique des mots ou des tweets comme nous souhaitons le faire.

En travaillant à partir du flux de données de Twitter, on reçoit des données en continu. Il est possible d'en changer la représentation et de les clusteriser au fur et à mesure pour pouvoir suivre la dynamique des clusters ainsi formés et de détecter les évènements qui correspondent à de brusques croissances de ces clusters ([2], [9], [5]). Les différentes approches font varier les critères d'association d'un tweet à un cluster ou les règles de construction de ces derniers.

Nous nous concentrons sur l'étape de description des clusters pour valider notre approche sur un ensemble de données statique. Nous prévoyons de nous rapprocher de ces travaux de clustering incrémental ultérieurement.

2.2 Description de clusters

Avoir des clusters de données n'est pas suffisant pour récupérer l'information qu'un cluster contient. Il faut la rendre compréhensible par un humain. Une explication est nécessaire pour comprendre rapidement les résultats d'un processus impliquant autant de données et dans le cadre de nos travaux, pour aider à comprendre quel genre d'évènement est en cours et comment il se manifeste.

Dans [13], des concepts sous la forme d'ensembles clos d'éléments sont recherchés et les clusters sont construits autour. Ces concepts servent finalement d'explication. Dans [3], les clusters descriptifs sont construits en même temps que leur description, les données étant définies par un ensemble d'attributs et un ensemble de tags. Les clusters sont calculés à base des attributs et sont décrits en utilisant les tags. Ces approches se basent toutefois sur des vecteurs binaires, représentant un tweet par un vecteur de la taille du vocabulaire où chaque 1 représente un mot présent dans le tweet. Elles ne peuvent pas tirer parti des espaces de plongement sémantique.

Dans [4], les auteurs proposent de décrire les clusters d'une partition en utilisant un ensemble de descripteurs en lien mais distinct de l'ensemble de données. Par exemple, expliquer une communauté d'utilisateurs de Twitter avec des hashtags employés par les utilisateurs de celle-ci. Ils montrent que c'est un problème difficile computationnellement voire impossible à résoudre mais qui peut être relaxé pour le rendre accessible. Nous ne pouvons pas ré-appliquer cette méthode, comme nous cherchons à expliquer les données directement depuis leur contenu.

L'abstraction et l'extraction de résumé sont des approches bénéficiant du développement des approches neuronales. En reformulant des phrases du corpus cible ou en extrayant les phrases saillantes, des modèles similaires à BERT

[16] génèrent des résumés pour des ensembles de documents. Ce sont toutefois des approches nécessitant d'une part de grands jeux de données annotées dont nous ne disposons pas, et qui sont encore difficilement explicables d'autre part.

Nouveauté. La contribution de nos travaux réside d'une part dans le fait d'utiliser le contenu même des tweets clusterisés dans un espace sémantique riche pour les décrire et d'autre part dans la stratégie mise en place pour choisir les descripteurs potentiels. Comme nous allons le détailler, nous proposons une méthode de sélection de candidats descripteurs basée sur une mesure de leur intérêt descriptif et une autre intégrant en plus un critère de couverture de l'espace de représentation sémantique du cluster décrit.

3 Approche proposée

Nous cherchons à caractériser un "grand" événement en décrivant ses différents sous-événements qui se manifestent sur Twitter. Pour ce faire, nous nous appuyons sur l'hypothèse distributionnelle [8] et son application à travers les plongements de documents [10]. Nous postulons que des tweets traitant d'un même sujet seront regroupés par un algorithme de clustering s'ils sont représentés dans un espace de plongement sémantique. En effet, selon cette hypothèse, les tweets traitant d'un même sujet ont un sens similaire et y sont plus proches. Nous présentons finalement les événements par la mosaïque des descriptions de leurs sous-événements que nous construisons.

Dans ce qui suit, nous utiliserons $\mathcal{C}$ pour désigner un clustering (une partition), C_i pour le i -ème cluster, w pour un descripteur (un mot) et d pour une description (un ensemble de mots).

3.1 Représentation des données et clustering

Représentation des données. Pour construire des vecteurs de documents avec les tweets, nous utilisons le plongement proposé dans [10] avec le modèle Doc2Vec. Cette approche permet de construire un espace de représentation sémantiquement fin et basé sur le discours en situation. La distance entre deux vecteurs y représente bien une forme de dissimilarité sémantique entre les tweets correspondants. Par ailleurs, ce modèle permet aussi de contrôler la dimension des vecteurs obtenus. En effet, leur taille a un grand impact à la fois sur la mesure de distance entre les vecteurs et sur le temps de calcul du clustering. Nous avons fixé la taille des vecteurs de documents représentant les tweets à 500, nombre déterminé empiriquement.

Clustering. Selon l'hypothèse distributionnelle [8], l'espace de représentation dans lequel nous plongeons les tweets est riche car directement capté au travers des usages qui sont faits des mots eux-mêmes. Par sa construction, la distance entre deux vecteurs de tweet y représente une forme de distance sémantique entre ces tweets. Nous la mesurons avec la distance euclidienne. Pour construire des clusters sémantiquement cohérents représentant des sous-événements, nous proposons donc d'utiliser l'algorithme K-Means. En effet, ce dernier construit des clusters com-

pacts au sein de l'espace de représentation dans lequel il travaille en minimisant la distance d'un point au centroïde de son cluster. Nous cherchons ainsi à obtenir les clusters les plus cohérents et centrés sur un objet.

3.2 Descriptions

Nous construisons une description pour un cluster de tweets dans le but de l'interpréter et d'avoir une image de ce qu'il concerne. Nous proposons de considérer les mots des tweets d'un cluster comme ses possibles tags de descriptions et de construire un modèle en Programmation Linéaire en Nombre Entier (PLNE) pour sélectionner les meilleurs mots pour décrire chaque cluster. Pour ce faire nous devons définir un critère indiquant si un mot est adapté pour décrire un cluster et les règles adéquates pour construire une description.

3.2.1 Score DF-IDF

Nous proposons de mesurer la pertinence d'un mot w par rapport à un cluster C_i par le score DF-IDF défini comme suit :

$$DFIDF(w, C_i) = \frac{N_w(C_i)}{|C_i|} \cdot \log\left(\frac{N(\mathcal{C})}{N_w(\mathcal{C})}\right) \quad (1)$$

où $N_w(C_i)$ est le nombre de tweets contenant w dans le cluster C_i , $|C_i|$ la taille de C_i , $N(\mathcal{C})$ le nombre total de tweets et $N_w(\mathcal{C})$ le nombre de tweets contenant w dans le clustering $\mathcal{C}$.

Le calcul de ce score est inspiré de [17]. Dans ces travaux, le DF-IDF donne des informations sur l'importance d'un mot pour discriminer le tweet dans lequel il apparaît des autres tweets, sur la base d'une fenêtre temporelle d'apparition (représentée dans la seconde opérande de la formule, l'IDF). Ici, nous en avons adapté la formule pour que le score aide à discriminer le tweet contenant le mot cible sur la base des clusters obtenus. Comme on peut le voir dans la formule 1, plus un terme est fréquent dans un cluster, plus élevé est son DF-IDF. Mais s'il est fréquent dans tout le corpus, son score diminue. L'idée de ce score est d'y compresser une forme de typicalité et de fréquence.

3.2.2 Sélection des descripteurs candidats

La sélection des mots candidats pour décrire des clusters est un point clef de notre approche. Nous présentons ici trois méthodes de sélection.

Plus hautes fréquences. Les 20 mots les plus fréquents dans chaque cluster sont sélectionnés pour constituer la liste initiale des candidats. On en fait l'ensemble final en supprimant leurs occurrences multiples.

Top DF-IDF. Les 20 mots avec les plus hauts scores DF-IDF sont sélectionnés depuis chaque cluster. Nous obtenons l'ensemble de candidats avec le même dédoublement. Avec cette méthode les mots sélectionnés sont les plus pertinents d'après le score DF-IDF.

Hybride DF-IDF/FPF. Cette méthode vise à sélectionner les mots pertinents d'après le score DF-IDF qui couvrent le plus la représentation des mots. Pour assurer cette couverture, nous utilisons l'algorithme Farthest Point First (FPF)

[7] qui vise à sélectionner des représentants (têtes) des points d'un ensemble des points. La première tête est le point le plus éloigné de tous les autres et tous les autres points ont ce point comme tête. À chaque itération, le point le plus éloigné de sa tête est sélectionné pour devenir une nouvelle tête. Chaque point plus proche de celle-ci que de sa précédente tête lui est alors rattaché.

Dans notre espace, les points représentent des mots et la distance représente l'éloignement sémantique. Par son fonctionnement, cet algorithme va donc choisir des mots qui en représentent d'autres au sens proche et ces représentants seront sémantiquement distants les uns des autres. Cet algorithme est appliqué dans chaque cluster pour sélectionner 20 têtes, en restreignant les mots possibles à ceux ayant un score DF-IDF au dessus de la médiane de leur cluster. Nous nous assurons ainsi de garantir une certaine pertinence des mots que peut proposer cette sélection hybride.

3.2.3 Construction des descriptions

Nous utilisons la PLNE pour formuler notre problème d'une manière déclarative. Étant donné un clustering $\mathcal{C} = \{C_1, C_2, \dots, C_k\}$ et un ensemble de candidats descripteurs $W = \{w_1, w_2, \dots, w_m\}$, le modèle prend en entrée :

- une matrice P de taille $k \times m$, où P_{ij} est le score DF-IDF du mot w_j dans le cluster C_i ;
- k matrices T^k de taille $|C_k| \times m$, où T_{ij}^k vaut 1 si le candidat w_j est présent dans le tweet i du cluster C_k .

Nous définissons comme variables une matrice binaire X de taille $k \times m$, où X_{ij} vaut 1 lorsque le mot w_j est utilisé pour décrire le cluster C_i . L'objectif du modèle est de maximiser la pertinence des descriptions construites, ce qui se traduit par la maximisation de la somme :

$$\sum_{i=1}^k \sum_{j=1}^m X_{ij} P_{ij} \quad (2)$$

Les exigences pour les descriptions sont exprimées par des contraintes du modèle :

- **Non-vide** Une description ne doit pas être vide :

$$\forall i \in [1, k], \sum_{j=1}^m X_{ij} \geq 1$$

De cette manière, nous garantissons d'avoir un support pour l'interprétation de chaque cluster ;

- **Non-recouvrement** Un mot ne peut être utilisé dans plusieurs descriptions :

$$\forall j \in [1, m], \sum_{i=1}^k X_{ij} \leq 1$$

Nous évitons ainsi de construire des descriptions trop similaires ;

- **Taille max** Une description doit contenir au plus t mots :

$$\forall i \in [1, k], \sum_{j=1}^m X_{ij} \leq t$$

Nous limitons la taille des descriptions pour ne pas compliquer leur interprétation ;

- **Score non-nul** Un mot dans une description doit avoir un score DF-IDF non nul dans le cluster correspondant :

$$\forall i \in [1, k], \forall j \in [1, m], P_{ij} = 0 \implies X_{ij} = 0$$

Du fait des contraintes non-vide et non-recouvrement, le modèle peut affecter un mot à un cluster même s'il n'y apparaît pas. Cela n'apportant rien pour l'interpréter, nous prévenons cette possibilité.

- **Couverture d'apparitions** La somme du nombre d'apparition des mots d'une description dans le cluster correspondant doit dépasser une proportion α de la taille du cluster :

$$\forall c \in [1, k], \sum_{j=1}^m \sum_{i=1}^{|C_c|} T_{ij}^c X_{cj} \geq \alpha |C_c|$$

Cette contrainte ne représente pas la couverture effective des tweets d'un cluster par les mots de sa description. En effet sur un cluster de deux tweets, si l'un contient deux mots descripteurs et l'autre aucun, une description contenant ces deux mots satisfera cette contrainte avec $\alpha = 1$ même si elle ne couvre effectivement que la moitié de ce cluster. Nous formulons une contrainte de couverture de cette manière pour qu'elle soit la plus riche possible tout en restant linéaire.

4 Expérimentations

Les expérimentations sont conduites pour répondre aux questions suivantes :

- Quelle est la meilleure manière de choisir les descripteurs potentiels que l'on peut affecter à un cluster ?
- Les descriptions construites sont-elles adéquates pour leur clusters ?

Nous apportons nos éléments de réponse à ces deux questions dans la section 4.2.

4.1 Paramétrage expérimental

4.1.1 Données et requête

Pour constituer un premier jeu de données sur lequel évaluer notre approche, nous avons ciblé l'évènement climatique marquant le plus récent en France au moment où nous avons lancé la collecte. C'est donc sur la tempête Alex, qui a frappé le sud-est et le nord-ouest du pays, que nous avons centré notre requête.

Les épisodes climatiques violents sont des cibles toutes indiquées pour vérifier nos hypothèses car ils touchent un territoire et tous les utilisateurs de Twitter qui s'y trouvent. D'autres types d'évènements peuvent en effet n'affecter ou n'intéresser que certaines communautés sur la plate-forme, limitant ainsi la variété du contenu que l'on peut en retirer.

En s'intéressant à une tempête, les messages récoltés seront diversifiés en contenu, forme et origine.

L'idéal pour construire une requête visant à capter des tweets pour caractériser un événement et ses manifestations serait de disposer d'une ontologie décrivant cette structuration. On disposerait alors de mots clefs pour construire plusieurs requêtes de différentes granularités.

Sans cette ontologie, nous travaillons à partir d'une liste de mots clefs restreinte se résumant dans notre cas à "tempête" et "alex". Nous ajoutons à la requête des variations de ces formes pour capturer les morphologies incorrectes qui peuvent apparaître sur Twitter.

4.1.2 Préparation des données

Les mots vides, ou stop words, et les tokens utilisés dans la requête sont supprimés des tweets collectés. Nous utilisons ensuite le `LefffLemmatizer`² à travers son implantation dans `spaCy` pour obtenir des lemmes et nous utilisons des expressions régulières pour remplacer les heures et les nombres avec respectivement les étiquettes 'matchedHour' et 'matchedNumber'. Comme le tokeniseur intégré au lemmatiseur s'est montré moins robuste sur nos tweets que sur des jeux de données classiques, nous avons aussi ajouté une expression régulière pour vérifier que la séparation des mots était correcte. Nous n'avons retenu que les lemmes d'au moins trois caractères et les tweets composés d'au moins deux lemmes.

4.1.3 Représentation des données

Pour représenter nos tweets avec des vecteurs, nous utilisons le plongement de document proposé dans [10]. Une fois préparées, nos données sont exploitées par ce modèle pour apprendre un espace de représentation dans lequel la distance entre deux points représente une distance sémantique entre les deux tweets correspondants.

Nous avons choisi le modèle de plongement Doc2Vec car il permet d'apprendre une représentation des documents en même temps qu'une représentation des mots qui les composent. Nous en avons utilisé l'architecture PV-DBow plutôt que sa contrepartie PV-DM. La première rend l'entraînement plus coûteux que la seconde mais fournit des vecteurs qui mènent à obtenir de meilleurs résultats pour les tâches qui en évaluent la qualité.

Le modèle apprend l'espace de représentation en s'entraînant à prédire le mot venant après une séquence de mots donnés. Dans notre cas cela consiste à extraire une séquence de mots du tweet en cours d'apprentissage, en en conservant l'ordre, pour prédire le dernier mot de cet extrait. La taille de cette séquence est contrôlée. Nous l'avons paramétrée pour que le modèle s'entraîne à prédire le dernier d'une séquence de trois mots. Nous l'avons déterminé empiriquement, une séquence trop longue se révélant inadaptée pour s'entraîner sur des tweets pouvant être plus petits qu'elle.

Nous avons utilisé l'implantation du modèle de Gensim 4.1.2³.

2. <https://spacy.io/universe/project/spacy-lefff>

3. <https://radimrehurek.com/gensim/>

4.1.4 Clustering

Nous avons utilisé l'algorithme K-Means pour clusteriser les tweets, en nous appuyant sur l'hypothèse que les tweets avec un sens proche seront proches dans l'espace de plongement. Nous l'avons paramétré avec $K=10$. Nous avons déterminé ce nombre empiriquement. Trop peu de clusters les font contenir des informations sur plusieurs sujets. Trop de clusters amène à séparer des informations qui devraient être regroupées, comme c'est le cas avec 20 clusters pour nos données. Ce paramétrage est expérimentalement apparu comme un juste compromis pour faire apparaître des clusters bien identifiés et d'autres qui pourraient être soit sous découpés soit réunis. Nous avons ainsi un échantillon de tous les cas qui peuvent nous intéresser. Nous initialisons le clustering avec la méthode K-Means++ [14], une méthode de sélection des centres de cluster initiaux semi-aléatoire qui permet d'optimiser le clustering. Par ailleurs nos expérimentations ont montré que travailler sur des vecteurs de score TF-IDF engendre des clusters déséquilibrés en taille tandis que travailler sur la base de plongements a permis de les harmoniser : avec un corpus de 19457 tweets, 7 clusters font entre 1000 et 3000 tweets, 2 autres moins de 700 et un dernier plus de 4700.

4.1.5 Problème de PLNE

Un framework déclaratif est utilisé pour présenter notre problème. Nous l'avons codé en utilisant l'interface Python 3.8.8⁴ du solver Gurobi 9.1.1.2⁵.

4.1.6 Métriques d'évaluations

À notre connaissance, il n'y a pas de métrique de référence qui serait adaptée à l'évaluation des descriptions que nous construisons. Par exemple, les différentes métriques ROUGE [12] offrent une évaluation en comparant un résumé généré à une référence, dite gold, produite par un humain, qui n'existe pas pour nos données. SUPERT [6] compare les phrases du résumé produit aux phrases saillantes du document. Notre description étant constituée seulement de mots, ce genre de métrique n'est pas approprié non plus. Nous proposons donc les deux critères suivants :

- **Importance Relative (IR).** L'importance relative d'un ensemble de mots d par rapport à un cluster C est définie en utilisant le score DF-IDF :

$$IR(d, C) = \sum_{w \in d} DFIDF(w, C) \quad (3)$$

Ceci permettra de mesurer la pertinence d'une description par rapport à un cluster.

- **Somme Pondérée des Distances (SDP).** Un autre objectif est de construire des descriptions de clusters sémantiquement cohérentes et distinctes. Nous proposons de mesurer la différence de deux descriptions d_i et d_j par :

$$SPD(d_i, d_j) = \frac{\sum_{w \in d_i} \sum_{w' \in d_j} dist(w, w')}{|d_i| |d_j|} \quad (4)$$

4. <https://www.python.org/downloads/release/python-388/>

5. <https://www.gurobi.com/documentation/9.1>

où $dist(w, w')$ est la distance dans l'espace de plongement entre les vecteurs représentant w et w' .

Une IR haute d'une description par rapport à un cluster indiquera que les mots de celles-ci sont adéquats pour identifier le cluster et le séparer des autres, au regard du score DF-IDF. On attend donc une IR d'une description plus élevée pour son cluster que pour les autres. Si une description a une IR similaire pour tous les clusters, elle les identifie et les sépare tous de la même manière et ne permet alors d'en distinguer aucun.

La SPD représente la distance entre deux descriptions. Plus elle est basse, plus elles sont proches dans l'espace de plongement des vecteurs de tweets et ainsi elles sont proches sémantiquement l'une de l'autre. On attend ainsi la SPD d'une description plus basse par rapport à elle-même que par rapport aux autres. Si toutes les SPD d'une description sont équivalentes, alors elle est aussi proche sémantiquement d'elle-même que des autres et il est difficile de déterminer ce qui la caractérise sur ce critère.

Comparer les descriptions avec ces métriques apporte donc des informations sur leur séparation et leur identification, et par extension sur celles des clusters concernés.

4.2 Résultats et analyses

Plusieurs exécutions du clustering avec un même paramétrage basés sur des plongements de documents donnant des clusters très similaires, l'échantillon de résultats présentés est représentatif de ceux que l'on peut obtenir avec notre méthode. Il faut noter que le score obtenu par le modèle d'optimisation construisant les descriptions varie, de manière non significative toutefois : d'un clustering à l'autre, le total des scores des descriptions obtenu varie en moyenne de moins de 0,08%. L'initialisation K-Means++ [14] sélectionne un premier centre de cluster aléatoirement parmi les données puis les centres suivants les plus éloignés les uns des autres pour optimiser le clustering. Nous avons exécuté un grand nombre de fois le clustering sur les mêmes données pour vérifier la robustesse des partitions construites. Les clusterings et les descriptions obtenus sont suffisamment similaires en terme de composition⁶ pour être représentatifs de l'ensemble des résultats. Nous en présentons donc un ici, sélectionné aléatoirement.

Nous donnons ici les résultats des sélections de candidats présentés dans la section 3.2.2. Nous fixons le paramètre α de la contrainte de couverture d'apparition présenté dans la section 3.2.3 à 30 pour des raisons de comparaisons. C'est le maximum qu'il est possible de fixer sur cette contrainte commun aux trois méthodes de sélection. Les clusters et leur description sont alignés dans le tableau de résultats. Le meilleur score d'une ligne de matrice apparaît en orange. Ainsi, une sélection de bonne qualité au regard de nos critères correspondra à un tableau avec une diagonale nord-ouest/sud-est fortement colorée avec des valeurs en orange et le reste du tableau faiblement coloré.

6. D'une exécution de clustering à l'autre, la composition des clusters est identique en moyenne à 99,5% et le score total du modèle construisant les descriptions varie en moyenne de 0,08%.

Les expérimentations sont réalisées sur une machine dotée d'un i7-11850H et de 32 Go de RAM. Le temps total d'exécution est inférieur à 2 minutes avec les sélections par les plus hautes fréquences et Tops DF-IDF. Avec la sélection hybride, on rajoute 1 minute 30. À noter que la majorité de ces temps est occupée par les écritures/lectures des fichiers de résultats.

4.2.1 Résultat sélection par les plus hautes fréquences

Cluster Id description	0	1	2	3	4	5	6	7	8	9
0	3.7625	0.4101	0.6088	0.3916	0.5083	0.4025	0.7407	0.8427	1.0482	0.489
1	0.6384	1.1374	0.4155	0.6609	0.7196	0.4025	0.2794	0.2187	0.2677	0.3943
2	0.0649	0.0663	0.407	0.0316	0.0809	0.0694	0.5158	0.1034	0.1454	0.0379
3	0.4414	0.0958	0.407	0.0316	0.0809	0.0694	0.5158	0.1034	0.1454	0.0379
4	0.4635	0.8924	0.5764	1.2289	1.1293	0.6666	0.46	0.3805	0.4049	0.9444
5	4.2155	0.6867	0.4911	0.5093	0.7322	1.2289	0.6809	0.4188	0.5048	0.5554
6	1.3401	0.6176	0.9346	0.9989	0.8777	1.2289	1.0043	1.307	1.1801	1.0643
7	1.2352	0.3813	1.051	0.2253	0.6228	1.2396	0.8427	3.3827	1.3322	0.2843
8	0.6688	0.4155	0.6964	0.1727	0.5684	0.641	0.4559	1.8767	1.3322	0.1795
9	0.6819	0.4655	0.3026	1.2289	0.5974	0.7689	0.5159	0.2897	0.1349	1.0643

(a) Matrice de l'importance relative (IR) des descriptions

Description Id	0	1	2	3	4	5	6	7	8	9
0	0.5695	0.7787	0.6878	0.6612	0.3793	0.5556	0.5246	0.5833	0.5618	0.5618
1	0.5695	0.5076	0.4223	0.4991	0.5009	0.6091	0.6982	0.6719	0.4533	0.4533
2	0.7787	0.5076	0.5076	0.679	0.6555	0.7046	0.8118	0.74	0.675	0.7642
3	0.6878	0.4223	0.679	0.5076	0.3126	0.6864	0.5659	0.7853	0.7921	0.2959
4	0.6612	0.4591	0.6555	0.3126	0.3748	0.6864	0.5659	0.7118	0.6964	0.4092
5	0.3793	0.5009	0.7046	0.6864	0.6832	1.2289	0.6179	0.6612	0.6779	0.583
6	0.5556	0.6091	0.8118	0.5659	0.5897	0.6179	0.468	0.5908	0.6678	0.5354
7	0.5246	0.6982	0.74	0.7853	0.7118	0.6612	0.5908	0.6678	0.4022	0.7262
8	0.5833	0.6719	0.675	0.7921	0.6964	0.6779	0.6678	0.4022	0.7262	0.7657
9	0.5618	0.4533	0.7642	0.2959	0.4092	0.583	0.5354	0.7262	0.7657	0.3053

(b) Matrice des sommes pondérées des distances (SPD) des descriptions

Id	Mots
0	'bretagne', 'violent', 'morbihan', 'météo', 'électricité', 'vendredi', 'fort', 'ouest', 'alerte', 'jeudi'
1	'matchedNumber', 'temps', 'nouveau', 'grand', 'contre', 'entre'
2	'ainsi', 'trois', 'bois', 'bout', 'vivement', 'dansaient', 'capucine', 'virant'
3	'faire', 'bien', 'voir', 'très', 'quand', 'trop', 'dire', 'beau', 'rien', 'ournée'
4	'comme', 'fait', 'sans', 'depuis', 'pouvoir', 'aussi', 'venir', 'merci'
5	'vent', '#tempête', 'nuit', 'heure', 'kilomètre', 'rafales', 'vents', '#bretagne', '#morbihan', 'côte'
6	'après', '#tempetealex', 'demain', 'cause', 'arriver', 'dégât', 'image', 'soir', 'vallée'
7	'alpes-maritimes', 'france', 'passage', 'vigilance', 'sinistre', 'personne', 'disparu', 'rouge', 'département', 'crue'
8	'alpesmaritimes', 'deux', 'direct', 'mort', 'corps', 'pompier', 'bilan', 'italie', 'orange', 'tousjours'
9	'aller', 'avant', 'calme', 'chez', 'pluie', 'tempête'

(c) Description des clusters

FIGURE 1 – Évaluation des descriptions obtenues avec la sélection de candidats par les plus hautes fréquences. Cette sélection propose 87 candidats dont 35, apparaissant en vert foncé, ne sont pas dupliqués dans leur liste initiale

Dans la figure 1a nous pouvons voir que même si la diagonale est fortement colorée, la cellule de meilleur score est dans 5 cas sur 10 située ailleurs sur la ligne. Les descriptions produites sont donc, du point de vue leur IR, adéquates pour leur cluster mais également pour d'autres. Les descriptions 0, 1, 4, 6 et 7 sont mêmes plus adéquates pour un autre cluster. On peut remarquer que les descriptions 2, 3, 5 et 8 ont une valeur d'IR pour leur cluster nettement plus élevée. En outre, on remarque que les lignes 1, 4 et 6 et 9 sont fortement colorées. Les descriptions correspondantes ont donc une IR similaire pour tous les clusters.

Dans la figure 1b, la meilleure valeur de la ligne est bien sur la diagonale de la matrice dans 8 cas sur 10. Dans la majorité des cas, les descriptions semblent donc plus proches sémantiquement d'elles-mêmes que d'autres descriptions. Par ailleurs, les lignes fortement colorées sont les 0, 6 et 9. Elles sont relativement peu nombreuses. Il semble donc que les descriptions soient distinctes sémantiquement les unes des autres. On constate des lignes où l'écart entre la meilleure valeur et les autres est net (1, 2, 3, 5 et 8) et d'autres plus

uniformes (0, 6 et 9). Les observations sur ces deux matrices semblent pointer vers les conclusions suivantes : soit il y a des clusters particulièrement faciles et d'autres particulièrement difficiles à distinguer, soit la sélection par fréquence n'offre pas de candidats adaptés pour pouvoir décrire les clusters de manière distincte.

Dans la figure 1c, on peut constater que moins de la moitié des mots à la disposition du modèle de PLNE n'étaient pas dupliqués dans la liste initiale des candidats par la fréquence. Il est possible d'interpréter cela de deux manières différentes : soit le clustering peine à séparer certains clusters, soit la fréquence seule n'est un pas un critère suffisant pour identifier des descripteurs distincts d'un cluster à l'autre. Du point de vue de l'interprétation, on remarque que certaines descriptions sont beaucoup plus compréhensibles que d'autres. Par exemple on voit clairement se dégager un sujet pour le cluster 5, les vents qui ont frappé les côtes de la Bretagne lors d'une nuit de la tempête, là où il est impossible d'attribuer un sujet en particulier à la description 3. Les descriptions 0, 5, 7 et 8 paraissent ainsi assez facilement interprétables, les autres ne semblant centrées sur rien en particulier.

4.2.2 Résultat sélection par les Tops DF-IDF

Id description	Cluster Id										
	0	1	2	3	4	5	6	7	8	9	
0	3.7628	0.4101	0.6088	0.3916	0.5083	0.7177	0.7407	0.8427	1.0482	0.489	
1	0.3886	0.4936	0.4958	0.747	0.7511	0.7779	0.3132	0.3642	0.4399	0.4258	
2	0.0917	0.2096	0.3661	0.1003	0.2291	0.0747	0.6342	0.4301	0.5152	0.1029	
3	0.2981	0.8142	0.3661	0.1003	0.2291	0.0747	0.6342	0.4301	0.5152	0.1029	
4	0.4218	1.131	0.7816	1.3004	1.320	0.5638	0.512	0.6415	0.7335	0.8475	
5	4.407	0.946	0.5667	0.5719	0.8897	1.0419	0.7195	0.4374	0.518	0.5943	
6	1.8286	0.6169	0.9092	0.9042	0.8464	1.061	1.0216	1.4498	1.4411	0.9264	
7	1.1049	0.2756	0.7668	0.1813	0.5214	1.2222	0.6732	2.7907	1.11	0.2335	
8	0.5277	0.3701	0.7693	0.111	0.5899	0.4122	0.5313	2.4536	0.729	0.1241	
9	0.8505	0.7084	0.4292	2.0592	0.9755	0.9627	0.7647	0.4009	0.1882	0.719	

(a) Matrice de l'importance relative (IR) des descriptions

Description												
Id	Id	0	1	2	3	4	5	6	7	8	9	
Id	description	0	0.5336	0.5336	0.7878	0.7132	0.6587	0.3985	0.4723	0.5402	0.6089	0.5943
1		0.5336	0.5336	0.5093	0.4456	0.4978	0.5251	0.5425	0.5655	0.5462	0.4784	
2		0.7878	0.5093	0.5336	0.6383	0.6427	0.7327	0.7742	0.6725	0.6126	0.749	
3		0.7132	0.4456	0.6383	0.6383	0.3934	0.6983	0.6435	0.7786	0.3069	0.2587	
4		0.6587	0.4978	0.6427	0.3934	0.4508	0.6885	0.5929	0.6383	0.655	0.4558	
5		0.3985	0.5251	0.7327	0.6983	0.6885	0.5929	0.5377	0.6924	0.7046	0.6087	
6		0.4723	0.5425	0.7742	0.6435	0.5929	0.5377	0.4806	0.5888	0.642	0.5739	
7		0.5402	0.5655	0.6725	0.7786	0.6383	0.6924	0.5888	0.3795	0.7512	0.7512	
8		0.6089	0.5462	0.6126	0.3069	0.655	0.7046	0.642	0.3617	0.7046	0.8175	
9		0.5943	0.4784	0.749	0.2587	0.4558	0.6087	0.5739	0.7512	0.8175	0.2808	

(b) Matrice des sommes pondérées des distances (SPD) des descriptions

Id	Mots
0	'vent', 'morbihan', 'violent', 'électricité', 'météo', 'ouest', 'fort', 'alerte', 'vendredi', 'jeudi'
1	'matchedNumber', 'temps', 'nouveau', 'grand', 'contre', '#france', 'entre', '#lemonde', 'place', 'suite'
2	'ainsi', 'bois', 'dansaient', 'capucine', 'virant', 'bout', 'vivement', 'trois', 'sinistré', 'village'
3	'voir', 'bien', 'quand', 'dire', 'beau', 'vouloir', 'rien', 'journée', 'prendre', 'bonjour'
4	'faire', 'comme', 'fait', 'sans', 'pouvoir', 'aussi', 'venir', 'jamais', 'vallée', '#alex06'
5	'#tempête', 'nuit', 'heure', 'kilomètre', 'rafale', 'vents', '#bretagne', '#morbihan', 'côte', 'terre'
6	'#tempêtealex', 'après', 'aller', 'bretagne', 'arriver', 'dégât', 'france', 'image', 'soir', 'premier'
7	'alpes-maritimes', 'vigilance', 'passage', 'personne', 'disparu', 'rouge', 'crue', 'département', 'recherchées', 'habitant'
8	'#alpesmaritimes', 'deux', 'direct', 'mort', 'corps', 'pompiers', 'bilan', 'italie', 'retrouvé', 'applegreen'
9	'depuis', 'avant', 'calme', 'demain', 'chez', 'trop', 'cause', 'pluie', 'tempête', 'dehors'

(c) Descriptions des clusters

FIGURE 2 – Évaluation des descriptions pour la sélection de candidat par les Top DF-IDF, avec 101 candidats dont 51 non dupliqués.

Dans la figure 2a nous pouvons voir la diagonale fortement colorée et la cellule de meilleur score d'une ligne dans 7 cas sur 10 sur la diagonale. Les descriptions produites sont donc, du point de vue leur IR, adéquates pour leur cluster et la majorité d'entre elles sont les plus adaptées pour leur

cluster. On peut remarquer que les descriptions 2, 3, 5 et 8 ont une valeur d'IR pour leur cluster qui se distingue nettement des autres. Dans 4 cas sur 10 donc, la description discrimine nettement mieux son cluster que les autres. On remarque d'autre part que les lignes 1, 4, 6, et 9 sont fortement colorées. Les descriptions sont donc dans 4 cas sur 10 similaires aux autres du point de vue IR.

Dans la figure 2b, la meilleure valeur de la ligne est sur la diagonale dans 6 cas sur 10. On peut remarquer que pour les descriptions 6, 7 et 9, la valeur de la diagonale est tout de même proche de la meilleure valeur de la ligne. Il semble donc que la sélection par DF-IDF propose des candidats qui permettent de construire des descriptions cohérentes mais qui peinent à se distinguer sémantiquement. Par ailleurs, on peut remarquer que les lignes 0, 6 et 9 sont fortement colorées. On constate ici des tendances similaires avec les résultats de la sélection par les plus hautes fréquences. Nous pouvons l'interpréter de deux manières différentes : soit la composante fréquentielle du calcul du score DF-IDF y est encore trop forte et cette sélection pose le même problème que la précédente, soit il y a des clusters plus difficiles à discerner des autres.

Dans la figure 2c, on constate que plus de la moitié des candidats initiaux n'étaient pas dupliqués. Les sujets des descriptions des clusters 0, 5, 6, 7 et 8 semblent se démarquer assez facilement alors que ceux des autres sont beaucoup moins clairs voire impossibles à identifier.

4.2.3 Résultat sélection hybride

Cluster Id	0	1	2	3	4	5	6	7	8	9
Id description	0	2.2925	0.3438	0.3741	0.3243	0.3897	0.3999	0.4423	0.6993	0.3419
1	0.5901	0.539	0.7665	0.8569	1.1598	0.3357	0.5115	0.6812	0.4923	
2	0.2562	0.6439	0.34	0.5215	0.6358	0.34	0.5215	1.0348	1.0925	0.2499
3	0.3958	0.6206	0.4111	1.1377	0.6343	0.5433	0.4317	0.2432	0.2422	0.7947
4	0.5731	0.9287	0.5976	1.1377	1.1377	0.7669	0.4401	0.4581	0.4584	0.7431
5	3.2166	0.9414	0.5286	0.549	0.7976	1.5453	0.5644	0.304	0.4037	0.5044
6	1.328	0.5612	0.8625	0.934	0.8703	1.0826	1.009	1.1369	1.207	1.1123
7	0.3262	0.1965	0.4959	0.2471	0.4124	0.4267	0.3843	1.434	1.1289	0.2623
8	0.484	0.4197	0.7192	0.2389	0.5704	0.3861	0.3665	1.4489	0.749	0.2293
9	0.303	0.5787	0.2732	0.6218	0.2186	0.4974	0.2837	0.103	1.2354	

(a) Matrice de l'importance relative (IR) des descriptions

Id	Description	Id									
	description	0	1	2	3	4	5	6	7	8	9
0		0.5011	0.4984	0.7308	0.6072	0.5431	0.3721	0.5503	0.6004	0.6369	0.5848
1	0.4984	0.4426	0.6068	0.4831	0.4817	0.5161	0.5092	0.519	0.5753	0.4683	
2	0.7308	0.6068	0.46	0.6409	0.6238	0.7639	0.639	0.4677	0.5284	0.6542	
3	0.6072	0.4831	0.6409	0.3048	0.4197	0.6225	0.4949	0.5798	0.692	0.3451	
4	0.5431	0.4817	0.6238	0.4197	0.439	0.5585	0.5277	0.5483	0.6209	0.3652	
5	0.3721	0.5161	0.7639	0.6225	0.5585	0.43	0.6011	0.6822	0.7201	0.6021	
6	0.5503	0.5092	0.639	0.4949	0.5277	0.6011	0.6155	0.5411	0.6512	0.499	
7	0.6004	0.519	0.4677	0.5798	0.5483	0.6822	0.5411	0.7201	0.3847	0.5782	
8	0.6369	0.5753	0.5284	0.692	0.6309	0.7201	0.6512	0.3847	0.7201	0.6783	
9	0.5848	0.4683	0.6542	0.2841	0.3652	0.6021	0.499	0.5782	0.6783	0.1152	

(b) Matrice des sommes pondérées des distances (SPD) des descriptions

Id	Mots
0	'électricité', '#tempêtes', 'pluie', 'tropical', 'dimanche', 'violent', 'alerte', 'inondation', 'jusqu', 'delta'
1	'temps', 'autre', 'nouveau', 'grand', 'place', 'rester', 'octobre', 'samedi', '#france', 'matchedHour'
2	'depuis', 'sinistré', '#tempêtealex', 'solidarité', 'parc', 'pont', 'trois', 'vallée', 'soutien', 'virant'
3	'après', 'falloir', 'hier', 'voir', 'très', 'alors', 'voici', 'parler', 'parce', 'plage'
4	'sans', 'jamais', 'faire', 'cette', 'jours', 'avis', 'aide', 'aussi', 'raison', 'face'
5	'vent', 'matchedNumber', 'côte', 'ouragan', 'heure', 'dépression', 'vers', 'entre', 'kilomètre', '#morbihan'
6	'demain', 'avant', 'image', 'cause', 'france', 'après', 'fermé', 'dégât', 'premier', 'jour'
7	'point', 'secours', 'nice', 'annoncer', 'crue', 'maison', 'région', 'disparu', 'emporté', 'épisode'
8	'emmanuel', 'macron', 'corps', 'rouge', 'moins', '#nice06', 'disparus', 'plusieurs', 'toujours', 'village'
9	'seul', 'comment', 'gros', 'calme', 'chez', 'fois', 'quoi', 'rien', 'ciel', 'dire'

(c) Descriptions des clusters

FIGURE 3 – Évaluation des descriptions pour la sélection de candidat hybride, avec 118 candidats dont 72 non dupliqués.

L'utilisation de l'algorithme FPF pour orienter la sélection des candidats permet de mieux couvrir la sémantique représentée par les vecteurs des mots des tweets d'un cluster. Ce gain de représentativité s'accompagne d'un autre effet. En choisissant des candidats à la périphérie de la représentation sémantique d'un cluster, on les rapproche naturellement des autres clusters même si cet effet est compensé par l'imposition d'un seuil de DF-IDF.

Nous pouvons le voir notamment dans la figure 3a. La meilleure valeur d'une ligne est sur la diagonale dans 7 cas sur 10. La matrice apparaît fortement colorée. Les lignes 1, 2, 4, 6, et 7 apparaissent presque uniformes. Au regard de ce critère, la sélection hybride semble proposer des candidats adéquats pour décrire les clusters mais pas pour différencier les descriptions. Nous pouvons remarquer toutefois que la valeur des diagonales des lignes 3, 5 et 8 se distinguent nettement des autres, ce qui semble indiquer que certains clusters sont plus faciles à différencier dans leur description. De manière générale les scores d'IR sont plus bas que pour les autres sélection, ce qui s'explique aussi par l'utilisation de FPF dans la sélection des candidats.

Cette tendance s'observe aussi dans la figure 3b. La meilleure valeur de la ligne se trouve sur la diagonale dans 8 cas sur 10 mais n'est nettement meilleure que dans 2 cas sur 8 (lignes 5 et 9). Comme pour l'IR, la SPD les descriptions obtenues avec la sélection hybride pourraient indiquer que les candidats de cette sélection sont adéquats pour décrire les clusters mais pas pour les isoler.

Dans la figure 3c, on peut voir qu'un peu plus d'un tiers des descripteurs initiaux de la sélection hybride sont dupliqués. C'est significativement moins qu'avec les autres méthodes. Le sujet des descriptions 2, 5, 6, 7 et 8 semble s'identifier assez facilement.

4.2.4 Comparaison de l'ensemble des méthodes

Nous comparons ici les trois approches présentées sur l'IR, la SPD et la couverture des données.

La figure 4 permet de voir qu'elles ont des performances différentes en terme d'IR mais similaires en terme de SPD. Nous voyons également apparaître que la sélection basée uniquement sur la fréquence est moins bonne que les deux autres au regard de ces critères. La sélection hybride et celle des Tops DF-IDF affichent des résultats similaires, l'hybride étant légèrement meilleure que les Tops DF-IDF.

En observant ces résultats, il apparaît que la fréquence seule ne suffit pas à identifier des descripteurs qui permettent de construire des descriptions appropriées pour leur cluster. Lorsqu'on s'appuie sur ce seul critère, on peut en effet sélectionner des mots très fréquents dans un cluster avec un score DF-IDF de fait plus bas que dans d'autres clusters. Ainsi, ce descripteur sera plutôt utilisé dans un autre cluster où son score sera plus élevé. De cette manière, l'ensemble des candidats présentés au modèle d'optimisation peut dans un cas non favorable compter une minorité de mots importants pour décrire les clusters. Les descriptions qui sont alors construites seront peu adéquates pour leur ensemble de tweets.

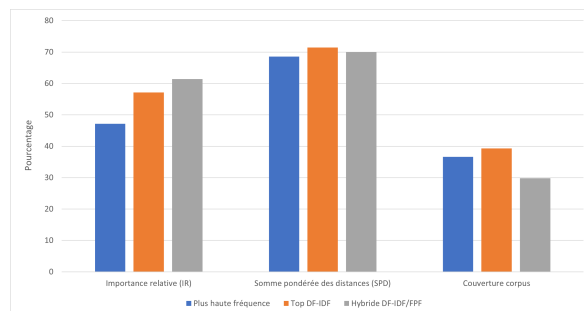


FIGURE 4 – Proportion des cas où le meilleur score d'une description est le meilleur pour le cluster qu'elle décrit et couverture du corpus par les descriptions

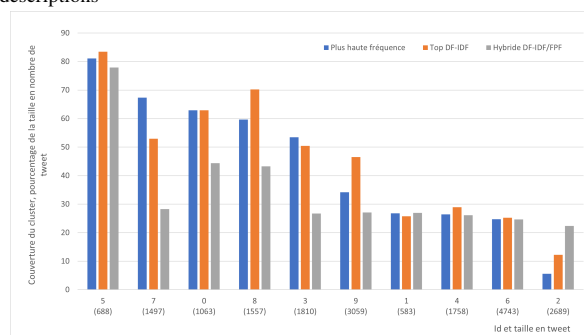


FIGURE 5 – Couverture des clusters par sélection de candidat

Nous pouvons voir, dans la figure 4 la couverture du corpus selon la sélection de candidat. Si la sélection par plus haute fréquence et des Tops DF-IDF ont des résultats similaires, respectivement 36,6% et 39,3%, la sélection hybride se montre moins performante au regard de critère, avec 29,9% de tweets du corpus couverts. En s'appuyant, même en partie, sur l'algorithme Farthest Point First pour identifier des candidats pour chaque cluster, on sélectionne des mots éloignés les uns des autres dans l'espace de plongement. Par construction les mots plus proches apparaissent souvent ensemble. De ce fait, les candidats sélectionnés par l'approche hybride apparaissent dans moins de tweets et la couverture du corpus par les descriptions obtenues en pâtit. Nous pouvons constater la même tendance, s'expliquant de la même manière, dans les couvertures par cluster de la figure 5. Nous pouvons aussi y constater que même si de manière générale l'approche hybride sélectionne des descripteurs moins couvrants, elle est globalement plus stable et s'en sort au moins aussi bien sinon mieux avec les clusters difficiles à couvrir, comme les 4, 6 et 2.

La couverture représentée dans les figures 4 et 5 n'est pas la même que celle imposée par les contraintes du modèle, voir section 3.2.3. Ici c'est la couverture réelle des descriptions qui est présentée, soit la proportion de tweets du corpus ou du cluster couverte. Ces deux mesures peuvent concorder, pour la couverture de l'ensemble du corpus par exemple (figure 4), mais aussi être complètement déconnectées. Nous pouvons le voir dans la figure 5 pour le cluster 2, dont seulement 5,6% des tweets sont couverts alors que c'est le cas pour 36,3% des apparitions des candidats dans ce cluster et

que la contrainte est respectée.

Les résultats des différentes sélections montrent des tendances communes. D'abord, le manque de variété au sein des candidats sélectionnés induit de moins bons résultats. Par ailleurs si une description peut poser problème à une approche et pas aux autres, certains clusters semblent de manière générale plus faciles ou plus difficiles à décrire. Par exemple, le cluster 5 est systématiquement bien identifié, tant en IR qu'en SPD. À l'inverse, le 4 pose toujours problème. Il apparaît donc qu'indépendamment des descripteurs potentiels choisis, certains clusters soient moins favorables à la construction de descriptions avec de bonnes IR et SPD. Le même problème s'observe avec la couverture, avec des clusters plus difficiles à couvrir que d'autres comme on peut le voir sur la figure 5.

Construire une description pour un cluster de tweets avec notre méthode comprend donc deux axes de difficulté : représenter ce sur quoi porte le contenu du cluster pour l'interpréter et couvrir ses tweets dans une proportion significative pour que l'interprétation soit robuste. Les résultats en IR et SPD permettent de se faire une idée de la difficulté sur le premier axe, la mesure de la couverture d'un cluster par sa description permettant de le faire pour l'autre. Intuitivement on pourrait se dire que ces deux difficultés sont liées, qu'il faut couvrir un cluster pour en capturer le sens. Les évaluations des clusters 5 et 6 semblent pointer dans cette direction. Pourtant le cluster 2, avec de bonnes IR et SPD apparaît particulièrement difficile à couvrir. Il n'est donc pas possible d'établir l'existence ou l'absence d'un lien entre ces deux axes.

Plusieurs questions se posent sur l'interprétabilité humaine des clusters et de leurs descriptions. Par exemple, les scores d'IR et de SPD de la description du cluster 5 sont toujours bons et quelle que soit la méthode. On y retrouve toujours certains mots : 'vent', '#morbihan', 'côte', 'kilomètre' et 'heure'. Dans les sélections par les plus hautes fréquences et Tops DF-IDF, on y trouve 'rafale'. Dans la sélection hybride, 'rafale' n'est pas présent mais 'ouragan' apparaît. À travers la construction des clusters, l'hypothèse distributionnelle semble se vérifier et permettre d'interpréter facilement ce genre de cluster. Les vecteurs de tweets contenant 'vent' et '#morbihan' semblent proches dans l'espace de plongement, suffisamment pour que ces mots soient systématiquement choisis pour décrire leur cluster.

En observant les descriptions des clusters qui sont consistantes en scores et en contenu avec les différentes sélections, nous pouvons construire une image de la tempête Alex. Au moins une des nuits lors de la tempête a été marquée par des vents très violents sur les côtes bretonnes dans le Morbihan (cluster 5). Par ailleurs les crues dans les alpes maritimes, en particulier dans la région de Nice, ont causé des dégâts et provoqué des disparitions à l'origine d'une mobilisation notable des pompiers (clusters 7 et 8).

Pour le cluster 4 en revanche, les choses sont moins claires. Ses trois descriptions présentées ont des scores indiquant qu'elles pourraient convenir à d'autres clusters et aucun objet particulier ne s'en dégage. Elles contiennent en ma-

jorité des mots outils : 'fait', 'après' ou encore 'comme'. Outre les pistes d'explications déjà évoquées, cela pourrait venir du fait que certains descripteurs portant une sémantique moins claire et présentant des ambiguïtés (l'ambiguïté n'étant par ailleurs pas traitée dans l'espace de plongement) ne peuvent décrire un sous-événement en particulier, tel qu'on le constate dans le cluster 5. En effet, les mots outils peuvent apparaître dans n'importe quel contexte, indépendamment du sujet abordé. Par construction, les vecteurs qui contiennent plusieurs de ces mots sont donc naturellement plus éloignés des zones de l'espace sémantique correspondant à un sujet particulier.

5 Conclusion et perspectives

Nous présentons ici deux méthodes de sélection de mots depuis des ensembles de tweets pour décrire au mieux ces ensembles, par les Tops DF-IDF et par la méthode hybride. Si la validation de ces résultats par l'expérimentation sur d'autres jeux de données et l'utilisation d'autres métriques d'évaluation sont nécessaires pour les confirmer, nos premières conclusions sont prometteuses. En effet, dans les cas où les clusters semblent concerner un sous-événement en particulier qui se distingue du reste des données, nos descriptions permettent de le représenter avec les mots sélectionnés. Ces derniers sont sémantiquement cohérents, entre eux et par rapport au cluster décrit. Nous envisageons plusieurs pistes d'amélioration. Tout d'abord l'utilisation de n-grammes plutôt que d'uni-grammes pour les composants des descriptions, notamment pour se défaire de l'ambiguïté de mots comme "faire". Nous souhaitons également mettre au point des contraintes sur les distances entre les mots et les documents dans l'espace de plongement sémantique. Leur élaboration devra passer par une étude poussée des liens entre les vecteurs de mots et de documents dans cet espace. Par ailleurs, nous souhaitons nous appuyer sur les descriptions pour identifier les clusters qui pourraient être re-découpés et ceux qui pourraient être fusionnés au regard de leur interprétation. C'est un axe qui devra passer par l'étude dans le détail de l'espace de plongement utilisé notamment pour comprendre comment identifier un cluster impossible à interpréter, comme un cluster de mot outil, d'un autre.

Références

- [1] J. Ansah, L. Liu, W. Kang, J. Liu, and J. Li. Leveraging burst in twitter network communities for event detection. *World Wide Web*, 2020.
- [2] H. Becker, M. Naaman, and L. Gravano. Beyond Trending Topics : Real-World Event Identification on Twitter. *ICWSM*, 2011.
- [3] T. Dao, C. Kuo, S. S. Ravi, C. Vrain, and I. Davidson. Descriptive clustering : ILP and CP formulations with applications. In *IJCAI 2018*, pages 1263–1269, 2018.
- [4] I. Davidson, A. Gourru, and S. Ravi. The Cluster Description Problem - Complexity Results, Formulations and Approximations. *NeurIPS*, 2018.

- [5] De Boom, C. and Van Canneyt, S. and Dhoedt, B. Semantics-driven event clustering in Twitter feeds. In *Proceedings of the 5th Workshop on Making Sense of Microposts*, 2015.
- [6] Y. Gao, W. Zhao, and S. Eger. SUPERT : Towards New Frontiers in Unsupervised Evaluation Metrics for Multi-Document Summarization. *ACL*, 2020.
- [7] T. Gonzalez. Clustering to minimize the maximum intercluster distance. *Theoretical Computer Science*, 1985.
- [8] Z. S. Harris. Distributional Structure. *WORD*, 1954.
- [9] M. Hasan, M. A. Orgun, and R. Schwitter. Twitter-News+ : A Framework for Real Time Event Detection from the Twitter Data Stream. In *Social Informatics*, Cham, 2016.
- [10] Q. V. Le and T. Mikolov. Distributed representations of sentences and documents. *CoRR*, 2014.
- [11] C. Li, A. Sun, and A. Datta. Twevent : Segment-based event detection from tweets. In *CIKM*, 2012.
- [12] C.-Y. Lin and F. Och. Looking for a few good metrics : Rouge and its evaluation. In *Ntcir workshop*, 2004.
- [13] A. Ouali, S. Loudni, Y. Lebbah, P. Boizumault, A. Zimmermann, and L. Loukil. Efficiently Finding Conceptual Clustering Models with Integer Linear Programming. In *IJCAI*, 2016.
- [14] E. Sherkat, J. Velcin, and E. E. Milios. Fast and Simple Deterministic Seeding of KMeans for Text Document Clustering. In *Experimental IR Meets Multilinguality, Multimodality, and Interaction*. 2018.
- [15] T. Tang and G. Hu. Detecting and Tracking Significant Events for Individuals on Twitter by Monitoring the Evolution of Twitter Followership Networks. *Information*, 2020.
- [16] B. C. Wallace, S. Saha, F. Soboczanski, and I. J. Marshall. Generating (Factual?) Narrative Summaries of RCTs : Experiments with Neural Multi-Document Summarization. *AMIA Summits on Translational Science Proceedings*, 2021.
- [17] J. Weng and B.-S. Lee. Event detection in twitter. In *ICWSM*, 2011.
- [18] Y. You, G. Huang, J. Cao, E. Chen, J. He, Y. Zhang, and L. Hu. GEAM : A General and Event-Related Aspects Model for Twitter Event Detection. In *WISE*, 2013.
- [19] D. Zhou, L. Chen, and Y. He. An unsupervised framework of exploring events on twitter : Filtering, extraction and categorization. In *AAAI*, 2015.

How COVID-19 is Changing Our Language: Detecting Semantic Shift in Twitter Word Embeddings

Yanzhu Guo¹, Christos Xypolopoulos^{1,2}, Michalis Vazirgiannis^{1,3}

¹ Ecole Polytechnique, LIX

² National Technical University of Athens (NTUA)

³ Athens University of Economics and Business (AUEB)

yanzhu.guo@polytechnique.edu

Résumé

Les mots sont des objets malléables, influencés par les événements reflétés dans des textes écrits. Située dans le contexte de COVID-19, notre recherche vise à détecter le changement sémantique dans le langage des médias sociaux provoqué par la crise sanitaire. À l'aide des données de grande taille liées au COVID-19 extraites de Twitter, nous entraînons des modèles de langue pour différentes périodes de temps avant et pendant l'épidémie. Nous comparons ces modèles de langue entraînés avec une série de mesures de similarité afin d'observer la variation sémantique. À l'aide d'une liste de mots saillants provenant de la mise à jour spéciale COVID-19 du dictionnaire anglais Oxford, nous déterminons l'approche de détection la plus adaptée à notre corpus Twitter. En nous appuyant sur cette approche, nous réalisons des études de cas sur un ensemble de mots sélectionnés par détection de sujets et visualisons l'évolution diachronique. Enfin, nous effectuons une analyse exploratoire sur des tweets français et obtenons des perspectives intéressantes sur des enjeux sociaux en France.

Mots-clés

Changement Sémantique, Twitter, COVID-19, Traitement du Langage Naturel.

Abstract

Words are malleable objects, influenced by events reflected in written texts. Situated in the global outbreak of COVID-19, our research aims at detecting semantic shift in social media language triggered by the health crisis. With COVID-19 related big data extracted from Twitter, we train word embeddings models for different time periods before and after the outbreak. We compare these trained embeddings with a series of different dissimilarity metrics in order to observe variation in semantics. Using a list of salient words from the COVID-19 special update of the Oxford English dictionary, we determine the detection approach most adapted to our Twitter corpus at hand. Drawing upon this approach, we carry out case studies on a set of words selected by topic detection and visualize the diachronic evolution. Finally, we perform an exploratory analysis on French tweets and gain interesting insights on societal issues in France.

Keywords

Semantic Shift, Twitter, COVID-19, Natural Language Processing.

1 Introduction

Les mots s'adaptent à l'environnement, leurs sens et leurs significations font l'objet de variations constantes, c'est-à-dire de glissements sémantiques. Comprendre comment ils changent à travers différents contextes et périodes de temps est crucial pour révéler le rôle du langage pendant l'évolution de la société.

Les glissements sémantiques peuvent être la conséquence de changements à long terme en raison de circonstances politiques, sociales, culturelles ou économiques. Un exemple est le mot "gay", dont le sens a changé au cours du 20e siècle, passant de la signification de "gai" et "joyeux" à celle d'"homosexualité" [25]. L'influence d'événements historiques ponctuels est tout aussi importante et peut entraîner des changements spectaculaires en un court laps de temps. Par exemple, l'événement tragique des attentats du 11 septembre 2001 a considérablement modifié l'interprétation générale du mot "terrorisme" [24].

La pandémie de COVID-19 a amené l'ensemble du discours mondial à se concentrer sur un seul sujet, cela s'était rarement produit auparavant. Les gens du monde entier ont dû adapter leur vocabulaire pour pouvoir parler de l'époque extraordinaire que nous vivons tous. Le glissement sémantique qui en résulte est illustré par le contenu généré par les utilisateurs sur les médias sociaux. En fait, l'un des facteurs contribuant à l'ampleur sans précédent du glissement sémantique déclenché par le COVID-19 et à son extraordinaire vitesse de propagation est le fait que la société humaine est connectée numériquement comme jamais auparavant. Les gens utilisent les médias sociaux non seulement pour rapporter les dernières nouvelles, mais aussi pour exprimer leurs opinions et leurs sentiments sur des événements du monde réel. Les utilisateurs montrent un intérêt particulier pour les situations d'urgence telles que COVID-19. Les données collectées sur des plateformes telles que Twitter peuvent constituer des ressources précieuses pour étudier l'effet de la pandémie sur le langage humain. Leur diversité et leur

compréhensibilité sont garanties par la nature inclusive des médias sociaux.

Dans cet article, nous explorons la stabilité sémantique des mots en calculant comment leurs significations, représentées par les embeddings de mots respectifs, ont été influencées par la pandémie. Nous collectons d'abord un corpus de tweets en anglais liés à COVID-19 ainsi qu'un corpus de référence de tweets hétérogènes en anglais postés entre janvier 2019 et décembre 2019. Nous générons des embeddings statiques ainsi que contextuels pour les deux corpus et calculons la stabilité des mots par deux métriques différentes basées sur les embeddings. En utilisant une liste de mots sélectionnés à partir de la mise à jour spéciales COVID-19 du dictionnaire anglais Oxford, nous évaluons les différentes approches et sélectionnons la plus efficace pour effectuer une analyse plus approfondie. Enfin, nous collectons un corpus de tweets en français liés à COVID-19 et réalisons une étude de cas pour la langue française.

2 Travaux antérieurs

Dans cette section, nous examinons la littérature connexe qui se divise en deux catégories : la détection du glissement sémantique et l'analyse de Twitter liée à COVID-19.

2.1 Détection du glissement sémantique

Alors que les premières études computationnelles du changement sémantique ont débuté par l'analyse des fréquences brutes des mots [11, 13] et des cooccurrences [10], des études plus récentes ont principalement fait appel à des embeddings neuronaux de mots. Les approches classiques utilisent généralement des embeddings de mots statiques et peuvent être classées en deux groupes : l'utilisation de mesures de stabilité basées sur la similarité cosinus avec alignement ou l'utilisation de mesures de stabilité basées sur le voisinage sans alignement.

Les approches nécessitant un alignement sont généralement séparées en deux étapes : d'abord générer des modèles d'intégration de mots distincts pour chaque corpus indépendamment, puis utiliser des approches mathématiques pour les aligner dans le même espace latent sous-jacent. La qualité de l'alignement est déterminante pour les résultats des comparaisons. Kulkarni et al. [15] ont calculé la transformation linéaire optimale entre l'espace d'embeddings de base et l'espace d'embeddings cible en résolvant un problème de moindres carrés de k voisins les plus proches. Hamilton et al. [9] ont également utilisé des transformations linéaires pour l'alignement mais n'ont considéré que les transformations orthogonales. Zhang et al. [29] ont réalisé l'alignement de manière similaire, en ajoutant l'utilisation de mots d'ancrage, dont la signification est censée rester stable entre les deux espaces d'embeddings.

Les approches basées sur le voisinage reposent sur l'hypothèse que les mots ayant des significations significativement divergentes d'un corpus à l'autre sont censés avoir un contexte différent et donc un ensemble de voisins différent dans chaque corpus. Hamilton et al. [8] mesurent les changements des plus proches voisins d'un mot dans le but de capturer les changements drastiques dans le sens prin-

cipal. Azarbyonad et al. [1] construisent un graphe pour chaque corpus avec les mots comme nœuds et les similarités entre eux comme arêtes. Ils calculent les similarités entre les voisinages d'un même mot dans différents graphes par des mesures de similarité basées sur le graphe. Gonen et al. [7] analysent directement l'espace de vocabulaire partagé des mots dans différents corpus, en considérant simplement les k plus proches voisins dans chaque corpus et en calculant la taille de l'intersection des deux listes. Cette méthode s'est avérée simple, interprétable et robuste.

Le succès récent des représentations de mots contextualisées telles que BERT [4] et ELMo [23] a ouvert de nouvelles portes aux chercheurs en détection de changements sémantiques. Hu et al. [12] construisent un embedding de sens distingué pour chaque sens d'un mot en s'appuyant sur des embeddings contextualisés profonds. En faisant correspondre les embeddings de chaque tranche de temps aux embeddings de sens construits, ils sont capables de suivre l'évolution du sens de chaque mot cible dans le temps. Giulianelli et al. [6] considèrent que chaque apparition d'un mot représente un sens différent. Ils déterminent comment les sens des mots varient dans le temps en affinant de manière incrémentée [14] sur des tranches de temps successives, puis en effectuant un regroupement K-means. Montariol et al. [20] améliorent l'évolutivité de l'approche précédente en fusionnant dynamiquement les clusters pendant la génération d'embeddings des mots. Martinc et al. [18] représentent directement le sens global d'un mot dans un corpus donné comme la moyenne de ses embeddings contextualisés et étudient leur évolution. Malgré la popularité croissante de l'utilisation d'embeddings contextualisés dans cette branche de la recherche, la récente tâche SemEval sur la détection non supervisée des changements sémantiques lexicaux [26] a montré que les méthodes d'embeddings statiques obtenaient les meilleures performances moyennes dans tous les corpus.

2.2 Analyse de Twitter liée à COVID-19

De nombreux articles s'intéressent à la réaction des utilisateurs de Twitter et d'autres médias sociaux face à cette crise sanitaire. Chen et al. [2] ont publié le premier dataset public de données Twitter sur le COVID-19. Cinelli et al. [3] ont traité de la diffusion de fausses informations concernant le COVID-19 sur Twitter. Ziems et al. [30] ont révélé l'origine et le mode de diffusion des comportements racistes en ligne pendant l'épidémie de COVID-19. Lopez et al. [17] ont permis de comprendre les réactions des gens aux politiques de COVID-19 en exploitant un ensemble de données Twitter multilingues. Müller et al. [21] ont publié COVID-Twitter-BERT, un modèle basé sur un transformateur pré-entraîné, avec un large éventail d'applications dans les tâches de NLP liées à COVID-19.

En ce qui concerne les glissements sémantiques, Tahmasbi et al. [28] ont entraîné des modèles Word2Vec hebdomadaires à partir de données Twitter collectées après l'épidémie et ont observé des glissements vers l'apparition d'injures plus sinophobes. Cependant, à notre connaissance, il n'y a pas eu d'étude systématique ou complète sur les glissements

sémantiques du langage des médias sociaux induits par le COVID-19.

3 Jeux de données et évaluation

Afin de suivre l'évolution des usages des mots, nous collectons deux corpus à grande échelle de tweets en anglais : un corpus de référence pré-COVID-19 avec les tweets publiés avant l'épidémie ainsi qu'un corpus lié à COVID-19 avec les tweets mentionnant explicitement la pandémie.

Corpus de référence pré-COVID-19 Pour construire un corpus de référence qui remonte à l'époque pré-COVID-19, nous téléchargeons le flux Twitter général saisi par l'équipe d'archivage : ¹, contenant des tweets diffusés de janvier 2019 à décembre 2019. Après avoir sélectionné uniquement les tweets en anglais et supprimé les tweets dupliqués avec l'outil open-source runiq ², nous obtenons un corpus de 127M tweets uniques.

Corpus lié à COVID-19 Les tweets utilisés pour construire notre corpus relatif à COVID-19 datent de mars 2020 à août 2020. Dans ce cas, nos filtres se concentrent sur les tweets qui incluent les hashtags "covid19" et "coronavirus". Grâce à l'API publique de streaming de Twitter, nous avons extrait les tweets en anglais marqués par les deux hashtags ci-dessus. Après les mêmes étapes de filtrage linguistique et de déduplication que pour le corpus précédent, nous obtenons un corpus de 53M de tweets uniques.

Évaluation Afin d'évaluer nos différentes approches de détection du glissement sémantique, nous sélectionnons une liste de mots dont le glissement sémantique lié à COVID-19 a été confirmé par des lexicographes experts du dictionnaire anglais Oxford. Le dictionnaire anglais Oxford ³ a récemment publié plusieurs mises à jour spéciales COVID-19 dédiées à la langue de COVID-19. Nous sélectionnons tous les mots avec de nouveaux sens ajoutés et tous les mots (sauf les mots d'arrêt) qui appartiennent à de nouvelles entrées de phrases composées. Notre étude se concentre sur le glissement sémantique des mots existants et laisse le sujet de l'innovation lexicale et des néologismes pour un travail futur. La liste complète des 37 mots se trouve en annexe A.

Comme démontré dans des travaux antérieurs [8] [7], les méthodes automatiques de détection des glissements sémantiques ne peuvent que fournir une liste candidate de mots susceptibles de subir des glissements sémantiques mais ne peuvent en aucun cas rendre compte des garanties. Les cas d'utilisation appropriés de notre cadre de détection consistent notamment à offrir aux lexicographes un premier aperçu avec une liste de mots sur lesquels mener une enquête plus approfondie, ainsi qu'à sensibiliser les fonctionnaires aux récits sociaux en cours afin de communiquer avec le grand public de manière adaptée et efficace. Par conséquent, nous choisissons le score de recall de la détection des glissements sémantiques comme métrique d'évaluation, conformément à la motivation sous-jacente des scénarios d'utilisation.

1. <https://archive.org/details/twitterstream>

2. <https://github.com/whitfin/runiq>

3. <https://www.oed.com>

4 Méthodologie

Comme l'a montré la tâche SemEval sur la détection non supervisée des changements sémantiques lexicaux [26], aucun modèle d'embeddings ni aucune métrique de stabilité ne permet d'obtenir des résultats optimaux pour tous les corpus. Les performances des différentes approches dépendent fortement des caractéristiques spécifiques des corpus analysés. Par conséquent, nous combinons les modèles d'intégration de mots les plus avancés avec les métriques de stabilité les plus largement appliquées, dans l'espoir de découvrir l'approche la plus adaptée à nos corpus.

4.1 Modèles d'embeddings

Nous expérimentons des modèles d'embeddings de mots statiques et contextualisés, en prenant word2vec [19] et BERT [4] comme représentants respectifs.

Word2vec Nous utilisons l'implémentation open source de Word2Vec dans le package gensim ⁴ pour générer des embeddings de mots. Nous supprimons les mots qui apparaissent moins de 10 fois et appliquons l'architecture Skipgram, en choisissant une taille de fenêtre de 4 et une dimensionnalité de 300. Nous entraînons deux modèles word2vec distincts, indépendamment l'un de l'autre, pour le corpus de référence pré-COVID-19 et le corpus connexe COVID-19. Le pipeline de prétraitement des tweets est détaillé dans l'annexe B.

BERT Pour la modélisation du langage avec BERT, nous utilisons BERTweet [22], le premier modèle de langage à grande échelle pré-entraîné pour les Tweets anglais. BERTweet est entraîné sur l'architecture RoBERTa [16], en utilisant la même configuration de modèle que BERT-base. Le corpus original utilisé pour le pré-entraînement de BERTweet est constitué de 850M de tweets anglais, contenant 845M de tweets diffusés de janvier 2012 à août 2019 et 5M de tweets liés à la pandémie COVID-19. Nous utilisons le *bertweet-covid19-base-uncased* ⁵ version du modèle disponible sur Hugging Face. Cette version est le résultat d'un pré-entraînement supplémentaire du modèle original sur un corpus de 23M de Tweets anglais COVID-19 pour 40 époques. Ce modèle répond bien à notre objectif car il a été entraîné de manière extensive sur les tweets pré-COVID-19 et les tweets liés à COVID-19, ce qui lui permet d'incorporer des sens de mots contextualisés pour les périodes pré- et post-pandémique.

4.2 Mesures de stabilité

Nous étudions à la fois les mesures de stabilité basées sur la similarité cosinus et celles basées sur le voisinage, qui s'avèrent plus performantes dans différents contextes : [8].

4.2.1 Mesure basé su la similarité cosinus

La similarité en cosinus est une méthode populaire en NLP pour estimer la similarité de deux vecteurs de mots. Cependant, pour les modèles **word2vec**, le problème de l'alignement est un élément clé de la comparaison d'embeddings de mots indépendants formés sur des corpus différents. En

4. <https://radimrehurek.com/gensim/>

5. <https://huggingface.co/vinai/bertweet-covid19-base-uncased>

	Word2vecCos	Word2vecNN	BertCos	BertNN
recall	0.78	0.73	0.19	0.22

TABLE 1 – Valeurs de recall de la détection du glissement sémantique pour les mots cibles.

Approche	Les 5 premiers mots détectés
Word2vecCos	cerb, vtm, ggd, pums, adria
Word2vecNN	redzone, cerb, wha, corona, ceba
BertCos	unionism, fisa, bandage, carona, hoses
BertNN	drywall, flintstone, spfl, corny, trav

TABLE 2 – Les 5 premiers mots détectés par chaque approche.

raison de l’invariance rotationnelle des fonctions de coût dans l’algorithme d’entraînement word2vec, les embeddings appris séparément sont placés dans des espaces latents différents. Cela n’affecte pas les similarités cosinus par paire au sein d’un même espace d’embeddings mais entrave la comparaison d’un même mot entre deux modèles différentes. La solution la plus recherchée est la méthode d’alignement des espaces vectoriels. En suivant Hamilton et al. [9], nous utilisons des Procrustes orthogonaux [27] pour aligner les embeddings de mots appris sur les mêmes axes de coordonnées. Nous faisons l’hypothèse simplificatrice que les espaces sont équivalents sous une rotation orthogonale. Plus précisément, nous définissons $W_{precovid} \in R^{d \times |V|}$ comme la matrice d’embeddings dans le modèle de référence pré-COVID-19 et $W_{covid} \in R^{d \times |V|}$ comme la matrice d’embeddings dans le modèle lié à COVID-19. Nous alignons W_{covid} sur $W_{precovid}$ tout en préservant les similarités cosinus en optimisant :

$$R = \underset{Q}{\operatorname{argmin}}_{Q^T Q = I} \|W_{covid}Q - W_{precovid}\|_F$$

Nous résolvons ce problème d’optimisation par une application de la décomposition en valeurs singulières et obtenons la meilleure transformation rotationnelle orthogonale R entre les deux espaces d’embeddings. Après avoir projeté les embeddings du modèle COVID-19 dans l’espace du modèle de référence avec R , nous pouvons calculer en toute sécurité les similarités en cosinus de tous les mots du vocabulaire partagé. La similarité en cosinus sert de mesure de la stabilité sémantique : plus la similarité en cosinus entre les embeddings d’un mot dans les deux espaces est élevée, plus sa stabilité est grande.

Comme pour les modèles **BERT**, la représentation de différents sens sémantiques est assurée par sa nature contextuelle. Nous n’avons pas besoin d’entraîner des modèles indépendants sur des corpus différents, il n’y a donc pas de problème d’alignement. Cependant, étant donné que l’architecture BERT consiste en 12 couches d’encodeurs retournant des sorties distinctes pour chaque instance de séquences d’entrée, il n’est pas simple d’obtenir un vecteur d’embedding unique pour chaque mot donné. Avec une approche similaire à celle de Martinc et al. [18], nous concaténons tous les tweets d’un même corpus et les séparons en séquences de 128 tokens. En

introduisant ces séquences dans le modèle par des batches de 32 séquences, nous générons des embeddings de séquences en additionnant les quatre dernières couches de sortie de l’encodeur. Nous séparons ensuite les embeddings de séquence en 128 sous-parties, chacune correspondant à l’un des 128 tokens de la séquence d’entrée. Maintenant, chaque token a un vecteur d’embedding différent pour chaque instance de contexte dans laquelle il est apparu. Pour chaque token du vocabulaire du corpus, nous prenons la moyenne de tous ses vecteurs d’intégration comme représentation globale de sa signification sémantique dans le corpus donné.

4.2.2 Mesure basé sur le voisinage

Plutôt que de tenter de projeter deux espaces d’embeddings dans un espace partagé, Gonen et al. [7] proposent de travailler dans l’espace de vocabulaire partagé. L’intuition sous-jacente est que les mots faisant l’objet d’un glissement sémantique sont susceptibles d’être interchangeables avec différents ensembles de mots, et donc d’avoir des voisins différents dans les deux espaces d’embeddings. Ceci donne lieu à un algorithme simple et efficace : chaque mot dans un corpus est représenté comme l’ensemble de ses k plus proches voisins (NN). La stabilité sémantique d’un mot à travers les corpus est simplement déterminée en considérant la taille de l’intersection des deux ensembles :

$$\operatorname{sim}NN^k(w) = |NN_{precovid}^k(w) \cap NN_{covid}^k(w)|$$

où $NN_i^k(w)$ est l’ensemble des k plus proches voisins du mot w dans l’espace i .

Pour calculer la liste des plus proches voisins, nous ne considérons que les mots du corpus dont la fréquence est supérieure au percentile 30% afin de filtrer le bruit présent dans les données générées par les utilisateurs de médias sociaux. À la suite de Gonen et al. [7], k est choisi à 1000 afin de retenir plus d’informations globales.

Bien que cette méthode ait été initialement proposée pour les embeddings statique de mots, elle peut être naturellement étendue aux embeddings contextuels en calculant les deux ensembles de k plus proches voisins dans le même espace d’embeddings contextuels au lieu des deux espaces d’embeddings statiques indépendants.

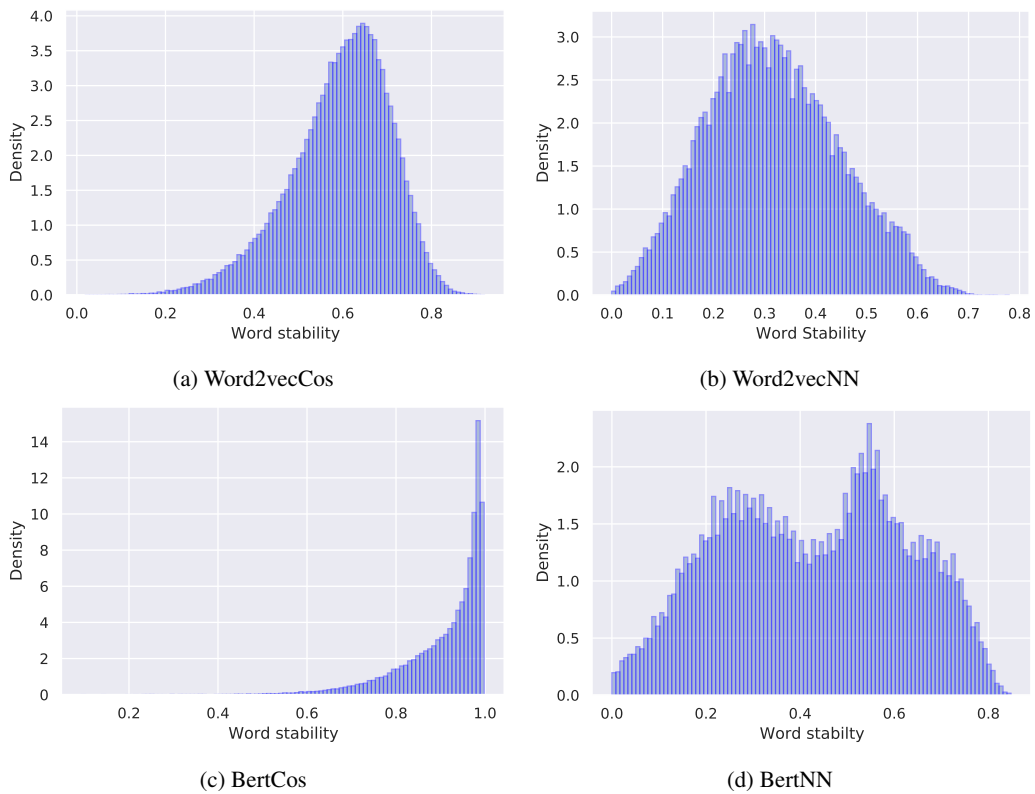


FIGURE 1 – Distribution du score de stabilité sémantique obtenu par chaque approche.

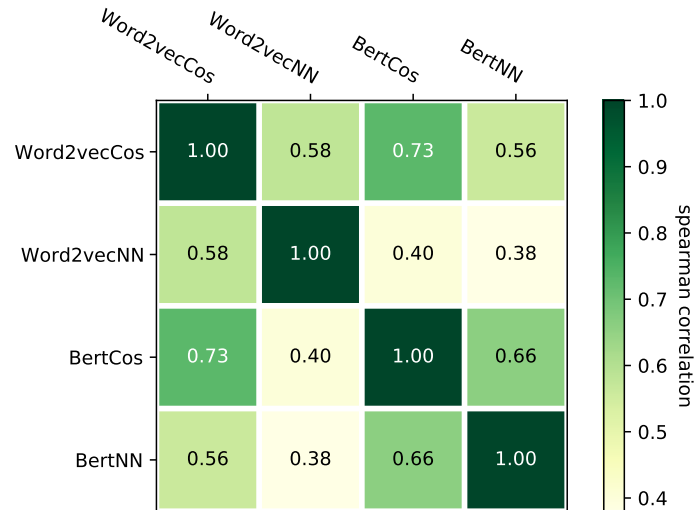


FIGURE 2 – Heatmap de la corrélation de Spearman entre les scores de stabilité des mots générés par différentes approches.

4.3 Approches étudiées

Dans la suite de ce document, nous ferons référence aux approches étudiées avec les noms suivants :

- **Word2vecCos** : combinaison d’embeddings word2vec avec la similarité cosinus.
- **Word2vecNN** : combinaison d’embeddings word2vec avec la similarité des plus proches voisins.
- **BertCos** : combinaison d’embeddings BERT avec la

similarité en cosinus.

- **BertNN** : combinaison d’embeddings BERT avec la similarité des plus proches voisins.

5 Résultats

Nous calculons les scores de stabilité d’une liste de 37 mots avec un glissement sémantique confirmé, comme mentionné dans 3. Le centre d’intérêt n’est pas les valeurs absolues des

Mot	S'éloigner	S'approcher
racism	sexism, homophobia	asians, sinophobia
hero	veteran, superman	frontliner, covidwarrior
quarantine	swineflu, flu	coranatine, corona
ai	math, data	ehealth, bloodtesting

TABLE 3 – Trajectoires des mots étudiés.

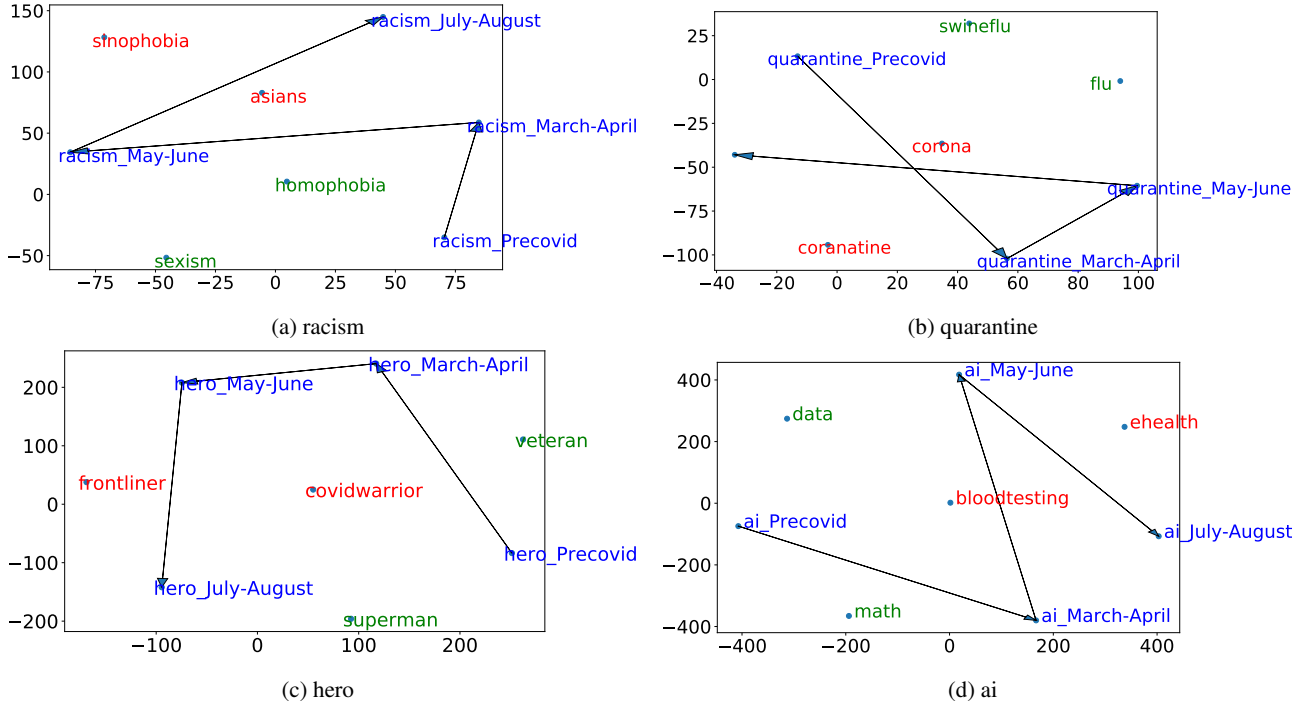


FIGURE 3 – Visualisation t-SNE des glissements sémantiques diachroniques dans les tweets anglais.

scores de stabilité eux-mêmes mais leurs positions relatives dans la distribution des scores de stabilité de l'ensemble du vocabulaire du corpus. Pour les mêmes raisons que dans 4.2.2, nous filtrons le bruit de la distribution en ignorant les mots à faible fréquence. Enfin, nous calculons le percentile du score de stabilité du mot cible dans l'ensemble de la distribution. En suivant la solution globale la plus performante de SemEval [26], nous définissons les mots dont le classement est inférieur au percentile 25% comme étant sémantiquement décalés. Les scores de recall et les distributions des scores de stabilité des mots cibles pour chaque approche sont respectivement présentés dans le tableau 1 et la figure 1. Le tableau 2 présente le top-5 des mots détectés par chaque approche. La corrélation de Spearman entre les scores générés par les différentes approches est présentée dans la figure 2.

Il est évident que les embeddings de mots statiques word2vec sont nettement plus performants que les embeddings de mots contextualisés BERT. Une explication possible est la caractéristique anisotrope des embeddings contextualisés. Par une série d'analyses géométriques, Ethayarajh [5] montre que dans toutes les couches de BERT, les

représentations de tous les mots occupent un cône étroit dans l'espace au lieu d'être distribuées partout. Cette observation se manifeste également dans la figure 1 : les valeurs de similarité en cosinus entre les embeddings BERT moyens dans différents corpus sont extrêmement concentrées vers 1, la valeur la plus élevée. Avec tous les vecteurs serrés dans un cône étroit et la similarité entre tous les mots extrêmement élevée, cette méthode est plus sensible au bruit, qui est couramment présent dans les données des médias sociaux. Nous argumentons que les embeddings word2vec sont plus adaptés à nos corpus.

Il convient de souligner que la corrélation de Spearman (illustrée dans la figure 2) entre les classements générés par word2vec et BERT est étonnamment élevée étant donné l'écart significatif entre leurs scores de recall. Nous en déduisons que BERT peut effectivement générer des classements de décalage sémantique raisonnables pour les mots sémantiquement significatifs, mais que leur sensibilité au bruit fait qu'une grande partie des mots sémantiquement insignifiants sont classés dans la partie supérieure du décalage sémantique. Ces mots bruyants peuvent être le résultat d'une orthographe non standardisée ou d'une utilisation intensive

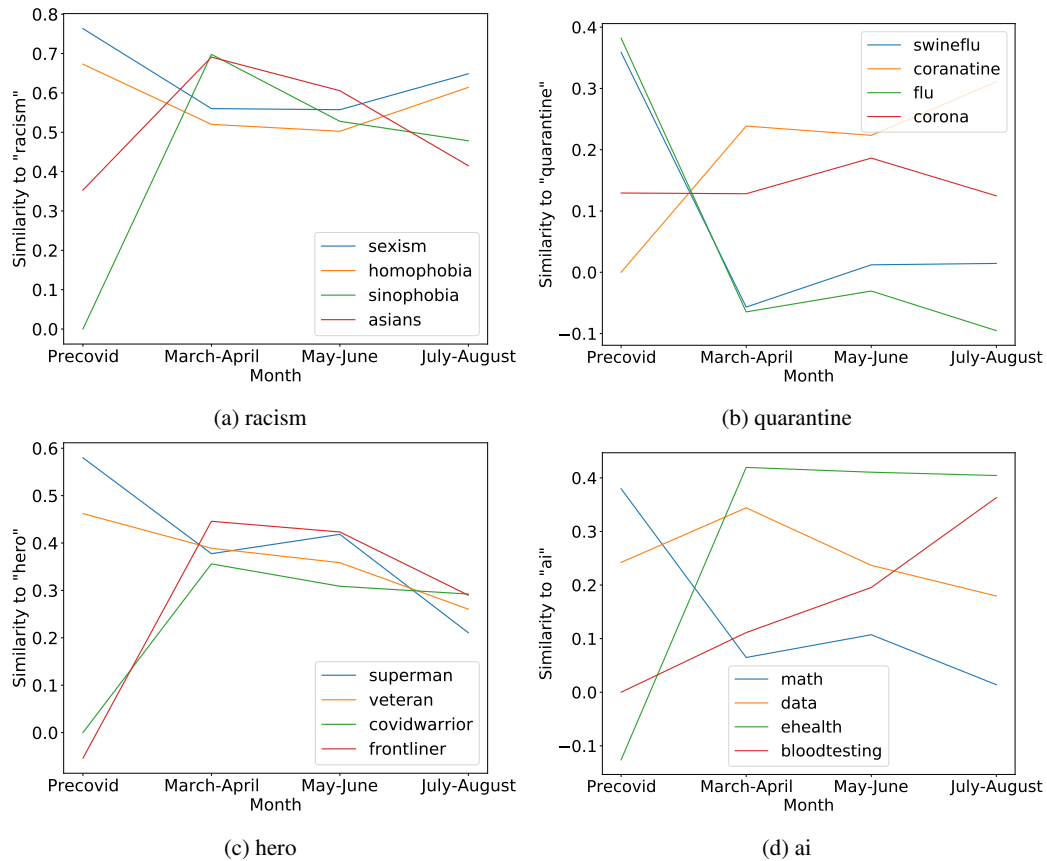


FIGURE 4 – Changements de similarité entre les mots sémantiquement décalés et leurs voisins.

d'abréviations et d'argots.

En ce qui concerne la comparaison entre les mesures de stabilité basées sur la similarité cosinus et les mesures de stabilité du plus proche voisin, leurs performances en termes de score de recall sont équivalentes, tandis que la corrélation entre leurs classements générés est relativement faible. Cela indique qu'elles sont toutes deux efficaces dans la tâche de détection des glissements sémantiques et que l'une peut compléter l'autre en détectant différents ensembles de mots.

La liste des cinq premiers mots détectés dans le tableau 2 confirme notre analyse. Les mots détectés à l'aide de word2vec sont plus significatifs que ceux détectés à l'aide de BERT, tandis que la métrique de similarité cosinus et la métrique des plus proches voisins se complètent. Les principaux mots détectés par word2vec sont principalement des acronymes d'organisations ou d'événements liés à COVID-19. Par exemple, "Cerb" est l'acronyme de "Canada Emergency Response Benefit", tandis que "adria" est le nom d'un tournoi de tennis où plusieurs joueurs ont été testés positifs. Les mots détectés avec BERT comprennent plus de bruit dû aux fautes d'orthographe et aux abréviations, comme "carona" qui est une forme erronée de "corona" et "trav" qui est le diminutif de "traveling".

6 Analyse diachronique

Nous utilisons des modèles word2vec dans ce qui suit car ils sont plus adaptés à nos corpus. Étant donné que la pandémie évolue rapidement, nous construisons des embeddings word2vec mensuels pour chacune des paires de mois *mars-avril*, *mai-juin* et *juillet-août*.

Nous choisissons un ensemble de mots clés pour analyser l'évolution diachronique (*c.-à.-d.*, comment la sémantique fluctue dans le temps). Nous n'utilisons pas directement les mots ayant les scores de stabilité les plus bas renvoyés par l'algorithme car nous voulons un ensemble de mots étroitement liés à des questions sociales d'actualité. À cette fin, nous effectuons une détection des sujets parmi les tweets liés à COVID-19, puis nous réalisons des études de cas sur les mots clés des sujets détectés. Une autre raison pour laquelle nous procédons à la détection des thèmes est que les mots présentant le plus grand changement sémantique global ne sont pas forcément ceux qui présentent les fluctuations mensuelles les plus importantes. Nous choisissons finalement quatre mots clés : "**racisme**", "**quarantaine**", "**héros**" et "**ai**".

Nous alignons les trois modèles word2vec mensuels liés à COVID-19 avec le modèle word2vec de référence en utilisant la méthode mentionnée dans 4.2.1. Nous visualisons les trajectoires des mots clés pour mieux comprendre leur évolution dans le temps. Le tableau 3 résume les glissements

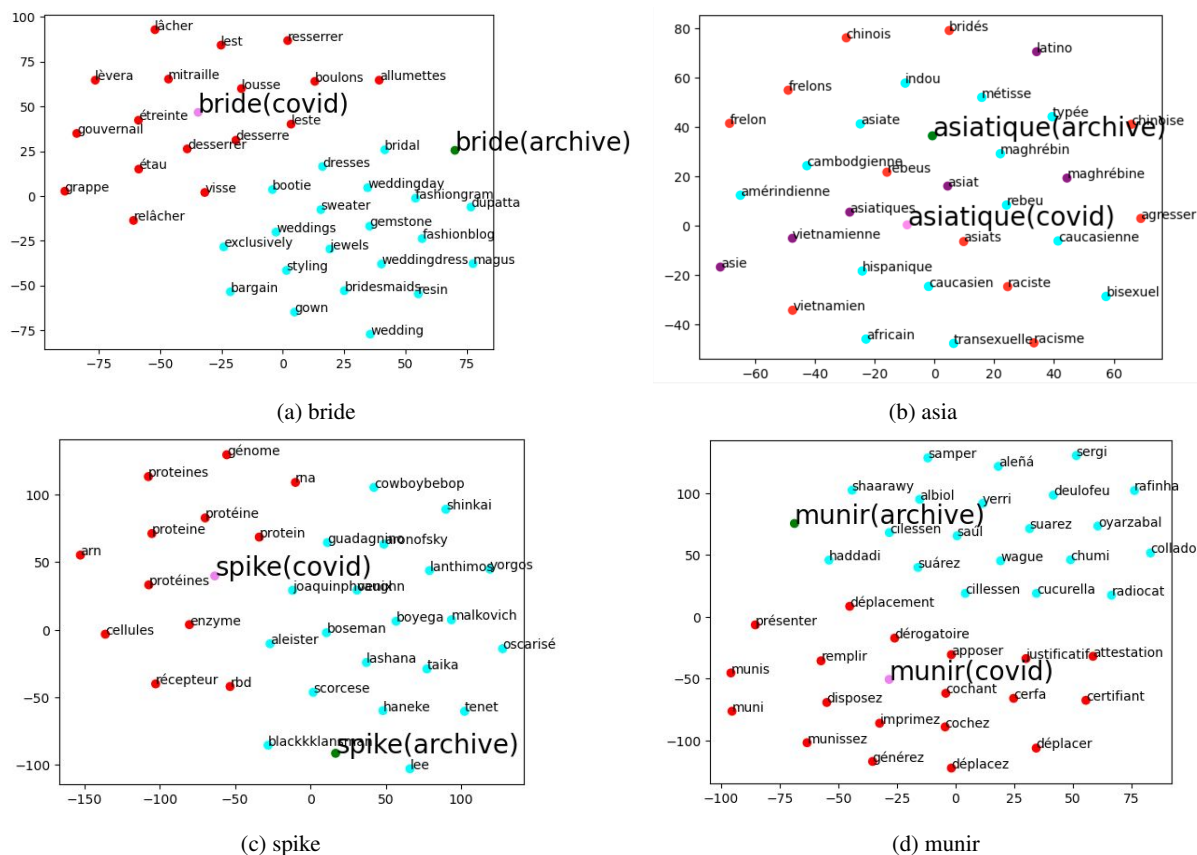


FIGURE 5 – Visualisation t-SNE des glissements sémantiques dans les tweets français.

sémantiques de l'ensemble des mots clés tandis que la figure 3 montre leurs trajectoires. Nous traçons la projection t-SNE bidimensionnelle des mots clés dans chacun des modèles alignés. Nous traçons également certains des mots environnants des mots clés dans les modèles alignés. La figure 4 illustre les changements de similarité cosinus entre le mot d'intérêt et les mots environnants.

Dans tous les cas, les illustrations de la trajectoire et la variation de la similarité cosinus démontrent que les mots d'intérêt ont changé de sens de manière significative après l'apparition de COVID-19 et montrent également une évolution sur des périodes mensuelles. Par exemple, nous voyons le mot "racisme" s'éloigner d'autres concepts généraux de discrimination et se rapprocher de mots exprimant explicitement la haine envers la communauté chinoise/asiatique. Cela coïncide avec le phénomène anti-asiatique mondial observable depuis le tout début de la pandémie. Il est intéressant de constater sur la figure 4 que la similitude entre les concepts de racisme et d'anti-Asiatique a atteint un pic en mars-avril et a commencé à diminuer légèrement en mai-juin et juillet-août, ce qui indique que les gens retrouvent lentement leur rationalité à mesure que le stade de COVID-19 progresse.

7 Étude de cas des tweets français

En suivant les mêmes procédures que pour l'anglais, nous avons construit un ensemble de données de référence pré-COVID-19 composé de 34M de tweets français et un ensemble de données lié au COVID-19 composé de 19M de tweets français. En utilisant l'approche de détection optimale décrite dans les sections précédentes, nous effectuons une analyse exploratoire pour la langue française. Dans la figure 5, nous montrons les résultats de la visualisation t-SNE pour quatre mots-clés. Nous représentons les mots voisins dans les espaces d'intégration respectifs avant et après l'épidémie. Les points bleus représentent les mots voisins de l'espace avant l'épidémie et les points rouges de l'espace après.

La figure 5 : (a) montre un exemple typique de glissement sémantique. Avant la pandémie, le mot "bride" dans les tweets signifiait principalement le mot anglais "bride", une femme qui est sur le point de se marier. Il se trouvait dans la même zone de l'espace d'embeddings que des mots tels que "weddingday", "bridesmaids", "jewels", etc. Cependant, après la pandémie, le mot est devenu plus étroitement associé aux politiques restrictives gouvernementales, ce qui signifie "contrôle renforcé". Il se rapproche de mots tels que "lâcher" et "resserrer" dans l'espace d'embeddings.

Comme le montre la figure 5 : (b), nous constatons que le mot "asiatique" était plus proche des noms d'autres groupes

ethniques avant l'épidémie, tels que "hispanique", "maghrébine" et "caucasienne". Nous faisons l'observation troublante qu'après l'épidémie, il s'est rapproché de mots à caractère violent tels que "racisme" et "agresser". Cette observation est cohérente avec la montée de l'agressivité envers les communautés asiatiques déclenchée par la pandémie de COVID-19. De tels résultats analytiques peuvent alerter le gouvernement et les décideurs politiques sur les enjeux sociaux. En effet, la société française a récemment été témoin d'un certain nombre d'incidents violents graves à l'encontre des Asiatiques, notamment la proposition largement retweetée de "poignarder tous les Asiatiques que vous rencontrez dans la rue".

8 Conclusion

Dans ce projet, nous avons effectué une analyse sémantique comparative sur des modèles d'embeddings de mots formés à partir de tweets postés avant ou pendant la pandémie mondiale COVID-19. Une telle période constitue un bon point de référence pour étudier les changements sémantiques induits par des événements d'urgence dans une société en crise.

Nos principales contributions sont les suivantes :

- (1) Nous avons montré que la pandémie COVID-19 a introduit des changements sémantiques notables dans le langage Twitter, en obtenant un ensemble de mots potentiellement décalés. Ces résultats peuvent éclairer les lexicographes, les décideurs politiques ainsi que le grand public.
- (2) Nous avons réalisé une étude complète sur l'adaptabilité de différentes approches de détection des glissements sémantiques aux corpus de médias sociaux. Nous sommes parvenus à la conclusion que les embeddings word2vec sont plus efficaces que BERT, tandis que les métriques de similarité cosinus et de plus proches voisins peuvent se compléter. Si les méthodes d'intégration et les mesures de stabilité que nous avons choisies sont parmi les plus représentatives, certaines d'entre elles n'ont jamais été combinées dans des recherches antérieures. Néanmoins, notre étude est la première à aborder le problème de la conformité des méthodes aux données des médias sociaux.
- (3) Nous avons démontré que les fluctuations sont continues de mars à août en visualisant les trajectoires d'un ensemble de mots-clés liés à COVID-19. Les résultats reflètent de manière vivante les questions sociétales en cours, attestant de la valeur de la détection des glissements sémantiques en sociologie.
- (4) Nous avons réalisé des études de cas pour la langue française, ce qui nous a permis d'obtenir des informations intéressantes sur les questions sociétales et de combler les lacunes de ce type d'analyse en français.
- (5) Nous avons construit une version anglaise et une version française des jeux de données Twitter lié à COVID-19. Chacun de ces jeux de données est accompagné d'un modèle d'embeddings de mots. Ces ressources sont utiles pour les recherches futures en matière d'analyse de Twitter.

Acknowledgements

Nous remercions la Chaire ANR AML/HELAS et le projet ANR XTCOVIF pour leur soutien de cette recherche.

Références

- [1] Hosein Azarbondy, Mostafa Dehghani, Kaspar Beelen, Alexandra Arkut, Maarten Marx, and Jaap Kamps. Words are malleable : Computing semantic shifts in political and media discourse. In *Proceedings of the 2017 ACM on Conference on Information and Knowledge Management*, pages 1509–1518, 2017.
- [2] Emily Chen, Kristina Lerman, and Emilio Ferrara. Covid-19 : The first public coronavirus twitter dataset. *arXiv preprint arXiv :2003.07372*, 2020.
- [3] Matteo Cinelli, Walter Quattrociocchi, Alessandro Galeazzi, Carlo Michele Valensise, Emanuele Brugnoli, Ana Lucia Schmidt, Paola Zola, Fabiana Zollo, and Antonio Scala. The covid-19 social media infodemic. *Scientific Reports*, 10(1) :1–10, 2020.
- [4] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert : Pre-training of deep bidirectional transformers for language understanding. *arXiv preprint arXiv :1810.04805*, 2018.
- [5] Kawin Ethayarajh. How contextual are contextualized word representations ? comparing the geometry of bert, elmo, and gpt-2 embeddings. In *Proceedings of the 2019 Conference on Empirical Methods in Natural Language Processing and the 9th International Joint Conference on Natural Language Processing (EMNLP-IJCNLP)*, pages 55–65, 2019.
- [6] Mario Giulianelli, Marco Del Tredici, and Raquel Fernández. Analysing lexical semantic change with contextualised word representations. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 3960–3973, 2020.
- [7] Hila Gonen, Ganesh Jawahar, Djamé Seddah, and Yoav Goldberg. Simple, interpretable and stable method for detecting words with usage change across corpora. In *Proceedings of the 58th Annual Meeting of the Association for Computational Linguistics*, pages 538–555, 2020.
- [8] William L Hamilton, Jure Leskovec, and Dan Jurafsky. Cultural shift or linguistic drift ? comparing two computational measures of semantic change. In *Proceedings of the Conference on Empirical Methods in Natural Language Processing. Conference on Empirical Methods in Natural Language Processing*, volume 2016, page 2116. NIH Public Access, 2016.
- [9] William L Hamilton, Jure Leskovec, and Dan Jurafsky. Diachronic word embeddings reveal statistical laws of semantic change. In *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 1 : Long Papers)*, pages 1489–1501, 2016.

- [10] Gerhard Heyer, Florian Holz, and Sven Teresniak. Change of topics over time-tracking topics by their change of meaning. *KDIR*, 9 :223–228, 2009.
- [11] Martin Hilpert and Stefan Th Gries. Assessing frequency changes in multistage diachronic corpora : Applications for historical corpus linguistics and the study of language acquisition. *Literary and Linguistic Computing*, 24(4) :385–401, 2009.
- [12] Renfen Hu, Shen Li, and Shichen Liang. Diachronic sense modeling with deep contextualized word embeddings : An ecological view. In *Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics*, pages 3899–3908, 2019.
- [13] Patrick Juola. The time course of language change. *Computers and the Humanities*, 37(1) :77–96, 2003.
- [14] Yoon Kim, Yi-I Chiu, Kentaro Hanaki, Darshan Hegde, and Slav Petrov. Temporal analysis of language through neural language models. In *Proceedings of the ACL 2014 Workshop on Language Technologies and Computational Social Science*, pages 61–65, 2014.
- [15] Vivek Kulkarni, Rami Al-Rfou, Bryan Perozzi, and Steven Skiena. Statistically significant detection of linguistic change. In *Proceedings of the 24th International Conference on World Wide Web, WWW '15*, page 625–635, Republic and Canton of Geneva, CHE, 2015. International World Wide Web Conferences Steering Committee.
- [16] Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. Roberta : A robustly optimized bert pretraining approach. *arXiv preprint arXiv :1907.11692*, 2019.
- [17] Christian E Lopez, Malolan Vasu, and Caleb Gallemore. Understanding the perception of covid-19 policies by mining a multilanguage twitter dataset. *arXiv preprint arXiv :2003.10359*, 2020.
- [18] Matej Martinc, Petra Kralj Novak, and Senja Pollak. Leveraging contextual embeddings for detecting diachronic semantic shift. In *Proceedings of The 12th Language Resources and Evaluation Conference*, pages 4811–4819, 2020.
- [19] Tomas Mikolov, Ilya Sutskever, Kai Chen, Greg Corrado, and Jeffrey Dean. Distributed representations of words and phrases and their compositionality. *arXiv preprint arXiv :1310.4546*, 2013.
- [20] Syrielle Montariol, Matej Martinc, and Lidia Pivovaro. Scalable and interpretable semantic change detection. In *Proceedings of the 2021 Conference of the North American Chapter of the Association for Computational Linguistics : Human Language Technologies*, pages 4642–4652, 2021.
- [21] Martin Müller, Marcel Salathé, and Per E Kummervold. Covid-twitter-bert : A natural language processing model to analyse covid-19 content on twitter. *arXiv preprint arXiv :2005.07503*, 2020.
- [22] Dat Quoc Nguyen, Thanh Vu, and Anh Tuan Nguyen. Bertweet : A pre-trained language model for english tweets. In *Proceedings of the 2020 Conference on Empirical Methods in Natural Language Processing : System Demonstrations*, pages 9–14, 2020.
- [23] Matthew E Peters, Mark Neumann, Mohit Iyyer, Matt Gardner, Christopher Clark, Kenton Lee, and Luke Zettlemoyer. Deep contextualized word representations. *arXiv preprint arXiv :1802.05365*, 2018.
- [24] Stephen D Reese and Seth C Lewis. Framing the war on terror : The internalization of policy in the us press. *Journalism*, 10(6) :777–797, 2009.
- [25] Justyna A. Robinson. A gay paper : why should sociolinguistics bother with semantics ? : Can sociolinguistic methods shed light on semantic variation and change in reference to the adjective gay ? *English Today*, 28(4) :38–54, 2012.
- [26] Dominik Schlechtweg, Barbara McGillivray, Simon Hengchen, Haim Dubossarsky, and Nina Tahmasebi. Semeval-2020 task 1 : Unsupervised lexical semantic change detection. In *Proceedings of the Fourteenth Workshop on Semantic Evaluation*, pages 1–23, 2020.
- [27] Peter H Schönemann. A generalized solution of the orthogonal procrustes problem. *Psychometrika*, 31(1) :1–10, 1966.
- [28] Fatemeh Tahmasbi, Leonard Schild, Chen Ling, Jeremy Blackburn, Gianluca Stringhini, Yang Zhang, and Savvas Zannettou. “go eat a bat, chang !” : On the emergence of sinophobic behavior on web communities in the face of covid-19. In *Proceedings of the Web Conference 2021*, pages 1122–1133, 2021.
- [29] Yating Zhang, Adam Jatowt, Sourav S Bhowmick, and Katsumi Tanaka. The past is not a foreign country : Detecting semantically similar terms across time. *IEEE Transactions on Knowledge and Data Engineering*, 28(10) :2793–2807, 2016.
- [30] Caleb Ziems, Bing He, Sandeep Soni, and Srikanth Kumar. Racism is a virus : Anti-asian hate and counterhate in social media during the covid-19 crisis. *arXiv preprint arXiv :2005.12423*, 2020.

A Liste des mots avec un glissement sémantique confirmé de l’Oxford English Dictionary

covid, coronavirus, corona, cv, infodemic, isolate, quarantine, distancing, comorbid, tracing, frontline, mers, zoom, elbow, bump, wfh, ppe, flatten, curve, cfr, spread, transmission, cytokine, facemask, mask, surgical, covering, reproductive, spike, shield, hydroxychloroquine, dexamethasone, bubble, stimulus, cpap, handgel, essential

B Pipeline de prétraitement des Tweet

Chaque tweet utilisé dans le processus de formation de nos modèles word2vec est prétraité avec les étapes suivantes :

- **Conversion en minuscules** : Toutes les lettres de chaque tweet sont converties en minuscules.
- **Suppression des URL et des mentions** : Nous supprimons toutes les mentions et URL apparaissant dans les tweets.
- **Normalisation des symboles emoji et des caractères d'émoticônes** : Nous remplaçons tous les symboles emoji et les émoticônes par leur correspondance dans une liste de représentations normalisées.
- **Suppression des tweets courts** : Les tweets extrêmement courts, c'est-à-dire les tweets comportant moins de 10 tokens au total, ne contiennent pas de contenu significatif dans la plupart des cas et sont donc supprimés.

Détection d'objets en temps réel : Entraînement de réseaux de neurones convolutifs sur images réelles et synthétiques

T. Zgheib¹, H. Borges¹, V. Feldman¹, T. Guntz¹, C. Di Loreto¹, O. Desmaison¹, F. Corduant¹

¹ KAIZEN Solutions, KZS LAB, 38330 Montbonnot-Saint-Martin

taline.zgheib@kaizen-solutions.net

Résumé

La détection des objets a prouvé son efficacité dans diverses applications industrielles. La réussite de l'entraînement d'un détecteur d'objets dépend largement de la disponibilité de jeux d'images réelles, qui sont souvent rares et dépendent de l'application visée. Cependant, il est possible de générer des images synthétiques pour l'entraînement via des modèles 3D des objets d'intérêt. Dans cet article, nous étudions les performances de deux architectures de réseaux neurones : MobileNetV2-SSD et ResNet50-SSD. Trois scénarios d'entraînement sont considérés : i) Entraînement ER sur des images réelles, ii) E1 sur des images synthétiques et iii) E2 sur des images synthétiques avec un fine-tuning sur des images réelles. Nos résultats montrent la baisse de la précision moyenne (mAP) lorsque l'entraînement est effectué sur des images synthétiques. Cette précision s'améliore lorsque l'entraînement est suivi d'un fine-tuning sur des images réelles.

Mots-clés

Détection d'objets, réseau de neurones convolutifs, smartphone, temps réel, images synthétiques.

Abstract

Object detection has been proven to be very effective in various industrial applications. Successfully training an object detector relies on real image datasets that are scarce and heavily dependent on the targeted application. However, creating synthetic image datasets for training is possible using 3D models of the objects of interest. Here, we study the performance of two Neural Network architectures : MobileNetV2-SSD and ResNet50-SSD. Three training scenarios are considered i) ER training on real images, ii) E1 training on synthetic images and iii) E2 training on synthetic images and finetuning on real images. Our results show a decrease in the mean average precision (mAP) when the models are trained on synthetic datasets compared to real image datasets. This precision improves when the training is finetuned with real images.

Keywords

Object detection, convolutional neural networks, smartphone, real time, synthetic images.

1 Introduction

La détection d'objet est une technique de vision par ordinateur qui consiste à identifier la présence et l'emplacement de toute instance d'une classe d'objets présente dans l'image (via une boîte englobante associée à une classe d'objets). Cette technique est devenue en quelques années un outil précieux dans des domaines très divers de l'industrie, comme par exemple : la classification et tri des déchets [32], la détection des défauts [31], etc. En effet, le secteur industriel peut bénéficier de la vision par ordinateur et particulièrement la détection d'objet pour automatiser plusieurs tâches. Ceci nécessite souvent le traitement d'images en temps réel sur des appareils embarqués ou mobiles qui sont plus limités en mémoire et en capacité de calcul [3].

Les algorithmes d'apprentissage automatique pour la détection d'objets sont largement utilisés dans l'apprentissage supervisé et peuvent résoudre de nombreux problèmes pratiques [20, 18]. Cependant, les réseaux neuronaux convolutifs (CNN) sont à l'heure actuelle plus performants dans le traitement d'images et la détection d'objets [34]. L'entraînement de modèles de réseaux de neurones est réalisé sur une base de données d'images annotées. Dans un contexte industriel, ceci pose beaucoup de contraintes liées à la disponibilité et l'accès aux données nécessaires. En effet, une base de données d'images réelles annotées est souvent inexistante, rare ou encore trop petite pour un entraînement efficace du CNN. De plus, la diversité des objets à détecter implique la génération d'une base de données spécifique pour chaque application. Ceci représente des contraintes temporelles et financières non négligeables pour l'industrie. Pour répondre à cette problématique, il est possible d'utiliser des techniques d'augmentation de données qui améliorent la taille et la qualité des ensembles de données d'entraînement de manière à pouvoir mieux entraîner les réseaux neuronaux profonds [30, 23]. Une seconde méthode consiste à générer automatiquement des données synthétiques annotées grâce aux modèles 3D des objets à détecter [25].

Nous étudions dans cet article les performances de deux détecteurs d'objets (MobileNetV2-SSD et ResNet50-SSD) entraînés sur des données réelles et synthétiques. Nous évaluons également l'influence de la qualité des données synthétiques générées sur la qualité et la performance des mo-

dèles choisis. Les modèles sont évalués sur des bases de données réelles et synthétiques que nous avons générées.

2 Travaux antérieurs

2.1 Architecture générale d'un détecteur d'objet

L'architecture typique d'un détecteur d'objet comprend trois parties : l'épine dorsale (*backbone*), le cou (*neck*) et la tête de détection (*detection head*). L'épine dorsale a pour rôle d'extraire les caractéristiques des images d'entrée. Les cartes de caractéristiques (appelées parfois cartes d'activations) générées sont ensuite introduites dans le cou du détecteur, responsable de l'extraction des caractéristiques pertinentes et de leurs fusions à multi-échelle. Les caractéristiques passent finalement par la tête de détection pour la classification finale et la régression des boîtes englobantes générées.

2.2 Extracteurs de caractéristiques

L'extraction de caractéristiques est une étape cruciale pour la détection d'objet. En effet, le choix de l'architecture utilisée a un impact sur la performance du modèle (vitesse d'inférence et précision). Les réseaux de neurones profonds peuvent apprendre et transmettre des caractéristiques complexes à plusieurs niveaux à partir de données d'entrée. Cependant, à mesure que les réseaux de neurones deviennent de plus en plus profonds, une saturation puis une dégradation de la précision se produisent. Ainsi des équipes de recherche [7] proposent les réseaux résiduels *ResNet* pour faciliter l'entraînement des réseaux de neurones très profonds et résoudre le problème de la dégradation. Ils introduisent la notion des connexions résiduelles entre les couches des réseaux neurones pour éviter le problème de disparition de gradients [7] (*vanishing gradient*). En effet, la rétropropagation du gradient dans un réseau très profond peut rendre le gradient infiniment petit. La connexion résiduelle proposée effectue des raccourcis qui permettent le passage des matrices d'activation d'une couche N à une couche plus profonde $N + k$. Cette cartographie d'identité et plus précisément les connexions résiduelles permettent la transmission du gradient à travers les couches les plus profondes, réduisant ainsi le problème de disparition de gradient [29].

Nos recherches s'intéressent aux réseaux de neurones applicables dans un contexte industriel où la vitesse d'inférence doit être proche du temps réel. Des réseaux récents ont été développés dans cette optique : Les *MobileNet* [9]. Ils offrent un compromis efficace entre vitesse et précision ainsi que la possibilité d'adapter la taille du réseau en fonction des besoins de l'utilisateur et de l'application proposée. L'architecture *MobileNet* introduit pour la première fois la notion de convolution séparable en profondeur qui vient remplacer la convolution classique dans les réseaux de neurones. Elle permet ainsi une diminution conséquente en coût de calcul sans dégrader la précision du modèle. En 2018, *MobileNetV2* [22] propose des blocs résiduels in-

versés (*inverted residual blocks*) et des goulots d'étranglements linéaires (*linear bottleneck*), offrant au réseau la possibilité d'apprendre sur des fonctions plus riches tout en réduisant la taille mémoire.

2.3 Détecteurs d'objets

Il existe deux approches pour la détection d'objet : l'approche à une passe et celle à deux passes. Les approches convolutives traditionnelles à deux passes (*two-stage detectors*, telles que R-CNN [5], SPP-net [6], FPN [12]) extraient les régions les plus susceptibles de contenir un objet (régions d'intérêt) [33, 35] lors d'une première passe. Ensuite, la seconde passe analyse ces régions et prédit à la fois la probabilité de présence d'un objet ainsi que sa classe d'appartenance. Ces méthodes sont certes très précises mais aussi fastidieuses et lentes. Des équipes ont alors développé des approches dites "en une passe" (*one stage detectors*, telles que SSD [14] et YOLO [21]), qui s'avèrent être plus rapides et plus adaptées pour la détection en temps réel. Ces modèles analysent l'image, localisent et identifient les objets en une seule passe [33, 35].

2.4 Entraînement sur des images synthétiques et réelles

L'entraînement des réseaux de neurones convolutifs nécessite l'utilisation d'une base de données synthétiques/réelles conséquente, or la taille de la base de données est un critère très important à considérer pour la performance du modèle. Cependant dans la pratique, ces bases de données sont le plus souvent rares ou bien ne répondent pas toujours aux exigences de taille et de qualité nécessaires pour un tel entraînement. Puisque générer un très grand nombre de données réelles engendrerait un sur-coût financier non négligeable, d'autres approches d'augmentation de données sont donc proposées pour créer des images synthétiques. Certaines méthodes reposent sur la modélisation du contexte visuel [4] et l'apprentissage automatique contradictoire [28] (*adversarial learning*) pour produire des images composites réalistes. D'autres utilisent la randomisation de domaine [24, 25] pour générer des images d'entraînement avec suffisamment de variations pour minimiser le sur-apprentissage.

Avoir une base de données de qualité est indispensable à l'entraînement des réseaux neurones. En effet, la littérature souligne l'importance d'utiliser un jeu de données équilibré où toutes les classes d'objets sont représentées de manière homogène. De même, la diversité de taille et de position des objets ainsi que le nombre de distracteurs impactent la qualité de l'apprentissage du modèle [24]. Par exemple, prendre en compte l'occlusion partielle des objets rend le modèle plus robuste au sur-apprentissage [1, 24]. Cependant, une étude montre que l'entraînement sur des images réelles reste nécessaire pour obtenir des performances acceptables [17].

3 Méthodologie

3.1 Architecture des modèles

Pour réaliser la détection d'objets, nous nous sommes appuyés sur le réseau de neurones Single Shot MultiBox Detector (SSD) [14]. SSD est un détecteur en une passe entièrement convolutif. Il permet une détection à multi-échelles facilitant ainsi la détection d'objets de tailles différentes [14]. SSD s'appuie sur l'architecture VGG-16 pour l'extraction de caractéristiques. Des couches supplémentaires convolutives succèdent le *backbone* et font partie du *neck* du détecteur. Elles ont pour rôle de réduire progressivement la taille des matrices d'activation. Enfin, la couche de la tête du détecteur prend en entrée des cartes de caractéristiques du neck et de l'extracteur. Dans notre application, on remplace VGG-16 par les réseaux MobileNetV2[22] et ResNet50 [7] pré-entraînés sur la base de données COCO [13]. Cette procédure est appelée "transfert d'apprentissage" et a pour avantage de réduire le temps d'exécution de l'algorithme et la quantité d'images utilisées pour la *fine-tuning* du modèle. Les modèles pré-entraînés sont fournis par le *TensorFlow 2 Détection Model Zoo*¹.

Le *neck* du modèle reste celui présenté dans l'article présentant le modèle SSD. A noter que ResNet50 désigne une architecture de réseau résiduel d'une profondeur de 50 couches.

3.2 Génération des images synthétiques et réelles

Pour entraîner les modèles de détection nous avons généré deux jeux de données réelles et synthétiques. Les deux jeux de données réelles sont composés respectivement de 100 (nommé R100) et 400 (R400) images. Ces données réelles sont des photos prises par l'équipe en conditions réelles (Figure 1b). Les deux jeux de données synthétiques SV1 et SV2 sont quant à eux respectivement composés de 2000 et 7500 images. Pour générer ces jeux de données synthétiques, nous avons adapté la technique décrite par [27, 26] qui proposent de créer des images synthétiques à partir de modèles 3D, d'images de fond réelles et en faisant varier plusieurs paramètres.

Les images synthétiques sont générées avec le moteur de jeu Unity3D en modifiant le fond et l'orientation des objets pour SV1 (Figure 1c) et l'inclinaison du fond, la lumière et la texture des objets pour SV2 (Figure 1d). Voir Table 1 pour plus de détails sur les différents paramètres choisis pour la génération de SV1 et SV2.

Les objets à détecter sont des briques Lego de trois formes différentes : 2x1, 2x2 et 4x1 (Figure 1a). Les images synthétiques sont générées en utilisant les modèles 3D des différents distracteurs et briques de Lego à positionner dans l'image. Nous avons choisi trois types de distracteurs : sphère, cube et cylindre. Leurs nombres par image varient aléatoirement et leurs textures sont modifiées en utilisant des images issues du jeu de données Flickr 30k. Le même jeu de données est utilisé pour modifier aléatoirement les

textures des briques Lego et les images de fond. De même, Unity nous permet aussi de varier plusieurs paramètres de façon aléatoire : le positionnement et orientation des objets, orientation des caméras, le nombre de sources lumineuses, etc. Tous ces paramétrages ont pour objectif de limiter le sur-apprentissage du modèle lors de l'entraînement. Les Figures 1c et 1d présentent deux images synthétiques type utilisées pour l'entraînement des modèles étudiés. Les annotations (boîtes englobantes et labels) des images synthétiques ainsi produites ont été générées automatiquement au format adapté pour le framework TensorFlow.

TABLE 1 – Récapitulatif des paramètres de génération des jeux de données d'images synthétiques SV1 et SV2

Paramètres	SV1	SV2
Nombre d'objets d'intérêts : 3 briques de Lego (e.g. 2x1,2x2,2x4)	X	X
Position et orientation aléatoire de l'objet d'intérêt dans la scène	X	X
Type de distracteurs : Sphère, cylindre, cube. Variation aléatoire de la position/rotation	X	X
Variation de la position (Z fixe) et de la rotation (X,Y,Z) de l'angle de la caméra	X	X
Variation du nombre de lumières (entre 0 et 12), leur intensité, portée et angle		X
Variation de la texture de l'objet d'intérêt en utilisant une image aléatoire de Flickr 30k		X
Inclinaison et translation de l'image de fond		X

3.3 Entraînement et évaluation des modèles

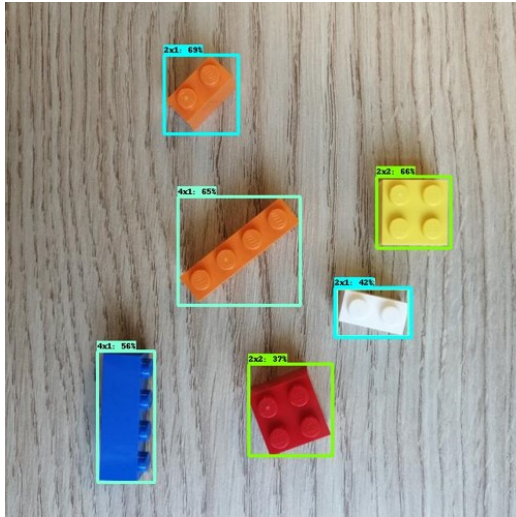
Pour le transfert de connaissance nous choisissons les modèles MobineNetV2-SSD (2.5 millions de paramètres) et ResNet50-SSD (31 millions de paramètres) pré-entraînés avec une résolution d'image de 640x640 pixels. Les performances des modèles choisis sont disponible dans la Table 2.

Pour chacun des modèles, nous proposons huit entraînements différents répartis sur trois scénarios (Table 3) :

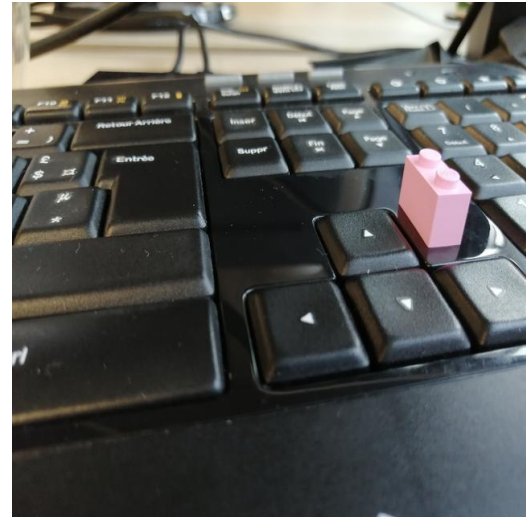
- Scénario de référence (ER) : Deux entraînements exclusivement sur des images réelles. Les jeux de données utilisés sont R100 et R400.
- Scénario E1 : Deux entraînements sur uniquement des images synthétiques. Les jeux de données utilisés sont SV1 et SV2.
- Scénario E2 : Quatre entraînements sur images synthétiques suivi d'un *fine-tuning* sur images réelles. Pour ce scénario quatre combinaisons de jeux de données sont utilisées : SV1+R100, SV2+R100, SV1+R400 et SV2+R400.

Les deux architectures sont optimisées par l'algorithme de descente de gradient stochastique (SGD, *stochastic gradient descent*) avec une taille de *batch* de 4 images, un *lear-*

1. shorturl.at/hmpKQ



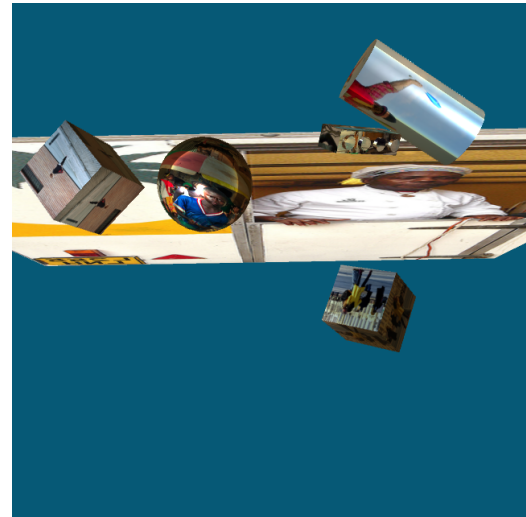
(a) Image réelle avec les différentes boîtes englobantes de 3 types de briques Lego (2x2, 2x1 et 4x1).



(b) Image réelle (jeu de données R100) prise par l'équipe contenant un Lego 2x1.



(c) Image synthétique (jeu de données SV1) générée avec notre approche contenant 2 distracteurs et un Lego 4x1.



(d) Image synthétique (jeu de données SV2) générée avec notre approche contenant 4 distracteurs et un Lego 4x1.

FIGURE 1 – Images réelles (a,b) et synthétiques (c, d) générées et utilisées pour l'entraînement des différents modèles étudiés.

ning rate de 0.01 et un total de 30 epochs pour les scénarios ER et E1. Pour le scénario E2 les modèles sont entraînés sur 20 epochs d'images synthétiques et 10 epochs d'images réelles avec respectivement un *learning rate* de 0.01 et de 0.001. La fonction perte utilisée pour la localisation est celle proposée par défaut, c'est-à-dire une fonction Smooth L1 (voir Hubert [10]). Tous les autres hyperparamètres des deux architectures sont ceux fournis par *TensorFlow 2 Détection Model Zoo*, nous conservons leur valeur par défaut.

L'entraînement et la validation des modèles sont réalisés sur une machine locale avec 32GB de RAM, un processeur Intel Xeon E5-1650 v2 et une carte Nvidia GTX1080. Les modèles sont validés sur 500 images réelles distinctes de celles utilisées pour les différents entraînements (R100 et

R400). La précision moyenne (mAP, *mean Average Precision*) est utilisée comme métrique d'évaluation des performances pour les différents scénarios.

TABLE 2 – Performances des modèles sur le Dataset COCO 2017 pour des images de dimensions 640x640 pixels (obtenus sur TFHub)

Modèle	Inférence (ms)	mAP	Nombre de paramètres (M)
MobileNetV2-SSD	39	28.2	2.5
ResNet50-SSD	46	34.3	31

TABLE 3 – Récapitulatif des scénarios d'entraînement pour MobileNetV2-SSD et ResNet50-SSD.

Scénarios	Jeux de données
(ER) : Entraînement de référence	R100 et R400
(E1) : Entraînement sur images synthétiques	SV1 et SV2
(E2) : Entraînement sur images synthétiques + fine-tuning sur images réelles	SV1+R100, SV2+R100, SV1+R400, SV2+R400

4 Résultats

La figure 2 présente la précision moyenne (mAP) de la détection pour les différentes analyses. Les résultats montrent que, tant pour MobileNetV2-SSD que pour ResNet50-SSD, la précision moyenne diminue lorsque l'on utilise des données entièrement synthétiques au lieu de données réelles. Pour Resnet50-SSD, la précision passe de 43% en moyenne pour des données réelles à 30% pour des données synthétiques. En revanche, pour Mobilenetv2-SSD, la baisse de la précision moyenne n'est que de 3 % (de 46%(scénario ER) à 41%(scénario E1)).

L'entraînement sur des données synthétiques suivi d'un ajustement fin sur des données réelles (scénario E2) augmente la précision de MobileNetV2-SSD à 47% (contre 41% pour scénario E1). Ainsi, l'utilisation de données réelles pour le *fine-tuning* augmente la mAP de MobileNetV2-SSD à des niveaux comparables à ceux de l'entraînement ER. De même, pour ResNet50-SSD, la mAP augmente de 31% pour le scénario E1 à 43% pour le scénario E2 (entraînement sur images synthétiques + *fine-tuning* sur images réelles). Notre résultat sur MobileNetv2-SSD et Resnet50-SSD est conforme à celui de Nowruzi *et al.* [17] qui suggèrent une augmentation de la précision moyenne quand l'entraînement sur des données synthétiques est suivi d'un *fine-tuning* sur les images réelles. De plus, Nowruzi *et al.* [17] montrent qu'en diminuant la proportion de données réelles utilisées pour le *fine-tuning*, nous dégradons la performance du modèle. Contrairement à cette dernière hypothèse, nous n'avons constaté aucun changement majeur dans les performances de notre modèle. L'ajustement sur 100 ou 400 images réelles n'a pas d'impact sur la précision de ResNet50 et augmente celle de MobileNet de seulement 3% (45% (SV2+R100) à 48% (SV2+R400)).

Notons que l'entraînement sur un plus grand jeu de données synthétiques pénalise la précision des modèles. En effet la mAP est diminuée de 39% (SV1) à 23% (SV2) pour Resnet50-SSD et de 45% (SV1) à 38% (SV2) pour MobilenetV2-SSD. Ceci est directement lié au choix des paramètres à faire varier lors de la création des données synthétiques. En effet, la complexité des variations des paramètres choisis pour la génération de SV2 par rapport à SV1 (Table 1) complique l'apprentissage et dégrade la qualité du jeu de données. Nous suspectons que l'incli-

naison du fond (Figure 1c) (source de zones uniformes de couleur aléatoire dans les images synthétiques), la texture ajoutée au Lego ainsi que les reflets lumineux ajoutent une trop grande variabilité des objets d'intérêts et rend l'apprentissage difficile. Toutefois nous obtenons des résultats surprenants pour Resnet50-SSD, dont la précision diminue de 44% à 42% pour un entraînement sur 400 images réelles. Des analyses supplémentaires devraient être menées pour déterminer les causes plausibles qui permettraient de justifier un tel résultat.

De façon générale, MobilenetV2-SSD a une meilleure précision que Resnet50-SSD pour les trois scénarios que nous avons examinés dans cette étude. Cette différence entre les deux modèles indique que MobilenetV2 a une meilleure capacité de généralisation aux images réelles que Resnet50. Ceci peut s'expliquer par le nombre de paramètres de MobilenetV2-SSD qui est inférieur à celui de ResNet50-SSD (Table 2), ce qui limite l'apprentissage de caractéristiques trop spécifiques aux images synthétiques lors de l'entraînement[3].

Les objets étudiés au cours de cette étude ont des caractéristiques semblables et relativement simples. Nous nous interrogeons ainsi sur la généralisation de nos résultats à différents types d'objets, aux textures plus complexes et aux formes irrégulières. En général, l'apparence réaliste d'un objet dépend de ses caractéristiques de bas niveau (forme, pose, texture, arrière-plan, *etc.*). L'étude de peng *et al.* [19] montre que la sensibilité à la pose, texture et l'arrière-plan de l'objet peut être surmontée lorsque l'entraînement du réseau de neurones (sur les données synthétiques) est suivi d'un *fine-tuning* sur des images réelles. Ainsi, si l'objet à détecter possède une texture complexe, il ne sera pas nécessaire de reproduire la complexité de cette texture lors de la création du jeu de données synthétiques. Cependant, si le réseau n'est pas affiné sur des images réelles, l'invariance des indices de bas niveau n'est plus valable. Néanmoins, si la texture est un facteur décisif dans la classification/détection de l'objet, une représentation réelle de la texture devrait être considérée. Par ailleurs, la forme de l'objet reste un critère important et indispensable pour la majorité des modèles qui se focalise sur les formes et les limites des objets pour l'apprentissage [16]. En conséquence, le résultat ne dépend pas de l'objet à détecter, mais *i)* du modèle choisi, *ii)* de la qualité de l'ensemble de données synthétiques utilisé et *iii)* du degré auquel l'objet est représenté. Il est alors préférable d'expérimenter avec le modèle de détection choisi afin d'identifier les aspects des objets qui méritent l'investissement nécessaire pour les représenter dans un jeu de données synthétiques de bonne qualité.

5 Conclusion et futurs travaux

L'approche proposée dans cet article permet de réaliser l'entraînement d'un détecteur d'objet sur trois différents type de briques Lego, avec très peu de données réelle annotées en notre possession. Nos premiers résultats montrent que nous pouvons entraîner un détecteur d'objet unique-

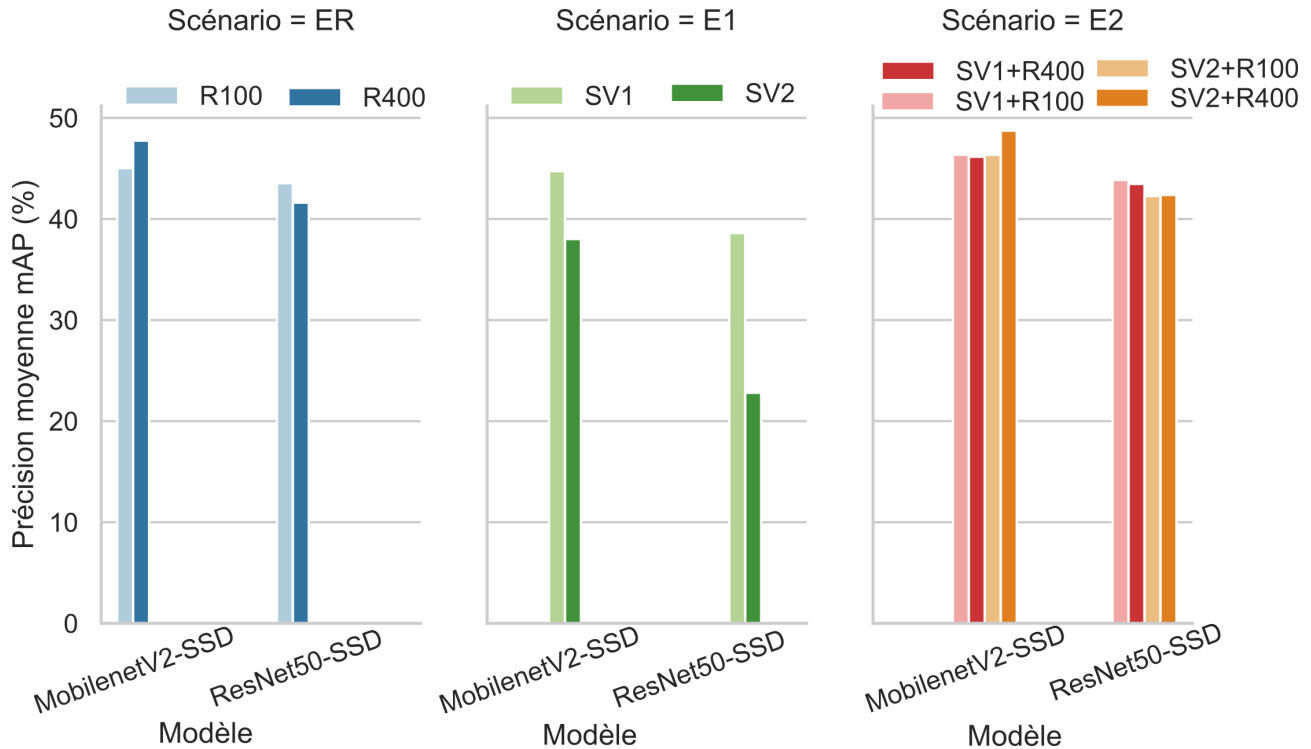


FIGURE 2 – Performance des modèles MobileNetV2-SSD et ResNet50-SSD pour les scénarios *i)* ER : entraînement sur images réelles, *ii)* E1 : entraînement sur images synthétiques et *iii)* E2 : entraînement sur images synthétiques suivi d'un *fine-tuning* sur images réelles. R100 et R400 sont les jeux de données d'images réelles avec respectivement 100 et 400 images. SV1 et SV2 sont les jeux de données synthétiques avec respectivement 2000 et 7500 images.

ment sur des données synthétiques et obtenir des performances satisfaisantes (39% pour MobileNetV2 et 45% pour ResNet-SSD, sur SV1), ce qui confirme le résultat de plusieurs autres études, comme Tremblay *et al.* [25]. Si nos résultats laissent à penser qu'utiliser seulement R100 ou R400 semble suffisant, nous pensons que l'utilisation d'une base de données hybride, combinant données réelles et synthétiques, est une approche très intéressante pour les petites équipes de recherche ou les industriels. En effet, la génération d'une telle base, par exemple 75% de données synthétiques et 25% de données réelles, est beaucoup moins chronophage et coûteuse. Cependant, de nouvelles expériences doivent être menées, cette fois-ci avec des objets à forme plus complexes que celle des Lego. Cela permettrait de confirmer la viabilité de cette approche à généraliser sur des cas plus proches des problématiques rencontrées en environnement industriel. Ces premiers résultats nous encourageant à poursuivre nos travaux en explorant plusieurs axes, notamment sur la qualité des données générées, l'optimisation du modèle et l'amélioration des performances de détection.

Tout d'abord, nous souhaitons étudier plus en détail l'impact du ratio données synthétiques/réelles dans une base de données hybride sur les performances des modèles de détection, de la même manière que cela est présenté dans les travaux de Nowruz *et al.* [17] et de Borkman *et al.* [2]. Cela nous permettrait de déterminer la proportion minimale de

données réelles nécessaires pour obtenir de bonnes performances.

Ensuite, pour expliquer les différences de performances entre les différentes stratégies de générations de données (entre SV1 et SV2), nous souhaitons mieux contrôler les variations de paramètres de ces stratégies. Une première piste est l'utilisation de la librairie Perception de Unity [2] pour générer des images synthétiques. En comprenant mieux l'influence de la variation des paramètres, nous espérons réduire l'écart de domaine entre les données réelles et synthétiques. Enfin, il serait intéressant d'étudier l'influence du photo-réalisme sur les performances de détection, dont l'efficacité reste une question ouverte dans la communauté [1, 15].

Concernant le modèle de détection, nous envisageons de tester les performances d'autres architectures potentiellement plus adaptées pour la détection d'objet en temps réel (par exemple, une architecture MobileNet-V3 [8] et YOLO-V5 [11]).

Références

- [1] Hassan Abu Alhaija, Siva Karthik Mustikovela, Lars Mescheder, Andreas Geiger, and Carsten Rother. Augmented reality meets computer vision : Efficient data generation for urban driving scenes. *International Journal of Computer Vision*, 126(9) :961–972, 2018.

- [2] Steve Borkman, Adam Crespi, Saurav Dhakad, Sujoy Ganguly, Jonathan Hogins, You-Cyuan Jhang, Mohsen Kamalzadeh, Bowen Li, Steven Leal, Pete Parisi, et al. Unity perception : Generate synthetic data for computer vision. *arXiv preprint arXiv :2107.04259*, 2021.
- [3] Julia Cohen, Carlos Crispim-Junior, Jean-Marc Chiappa, and Laure Tougne. Mobilenet ssd : étude d'un détecteur d'objets embarquable entraîné sans images réelles. In *ORASIS 2021*, 2021.
- [4] Nikita Dvornik, Julien Mairal, and Cordelia Schmid. Modeling visual context is key to augmenting object detection datasets. In *Proceedings of the European Conference on Computer Vision (ECCV)*, pages 364–380, 2018.
- [5] Ross Girshick, Jeff Donahue, Trevor Darrell, and Jitendra Malik. Rich feature hierarchies for accurate object detection and semantic segmentation. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 580–587, 2014.
- [6] Kaiming He, X. Zhang, Shaoqing Ren, and Jian Sun. Spatial pyramid pooling in deep convolutional networks for visual recognition. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 37 :1904–1916, 2015.
- [7] Kaiming He, Xiangyu Zhang, Shaoqing Ren, and Jian Sun. Deep residual learning for image recognition. In *2016 IEEE Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 770–778, 2016.
- [8] Andrew Howard, Mark Sandler, Grace Chu, Liang-Chieh Chen, Bo Chen, Mingxing Tan, Weijun Wang, Yukun Zhu, Ruoming Pang, Vijay Vasudevan, et al. Searching for mobilenetv3. In *Proceedings of the IEEE/CVF International Conference on Computer Vision*, pages 1314–1324, 2019.
- [9] Andrew G Howard, Menglong Zhu, Bo Chen, Dmitry Kalenichenko, Weijun Wang, Tobias Weyand, Marco Andreetto, and Hartwig Adam. Mobilenets : Efficient convolutional neural networks for mobile vision applications. *arXiv preprint arXiv :1704.04861*, 2017.
- [10] Peter J Huber. Robust estimation of a location parameter. In *Breakthroughs in statistics*, page 492–518. Springer, 1992.
- [11] Glenn Jocher, Alex Stoken, Jirka Borovec, A Chaurasia, and L Changyu. ultralytics/yolov5. *Github Repository*, YOLOv5, 2020.
- [12] Tsung-Yi Lin, Piotr Dollár, Ross Girshick, Kaiming He, Bharath Hariharan, and Serge Belongie. Feature pyramid networks for object detection. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 2117–2125, 2017.
- [13] Tsung-Yi Lin, Michael Maire, Serge Belongie, James Hays, Pietro Perona, Deva Ramanan, Piotr Dollár, and C Lawrence Zitnick. Microsoft coco : Common objects in context. In *European conference on computer vision*, pages 740–755. Springer, 2014.
- [14] Wei Liu, Dragomir Anguelov, Dumitru Erhan, Christian Szegedy, Scott Reed, Cheng-Yang Fu, and Alexander C Berg. Ssd : Single shot multibox detector. In *European conference on computer vision*, pages 21–37. Springer, 2016.
- [15] Nikolaus Mayer, Eddy Ilg, Philipp Fischer, Caner Hazirbas, Daniel Cremers, Alexey Dosovitskiy, and Thomas Brox. What makes good synthetic training data for learning disparity and optical flow estimation ? *International Journal of Computer Vision*, 126(9) :942–960, 2018.
- [16] Sergey I Nikolenko et al. *Synthetic data for deep learning*. Springer, 2021.
- [17] Farzan Erlik Nowruzi, Prince Kapoor, Dhanvin Kohli, Fahed Al Hassanat, Robert Laganier, and Julien Rebut. How much real data do we actually need : Analyzing object detection performance using synthetic and real data. *arXiv preprint arXiv :1907.07061*, 2019.
- [18] Constantine Papageorgiou and Tomaso Poggio. A trainable system for object detection. *International journal of computer vision*, 38(1) :15–33, 2000.
- [19] Xingchao Peng, Baochen Sun, Karim Ali, and Kate Saenko. Learning deep object detectors from 3d models. In *2015 IEEE International Conference on Computer Vision (ICCV)*, pages 1278–1286, 2015.
- [20] Dominik Maximilián Ramík, Christophe Sabourin, Ramon Moreno, and Kurosh Madani. A machine learning based intelligent vision system for autonomous object detection and recognition. *Applied intelligence*, 40(2) :358–375, 2014.
- [21] Joseph Redmon, Santosh Divvala, Ross Girshick, and Ali 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.
- [22] Mark Sandler, Andrew Howard, Menglong Zhu, Andrey Zhmoginov, and Liang-Chieh Chen. Mobilenetv2 : Inverted residuals and linear bottlenecks. In *Proceedings of the IEEE conference on computer vision and pattern recognition*, pages 4510–4520, 2018.
- [23] Uday Shankar Shanthamallu, Andreas Spanias, Cihan Tepedelenlioglu, and Mike Stanley. A brief survey of machine learning methods and their sensor and iot applications. In *2017 8th International Conference on Information, Intelligence, Systems & Applications (IISA)*, pages 1–8. IEEE, 2017.
- [24] Josh Tobin, Rachel Fong, Alex Ray, Jonas Schneider, Wojciech Zaremba, and Pieter Abbeel. Domain randomization for transferring deep neural networks from simulation to the real world. In *2017 IEEE/RSJ international conference on intelligent robots and systems (IROS)*, pages 23–30. IEEE, 2017.
- [25] Jonathan Tremblay, Aayush Prakash, David Acuna, Mark Brophy, Varun Jampani, Cem Anil, Thang To,

- Eric Cameracci, Shaad Boochoon, and Stan Birchfield. Training deep networks with synthetic data : Bridging the reality gap by domain randomization. In *Proceedings of the IEEE conference on computer vision and pattern recognition workshops*, pages 969–977, 2018.
- [26] Jonathan Tremblay, Thang To, and Stan Birchfield. Falling things : A synthetic dataset for 3d object detection and pose estimation. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition Workshops*, pages 2038–2041, 2018.
- [27] Jonathan Tremblay, Thang To, Artem Molchanov, Stephen Tyree, Jan Kautz, and Stan Birchfield. Synthetically trained neural networks for learning human-readable plans from real-world demonstrations. In *2018 IEEE International Conference on Robotics and Automation (ICRA)*, pages 5659–5666. IEEE, 2018.
- [28] Shashank Tripathi, Siddhartha Chandra, Amit Agrawal, Ambrish Tyagi, James M Rehg, and Visesh Chari. Learning to generate synthetic data via compositing. In *Proceedings of the IEEE/CVF Conference on Computer Vision and Pattern Recognition*, pages 461–470, 2019.
- [29] Andreas Veit, Michael J Wilber, and Serge Belongie. Residual networks behave like ensembles of relatively shallow networks. *Advances in neural information processing systems*, 29, 2016.
- [30] Sebastien C Wong, Adam Gatt, Victor Stamatescu, and Mark D McDonnell. Understanding data augmentation for classification : when to warp? In *2016 international conference on digital image computing : techniques and applications (DICTA)*, pages 1–6. IEEE, 2016.
- [31] Jing Yang, Shaobo Li, Zheng Wang, Hao Dong, Jun Wang, and Shihao Tang. Using deep learning to detect defects in manufacturing : a comprehensive survey and current challenges. *Materials*, 13(24) :5755, 2020.
- [32] Qiang Zhang, Xujuan Zhang, Xiaojun Mu, Zhihe Wang, Ran Tian, Xiangwen Wang, and Xueyan Liu. Recyclable waste image recognition based on deep learning. *Resources, Conservation and Recycling*, 171 :105636, 2021.
- [33] Zhong-Qiu Zhao, Peng Zheng, Shou-tao Xu, and Xindong Wu. Object detection with deep learning : A review. *IEEE transactions on neural networks and learning systems*, 30(11) :3212–3232, 2019.
- [34] Haiying Zhou, Xiangyu Yu, Ahmad Alhaskawi, Yanzhao Dong, Zewei Wang, Qianjun Jin, Xianliang Hu, Zongyu Liu, Vishnu Goutham Kota, Mohamed Hassan Abdulla Hasan Abdulla, et al. A deep learning approach for medical waste classification. *Scientific reports*, 12(1) :1–9, 2022.
- [35] Zhengxia Zou, Zhenwei Shi, Yuhong Guo, and Jieping Ye. Object detection in 20 years : A survey. *arXiv preprint arXiv :1905.05055*, 2019.

Détection d'objets urbains dans des nuages de points LiDAR

Y. Zegaoui¹, P. Borianne², M. Chaumont³, G. Subsol³, A. Seriai¹, M. Derras¹

¹ Berger-Levrault

² CIRAD ³ LIRMM

{younes.zegaoui, abderrahmane.seriai, mustapha.derras}@berger-levrault.com, {marc.chaumont, gerard.subsol}@lirmm.fr, philippe.borianne@cirad.fr

Résumé

Les travaux en apprentissage profond sur les données 3D sont depuis peu très étudiés. Leurs applications en milieu urbain permettraient de faciliter grandement la gestion des objets urbains par les agglomérations. Nous proposons une méthode originale pour la détection des objets urbains dans des nuages de points 3D issus d'acquisitions LiDAR.

Mots-clés

LiDAR, apprentissage profond, nuages de points 3D, convolution 3D, détection d'objets.

Abstract

Deep learning based research on 3D data has recently gained a lot of interest. Their application in urban areas would greatly facilitate the monitoring of urban objects by city managers. We propose an original method for the detection of urban objects in 3D point clouds from LiDAR acquisitions.

Keywords

LiDAR, deep-learning, 3D points cloud, 3D convolution, object detection.

1 Introduction

1.1 Gestion technique des équipements urbains

Le groupe Berger-Levrault, éditeur historique de logiciels administratifs pluriels, compte parmi ses clients un nombre important de collectivités territoriales avec des produits adaptés à leur diverses missions. Dans notre cas, nous nous intéressons particulièrement à la gestion technique des équipements en milieu urbain et notamment les objets urbains.

Nous désignons comme objet urbain tout objet disposé dans l'espace public et lié à un service offert par la collectivité territoriale. Cette définition inclut une grande diversité de cas, parmi lesquels des objets fixes tels que les lampadaires, des objets mobiles comme les poubelles et même des objets vivants comme les arbres urbains.

La mission des gestionnaires de ville est alors d'être en mesure de surveiller l'ensemble des objets urbains, de suivre

l'évolution de leurs statuts au fil du temps ainsi que de prévenir les interactions problématiques entre objets urbains, comme par exemple une chute de branche d'arbres sur des câbles. La difficulté de cette tâche est d'autant plus grande dans les agglomérations urbaines où le nombre important d'objets urbains constitue le défi principal. Il serait alors intéressant d'automatiser en partie la chaîne de traitement et de suivi des objets urbains et en particulier leur recensement.

1.2 Reconnaissance d'objets dans des nuages de points LiDAR

Nos travaux s'inscrivent dans la continuité des recherches sur la télédétection en milieu urbain, réalisées en collaboration avec le laboratoire d'informatique de robotique et de micro-électronique de Montpellier [10, 11]. Les développements récents dans le domaine de la reconnaissance d'objets dans les nuages de points, notamment basés sur l'apprentissage profond, nous encourage à expérimenter la technologie LiDAR terrestre comme dispositif d'acquisition des objets urbains.

En effet, les nuages de points issus d'acquisition LiDAR comptent habituellement plusieurs millions de points et peuvent recouvrir plusieurs kilomètres. De plus la manipulation de données 3D n'étant pas une tâche triviale, un traitement automatique de ceux-ci afin d'extraire les éléments d'intérêt, ici les objets urbains, s'avère une nécessité. Ainsi, nous comptons nous baser sur les avancées récentes afin de développer une méthode capable de réaliser une détection d'objets, c'est-à-dire prédire des boîtes englobantes autour des objets (voir figure 1 b), ce qui est différent de la segmentation sémantique, qui consiste à associer une classe à chaque point du nuage (voir figure 1 a).

Nous présentons en section 2 une revue des travaux connexes, en section 3 nous détaillons notre contribution et nous présentons les diverses expérimentations réalisées autour de celle-ci en section 4.

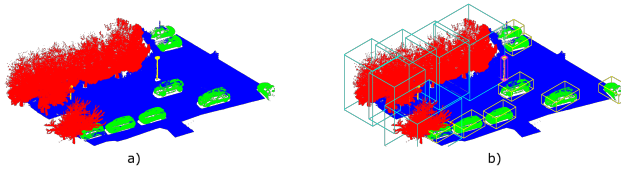


FIGURE 1 – Illustration de la tâche de segmentation sémantique (a) où une classe, représentée par une couleur, est associée à chaque point, et de la tâche de détection d'objets (b), où des boîtes englobantes sont définies autour des objets présents.

2 Travaux connexes

2.1 Apprentissage profond à partir de nuages de points 3D

2.1.1 Approches par structuration

Les premières tentatives de généralisation des réseaux *CNN* aux nuages de points 3D s'appuient sur une structuration préalable des nuages de points. Ainsi le réseau *VoxNet* [8] propose de discrétiser le nuage de points sous forme de voxels avant de le passer en entrée d'un *CNN*, dont l'architecture est légèrement modifiée afin de prendre en compte la dimension supplémentaire. Le réseau *MVCNN* [20] propose quant à lui de simuler des photos, appelées vues, prises tout autour d'un nuage de points et de passer celles-ci en entrée d'un *CNN* classique traitant des images 2D et de fusionner les différentes sorties entre elles.

Bien que ces méthodes par structuration aient permis d'adapter les méthodes de *deep-learning* aux nuages de points et d'obtenir les premiers résultats, elles comportent des limites importantes. Elles ont globalement pour conséquence de discrétiser la forme des nuages de points, causant ainsi une perte d'informations. De plus, les approches par multi-vues sont peu adaptées aux grands nuages de points avec beaucoup d'occultations et nécessitent souvent un maillage des nuages de points. Les approches volumiques souffrent quant à elles d'un coût en mémoire accrue causée par l'utilisation de convolutions volumiques. A noter sur ce point que les convolutions éparses (*sparse convolutions*) [24] permettent de réduire considérablement l'impact de ce défaut.

2.1.2 PointNet et ses successeurs

Le réseau *PointNet* [14] a introduit une architecture capable de réaliser l'apprentissage directement à partir du nuage de points sans transformation intermédiaire. Cette architecture repose en grande partie sur l'utilisation de perceptrons multi-couches (*multi-layer perceptrons*). Ceux-ci sont appliqués indépendamment sur chaque point du nuage en prenant uniquement les coordonnées (x, y, z) des points en entrée, permettant ainsi d'obtenir des vecteurs de *caractéristiques ponctuelles* pour chaque point, en rouge dans la figure 2. Une opération d'agrégation maximale est ensuite appliquée sur ces vecteurs afin d'obtenir le vecteur de caractéristique global, en vert dans la figure 2.

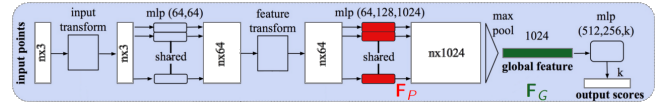


FIGURE 2 – Schéma de l'architecture du réseau *PointNet* [14], avec en rouge les vecteurs de *caractéristiques ponctuelles* et en vert le vecteur de *caractéristiques globales*.

La principale limite du réseau *PointNet* [14] est le manque de prise en compte du voisinage lors du calcul des *caractéristiques ponctuelles*, en effet celles-ci sont calculées pour chaque point indépendamment des autres or l'utilisation du voisinage est importante dans certaines applications comme la segmentation sémantique. Plusieurs méthodes ont alors proposé d'étendre *PointNet* [14] afin d'y ajouter des mécanismes permettant l'extraction de caractéristiques locales.

Le réseau *PointNet++* [15] propose d'appliquer le réseau *PointNet* [14] localement à plusieurs échelles en sous-échantillonnant à chaque étape le nuage de points. Le réseau *RandLaNet* [4] reprend cette idée en choisissant un sous-échantillonnage aléatoire plutôt que celui basé sur le *farthest point sampling (FPS)* et en le combinant avec un module d'attention. Le réseau *DGCNN* [22] propose d'utiliser un graphe dynamique afin de lier les *caractéristiques ponctuelles* entre elles. *SPG* [6] propose de partitionner le nuage de points en *super-points* liés entre eux dans une structure de graphe. *KP-Conv* [21] et *ConvPoint* [1] proposent d'étendre l'opération de convolution aux nuages de points 3D en définissant des noyaux de convolutions à partir de points.

2.2 Détection d'objets dans des nuages de points

Les méthodes post-*PointNet* [14] donnent de très bons résultats pour la tâche de segmentation sémantique, notamment dans le cas de nuages de points en extérieur. Néanmoins, ceux-ci ne permettent pas de réaliser la tâche de détection d'objets car elles ne comportent pas de notion de groupement de points : les traitements sont effectués uniquement au niveau des points (*caractéristiques ponctuelles*) et au niveau du nuage entier (*caractéristiques globales*). Ainsi, les méthodes de détection d'objets dans des nuages de points ont pour caractéristique d'ajouter un niveau intermédiaire représentant des groupements de points qui peuvent, ou non, contenir un objet. Dans la suite nous appelons ces groupements de points *régions* et nous distinguons plusieurs familles de méthodes selon leur manière de définir les *régions*.

2.2.1 Approches par vue d'oiseau

Les approches par vue d'oiseau s'appuient sur un concept simple : si nous observons un nuage de points 3D selon l'axe vertical, il est tentant de schématiser celui-ci sous la forme d'une image aérienne. Ainsi le nuage de points est partitionné selon une grille régulière horizontale et pour chaque cellule de cette grille, des *régions* sont définies en suivant le modèle des ancres (*anchor boxes*) utilisées dans

la détection d'objets dans des images [16]. Il est alors possible d'entraîner un réseau *CNN* à détecter des objets à partir de cette grille.

Le réseau *AVOD* [5] propose de combiner une grille horizontale (image aérienne) et une grille verticale (image frontale) et d'utiliser un premier *CNN* pour détecter les objets présents à partir de chacune des deux grilles, puis un second *CNN* afin de fusionner les prédictions. Les réseaux *VoxelNet* [26] et *PointPillars* [7] proposent d'utiliser *PointNet* [14] afin d'encoder la valeur des cellules de la grille horizontale en fonction des points qu'elles contiennent.

2.2.2 Projection d'images

Les approches par projection d'images utilisent, en plus du nuage de points, une photo rigoureusement étalonnée par rapport au nuage de points. Il est alors possible de faire le lien entre les pixels de cette photo et l'espace géométrique dans lequel est défini le nuage de points. La méthodologie est alors simple : un détecteur d'objets est appliqué sur la photo et permet d'obtenir des boîtes englobantes 2D. Ces boîtes sont ensuite projetées dans l'espace géométrique du nuage afin d'obtenir des boîtes englobantes 3D : il s'agit des *régions*. Enfin, les *régions* sont passés en entrée d'un réseau *PointNet* [14] afin de détecter les objets qui s'y trouvent.

Le réseau *FrustumNet* [13] est la première méthode introduisant cette méthodologie. Celui-ci a été plusieurs fois repris et amélioré, par exemple en appliquant le réseau *PointNet* plusieurs fois pour chaque *région* en "glissant" selon l'axe horizontal [23].

2.2.3 Génération de régions à partir des points

Les méthodes de génération de *régions* consistent à travailler directement sur les points à la manière de *PointNet*. Ainsi le réseau *PointRCNN* [19] propose dans un premier temps de segmenter le nuage de points entre les points de premier plan, appartenant à un objet, et les points d'arrière-plan, appartenant au fond. Dans un deuxième temps, des *régions* sont définies autour des points de premier plan et un réseau *PointNet* [14] est appliqué afin d'y détecter les objets présents. Le réseau *PointRGCN* [25] propose d'ajouter un mécanisme de *graph-convolution* afin de lier les *régions* entre elles et d'exploiter les relations pertinentes pour la détection d'objets. Le réseau *VoteNet* [12] propose de faire "voter" les points du premier plan, c'est à dire de prédire pour chacun d'entre eux un nouveau point censé être le centre de l'objet le plus proche. Les *régions* sont alors obtenues en groupant les *votes* les plus proches.

2.3 Bilan

Cette revue de l'état de l'art nous permet de présenter les familles de méthodes permettant d'obtenir de bons résultats sur les tâches de détection d'objets dans des nuages de points 3D en extérieur. Cependant aucune de ces méthodes n'a, à notre connaissance, été appliquée au problème de la détection d'objets urbains. De plus nous pouvons objecter que les méthodes par vue d'oiseau, bien que très performante dans le cas des voitures autonomes, ne sont pas bien adaptées aux problèmes des objets urbains où il est fréquent de retrouver des objets 'superposés' les uns sur les autres

(par exemple des lampadaires et des arbres ou bien des poubelles et des abri-bus). Quant aux méthodes par projection d'images, la nécessité d'utiliser une caméra peut être un frein dans la mesure où celle-ci donne une moins grande résolution pour les objets les plus éloignés de la caméra comme les arbres ou les caténaires.

Nous nous donnons alors comme objectif d'étendre les méthodes de génération de *régions* en unifiant leur processus d'apprentissage au lieu d'utiliser deux sous réseau entraînés séquentiellement, ainsi qu'en définissant des *régions* de formes arbitraires au lieu d'utiliser des boîtes englobantes, afin d'être plus proche de la forme des objets urbains.

3 L'architecture RPC

Dans cette section nous présentons notre architecture générale pour la détection d'objets urbains dans des nuages de points 3D. Celle-ci se nomme *RPC* pour réseau avec proposition de *cluster*. En effet, notre architecture se base sur la définition de *régions* de formes arbitraires appelées *clusters*. Celle-ci peut être entraînée de bout-en-bout et prend en entrée directement les points du nuage.

Notation : Dans la suite nous désignons un point par $\mathbf{P} \in \mathbb{R}^{d+3}$, qui peut être décomposé en sa partie géométrique $\mathbf{P}_g \in \mathbb{R}^3 = (x, y, z)$ et sa partie caractéristique $\mathbf{P}_f \in \mathbb{R}^d$, qui peut être la couleur du point, son intensité ou bien tout autre vecteur de caractéristiques associé au point. Nous notons un nuage de N points comme $S = \{\mathbf{P}_i\}_1^N$.

3.1 Opération de cluster-convolution

Notre architecture se base sur la définition de l'opération originale de *cluster-convolution*. C'est cette opération qui permet d'extraire les informations nécessaires à la génération des *régions* pour la détection d'objets. Cette opération prend en entrée un nuage de points S_l et en extrait des caractéristiques pour aboutir à un nuage de points S_{l+1} plus petit mais plus "profond". Elle peut être formalisée comme ceci :

$$ClConv : \begin{cases} \mathbb{R}^{N_l \times (3+d_l)} \rightarrow \mathbb{R}^{N_{l+1} \times (3+d_{l+1})} \\ S_l \rightarrow S_{l+1} \end{cases} \quad (1)$$

où $d_{l+1} > d_l$ et $N_l > N_{l+1}$. L'opération de *cluster-convolution* est schématisée dans la figure 3. Cette opération peut être appliquée successivement en diminuant à chaque fois la taille du nuage et en augmentant la taille du vecteur de caractéristiques extraites pour chaque point.

Cette opération se décompose en 2 couches de calcul :

- la couche de *clustering*, c'est elle qui permet de diminuer la taille du nuage de points en créant des groupes de points appelés *clusters*
- la couche de *graph-convolution*, c'est elle qui permet d'extraire de nouvelles caractéristiques en exploitant les liens entre les points appartenant à un même *cluster*.

3.1.1 La couche de clustering

Le but de la couche de *clustering*, schématisée en figure 4, est de créer des groupements de points cohérents à partir

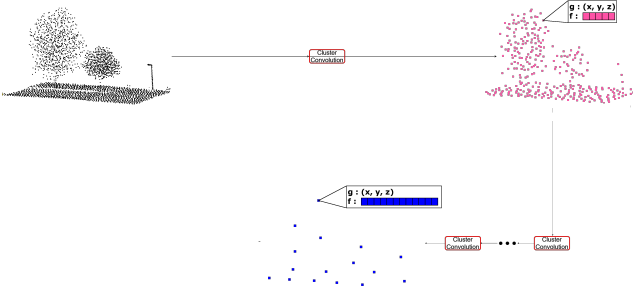


FIGURE 3 – Schéma de l'opération de *cluster-convolution* : un nuage de points est passé en entrée de celle-ci, obtenant en sortie un nuage de points de plus petite taille avec un nouveau vecteur de caractéristiques pour chaque point. Cette opération peut être appliquée successivement impliquant à chaque fois un nuage de plus petite taille avec des vecteurs de caractéristiques de dimension supérieur.

d'un nuage de points, ces groupements se basent sur la définition d'un nombre fixe de *clusters*, notés K et formant une partition du nuage de points original :

$$\bigcup_i K_i = S \quad (2)$$

$$\bigcap_i K_i = \emptyset \quad (3)$$

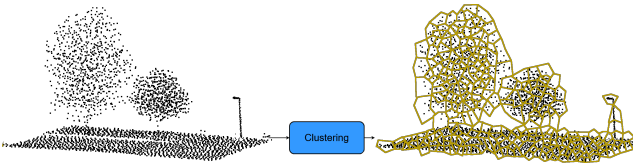


FIGURE 4 – Schéma de la couche de *clustering* avec le nuage de points en entrée qui est partitionnée en plusieurs *clusters* par l'application de la couche.

La cohérence du *clustering* est assuré par le fait que les points regroupés entre eux doivent être aussi similaires que possible aussi bien au niveau de leurs parties géométriques, que de leurs parties caractéristiques. Afin de vérifier au mieux cette propriété, nous nous sommes basés sur l'algorithme du *k-means* afin d'implémenter la couche de *clustering*. La différence principale avec une implémentation classique du *k-means* est la fonction de distance utilisée. En effet, pour pouvoir équilibrer les parties géométriques et caractéristiques dans le calcul des *clusters*, nous avons expérimenté avec plusieurs fonctions de distance.

Distance euclidienne Le choix naïf consiste à utiliser une distance euclidienne :

$$D_e(\mathbf{P}_1, \mathbf{P}_2) = \|\mathbf{P}_1, \mathbf{P}_2\|$$

où $\|\cdot\|$ représente la norme euclidienne. Ici aucune distinction entre partie géométrique et caractéristique n'est faite, ce qui a pour conséquence de déséquilibrer la distance. En effet, mis à part lors de la première étape de *cluster-convolution*, la partie caractéristique des points compte plus

de composantes que la partie géométrique et aura donc tendance à s'imposer uniformément. C'est pourquoi nous avons expérimenté plusieurs autres fonctions de distance, permettant d'équilibrer les deux parties.

Distance normalisée Afin d'équilibrer le calcul de distance, une approche peut être de normaliser les parties géométriques et caractéristiques des points en les considérant comme indépendantes :

$$D_{norm}(\mathbf{P}_1, \mathbf{P}_2) = \frac{\|\mathbf{P}_{g_1} - \mathbf{P}_{g_2}\|}{\sigma_g} + \frac{\|\mathbf{P}_{f_1} - \mathbf{P}_{f_2}\|}{\sigma_f}$$

où σ est la variance des séries associées.

Distance pondérée Cette approche consiste à calculer séparément les distances des parties géométriques et caractéristiques, puis de faire la moyenne pondérée de celles-ci :

$$D_\lambda(\mathbf{P}_1, \mathbf{P}_2) = \lambda D_e(\mathbf{P}_{g_1}, \mathbf{P}_{g_2}) + (1 - \lambda) D_e(\mathbf{P}_{f_1}, \mathbf{P}_{f_2})$$

où λ est le coefficient de pondération.

Distance maximale En s'inspirant du processus d'apprentissage de *PointNet* [14], nous avons défini la distance par composante maximale. Cette distance consiste à calculer une distance euclidienne sur les parties géométrique, augmentée du carré de la différence entre les composantes maximales des parties caractéristiques :

$$D_{max}(\mathbf{P}_1, \mathbf{P}_2) = \lambda D_e(\mathbf{P}_{g_1}, \mathbf{P}_{g_2}) + (1 - \lambda) \left(\max_{\mathbf{P}_{f_1}} - \max_{\mathbf{P}_{f_2}} \right)^2$$

où $\max_{\mathbf{P}_f}$ est la composante maximale du vecteur caractéristique $\mathbf{P}_f$.

Expérimentalement, nous avons déduit que les distances maximales et normalisées permettent un meilleur équilibre entre les parties géométriques et caractéristiques, plus de détails est donné en section 4.

3.1.2 La couche de *graph-convolution*

L'objectif de la couche de *graph-convolution* est d'extraire des caractéristiques pertinentes pour la détection d'objets à partir des points du nuage. Cette opération est appliquée à la suite de la couche de *clustering* et prend en entrée les *clusters* de points. Le but est alors de représenter chaque *cluster* par un unique point dont la partie caractéristique prend en compte la répartition des points dans le *cluster* :

$$G_{conv} : K^l \rightarrow \mathbf{P}^{l+1}$$

où $\mathbf{P}^{l+1}$ est un nouveau point n'appartenant pas au nuage de point original S .

La définition de l'opération de *graph-convolution* diffère donc des opérations calculées par *PointNet* et *DGCNN* (voir comparatif dans la figure 5).

La fonction calculée par *PointNet* [14] est telle que :

$$PointNet : S \rightarrow \mathbf{F}$$

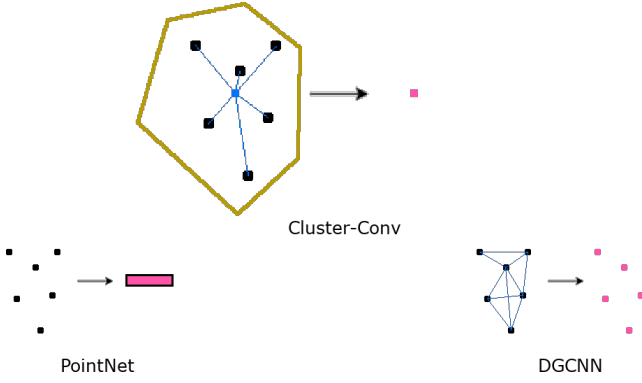


FIGURE 5 – Comparaison entre la couche de *graph-convolution* et les fonctions calculées par *PointNet* [14] et *DGCNN* [22]

où $\mathbf{F}$ est un vecteur de caractéristique non associé à une partie géométrique, ce qui empêche d'appliquer cette opération de manière successive.

La fonction de *convolution* calculée par *DGCNN* est telle que :

$$DGCNN : \{\mathbf{P}_i\}_1^N \in \mathbb{R}^{3+d_l} \rightarrow \{\mathbf{P}_i\}_1^N \in \mathbb{R}^{3+d_{l+1}}$$

Cette opération permet de calculer des nouvelles caractéristiques pour chaque point mais sans notion de groupes et donc ne permettant pas de restreindre le nombre de points à chaque application.

La procédure de calcul de l'opération de *graph-convolution* est décrite dans la figure 6, celle-ci est décomposée en plusieurs étapes :

- la couche prend en entrée les *clusters* de points précédemment définis
- pour chaque *cluster*, la moyenne des parties géométriques est calculée afin d'obtenir la partie géométrique du résultat $\mathbf{P}_g^{l+1}$
- le vecteur de caractéristique maximal $\mathbf{F}$ est calculé, puis est concaténé à chacune des parties caractéristiques du *cluster* et passé en entrée d'un perceptron multi-couches
- une opération d'agrégation maximale est appliquée sur la sortie du perceptron multi-couches afin d'obtenir la partie caractéristique du résultat $\mathbf{P}_f^{l+1}$

L'opération de *cluster-convolution* dépend ainsi de deux paramètres principaux :

1. le nombre de *clusters* k (décroissant au fil des applications successives)
2. le nombre de caractéristiques extraites (croissant au fil des applications)

3.2 Prédiction de boîtes englobantes

Après la dernière application de *graph-convolution*, nous obtenons un nuage de points S^l de taille bien inférieure au nuage de points en entrée S . Nous pouvons alors calculer une régression à partir des parties géométriques des

points de S^l vers des vecteurs représentant des boîtes englobantes. Les boîtes englobantes sont ainsi définies :

$$\mathbf{B} = (x, y, z, h, w, d, \theta)$$

où (x, y, z) est le centre de la boîte englobante, h sa hauteur, w sa largeur, d sa profondeur et θ son angle d'orientation par rapport à l'axe vertical.

La cible de cette opération de régression peut être déterminée facilement à partir des points $\mathbf{P}_i \in S^l$, en effet comme ceux-ci sont obtenus par application successive de *cluster-convolution*, il suffit de garder en mémoire l'affectation des *clusters* à chaque application de *cluster-convolution* pour pouvoir remonter au nuage de points original. Ainsi, nous pouvons à partir d'un point $\mathbf{P} \in S$ remonter à un sous-ensemble de S , c'est ce sous-ensemble que nous utilisons comme *région* pour la détection d'objets.

Nous considérons que chaque *région* peut au maximum contenir un seul objet que nous appelons l'objet majoritaire, noté O_{maj} . Cet objet est celui dont la proportion de points appartenant à la *région* par rapport à l'objet entier est maximal. La *région* est considérée comme positive si cette proportion est supérieure au seuil τ et négative dans le cas contraire.

La fonction de coût est alors définie en fonction de l'objet majoritaire pour chaque *région* :

$$\begin{aligned} \mathcal{L}(S^l, \mathbf{B}^*) &= \frac{\alpha}{N_l} \sum_i^{N_l} \mathcal{L}_{cls}(\mathbf{B}_i, \mathbf{B}_{O_{maj}}^*) \\ &+ \frac{\beta}{N_l} \sum_i^{N_l} \mathcal{L}_{reg}(\mathbf{B}_i, \mathbf{B}_{O_{maj}}^*) \end{aligned}$$

c'est-à-dire une somme pondérée d'une fonction de classification, $\mathcal{L}_{cls}$, pour déterminer la classe des boîtes englobantes, et d'une fonction de régression, $\mathcal{L}_{reg}$, pour déterminer leurs géométries. Les boîtes englobantes vérité terrain sont notées $\mathbf{B}^*$.

La fonction $\mathcal{L}_{cls}$ est une fonction d'entropie croisée classique :

$$\mathcal{L}_{cls}(x, y) = - \sum_{C=1}^{N_C} y_C \log(x_C)$$

où y_C vaut 1 quand la classe C est la classe attendue et 0 sinon.

La fonction $\mathcal{L}_{reg}$ est une fonction de Huber :

$$\mathcal{L}_{Huber}(x, y, \delta) = \begin{cases} \frac{1}{2}(x - y)^2, & \text{si } |x - y| \leq \delta \\ \delta(|x - y|) - \frac{1}{2}\delta, & \text{sinon} \end{cases}$$

où $|x - y|$ est la valeur absolue de $x - y$.

4 Résultats expérimentaux

Dans cette section, nous présentons les résultats que nous avons obtenus sur la tâche de détection d'objets urbains dans des nuages de points 3D à partir de notre architecture *RPC*.

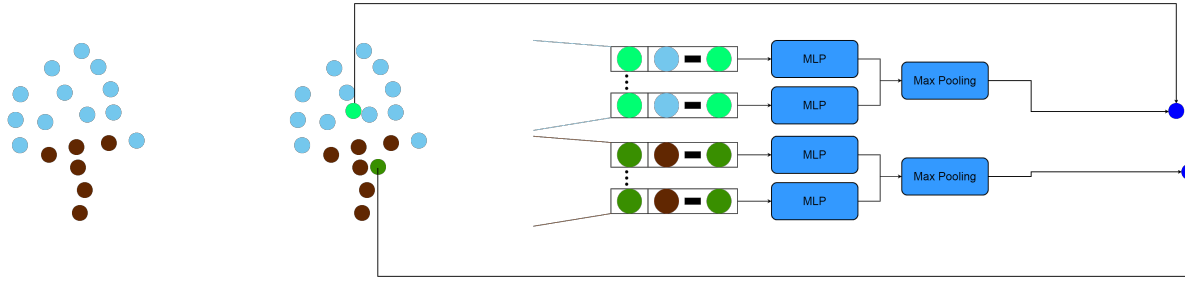


FIGURE 6 – Schéma du calcul de la couche de *graph-convolution*. Les *clusters* sont passés en entrée afin de calculer la partie géométrique du résultat, puis des perceptrons multi-couches sont suivi d’une agrégation maximale sont appliqués dans chaque *cluster* afin d’obtenir la partie caractéristique du résultat.

	Classe	poteau	personne	voiture	arbre
entraîn- ement	Paris	11	17	10	28
	Lille A	12	0	24	15
validation	Lille B	7	1	29	16
	Total	30	18	63	59

TABLE 1 – Bilan de la répartition des objets dans le jeu de données utilisé.

4.1 Jeu de données utilisé

La segmentation sémantique dans les nuages de points en extérieur est une tâche bien étudiée avec plusieurs jeux de données disponibles [18, 3, 9]. Cependant il n’en est pas de même pour la détection d’objets où il n’existe, à notre connaissance, pas de jeux de données publics, mis à part ceux orientés voitures autonomes qui ont leurs propres contraintes [2].

Nous avons alors décidé de nous baser sur le jeu de données *ParisLille* [17] car bien qu’initialement développé pour la tâche de segmentation sémantique, celui-ci possède des annotations au niveau instance. Nous utilisons alors une méthode semi-automatique afin de générer les boîtes englobantes associés aux objets et de les ajuster autour d’eux.

Le jeu de données qui en résulte est composé de 3 scènes issues d’acquisitions *LiDAR* différentes, dont 2 sont utilisés comme base d’entraînement, *Paris* et *Lille A*, et une pour la validation, *Lille B*. Nous comptons au total 170 objets répartis en 4 classes. Le détail de la répartition des objets est présenté dans le tableau 1.

Comme ces scènes comptent plusieurs millions de points chacune et s’étendent sur plusieurs dizaines de mètres, nous effectuons un pré-traitement afin de pouvoir les passer en entrée de notre architecture. Les nuages de points sont découpés en blocs de dimension égaux et sous-échantillonnés à 2048 points.

4.2 Implémentation

Le modèle que nous utilisons est basé sur notre architecture *RPC*, celui-ci se compose de 3 opérations de *cluster-convolution* pour la génération des *régions*, suivi de perceptrons multi-couches pour la prédiction des boîtes englobantes. Les détails du modèle sont présentés dans la figure

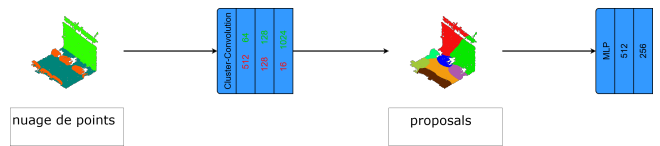


FIGURE 7 – Architecture du modèle utilisé pour nos expériences, avec en rouge le nombre de *clusters* et en vert le nombre de caractéristiques extraites pour chaque application de *cluster-convolution*.

	Moyenne	Signalisation/Poteau	Voiture	Arbre
Segmentation sémantique + clustering hiérarchique	0.29	0.11	0.64	0.11
<i>RPC</i> (notre architecture)	0.21	0	0.52	0.10
<i>VoteNet</i>	0.04	0	0.09	0.03

TABLE 2 – Résultats, en termes de précision moyenne (*average precision AP*), sur la scène de validation.

7. Le modèle est implémenté à l’aide du *framework TensorFlow*. Afin d’être dérivable, le nombre d’itération est fixé à 10 lors du calcul du *k-means* de la couche de *clustering*. La valeur du seuil τ pour la définition de l’objet majoritaire est fixé à 0,5 et nous prenons $\beta = 10 \times \alpha$ dans le calcul de la fonction de coût. Les exemples de *régions* positives et négatives sont équilibrés lors de l’apprentissage et un algorithme de *non maximum suppression (NMS)* est appliqué afin de filtrer les boîtes prédites.

4.3 Résultats

Nous entraînons notre modèle pendant 200 epochs¹ sur notre base d’entraînement et nous évaluons celui-ci sur notre jeu de validation avec la métrique *AP*. Les résultats sont reportés dans le tableau 2 et des visualisations de prédictions du réseau sont illustrées dans les figures 8 à 11.

Comme il n’existe pas de *benchmark* pour la détection d’objets urbains, nous utilisons comme comparatif une méthode en deux étapes : segmentation sémantique suivi de *clustering* adaptatif. Cependant, comme le *clustering* adap-

1. Le temps d’entraînement est de 8 heures environ avec un GPU Titan Maxwell.

tatif est optimisé directement sur le jeu de validation, les résultats sont biaisés en la faveur de la méthode en deux étapes. En ce qui concerne l'évaluation de la méthode *VoteNet* [12], celle-ci étant optimisée pour des scènes en intérieur, nous avons remarqué que le réseau essayait d'apprendre une répartition spatiale des objets dans l'espace, or celle-ci est totalement aléatoire dû à la pré-segmentation en blocs, biaisant ainsi fortement les résultats en défaveur de *VoteNet*. Globalement, les résultats demeurent plutôt bas, témoignant ainsi de la difficulté de la tâche, mis à part pour la classe voiture. Les performances de notre architecture sont proches de celles de la méthode en deux étapes, à part pour la classe signalisation et nous pouvons nous attendre à ce qu'elles les dépassent en s'appuyant sur un jeu d'entraînement plus grand.

vérité terrain

prédiction

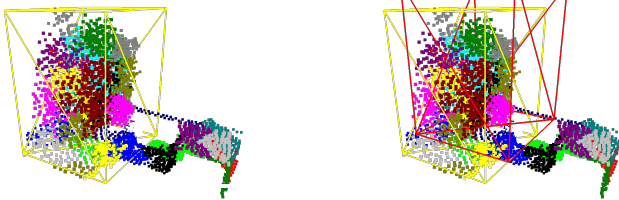


FIGURE 8 – Visualisation d'un exemple positif de prédiction d'un arbre par notre modèle. Chaque couleur correspond à une unique *région*, les boîtes englobantes jaunes sont la vérité terrain, les rouges sont les prédictions du réseau.

Dans la figure 8 nous pouvons voir que bien que l'arbre soit décomposé en plusieurs *régions* différentes, le réseau est tout de même capable de prédire une boîte englobante valide à partir d'une des *régions* située au centre de l'arbre.

vérité terrain

prédiction

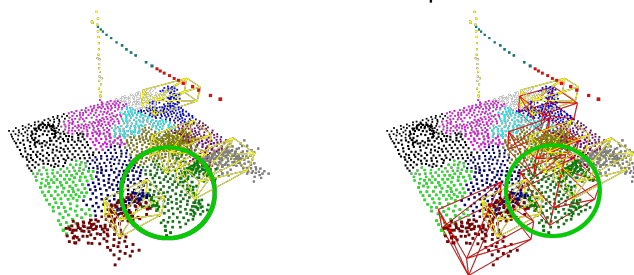


FIGURE 9 – Visualisation d'un exemple de prédiction négative du réseau dû à un chevauchement de 2 voitures par la *région*. Chaque couleur correspond à une unique *région*, les boîtes englobantes jaunes sont la vérité terrain, les rouges sont les prédictions du réseau.

Dans la figure 9, nous pouvons distinguer que la *région* est 'à cheval' sur deux voitures entraînant ainsi une erreur de

vérité terrain

prédiction

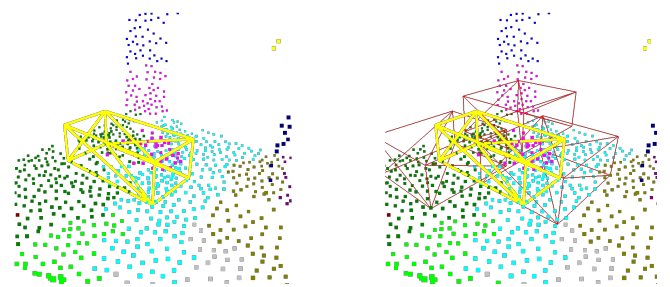


FIGURE 10 – Visualisation d'un exemple de prédiction négative du réseau dû à une voiture, chevauchant plusieurs *régions*. Chaque couleur correspond à une unique *région*, les boîtes englobantes jaunes sont la vérité terrain, les rouges sont les prédictions du réseau.

vérité terrain

prédiction

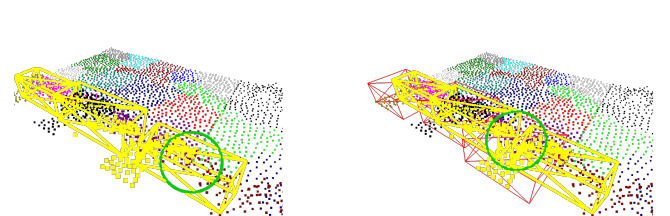


FIGURE 11 – Visualisation d'un exemple de prédiction négative du réseau dû à une voiture chevauchant deux blocs différents. Chaque couleur correspond à une unique *région*, les boîtes englobantes jaunes sont la vérité terrain, les rouges sont les prédictions du réseau.

prédiction du réseau avec une boîte englobante se situant entre les deux voitures. De même dans la figure 10, où cette fois un même objet, une voiture, est recoupé par plusieurs *régions* dont aucune n'est suffisamment proche de la forme de la voiture, impliquant ainsi une erreur de prédiction du réseau avec 3 boîtes englobantes prédites dont aucune ne recoupe correctement la voiture. Un autre type d'erreur est dû à la segmentation de la scène en blocs, ainsi dans la figure 11, un même objet est 'à cheval' sur deux blocs causant une erreur de prédiction avec une boîte englobante décalée par rapport à la voiture.

Nous pouvons alors formuler l'hypothèse suivante : les mauvaises prédictions du réseau sont en grande partie causées par des mauvaises répartition de *régions* et inversement, qu'à partir d'une bonne répartition des *régions*, les prédictions du réseau ont tendance à être valide.

4.4 Expériences supplémentaires

Afin d'évaluer notre hypothèse nous proposons un indice de qualité des *régions* mesurant à quel point une *région* recoupe un unique objet. Cet indice est défini pour une *région*

R comme ceci :

$$qualite = \frac{Card(O_{maj} \cap R)}{Card(O_{maj} \cup R)}$$

c'est-à-dire la proportion de points appartenant à l'objet majoritaire par rapport au nombre total de points de l'objet majoritaire et de la *région*.

Nous calculons alors cet indicateur pour chaque *région* du réseau sur le jeu de validation et nous le mettons en lien avec la validité de la boîte englobante prédite à partir de cette même *région*. Cette prédiction est évaluée par l'intersection sur l'union (*intersection over union* ou *IoU*) entre la boîte prédite et la boîte vérité terrain de l'objet majoritaire. Les résultats sont reportés dans le diagramme de la figure 12.

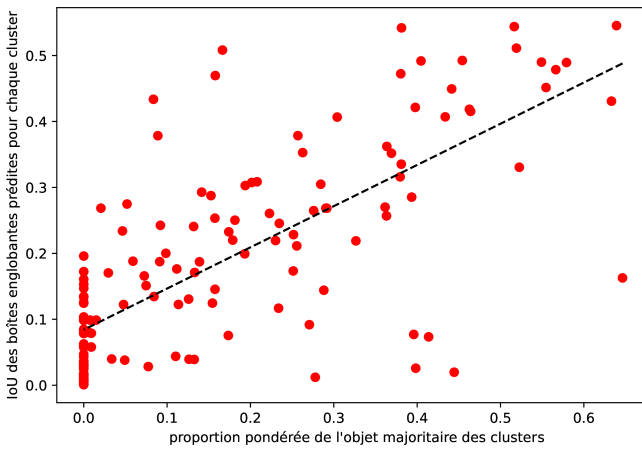


FIGURE 12 – Diagramme de l'intersection sur l'union des prédictions du réseau en fonction de la proportion d'objet majoritaire dans la *région* associée. Chaque point représente une unique *région*.

Nous remarquons alors qu'il existe bel et bien un lien entre ces deux grandeurs, qui peut par exemple être évalué avec un indice de corrélation de Pearson de 0.73, d'autant plus que les points *outliers* correspondent souvent à des erreurs de classification du réseau *i.e.* le réseau a reconnu la présence d'un objet mais pas sa classe.

A partir d'une répartition des *régions* correcte, le réseau a alors tendance à prédire des boîtes englobantes valides. Pour aller plus loin, nous proposons d'étudier plus en profondeur les paramètres qui régissent la répartition des *régions*.

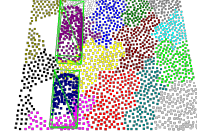
4.4.1 Équilibrage de la fonction de distance

La répartition des *régions* dépend essentiellement de l'application de la couche de *clustering*. L'élément central de la couche de *clustering* est la fonction de distance calculée par l'algorithme du *k-means* (voir section 3 pour plus détails sur celui-ci). L'équilibrage entre partie géométrique et partie caractéristique joue ainsi un rôle crucial dans le calcul de cette distance.

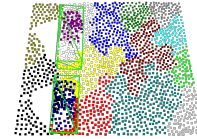
Ce rôle peut être montré en fixant l'initialisation du *k-means* pour un même bloc et en faisant varier uniquement le coefficient λ (plus celui-ci est élevé plus la partie géométrique

est importante et inversement). Une visualisation des répartitions ainsi obtenues pour plusieurs valeurs de λ est en figure 13 et l'évolution de l'indice de qualité des 2 *régions* qui recoupent le mieux les 2 voitures est présentée dans le diagramme 14.

lambda :
0.9



lambda :
0.5



lambda :
0.1

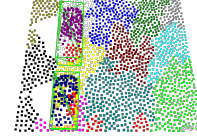


FIGURE 13 – Visualisation de la répartition des *régions* dans un même blocs, selon le coefficient λ de la fonction de distance.

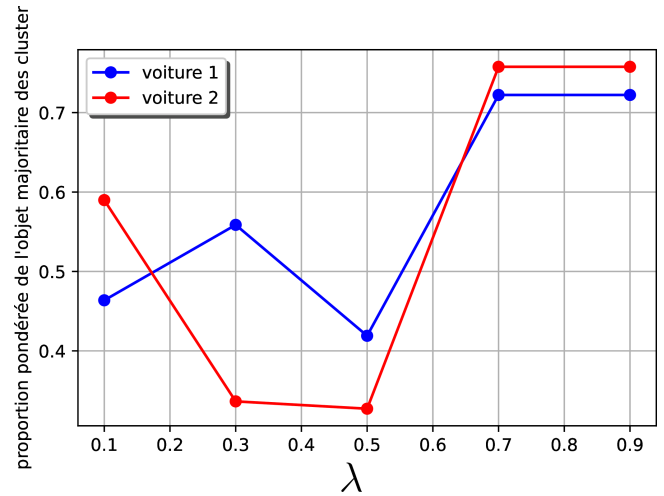


FIGURE 14 – Diagramme de la variation de l'indice de qualité de deux *régions* recoupant une voiture différente chacun, selon la valeur du coefficient λ de la fonction de distance.

Nous remarquons ainsi que cette répartition varie grandement selon la valeur de λ , avec des *régions* de bonne qualité pour des valeurs de λ supérieur à 0,7 et inférieur à 0,2 et de mauvaise qualité entre ces 2 valeurs. Néanmoins, il est difficile de tirer des conclusions globales, en effet les effets de la variation de λ sur la répartition des *régions* est clairement non uniforme, voir chaotique, et varie grandement d'un bloc à l'autre, ou selon la nature des objets présents.

4.4.2 Nombre de régions

L'autre paramètre régissant la répartition des *régions* est leur nombre, noté k , c'est-à-dire le nombre de *clusters* calculé lors de la dernière application de *cluster-convolution*. Nous fixons l'initialisation du *k-means* et faisons varier la valeur k . A chaque valeur de k nous calculons la valeur moyenne de qualité des *régions*, reportée dans le diagramme 16, et des exemples de visualisation sont présentés en figure 15.

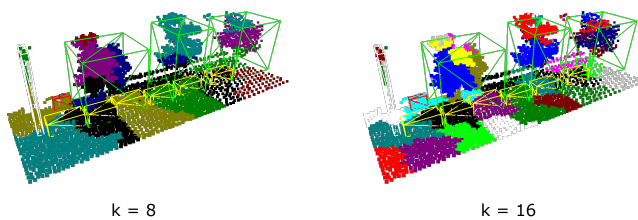


FIGURE 15 – Visualisation de la répartition des *régions* selon la valeur de k .

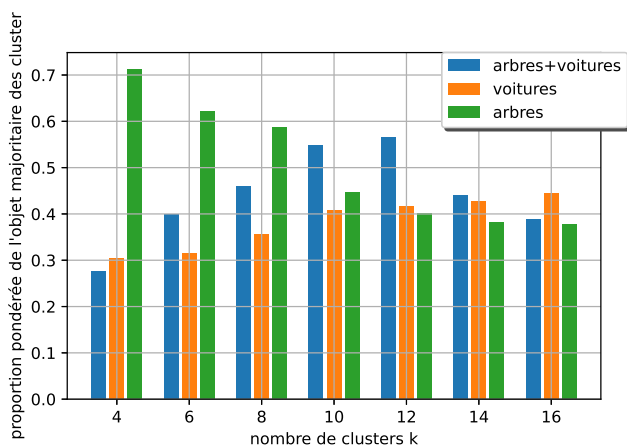


FIGURE 16 – Diagramme de la variation de l'indice de qualité des *régions* selon la valeur de k et la classe des objets présents en majorité dans les blocs. Comme nous pouvons le constater dans le diagramme, la valeur de k exerce une influence directe sur la qualité des *régions*. Ainsi nous pouvons voir que celle-ci est liée à la taille des objets présents dans les nuages de points. Dans des blocs contenant en majorité des petits objets, comme des voitures, des grandes valeurs de k donnent les meilleurs résultats, et inversement pour les blocs contenant majoritairement des grands objets, comme des arbres, des petites valeurs de k donnent les meilleurs résultats. Il est alors difficile de choisir a priori, une valeur de k qui serait adapté à tous les blocs. C'est pourquoi nous fixons expérimentalement celle-ci à 8.

5 Conclusion

Nous avons présenté nos travaux portant sur la détection d'objets par apprentissage profond, ceux-ci ont abouti au développement d'une nouvelle architecture permettant la

détection d'objets urbains dans des nuages de points *LiDAR*. Notre architecture se base sur l'opération originale de *cluster-convolution* permettant de simultanément extraire des caractéristiques à partir des nuages de points et de réduire leurs tailles. Nous avons montré que la répartition des *régions* générées par notre architecture exerce une influence majeure sur ses performances et comment celle-ci repose en partie sur les hyper-paramètres λ et k . L'importance de ces hyper-paramètres constitue une limite de notre architecture, dans la mesure où il serait préférable d'intégrer ceux-ci dans le réseau afin qu'ils soient optimisés lors du processus d'apprentissage. Une autre limite provient de la petite taille de notre jeu de données utilisées lors de nos évaluations expérimentales.

5.1 Perspectives

Afin de poursuivre nos travaux nous comptons utiliser des méthodes d'apprentissage semi-supervisé, comme le *deep-embedding clustering (DEC)* dans notre couche de *clustering* au lieu de l'algorithme du *k-means* dans le but de limiter l'influence des hyper-paramètres et d'accroître la robustesse de notre architecture. Nous comptons également adapter des jeux de données plus volumineux à la tâche de détection d'objets urbains voir d'annoter nos propres données *LiDAR* pour augmenter la portée de nos expériences et tester la généralité de notre approche. Enfin, les approches récentes par mécanismes d'attention pourraient être facilement adapté au problème de la détection d'objets urbains et ainsi être comparé à notre approche.

Références

- [1] Alexandre Boulch. ConvPoint : Continuous convolutions for point cloud processing. *Computers and Graphics (Pergamon)*, 88 :24–34, 2020.
- [2] Andreas Geiger, Philip Lenz, and Raquel Urtasun. Are we ready for autonomous driving? the KITTI vision benchmark suite. In *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pages 3354–3361, 2012.
- [3] T. Hackel, N. Savinov, L. Ladicky, J. D. Wegner, K. Schindler, and M. Pollefeys. Semantic3D.Net : a New Large-Scale Point Cloud Classification Benchmark. In *ISPRS Annals of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, volume 4, pages 91–98, 2017.
- [4] Qingyong Hu, Bo Yang, Linhai Xie, Stefano Rosa, Yulan Guo, Zhihua Wang, Niki Trigoni, and Andrew Markham. Randla-net : Efficient semantic segmentation of large-scale point clouds. In *2020 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR)*, pages 11105–11114, 2020.
- [5] Jason Ku, Melissa Mozifian, Jungwook Lee, Ali Hakeem, and Steven L. Waslander. Joint 3D Proposal Generation and Object Detection from View Aggregation. *IEEE International Conference on Intelligent Robots and Systems*, pages 5750–5757, 2018.

- [6] Loic Landrieu and Martin Simonovsky. Large-Scale Point Cloud Semantic Segmentation with Superpoint Graphs. *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pages 4558–4567, 2018.
- [7] Alex H. Lang, Sourabh Vora, Holger Caesar, Lubing Zhou, Jiong Yang, and Oscar Beijbom. Pointpillars : Fast encoders for object detection from point clouds. *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, 2019-June :12689–12697, 2019.
- [8] Daniel Maturana and Sebastian Scherer. Vox-Net : A 3D Convolutional Neural Network for real-time object recognition. *IEEE International Conference on Intelligent Robots and Systems*, 2015-December :922–928, 2015.
- [9] Daniel Munoz, J. Andrew Bagnell, Nicolas Vandapel, and Martial Hebert. Contextual classification with functional max-margin markov networks. In *2009 IEEE Computer Society Conference on Computer Vision and Pattern Recognition Workshops, CVPR Workshops 2009*, volume 2009 IEEE Computer Society Conference on Computer Vision and Pattern Recognition, pages 975–982, 2009.
- [10] Jérôme Pasquet. *Modélisation, détection et classification d'objets urbains à partir d'images photographiques aériennes*. PhD thesis, Université de Montpellier, 2016.
- [11] Lionel Pibre. *Localisation d'objets urbains à partir de sources multiples dont des images aériennes*. PhD thesis, Université de Montpellier, 2018.
- [12] Charles R. Qi, Or Litany, Kaiming He, and Leonidas Guibas. Deep hough voting for 3D object detection in point clouds. In *Proceedings of the IEEE International Conference on Computer Vision*, volume 2019-October, pages 9276–9285, 2019.
- [13] Charles R. Qi, Wei Liu, Chenxia Wu, Hao Su, and Leonidas J. Guibas. Frustum PointNets for 3D Object Detection from RGB-D Data. *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pages 918–927, 2018.
- [14] Charles R. Qi, Hao Su, Kaichun Mo, and Leonidas J. Guibas. PointNet : Deep learning on point sets for 3D classification and segmentation. In *Proceedings - 30th IEEE Conference on Computer Vision and Pattern Recognition, CVPR 2017*, volume 2017-January, pages 77–85, 2017.
- [15] Charles R. Qi, Li Yi, Hao Su, and Leonidas J. Guibas. PointNet++ : Deep hierarchical feature learning on point sets in a metric space. In *Advances in Neural Information Processing Systems*, volume 2017-December, pages 5100–5109, 2017.
- [16] Shaoqing Ren, Kaiming He, Ross Girshick, and Jian Sun. Faster R-CNN : Towards Real-Time Object Detection with Region Proposal Networks. *IEEE Transactions on Pattern Analysis and Machine Intelligence*, 39(6) :1137–1149, 2017.
- [17] Xavier Roynard, Jean Emmanuel Deschaud, and François Goulette. Paris-Lille-3D : A large and high-quality ground-truth urban point cloud dataset for automatic segmentation and classification. *International Journal of Robotics Research*, 37(6) :545–557, 2018.
- [18] Andrés Serna, Beatriz Marcotegui, François Goulette, and Jean Emmanuel Deschaud. Paris-rue-madame database : A 3D mobile laser scanner dataset for benchmarking urban detection, segmentation and classification methods. In *ICPRAM 2014 - Proceedings of the 3rd International Conference on Pattern Recognition Applications and Methods*, pages 819–824, 2014.
- [19] Shaoshuai Shi, Xiaogang Wang, and Hongsheng Li. PointRCNN : 3D object proposal generation and detection from point cloud. *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, 2019-June :770–779, 2019.
- [20] Hang Su, Subhransu Maji, Evangelos Kalogerakis, and Erik Learned-Miller. Multi-view convolutional neural networks for 3D shape recognition. *Proceedings of the IEEE International Conference on Computer Vision*, 2015 International Conference on Computer Vision, ICCV 2015 :945–953, 2015.
- [21] Hugues Thomas, Charles R. Qi, Jean Emmanuel Deschaud, Beatriz Marcotegui, François Goulette, and Leonidas Guibas. KPConv : Flexible and deformable convolution for point clouds. *Proceedings of the IEEE International Conference on Computer Vision*, 2019-October :6410–6419, 2019.
- [22] Yue Wang, Yongbin Sun, Ziwei Liu, Sanjay E. Sarma, Michael M. Bronstein, and Justin M. Solomon. Dynamic graph Cnn for learning on point clouds. *ACM Transactions on Graphics*, 38(5), 2019.
- [23] Zhixin Wang and Kui Jia. Frustum ConvNet : Sliding Frustums to Aggregate Local Point-Wise Features for Amodal. *IEEE International Conference on Intelligent Robots and Systems*, pages 1742–1749, 2019.
- [24] Yan Yan, Yuxing Mao, and Bo Li. Second : Sparsely embedded convolutional detection. *Sensors (Switzerland)*, 18(10), 2018.
- [25] Jesus Zarzar, Silvio Giancola, and Bernard Ghanem. PointRGCN : Graph Convolution Networks for 3D Vehicles Detection Refinement. *ArXiv*, abs/1911.1, 2019.
- [26] Yin Zhou and Oncel Tuzel. VoxelNet : End-to-End Learning for Point Cloud Based 3D Object Detection. *Proceedings of the IEEE Computer Society Conference on Computer Vision and Pattern Recognition*, pages 4490–4499, 2018.

Modèle et jeu de données pour la détection multi-spectrale de feux de forêt à bord de satellites

H. Farhat^{1,2}, L. Daniel¹, M. Benguigui^{1,3}, A. Girard¹

¹ Institut de Recherche Technologique Saint-Exupéry

² AViSTO

³ Activeeon

Contact : houssem.farhat@irt-saintexupery.com

Résumé

Le nombre de feux de forêts risque d'augmenter de +50% d'ici à 2100. Dans cet article, nous nous intéressons à la télédétection de ces feux à bord d'un satellite afin de lever des alertes précoces directement depuis l'espace. Dans ce cadre, nous avons entraîné un réseau de neurones UNet-MobileNetV3 à segmenter des images multispectrales de Sentinel-2, que nous avons annotées semi-automatiquement puis vérifiées manuellement. Nous avons ensuite déployé ce réseau sur un GPU embarquable sur un satellite en orbite basse. Les données et le réseau entraîné sont disponibles ici : www.irt-saintexupery.com/firesat

Mots-clés

Détection feux de forêts, Imagerie satellite, IA embarquée, Surveillance permanente.

Abstract

The number of wildfires will likely increase by 50% before the end of the century. In this article, we focus on detecting wildfires onboard satellites to raise early warnings directly from space. In this context, we trained a neural network UNet-MobileNetV3 to segment multispectral Sentinel-2 images, which we annotated semi-automatically and verified manually. We then deployed this network to a GPU that can be embedded into a low-orbit satellite. We provide our dataset and the trained network here : www.irt-saintexupery.com/firesat

Keywords

Wildfire detection, Satellite imagery, Embedded AI, Permanent monitoring.

1 Introduction

1.1 Gestion intégrée des forêts : la prévention

La biodiversité des forêts est riche mais en proie aux activités anthropiques. Entre 2001 et 2015, c'est plus de 5 fois la superficie de la France qui a été déforestée sur Terre[15]. La gestion des forêts est intimement liée à la localisation de l'agriculture et appelle à une optimisation conjointe globale[7].

En effet, la déforestation est due à 27% à l'expansion permanente de l'agriculture, des mines, et des infrastructures liées à l'énergie, portion qu'il convient de distinguer de la déforestation temporaire composée de l'exploitation forestière (26%), de la jachère longue (24%) et des feux de forêt (23%); l'urbanisation (<1%) n'a que peu participé à la perte nette d'arbres à l'échelle mondiale. Les effets de cette perte forestière sur la biodiversité se combinent avec des effets plus complexes et moins visibles. Ainsi, [40] insiste sur la nécessité de réguler le transport des végétaux pour éviter la diffusion de nuisibles non-endémiques dont les effets s'opèrent sur des dizaines d'années : 30 ans ont suffi à un champignon pour réduire de 99% la surface forestière des noisetiers des Appalaches, imposant aux espèces qui en dépendent de s'adapter ou disparaître. Enfin, la perte de biodiversité ne se résume pas à la perte d'une essence ou espèce : au niveau individuel, chaque arbre centenaire tué est une perte génétique *inestimable*[13].

Le nombre de feux de forêts devrait augmenter de +50% d'ici à 2100[9]. Leur détection précoce et leur suivi permettra d'anticiper leur progression et de coordonner les moyens d'actions. Les vigies, les patrouilles, et les appels téléphoniques des citoyens peuvent suffire à détecter rapidement et exhaustivement les feux dans les zones habitées, mais la surveillance permanente de larges zones inhabitées nécessite d'autres moyens d'observation tels que les satellites. Cartographier précisément ces feux à une échelle spatiale et temporelle globale est aujourd'hui difficile [38], mais nécessaire pour alimenter les modèles prospectivistes qui éclaireront nos décideurs sur les politiques d'occupation des sols. Notons pour être complet que tous les feux n'appellent pas forcément de réaction précoce. Ainsi les incendies dans des zones tampons permettent de protéger la forêt de 95% des feux de savane en Afrique centrale[14] et ne doivent pas être éteints. Certaines essences comme les séquoias géants profitent de petits feux[29]. Un rapport de mars 2021 des nations unies recommande un renversement de notre politique globale de gestion des feux, en privilégiant la prévention des feux à leur suppression[37], incluant notamment l'écobuage et la prévention des incivilités[38].

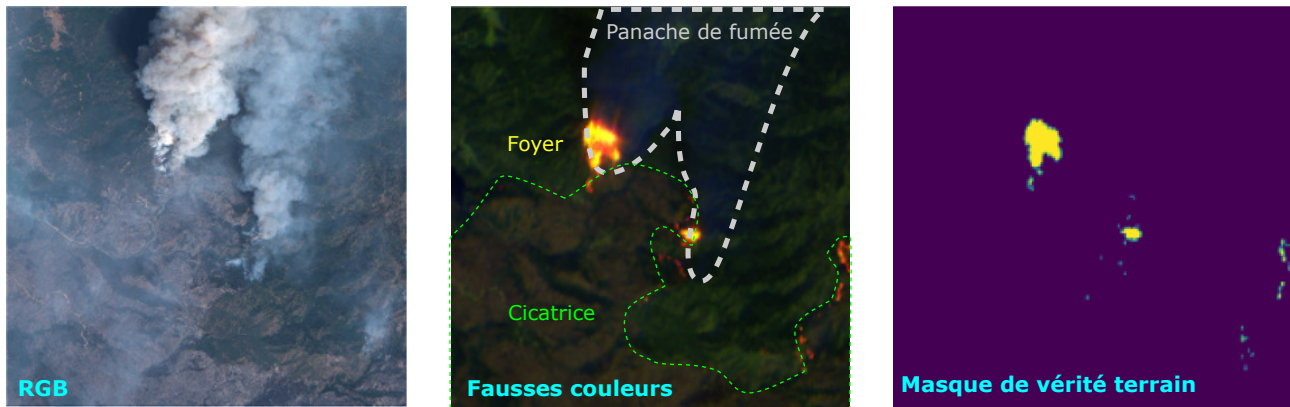


FIGURE 1 – Image Sentinel-2 en RGB (rouge=B04, vert=B03, bleu=B02) à gauche, fausses couleurs (rouge=B12, vert=B11, bleu=B04) au milieu, et le masque vérité terrain des foyers à droite ; l'image en fausse couleur présente un panache de fumée en gris, un foyer en jaune et une cicatrice en vert.

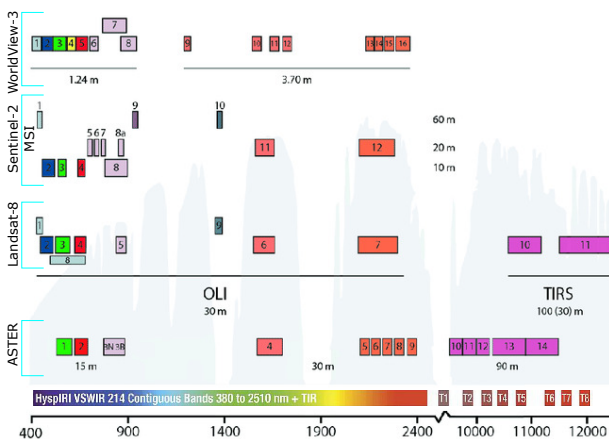


FIGURE 2 – Combinaison de deux figures extraites de [12, 11] comparant les différentes bandes spectrales de satellites d'observation. Dans cet article nous avons utilisé principalement les bandes 11 et 12 de Sentinel-2 MSI. Le fond de la figure représente la transmittance atmosphérique estivale aux latitudes moyennes.

1.2 Surveillance spatiale des incendies

Une zone d'incendie se décompose en trois objets, comme indiqué sur la figure 1 : les foyers, aussi appelés points chauds, le panache de fumée, et la cicatrice, aussi appelée zone brûlée. Cet article se focalise sur la *détection des foyers* à partir d'images satellitaires acquises par l'instrument multi-spectral de Sentinel-2. Cet instrument observe la Terre de 442 à 2202nm et nous utilisons les bandes spectrales rouge (B4) et proche infrarouge (B11 et B12) dans lesquelles les foyers sont les mieux observables. La figure 2 compare ces bandes à celles d'autres satellites. Les images acquises par Sentinel-2 sont disponibles sur les plateformes DIAS (Data and Information Access Service) financées par la commission européenne et couramment utilisées par la communauté scientifique de l'observation de la Terre.

Les techniques d'apprentissage automatique, en particulier

l'apprentissage profond de réseaux de neurones convolutifs (CNN), sont déjà employées pour analyser des images satellitaires hyper-spectrales[2, 32], notamment pour détecter les feux de forêt actifs [34, 43]. L'augmentation progressive de la disponibilité et de la précision des données multi-spectrales fournies par les séries de satellites MODIS, Landsat-8 et Sentinel-2, a créé de nouvelles opportunités d'études scientifiques sur la gestion des incendies[24, 18, 35, 42, 31, 44]. Des approches basées sur l'apprentissage automatique ont été introduites dans [19, 6] afin d'identifier les zones brûlées. [17] propose une approche basée sur l'apprentissage supervisé pour estimer automatiquement la sévérité des dégâts des zones affectées après l'extinction d'un feu de forêt en utilisant deux réseaux de neurones de type Unet. Le premier effectue une classification binaire de chaque pixel pour déterminer l'appartenance à une zone brûlée. Le second applique une régression pour estimer le niveau de gravité. Cette estimation est cruciale pour rapidement évaluer les pertes financières et efficacement mettre en place la restauration des zones affectées. [16] utilise des CNN afin d'améliorer la résolution des bandes infrarouges à ondes courtes (SWIR) et obtenir des cartes plus détaillées des feux de forêt actifs.

La fumée est un signal précoce de combustion de la biomasse. Elle joue donc un rôle important dans la détection précoce des incendies de forêt et la limitation de leurs impacts. Malheureusement, la détection des fumées en imagerie satellitaire demeure difficile, de par la diversité des formes, des couleurs, des étendues et des chevauchements spectraux de celles-ci [26, 46, 25]. Il est donc difficile de distinguer la fumée des autres types de textures, telles que les nuages, la poussière, la brume. Ainsi, le CNN SmokeNet[4] intègre l'information spatiale via une channel-wise-attention afin d'améliorer la représentation des caractéristiques pour la classification d'image avec ou sans fumée.

La détection de foyers a reçu davantage d'attention : [41, 30, 23] proposent des algorithmes basés sur des approches statistiques avec un seuil fixe. Leurs données sont multi-

spectrales avec une bande infrarouge sensible au feu (bande 7) et six bandes dans le visible et le proche infrarouge (1-6); ces données sont fournies par l'instrument Operational Land Imager (OLI) de Landsat-8. Sur le même concept, [22] propose l'algorithme AFD-S2 pour détecter les foyers sur des données Sentinel-2. Cet algorithme reconnaît les petits points de feu actifs avec une erreur moyenne de commission (CE) et une erreur moyenne d'omission (OE) sur 8 images Sentinel-2 autour de 0,14 et 0,04 respectivement. Cependant, ces approches à seuil fixe ont une précision limitée en raison de la complexité de la mise en place des seuils. Ainsi, [27] propose une petite architecture de réseau de neurones entièrement connectés avec 1 couche d'entrée, 3 couches cachées et 1 couche de sortie. Leur modèle a été entraîné en exploitant un sous-ensemble de 4 bandes (1, 5, 6 et 7) de Landsat-8. Ce choix de bandes est basé sur les analyses faites par [41] pour déterminer la forte corrélation de ces bandes avec les applications de détection d'incendie. Ensuite, [27] a comparé leur modèle avec l'algorithme à seuil fixe proposé par [41]. [33] propose trois architectures différentes basées sur l'architecture UNet avec des données de Landsat8. Le premier UNet est entraîné sur les 10 canaux de landsat8. Le deuxième est entraîné sur les 3 canaux B7, B6 et B4. Le dernier, aussi entraîné avec ces 3 bandes, est un UNet réduit appelé UNet-Light. Leur vérité terrain a été construite à partir des algorithmes [41, 30, 23], et l'architecture UNet-Light a obtenu le meilleur score IoU de 89%. Des réseaux de neurones à base de convolutions de type DeepLabV3 et DCPA+HRNetV2, ont été utilisés [45] pour segmenter sémantiquement et binairement (feu ou non-feu pour chaque pixel) des images Sentinel-2 en utilisant les bandes B12, B11 et B04. C'est aussi ce sous-ensemble de bandes que nous avons choisi l'année dernière lorsque nous avons commencé notre étude; cependant, nous avons choisi une architecture de CNN efficace afin de l'embarquer dans un satellite qui transmettrait ses alertes incendie aux intervenants sans passer par des infrastructures au sol.

1.3 Nos contributions

Cet article présente les contributions suivantes :

- Un jeu d'images multispectrales de 13 bandes spectrales provenant du satellite Sentinel-2, dont la résolution spatiale au sol varie de 40 à 80m, et où les foyers ont été pré-annotés automatiquement puis corrigés manuellement.
- L'entraînement d'un CNN embarquable référent sur ce jeu d'images atteignant une IoU de 94%.
- Une implémentation C++/TensorRT du segmenteur de foyers sur Nvidia Jetson Nano et NX, et donc compatible du Jetson TX2i embarquable à bord de satellites en orbite basse.

1.4 Plan de l'article

Cet article décrit la collecte et l'annotation des images puis, dans un second temps, l'entraînement, le déploiement, et l'évaluation du modèle utilisé. Une évaluation succincte de l'impact environnemental de cette étude est exposée après la conclusion.



FIGURE 3 – Géo-localisation de nos 90 scènes.

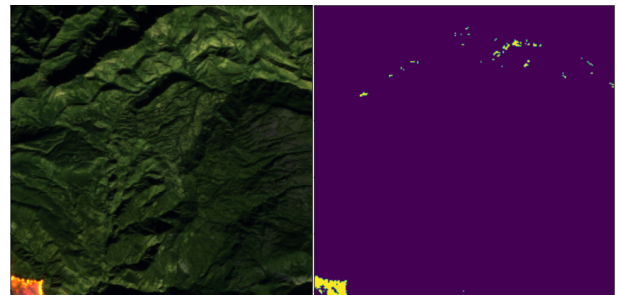


FIGURE 4 – Image en fausses couleurs, à gauche, où la formule de Pierre Markuse[28] génère un masque, à droite, contenant de nombreux faux positifs sur la moitié haute.

2 Expérimentation

2.1 Collection des données

Dans cette section, nous présentons comment nous avons utilisé le Web Coverage Service (WCS) de l'API OGC Sentinel-Hub[6, 18] pour télécharger les scènes satellitaires. Nous détaillons ensuite comment nous avons pré-annoté ces scènes avec le script de Pierre Markuse, puis finalement créé la base de données annotées qui a servi à notre expérimentation.

2.1.1 Sélection des incendies

A partir de divers moteurs de recherche, nous avons identifié le lieu et la date d'une liste d'incendies, comme le montre la carte de la Figure 3. Chaque incendie a ensuite été recherché sur le portail Sentinel-Hub, et pour chaque image identifiant l'incendie, l'aperçu de la zone au format GéoTIFF a été manuellement téléchargé, et ce à plusieurs moments. Au total, nous avons téléchargé 90 aperçus correspondant à des scènes acquises entre 2016 et 2020 par l'instrument MSI des satellites Sentinel-2a et 2b.

2.1.2 Téléchargement des scènes L1C

Ensuite, nous avons automatisé le téléchargement de la scène correspondant à chacun des aperçus géolocalisés; le but étant d'embarquer nos algorithmes sur satellite, nous avons choisi le niveau de traitement L1C qui est le plus proche de la sortie brute de l'instrument MSI [10]. Notre automatisation, programmée en Python, utilise le service WCS de l'API OGC. Les contraintes de WCS imposent une limite haute d'environ 5000x5000 pixels sur la taille

des scènes téléchargeables. Notre programme Python génère donc des requêtes WCS où le GSD (Ground Sampling Distance, en mètre par pixel) est augmenté itérativement jusqu'à ce que la zone d'intérêt soit suffisamment petite pour être téléchargeable ; le GSD de nos scènes varie ainsi entre 40m et 80m, et une possible conséquence est que notre CNN soit moins sensible au changement de GSD et se focalise davantage sur la dimension spectrale. A noter que WCS sur-échantillonne les bandes dont le GSD est plus faible que celui requêté : ainsi, au sein d'une même scène, toutes les bandes ont le même GSD. Au final, chaque scène téléchargée est composée des 13 bandes spectrales de MSI, de B01 à B12, voir figure 2, enregistrée en TIFF où les pixels sont codés en entier-16 bits.

2.1.3 Présentation du script de Pierre Markuse

Le script de Pierre Markuse[28] est un programme Javascript créé pour visualiser les feux actifs sur la plateforme de Sentinel-hub. Il est basé sur une formule simple, qui utilise les bandes B11 et B12 pour faire ressortir les feux. Cette formule, appliquée à chaque pixel, consiste à sommer les valeurs de la bande B11 et B12 puis classifie un pixel comme feu si et seulement si ce résultat est strictement supérieur à 1.

$$\text{feu} \Leftrightarrow B11 + B12 > 1 \quad (1)$$

2.1.4 Annotation des scènes

Pour pré-annoter ces scènes, nous utilisons la formule 1 de Pierre Markuse qui est conservative : elle permet de classifier correctement quasiment tous les pixels qui représentent un feu, mais avec un taux de faux positifs assez élevé comme illustré dans la figure 4. Nous avons appliqué en Python cette formule sur les scènes après les avoir converties en flottant avec des valeurs comprises entre 0 et 1. Une fois le masque binaire généré, nous l'avons ouvert dans un calque semi-transparent d'un logiciel de retouche raster, superposé à un calque contenant la scène en fausses couleurs (canal rouge=B12, canal vert=B11 et canal bleu=B04) et un calque contenant la scène en RGB (canal rouge=B04, canal vert=B03 et canal bleu=B02), voir figure 1. Jongler entre ces scènes nous permet de mieux estimer la classification de chaque pixel, notamment en regardant la fumée et les cicatrices visibles en RGB. Une fois ces pré-annotations corrigées, elles sont enregistrées au format TIFF où chaque pixel est un entier-non-signé 8-bits avec les valeurs possibles 0=non-feu et 255=feu ; ces masques sont ainsi notre vérité-terrain.

2.2 Entraînement du CNN de segmentation

2.2.1 Génération d'images à partir des scènes

Les scènes annotées doivent être tuilées avant l'entraînement pour des raisons de dépassement mémoire GPU et de diversité par batch. Le tuilage consiste à découper la scène en images de 256x256 avec une zone de recouvrement de 128 pixels. Si l'image contient plus de 15 pixels de feu, au regard de la vérité terrain, alors elle est ajoutée au *dataset*, qui est notre jeu de données servant à entraîner, valider et tester les modèles. Si elle contient entre 0 et 14 pixels, alors elle est rejetée car considérée comme trop peu sûr

	nombre de scènes	nombre d'images	ratio pixels feu/non-feu
dataset	90	2244 feu + 1104 non-feu	0.162%
trainset	72	2742	0.143%
valset	9	312	0.360%
testset	9	294	0.136%

TABLE 1 – Déséquilibre des classes feu/non-feu du dataset

qu'il y ait un feu, et pas assez sûr qu'il n'y en ait pas. Si elle contient 0 pixel de feu, alors elle est ajoutée avec une probabilité de 0.05 au dataset afin d'augmenter la diversité de terrain sans trop accentuer le déséquilibre entre les classes feu et non-feu. Même avec cette stratégie, le dataset obtenu contient seulement 0.162% de pixels de feux. Le partitionnement du dataset en trois sous-ensembles a été réalisé de manière aléatoire au niveau des scènes dans les proportions indiquées dans la table 1, assurant ainsi qu'aucune donnée d'entraînement ne soit utilisée pour la validation ou le test.

2.2.2 Description du modèle UNetMobileNetV3

U-Net[39] est un CNN développé pour segmenter des images biomédicales, mais aussi utilisé pour segmentation de nuages à bord de satellite[5]. Son architecture est composée d'un encodeur et d'un décodeur. Le but de l'encodeur est de capturer le contexte tandis que le rôle du décodeur est de calculer le masque de segmentation. Pour des raisons d'efficacité, nous avons remplacé l'encodeur original par MobileNet-v3-small[21], en prenant le soin de relier ce nouvel encodeur au décodeur. Le décodeur a été simplifié à 5 blocs de déconvolution puis une couche de Convolution et une activation Sigmoid permettant d'aboutir à un masque ayant la même taille que l'image d'entrée. Les 4 premiers des 5 blocs du décodeur sont composés d'une couche de Convolution Transposée, une couche de BatchNormalisation, et une couche de Dropout suivie d'une activation Relu. La sortie de chacun de ces 4 blocs est concaténée avec celles de l'encodeur. Ainsi, notre UNetMobileNetV3 possède 4 693 553 paramètres dont 1 529 968 pour l'encodeur mobilenetv3-small et 3 163 585 pour notre décodeur.

2.2.3 Procédure d'entraînement

Nous avons entraîné notre UNetMobileNetV3 sur deux GPU Nvidia GTX 1080 Ti avec les hyperparamètres suivants :

- Fonction de perte : Dice + Jaccard + Focal
- Poids des classes : feu=400, non-feu=1
- Initialisation des poids : encodeur mobilenetv3-small=ImageNet, decoder=aléatoire
- Gel des poids : non
- Taille du Batch : 128, soit 64 par GPU
- Optimiseur : Adam
- Learning rate initial : 0.01
- Strategies : Earlystopping + ReduceOnPlateau
- Max Epochs : 1000, mais s'arrête à ~400 epochs

La figure 5 montre que l'entraînement converge en moins de 400 epochs et présente un léger sur-apprentissage.

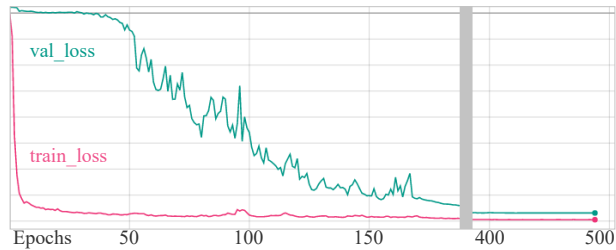


FIGURE 5 – Valeur des fonctions de perte au fil des epochs

	Puissance	Débit en MPixel/s	6711 km ² GSD=40m
Jetson Nano	<10W	1.87	2.236 s
Jetson NX	<20W	8.99	0.466 s

TABLE 2 – Vitesse du segmenteur de feu embarqué sur Jetson et évalué sur 52 batchs de 64 imageries de 256² pixels.

2.3 Déploiement sur électronique embarquée

Nous avons converti notre modèle tf.keras, représentant 2.98 milliards de MAC (Multiplieur-accumulateur), vers onnx puis nous l'avons compilé avec TensorRT entièrement en C++. Nous avons évalué ses performances sur les Jetson Nano et NX comme indiqué à la table 2. Afin de mettre ces performances en perspective, nous avons calculé qu'il est capable de segmenter en moins de 2 secondes et à un GSD de 40m les 5000km² du département de la Loire.

2.4 Évaluation des résultats

Pour mesurer la qualité de notre modèle, nous avons utilisé les métriques de la table 3 calculées à partir des masques générés soit par TF.keras en float-32, soit par TensorRT avec une quantification en float-16 lors du déploiement sur les deux Jetson. A partir du nombre de pixels de ces masques évalués comme vrais ou faux positifs (TP ou FP) et vrais ou faux négatifs (TN ou FN), nous définissons les quatre métriques suivantes : Rappel=TP/(TP+FN), Précision=TP/(TP+FP), F1-score=2TP/(2TP+FP+FN), IoU=TP/(TP+FP+FN). Ces résultats sont illustrés par la figure 6. Notre modèle possède une bonne capacité de généralisation en atteignant plus de 94% d'IoU sur le jeu de test, composé de scènes Sentinel-2 utilisées ni dans le trainset ni dans le valset; cette capacité est notamment due aux fonctions de pertes dédiées à compenser le fort déséquilibre des classes (ratio de pixels feu/non-feu=0.162%). La quantification réalisée par TensorRT introduit un léger biais en manquant moins de pixel feu mais en introduisant plus de faux positifs, ce qui augmente le Rappel et baisse la Précision.

3 Discussion & conclusion

Notre base de scènes Sentinel-2 de feux de forêts *annotées* semble être la première librement disponible, et nous espérons qu'elle servira de référence pour comparer différentes approches algorithmiques, telles que notre segmenteur qui semble aussi être le premier à être embarquable

Model	Rappel	Preci.	F1-score	IoU
TensorRT (test)	0.990	0.949	0.969	0.940
TF.keras (test)	0.973	0.975	0.974	0.950
TF.keras (val)	0.961	0.987	0.974	0.949
TF.keras (train)	0.998	0.999	0.999	0.997
P. Markuse (test)	1.000	0.871	0.931	0.871

TABLE 3 – Evaluation de Unet-mobilenet-v3 sur le testset, trainset et valset en TensorFlow.Keras, et en TensorRT sur Nvidia Jetson NX et Nano.

dans un satellite. A noter que ce segmenteur, entraîné sur des images avec un GSD de 40 à 80m, permet d'être embarqué sur des satellites ayant une orbite elliptique. Comme les satellites Sentinel ne peuvent pas être modifiés pour exécuter nos algorithmes, nous espérons que de prochaines versions embarqueront des électroniques le permettant. Concernant l'algorithme, son léger sur-apprentissage (valset IoU=94.9% < trainset IoU=99.7%) dû au manque de données d'entraînement et de régularisation fera l'objet d'un prochain travail, tout comme l'évaluation de notre segmenteur sur des images sur-résolues, voire acquises par d'autres satellites.

Déclaration de conflit d'intérêt

Les auteurs déclarent que le contenu de cet article n'a pas été influencé par d'autres conflits d'intérêt que ceux engendrés par une vie exposée aux incendies, sécheresses, déforestation, et effondrement de la biodiversité en général.

Remerciements

Ces travaux sont menés dans le projet CIAR (Chaîne Image Autonome et Réactive) de l'Institut de Recherche Technologique Saint-Exupéry. Les auteurs remercient les partenaires industriels et académiques du projet : Activeeon[1], Avisto, ELSYS Design, Geo4i, Inria, LEAT/CNRS, MyDataModels, Thales Alenia Space et TwinswHeel.

Impact environnemental

La part du numérique est projeté à plus de 14% des émissions globales de CO₂ en 2040 et, d'après [36], l'IA en représenterait une part croissante. Donc, bien que l'objectif principal de notre travail soit d'embarquer à bord de futurs satellites un détecteur de feux de forêts efficace, nous nous interrogeons sur l'impact environnemental de notre étude, comme recommandé par [20], et sur les moyens de les mesurer[3] et de les minimiser. Des pistes sont proposées dans la figure 7. Nous estimons que l'ensemble de nos travaux représentent environ 6,5 tCO₂ en se basant sur l'empreinte carbone d'un français de 9 tCO₂/an ayant travaillé 9 mois-homme sur ce projet. Sans compter l'impact de la fabrication, maintenance et l'impossible "recyclage" du matériel, et en ignorant tout sauf la consommation des deux GPUs Nvidia GTX 1080Ti, nous estimons que les 140 heures d'entraînement représentent 35 kWh, ce qui équivaut à moins de 2 kgCO₂ si le mix énergétique français est

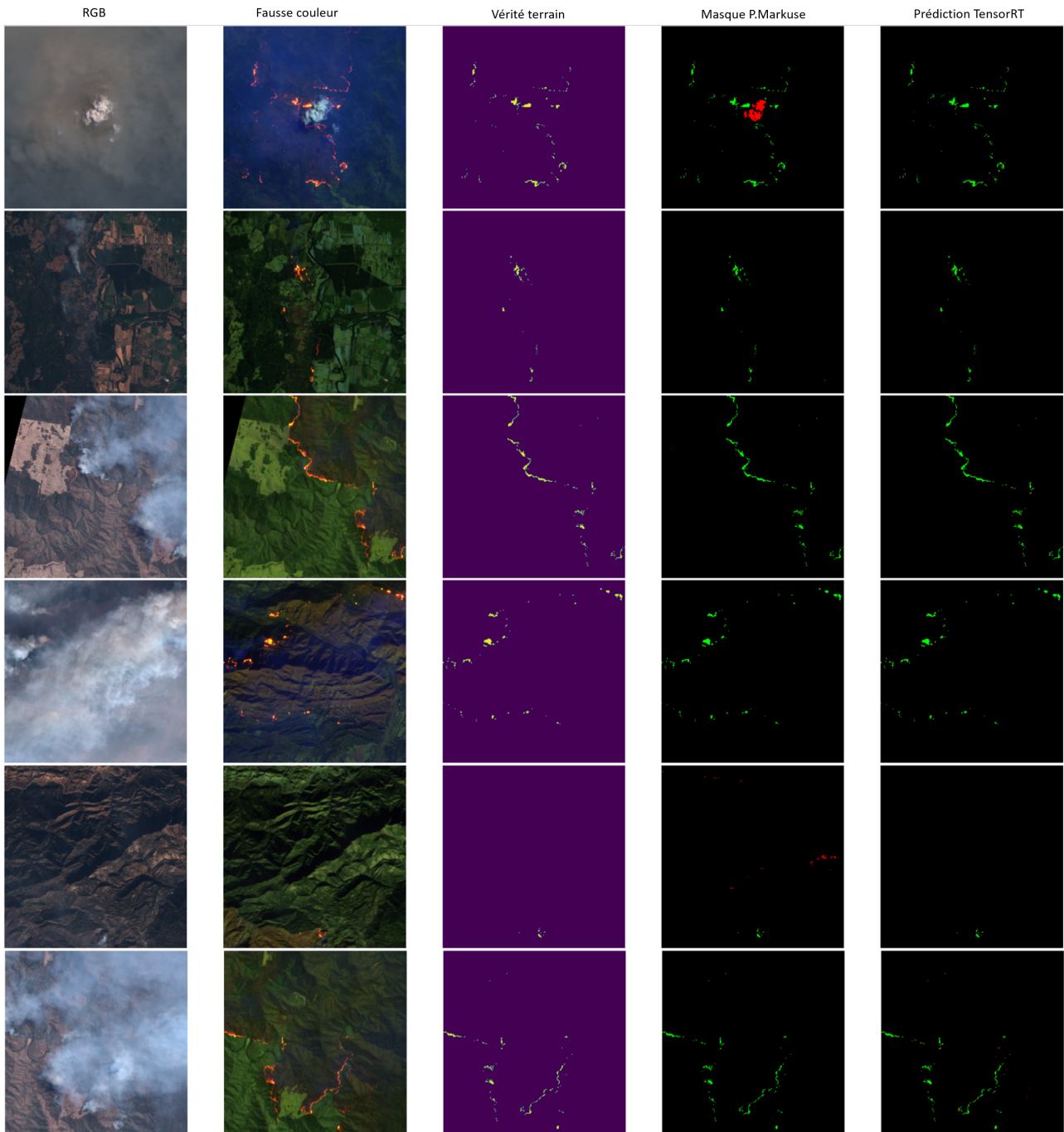


FIGURE 6 – Exemples de prédictions du modèle Unet-mobileNetV3 sur six imagettes appartenant au testset. Images de gauche à droite : image en RGB, image en fausse couleur donnée en entrée du modèle, vérité terrain, masque de Pierre Markuse, et masque prédit par le modèle TensorRT embarqué sur Jetson. Légende des couleurs des masques : les pixels noir, vert et Rouge représentent respectivement les vrais négatifs, les vrais positifs et les erreurs de classifications (faux positifs et faux négatives).



FIGURE 7 – Éco-actions pour les (neuro)scientifiques[36]

d'environ 50 gCO₂/kWh[8].

Références

- [1] ProActive: Job Scheduling & Workload Automation technologies. [Online; accessed 16. Mar. 2022].
- [2] Adrian Albert, Jasleen Kaur, and Marta Gonzalez. Using convolutional networks and satellite imagery to identify patterns in urban environments at a large scale. *arXiv*, Apr 2017.
- [3] Lasse F. Wolff Anthony, Benjamin Kanding, and Raghavendra Selvan. Carbontracker : Tracking and Predicting the Carbon Footprint of Training Deep Learning Models. *arXiv*, Jul 2020.
- [4] Rui Ba, Chen Chen, Jing Yuan, Weiguo Song, and Siuming Lo. SmokeNet : Satellite Smoke Scene Detection Using Convolutional Neural Network with Spatial and Channel-Wise Attention. *Remote Sens.*, 11(14) :1702, Jul 2019.
- [5] Gaetan Bahl, Lionel Daniel, Matthieu Moretti, and Florent Lafarge. Low-power neural networks for semantic segmentation of satellite images. In *Proceedings of the IEEE/CVF International Conference on Computer Vision (ICCV) Workshops*, Oct 2019.
- [6] Yifang Ban, Puzhao Zhang, Andrea Nascetti, Alexandre R. Bevington, and Michael A. Wulder. Near Real-Time Wildfire Progression Monitoring with Sentinel-1 SAR Time Series and Deep Learning. *Sci. Rep.*, 10(1322) :1–15, Jan 2020.
- [7] Robert M. Beyer, Fangyuan Hua, Philip A. Martin, Andrea Manica, and Tim Rademacher. Relocating croplands could drastically reduce the environmental impacts of global food production. *Commun. Earth Environ.*, 3(49) :1–11, Mar 2022.
- [8] EDF, Feb 2022. [Online; accessed 15. Mar. 2022].
- [9] Spreading like Wildfire : The Rising Threat of Extraordinary Landscape Fires, Feb 2022. [Online; accessed 15. Mar. 2022].
- [10] Products and Algorithms – Sentinel-2 MSI Technical Guide – Sentinel Online - Sentinel Online, Mar 2022. [Online; accessed 16. Mar. 2022].
- [11] Comparison between ASTER, Landsat 8, Sentinel-2, and WorldView-3 bands (Figure 4.), Mar 2022. [Online; accessed 15. Mar. 2022].
- [12] Landsat-9 - Search EO Satellite Missions, Mar 2022. [Online; accessed 15. Mar. 2022].
- [13] Charles H. Cannon, Gianluca Piovesan, and Sergi Munné-Bosch. Old and ancient trees are life history lottery winners and vital evolutionary resources for long-term adaptive capacity. *Nat. Plants*, 8 :136–145, Feb 2022.
- [14] Anabelle W. Cardoso, Imma Oliveras, Katharine A. Abernethy, Kathryn J. Jeffery, Sarah Glover, David Lehmann, Josué Edzang Ndong, Lee J. T. White, William J. Bond, and Yadvinder Malhi. A distinct ecotonal tree community exists at central African forest–savanna transitions. *J. Ecol.*, 109(3) :1170–1183, Mar 2021.
- [15] Philip G. Curtis, Christy M. Slay, Nancy L. Harris, Alexandra Tyukavina, and Matthew C. Hansen. Classifying drivers of global forest loss. *Science*, 361(6407) :1108–1111, Sep 2018.
- [16] D. A. G. Dell’Aglia, M. Gargiulo, A. Iodice, D. Riccio, and G. Ruello. Active Fire Detection in Multispectral Super-Resolved Sentinel-2 Images by Means of Sam-Based Approach. In *2019 IEEE 5th International forum on Research and Technology for Society and Industry (RTSI)*, pages 124–127. IEEE, Sep 2019.
- [17] Alessandro Farasin, Luca Colomba, and Paolo Garza. Double-Step U-Net : A Deep Learning-Based Approach for the Estimation of Wildfire Damage Severity through Sentinel-2 Satellite Data. *Appl. Sci.*, 10(12) :4332, Jun 2020.
- [18] Federico Filippini. Exploitation of Sentinel-2 Time Series to Map Burned Areas at the National Level : A Case Study on the 2017 Italy Wildfires. *Remote Sens.*, 11(6) :622, Mar 2019.
- [19] Louis Giglio, Luigi Boschetti, David P. Roy, Michael L. Humber, and Christopher O. Justice. The Collection 6 MODIS burned area mapping algorithm and product. *Remote Sens. Environ.*, 217 :72–85, Nov 2018.
- [20] Peter Henderson, Jieru Hu, Joshua Romoff, Emma Brunskill, Dan Jurafsky, and Joelle Pineau. Towards the Systematic Reporting of the Energy and Carbon Footprints of Machine Learning. *arXiv*, Jan 2020.
- [21] Andrew Howard, Mark Sandler, Grace Chu, Liang-Chieh Chen, Bo Chen, Mingxing Tan, Weijun Wang, Yukun Zhu, Ruoming Pang, Vijay Vasudevan, Quoc V. Le, and Hartwig Adam. Searching for MobileNetV3, 2019. [Online; accessed 16. Mar. 2022].
- [22] Xikun Hu, Yifang Ban, and Andrea Nascetti. Sentinel-2 MSI data for active fire detection in major fire-prone biomes : A multi-criteria approach. *DIVA*, 101, 2021.

- [23] S. S. Kumar and D. P. Roy. Global operational land imager Landsat-8 reflectance-based active fire detection algorithm. *Int. J. Digital Earth*, 11(2) :154–178, Feb 2018.
- [24] Rosa Lasaponara, Biagio Tucci, and Luciana Ghermandi. On the Use of Satellite Sentinel 2 Data for Automatic Mapping of Burnt Areas and Burn Severity. *Sustainability*, 10(11) :3889, Oct 2018.
- [25] Xiaolian Li, Weiguo Song, Liping Lian, and Xiaoge Wei. Forest Fire Smoke Detection Using Back-Propagation Neural Network Based on MODIS Data. *Remote Sens.*, 7(4) :4473–4498, Apr 2015.
- [26] Zhanqing Li, A. Khananian, R. H. Fraser, and J. Cihlar. Automatic detection of fire smoke using artificial neural networks and threshold approaches applied to AVHRR imagery. *IEEE Trans. Geosci. Remote Sens.*, 39(9) :1859–1870, Sep 2001.
- [27] Z. Liu, K. Wu, R. Jiang, and H. Zhang. A SIMPLE ARTIFICIAL NEURAL NETWORK FOR FIRE DETECTION USING LANDSAT-8 DATA. *ISPRS - International Archives of the Photogrammetry, Remote Sensing and Spatial Information Sciences*, XLIII-B3-2020 :447–452, Aug 2020.
- [28] Pierre Markuse. Visualizing Wildfires and Burn Scars with the Sentinel Hub EO Browser V2, April 2018. [Online; accessed 14. Mar. 2022].
- [29] Marc D. Meyer and Hugh D. Safford. Giant Sequoia Regeneration in Groves Exposed to Wildfire and Retention Harvest. *Fire Ecol.*, 7(2) :2–16, Aug 2011.
- [30] Sam W. Murphy, Carlos Roberto de Souza Filho, Rob Wright, Giovanni Sabatino, and Rosa Correa Pabon. HOTMAP : Global hot target detection at moderate spatial resolution. *Remote Sens. Environ.*, 177 :78–88, May 2016.
- [31] Fiona Ngadze, Kudzai Shaun Mpakairi, Blessing Kavhu, Henry Ndaimani, and Monalisa Shingirayi Maremba. Exploring the utility of Sentinel-2 MSI and Landsat 8 OLI in burned area mapping for a heterogenous savannah landscape. *PLoS One*, 15(5) :e0232962, May 2020.
- [32] Taekjun Oh, Myung Jin Chung, and Hyun Myung. Accurate Localization in Urban Environments Using Fault Detection of GPS and Multi-sensor Fusion. In *Robot Intelligence Technology and Applications 4*, pages 43–53. Springer, Cham, Switzerland, Jul 2016.
- [33] Gabriel Henrique de Almeida Pereira, André Minoro Fusioka, Bogdan Tomoyuki Nassu, and Rodrigo Minetto. Active Fire Detection in Landsat-8 Imagery : a Large-Scale Dataset and a Deep-Learning Study. *arXiv*, Jan 2021.
- [34] Alessandro Piscini and Stefania Amici. Fire detection from hyperspectral data using neural network approach. *ResearchGate*, page 963721, Oct 2015.
- [35] Luca Pulvirenti, Giuseppe Squicciarino, Elisabetta Fiori, Paolo Fiorucci, Luca Ferraris, Dario Negro, Andrea Gollini, Massimiliano Severino, and Silvia Puca. An Automatic Processing Chain for Near Real-Time Mapping of Burned Forest Areas Using Sentinel-2 Data. *Remote Sensing*, 12(4) :674, Feb 2020.
- [36] Charlotte L. Rae, Martin Farley, Kate J. Jeffery, and Anne E. Urai. Climate crisis and ecological emergency : Why they concern (neuro)scientists, and what we can do. *Brain Neurosci. Adv.*, 6 :23982128221075430, Jan 2022.
- [37] François-Nicolas Robinne. UNFF16 background paper : Impacts of disasters on forests, in particular forest fires. *ResearchGate*, Apr 2021.
- [38] Marcos Rodrigues, María Zúñiga-Antón, Fermín Alcasena, Pere Gelabert, and Cristina Vega-Garcia. Integrating geospatial wildfire models to delineate landscape management zones and inform decision-making in Mediterranean areas. *Saf. Sci.*, 147 :105616, Mar 2022.
- [39] Olaf Ronneberger, Philipp Fischer, and Thomas Brox. U-Net : Convolutional Networks for Biomedical Image Segmentation. In *Medical Image Computing and Computer-Assisted Intervention – MICCAI 2015*, pages 234–241. Springer, Cham, Switzerland, Nov 2015.
- [40] B. Roy, H. M. Alexander, J. Davidson, F. Campbell, J. Burdon, R. Snieszko, and C. Brasier. Increasing forest loss worldwide from invasive pests requires new trade regulations. *undefined*, 2014.
- [41] Wilfrid Schroeder, Patricia Oliva, Louis Giglio, Brad Quayle, Eckehard Lorenz, and Fabiano Morelli. Active fire detection using Landsat-8/OLI data. *Remote Sens. Environ.*, 185 :210–220, Nov 2016.
- [42] Daniela Smiraglia, Federico Filipponi, Stefania Mandrone, Antonella Tornato, and Andrea Taramelli. Agreement Index for Burned Area Mapping : Integration of Multiple Spectral Indices Using Sentinel-2 Satellite Images. *Remote Sens.*, 12(11) :1862, Jun 2020.
- [43] Nguyen Thanh Toan, Phan Thanh Cong, Nguyen Quoc Viet Hung, and Jun Jo. A deep learning approach for early wildfire detection from hyperspectral satellite images. In *2019 7th International Conference on Robot Intelligence Technology and Applications (RiTA)*, pages 38–45. IEEE, Nov 2019.
- [44] Daan van Dijk, Sorosh Shoaie, Thijs van Leeuwen, and Sander Veraverbeke. Spectral signature analysis of false positive burned area detection from agricultural harvests using Sentinel-2 data. *Int. J. Appl. Earth Obs. Geoinf.*, 97 :102296, May 2021.
- [45] Qi Zhang, Linlin Ge, Ruiheng Zhang, Graciela Isabel Metternicht, Chang Liu, and Zheyuan Du. Towards a Deep-Learning-Based Framework of Sentinel-2 Imagery for Automated Active Fire Detection. *Remote Sens.*, 13(23) :4790, Nov 2021.
- [46] Tom X.-P. Zhao, Steve Ackerman, and Wei Guo. Dust and Smoke Detection for Multi-Channel Imagers. *Remote Sens.*, 2(10) :2347–2368, Oct 2010.

