



L'intelligence artificielle comme facilitateur de la prise de décision

Andrés Romero

► To cite this version:

Andrés Romero. L'intelligence artificielle comme facilitateur de la prise de décision. Gestion et management. 2020. dumas-03000582

HAL Id: dumas-03000582

<https://dumas.ccsd.cnrs.fr/dumas-03000582>

Submitted on 19 Mar 2021

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



L'Intelligence artificielle comme facilitateur de la prise de décision

Présenté par : ROMERO Andrés

Entreprise d'accueil : ProLeads

**9 Rue Marcel Chabloz
38400 Saint-Martin-d'Hères**

Date de stage : du 06/04/2020 au 30/09/2020

Tuteur entreprise : ZABONI Rodolphe

Tuteur universitaire : PEREA Céline

**Master 2 Pro. FI
Master management des systèmes d'information
Parcours intelligence des données
2019 - 2020**



Mémoire de stage/ de recherche

L'Intelligence artificielle comme facilitateur de la prise de décision



Présenté par : ROMERO Andrés

Entreprise d'accueil : ProLeads

**9 Rue Marcel Chabloz
38400 Saint-Martin-d'Hères**

Date de stage : du 06/04/2020 au 30/09/2020

Tuteur entreprise : ZABONI Rodolphe

Tuteur universitaire : PEREA Céline

**Master 2 Pro. FI
Master management des systèmes d'information
Parcours intelligence des données
2019 - 2020**



Avertissement :

Grenoble IAE, au sein de l'Université Grenoble Alpes, n'entend donner aucune approbation ni improbation aux opinions émises dans les mémoires des candidats aux masters en alternance : ces opinions doivent être considérées comme propres à leur auteur.

Tenant compte de la confidentialité des informations ayant trait à telle ou telle entreprise, une éventuelle diffusion relève de la seule responsabilité de l'auteur et ne peut être faite sans son accord.

RÉSUMÉ

La surcharge informationnelle à laquelle nous sommes exposés aujourd'hui, est en partie le résultat des développements technologiques comme le Web 2.0. , les téléphones intelligents et les appareils connectés. En parallèle, on a vu croître la dépendance des personnes à l'égard des médias numériques pour la communication interpersonnelle, l'acquisition et la diffusion d'informations dans un contexte privé et/ou professionnel. 74 % des managers souffrent de surinformation ayant un impact direct sur leurs processus de décision et trois Français sur quatre disent avoir été exposés à des fakenews. Cependant, en raison de ressources cognitives limitées de l'humain, ce rapport propose l'exploration d'une solution utilisant techniques d'intelligence artificielle (AI) qui facilite le processus de prise de décision en diagnostiquant les manipulations d'informations. Basé sur mon expérience en entreprise dans la mise en œuvre de la procédure du web mining et la revue de littérature, il est proposé une méthodologie pour l'identification des manipulations d'informations capable d'analyser simultanément de grandes quantités de texte :

MOTS CLÉS : IA, NLP, surcharge informationnelle , fakenews, prise de décision

SOMMAIRE

CHAPITRE 1 – L'INFORMATION, UNE RESSOURCE FIABLE ?.....	11
I. Clarifications conceptuelles	11
II. Les manipulations de l'information	14
CHAPITRE 2 – PRISE DE DECISION DANS UN ENVIRONNEMENT DE MANIPULATION D'INFORMATIONS.....	19
I. Prise de décision et rationalité	19
II. Influencer le processus de prise de décision	19
CHAPITRE 3 – DETECTER LES MANIPULATIONS DE L'INFORMATION.....	22
I. Les Décodeurs du <i>Monde</i>	22
II. Hoaxbuster.....	22
CHAPITRE 4 – L'INTELLIGENCE ARTIFICIELLE COMME OUTIL D'AIDE A LA DECISION	23
I. NLP : linguistique informatique	23
II. Web Mining et la prise de décision automatique.....	25
CHAPITRE 5 – PROBLEMATISATION	31
CHAPITRE 6 – STAGE CHEZ PROLEADS	33
I. Présentation de l'entreprise	33
II. Description des missions	33
III. Mise en œuvre de missions	33
CHAPITRE 7 – EXPLORATION DE LA PRISE DE DECISION DANS UN CONTEXTE DE MANIPULATIONS DE L'INFORMATION	37
I. Choix du sujet	37
II. Choix de la méthodologie	37

INTRODUCTION

Une épidémie à laquelle la société, dite numérique, est de plus en plus exposée, est en train de devenir un sujet préoccupant pour les citoyens, mais impacte également la prise de décision en entreprise : il s'agit de *l'infobésité*. Ce phénomène récent, décrit par Sauvajol-Rialland comme « la pathologie de la surcharge informationnelle » [5], est en partie le résultat des développements technologiques que nous avons connus au cours des dernières décennies avec l'émergence du Web 2.0.¹, des téléphones intelligents, des appareils connectés, etc. [6]. En parallèle, on a vu croître la dépendance des personnes à l'égard des médias numériques pour la communication interpersonnelle, l'acquisition et la diffusion d'informations dans un contexte privé et/ou professionnel [7].

Dans le domaine professionnel, 74 % des managers souffrent de surinformation et d'un sentiment d'urgence généralisé [8]. Bien que la participation à des communautés en ligne (LinkedIn, Twitter, etc.) soit considérée comme un élément permettant d'accélérer le processus de décision dans l'activité managériale [9], la surinformation pourrait impacter négativement la prise de décision. Car paradoxalement, si l'un des rôles fondamentaux de l'information est de faciliter ce processus la surabondance d'informations n'ajoute que de la complexité au processus [11].

Cependant, en raison de ressources cognitives limitées, la précision de la détection des *éléments trompeurs* par l'être humain est d'environ 54% [14] : c'est tout juste un peu mieux que de jouer au pile ou face pour décider si oui ou non l'information est juste. Cela est d'autant plus difficile que la quantité de contenu disponible sur le Web est immense, et que la désinformation n'est pas un phénomène rare : 74% de français considèrent avoir déjà été exposés à des *fake news*[15].

En somme, pour pouvoir faire face à la masse de données potentiellement mobilisables pour prendre des décisions en entreprise, tout en évitant les biais cognitifs et autres limites du cerveau humain, il semble intéressant de développer des outils d'aide à la décision pour l'identification et le diagnostic automatique des manipulations de l'information numérique.

Ainsi, ce rapport s'intéresse à la question suivante :

Quelle solution informatique pour faciliter la prise de décision dans un environnement de manipulations de l'information ?

¹Le Web 2.0 se caractérise par la simplicité et l'interactivité. Les utilisateurs peuvent contribuer, échanger et collaborer sous différentes formes sans avoir de grandes connaissances techniques. [1]

L'exploration de cette problématique sera faite en utilisant mon expérience chez ProLeads ainsi qu'une revue bibliographique. Dans un premier temps, nous allons faire un constat de la qualité et la qualité des informations sur le Web. Ensuite, nous analyserons les limites dans les processus informationnels en entreprise. Enfin, nous allons étudier l'adéquation entre les caractéristiques des technologies de l'intelligence artificielle et la diminution de la complexité du processus de décision en entreprise.

PARTIE 1 :

-

PARTIE THEORIQUE

CHAPITRE 1 – L'INFORMATION, UNE RESSOURCE FIABLE ?

I. CLARIFICATIONS CONCEPTUELLES

A. Définition

La définition de l'information en tant que concept a fait l'objet de plusieurs controverses et diffère selon la discipline qui l'analyse [1]. Le dictionnaire Larousse présente dans sa liste de définitions sur l'information, la suivante:

Tout événement, tout fait, tout jugement porté à la connaissance d'un public plus ou moins large, sous forme d'images, de textes, de discours, de sons.[2]

Il faut noter que sa racine vient du verbe latin « *informare* » qui signifie « donner forme à » ou « se former une idée de »[3]. En outre, l'Association des professionnels de l'information et de la documentation la définit comme:

Élément de connaissance susceptible d'être représenté à l'aide de convention pour être conservé, traité ou communiqué. [l'information] se caractérise par un contenu (signifiant), un signifié et une forme [4].

Ainsi, l'information est étroitement liée aux processus cognitifs par lesquels seront traitées les données, qui sont quant à elles les éléments primaires, avant toute mise en contexte ou appréhension par un cerveau humain. En d'autres termes et comme elle est définie dans le domaine de la gestion de connaissances, l'information constitue « une donnée à laquelle un sens (ou une interprétation) a été ajouté » [5]. Cette approche place la connaissance dans la suite du processus, dans le sens où il s'agit d'une donnée qui a été intégrée pour la mise en œuvre d'une tâche par un sujet.

Dans ce rapport, nous utiliserons l'approche de la gestion des connaissances telle qu'elle est décrite dans l'illustration 1 pour faire référence aux concepts de *donnée*, d'*information* et de *connaissance*.

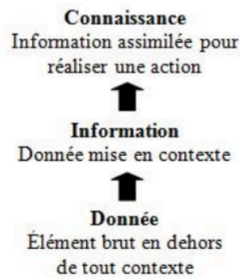


Figure 1 : modèle hiérarchique de la connaissance[41].

B. Repères historiques

L'information a accompagné l'évolution de l'humanité depuis son plus jeune âge (Illustration 2). Au départ, l'humanité utilisait des méthodes non verbales, comme les gestes, pour transmettre des informations. Plus tard, les dessins rupestres sont apparus comme les premières formes de communication, de stockage et de consultation des informations. Plus tard, le développement de différentes langues permettraient des échanges d'information plus efficace.

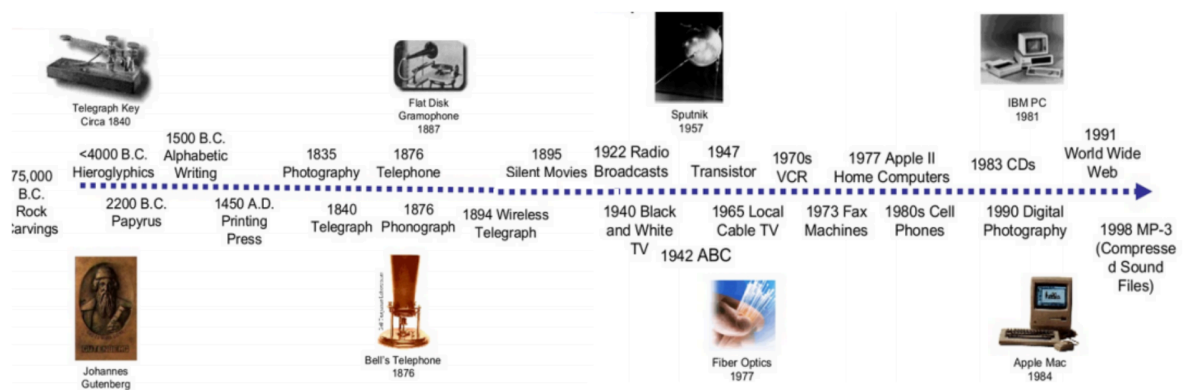


Figure 2 : perspective historique de l'information [6].

L'arrivée de la machine à imprimer au milieu du XVe siècle a permis la diffusion massive d'informations provenant de la même source par l'impression et la distribution de manuscrits. À partir du XIXe siècle, l'information peut être diffusée à un large public sous d'autres formes que le texte, par exemple en images grâce à la cinématographie ou en sons avec le phonographe [6].

Le XXe siècle, apporte beaucoup de progrès technologiques, dont l'ordinateur dans les années 70 et l'internet dans les années 90. En fait, au cours des 50 dernières années, les technologies de l'information sont devenues de plus en plus puissantes et interconnectées.

- La puissance consiste dans la capacité de traitement des ordinateurs, qui s'est accrue au fil du temps. Ainsi, en moyenne, tous les deux ans, la densité des transistors a doublé dans les microprocesseurs entre 1971 et 2001.

- L'interconnectivité, attestée par la vitesse accrue de la connexion au réseau et la facilité d'accès aux ordinateurs, font de l'internet aujourd'hui un élément essentiel de notre vie quotidienne [7].

L'ère de l'internet au début du XXI^e siècle (Illustration 3) commence par une « phase contenu » dans laquelle il est possible d'envoyer des e-mails dans lesquels peuvent être transmises des informations sous différents formats (texte, images et/ou sons). Ensuite arrive la « phase service », qui met l'accent sur la productivité et le commerce électronique. Puis, on passe à une « phase de l'internaute » où les gens occupent une place centrale et une explosion de plateformes de réseaux sociaux est générée. Enfin, on atteint la phase actuelle de l'« internet des objets » où les appareils électroniques communiquent entre eux pour effectuer une série d'activités préprogrammées ou dirigées en temps réel.

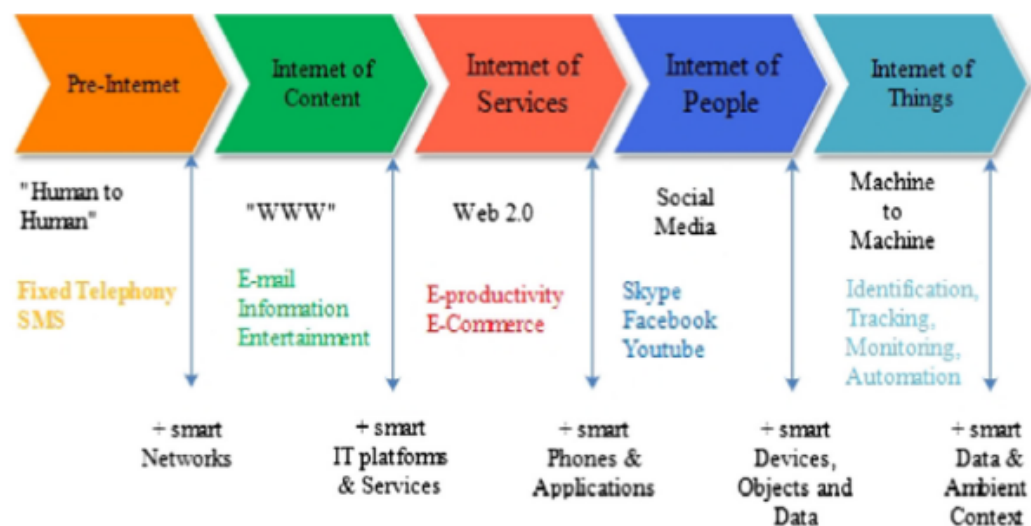


Figure 3 : les périodes de l'ère de l'internet [8].

Aujourd'hui, la communauté scientifique s'intéresse à l'intégration de l'intelligence artificielle (IA) dans des dispositifs interconnectés afin que les machines elles-mêmes puissent prendre les décisions nécessaires et agir sans intervention humaine, ce qui nous amènerait à une « phase de l'internet des choses propulsées par l'intelligence artificielle » [8].

C. L'information, aujourd'hui

Nous vivons actuellement une période inédite de surcharge informationnelle : en effet, dans l'histoire de l'humanité, il n'y aurait jamais eu dans notre environnement autant d'informations accessibles et à autant de personnes. De manière concomitante, le nombre d'informations échangées entre les individus s'est accru, tant dans le cadre de leur vie privée que dans le monde du travail. A titre d'illustration, on peut noter que chaque minute, sont téléversés 100 heures de vidéo

sur la plateforme YouTube, des centaines de millions de messages sont envoyés sur Twitter par jour, et des dizaines de millions de photos sont présentes sur Instagram. Ces évolutions peuvent s'expliquer par les progrès récents des technologies numériques, ayant amené à l'explosion de la quantité d'informations échangée sur Internet – on a atteint 40ZB d'informations digitales stockées en 2020. Cela a généré la nécessité de traiter cette masse, expliquant le développement des professions du Big Data.

Mais à cette quantité immense d'informations est associée la problématique de la qualité de l'information : quelle valeur ? Quelle fiabilité ? Quelle confiance accorder aux informations disponibles sur le Web ? Ces questions sont d'autant plus justifiées qu'il existe des phénomènes avérés de manipulation de l'information.

II. LES MANIPULATIONS DE L'INFORMATION

La manipulation peut être définie comme «obtenir de quelqu'un qu'il fasse quelque chose dont il aurait préféré se dispenser – et qu'il n'aurait pas fait à la suite d'une simple demande » [9]. La recherche en psychologie sociale a montré que nos comportements ne sont pas toujours sous notre contrôle, et que sans le savoir, nous pouvons être manipulés au quotidien par des facteurs apparemment inoffensifs qui nous amènent à prendre des décisions que nous n'aurions jamais prises spontanément [10]. Ces facteurs, qui semblent à première vue anodins, comme le fait de demander l'heure avant de demander de l'argent dans la rue pour augmenter ses chances de recevoir effectivement cet argent, sont considérés comme des techniques de manipulation lorsque la personne qui les utilise est consciente des effets qu'ils vont produire.



Figure 4 : publicité en faveur du tabac aux États-Unis qui présente (a) des avis d'experts et (b) des chiffres précis. [42]

Les techniques de manipulation sont par ailleurs utilisées dans le monde des affaires, notamment dans le marketing et le commerce. Par exemple, dans l'illustration 4, on associe l'image de professionnels de santé aux cigarettes, pour faire penser que la communauté médicale soutient leur consommation et donc qu'elles sont des produits favorables à la santé.

Mais ces techniques peuvent aussi contribuer à des causes sociales ou qui servent l'intérêt général. Par exemple, elles peuvent permettre de faire passer des messages de prévention routière comme dans l'illustration 5 [9] : grâce à une image de véhicule accidenté associée à un message ironique, on montre que minimiser les risques d'un comportement (ici, celui de rouler à vive allure) peut avoir de graves conséquences : la mort d'une personne. Les informations diffusées dans cette affiche visent à générer des émotions (culpabilité, peur, etc.) qui sont supposées inciter les personnes à adopter un comportement plus raisonnable, celui de conduire à vitesse modérée, tout en ne divulguant pas explicitement cette intention. C'est pourquoi on peut catégoriser cela dans la manipulation d'information.



Figure 5 : publicité au service la sécurité routière [40]

Dans l'histoire, les techniques de manipulation ont fait partie de la stratégie militaire qui contribuent à déstabiliser l'adversaire ou influencer l'opinion publique en utilisant toujours les dernières avancées technologiques en matière de communication afin de maximiser son pouvoir de persuasion. Ainsi, les principaux moyens de propagation de l'information pendant la première guerre mondiale (1914-1918) était la presse écrite, au cours de la seconde guerre mondiale (1939-1945), la radio et le cinéma, après 1945, la télévision, et aujourd'hui les médias numériques [11].

La tromperie (*deception* en anglais) sera plus particulièrement explorée dans la partie pratique de ce rapport. Voici un exemple de la manière dont elle peut se concrétiser autour de la controverse sur la ré-autorisation du glyphosate pour cinq ans par les autorités réglementaires européennes. En 2017, le journal *The Guardian* a révélé que plusieurs pages du rapport produit par l'Institut fédéral

allemand pour l'évaluation des risques (BfR) provenaient du groupe *Glyphosate Task Force*, un organisme qui réunissait différents producteurs de l'herbicide et qui défendait la position selon laquelle la substance n'était pas cancérogène [26].

La tromperie est identifiée dans cet exemple par le fait que le rapport BfR transmet intentionnellement des informations parce que les experts BfR ont volontairement ignoré les risques détectés de l'herbicide dans leurs diagnostics et ont décidé de communiquer des diagnostics modifiés, créant ainsi une fausse conclusion des régulateurs européens [27].

A. Les manipulations à l'ère numérique

Dans l'ère de l'internet, il existe une longue liste de termes qui font référence aux manipulations de l'information, tels que *fake-news*, post-vérité ou *information warfare* qui sont mélangés à des notions classiques telles que la propagande et la désinformation (Tableau 1). Dans ce rapport, nous utilisons l'expression « manipulations de l'information » qui désigne ces phénomènes qui se caractérisent par « l'utilisation des technologies de l'information pour influencer de manière cachée le processus décisionnel d'une autre personne, en ciblant et en exploitant les vulnérabilités du processus décisionnel »[12].

Tableau 1 : définition des concepts liés aux manipulations de l'information

Concept	Définition
<i>fake-news</i>	Traduite par « fausses informations » ce sont des informations falsifiées, contrefaites ou créées de toutes pièces. » [13]
post-vérité	Se réfère à un phénomène de distorsion où les faits objectifs ont moins d'influence sur la formation de l'opinion publique que fait l'appel à l'émotion et aux convictions personnelles. [14]
tromperie	« information transmise intentionnellement pour créer une fausse impression ou conclusion » [25]
guerre cognitive	Manière d'utiliser la connaissance dans un but conflictuel. Ce concept généralise l' <i>information warfare</i> qui est limité à la manière de leurrer l'adversaire en termes de commandement. [15]
propagande	« Tentative d'influencer l'opinion et la conduite de la société de

	telle sorte que les personnes adoptent une opinion et une conduite déterminée » [16]. Le plus souvent, la propagande se situe dans un contexte politique.
désinformation	« informations dont on peut vérifier qu'elles sont fausses ou trompeuses, qui sont créées, présentées et diffusées dans un but lucratif ou dans l'intention délibérée de tromper le public et qui sont susceptibles de causer un préjudice public » Elle se distingue de la 'mésinformation' (<i>misinformation</i>), qui n'est pas intentionnelle. [13]

À l'heure actuelle, les technologies de l'information qui sont particulièrement adaptées aux influences manipulatrices car elles présentent 3 caractéristiques [12] :

1. La surveillance numérique omniprésente : les actions que nous entreprenons laissent une trace numérique qui peut être stockée dans des bases des données détaillées qui peuvent être analysées. Par exemple, en surveillant les messages, les photos, les interactions et l'activité sur Internet en temps réel, Facebook peut déterminer des différents états d'âme des adolescents [17].
2. Les plateformes numériques (sites web et applications pour smartphones) sont dynamiques, interactives, intrusives et adaptatives. Ils se configurent en temps réel, en utilisant les informations stockées et, en même temps, apprennent en interagissant avec elles. Il existe une demande d'attention intrusive sous la forme de messages automatiques (e-mails, SMS, *push-notifications*, etc.) qui apparaissent aux moments précis où ils sont le plus susceptibles de nous tenter.
3. *L'invisibilité technologique* qui est le fait que nous voyons, entendons ou percevons les informations transmises par les technologies, comme si ces dernières étaient invisibles. Absorbé par l'information, on perd la conscience qu'on est en train de regarder un écran

Enfin, la surcharge d'informations accroît nos vulnérabilités au détriment de notre autonomie au sein d'une société qui souffre d'infobésité. Car exposés à un océan d'informations diverses, il devient difficile de détecter les manipulations.

B. Conséquences sur la société et l'entreprise

Des phénomènes de manipulations de l'information peuvent se rencontrer dans divers contextes :

1. Des activités personnelles privées : par exemple Facebook cible les publicités sur les adolescents à des moments où ils sont perçus comme particulièrement vulnérables à l'influence pour modifier leurs comportements d'achat. [17]
2. La sphère politique : le cas de Cambridge Analytica et sa participation à la manipulation des processus politiques démocratiques [18]
3. L'entreprise : des plateformes comme Uber utilisent des stratégies pour influencer le comportement des travailleurs [19])

On constate à travers d'exemples que les manipulations de l'information agissent fondamentalement sur nos prises de décisions.

CHAPITRE 2 – PRISE DE DECISION DANS UN ENVIRONNEMENT DE MANIPULATION D'INFORMATIONS

I. PRISE DE DECISION ET RATIONALITE

La prise de décision est un « processus complexe d'identification et de résolution de problèmes » [20]. Elle est considérée comme complexe en raison d'un grand nombre d'éléments (information, instinct, personnalité, émotions, habitudes, contexte social et/ou économique, etc.) qui interagissent les uns avec les autres durant ce processus et peuvent évoluer dans le temps. En même temps, la cognition humaine est limitée : même si nous avons la capacité technologique d'obtenir toutes les informations sur un problème, nous ne serions pas en mesure de trouver la solution optimale en raison de défauts dans notre raisonnement. Cela est la principale proposition de la théorie de la rationalité limitée proposée par Herbert Simon en 1947 [21].

II. INFLUENCER LE PROCESSUS DE PRISE DE DECISION

Cette théorie s'oppose au modèle de rationalité parfaite où « le décideur possède toute la connaissance requise, pour prendre la meilleure décision possible en fonction de ses objectifs ». Cependant, il est très difficile de prendre des décisions totalement rationnelles car nos ressources cognitives pour le traitement de l'information sont limitées, surtout lorsque les problèmes sont complexes. La rationalité limitée apparaît comme une proposition alternative : le décideur ne cherchera pas la solution optimale, mais une solution qui satisfasse ses aspirations. Le modèle de la rationalité limitée soutient que le décideur est soumis à un certain nombre de limitations telles que [22]:

1. Une connaissance limitée de l'environnement puisque seule une partie de la globalité des informations est connue.
2. La soumission aux émotions
3. Impossibilité d'anticiper et d'envisager toutes les options possibles dans une situation
4. Limitation du temps pour prendre la décision
5. Identification et maîtrise limitées du nombre de variables dans l'environnement
6. Capacité d'attention insuffisante pour traiter toutes les informations disponibles
7. Incapacité à analyser toutes les informations simultanément

Nous allons voir comment ces limites sont exploitées à travers des exemples présentés à la fin de la section précédente, afin de se représenter comment le processus décisionnel est modifié pour favoriser des intérêts privés.

A. Facebook et les adolescents

En mai 2017, un journal australien a montré au public les preuves de la stratégie de Facebook qui décrivait comment les annonceurs pouvaient utiliser la plateforme pour faire de la publicité auprès des adolescents en période de vulnérabilité [17]. En surveillant les messages, les photos, les interactions et l'activité sur Internet en temps réel, Facebook peut déterminer les différents états d'âme. C'est un exemple clair de l'exploitation de la vulnérabilité des émotions pour manipuler le processus de décision des gens : l'objectif était les adolescents achètent des produits qui, dans une situation émotionnelle différente, ils n'auraient pas eu l'intention d'acheter.

B. Cambridge Analytica et les élections

En mars 2018, il a été révélé que le groupe commercial Cambridge Analytica (CA) a utilisé une quantité massive de données provenant des utilisateurs de Facebook pour générer du profilage psychologique et envoyer des messages politiques ciblés dans le but d'influencer le processus démocratique lors de plus de 200 élections dans le monde.

Cela a été réalisé en analysant les « likes » sur Facebook avec lesquels l'entreprise pouvait déduire des caractéristiques personnelles telles que le sexe, l'orientation sexuelle, l'origine ethnique, la religion, les opinions politiques, l'état des relations, la consommation de substances et les traits psychologiques. Grâce à ce profilage, CA a ensuite eu la capacité de personnaliser des messages et les transmettre sous forme d'email ou de publicité sur les réseaux sociaux [23].

Ce cas montre comment les limitations humaines comme la connaissance limitée de l'environnement, l'impossibilité d'envisager toutes les options dans des élections démocratiques et la soumission aux émotions peuvent être exploitées pour servir un intérêt privé.

C. Uber et le comportement des travailleurs

Dans la gestion algorithmique d'entreprises sont pensées des méthodes d'altération des processus de décision des employés. Chez Uber, 3 ont été identifiées.

- A. Uber bombarde les conducteurs de textes, d'e-mails, de pop-ups d'informations qui renseignent sur les événements et horaires susceptibles de favoriser l'utilisation des services d'Uber par les gens du secteur, par exemple les événements sportives et les heures de fermeture de bars. Cette information est présentée comme prometteuse (forte demande de clients donc profits élevés) ce qui devrait modifier le processus de décision

du conducteur puisqu'elle l'incite à se rendre dans des endroits où il n'avait pas initialement l'intention d'aller. Dans ce cas, l'entreprise pousse l'employé à fournir un service pour des récompenses très incertaines, en exploitant la vulnérabilité du manque de capacité à analyser simultanément toutes les informations.

- B. Dans l'objectif d'allonger le temps de travail des conducteurs, lorsque ceux-ci décident de se déconnecter une autre méthode est employée par Uber : ils peuvent, dans certains cas, recevoir des messages tels que :

Êtes-vous sûr de vouloir vous déconnecter ? La demande est très forte dans votre région. Gagnez plus d'argent, n'arrêtez pas maintenant !

vous êtes à 10€ de gagner 330€ !

Ce type de stratégie tente d'exploiter la vulnérabilité du décideur, en présentant des objectifs qui paraissent facilement atteignables, générant ainsi des émotions.

- C. Enfin, la dernière méthode d'Uber exploite la vulnérabilité de la limitation du temps et de l'attention au moment de prendre une décision. Avant la fin d'un trajet, Uber indique automatiquement la demande de trajet suivante, rendant plus difficile de refuser que de faire un autre trajet. Cette stratégie est largement utilisée dans d'autres plateformes telles que Youtube et Netflix pour regarder les vidéos suivantes[23].

Les phénomènes présentés précédemment viennent montrer comment les technologies de l'information peuvent aujourd'hui être un moyen d'exploiter les vulnérabilités dans la prise de décision pour favoriser un intérêt particulier même au prix de notre santé (risque d'accident du chauffeur Uber en raison de la fatigue) ou de nos droits fondamentaux.

CHAPITRE 3 – DETECTER LES MANIPULATIONS DE L'INFORMATION

Différents disciplines tels que la philosophie, la psychologie, la sociologie, la criminologie et l'anthropologie se sont intéressés à la détection des manipulations de l'information. Les études dans ces disciplines ont montré qu'en raison de ressources cognitives limitées, la précision de la détection des éléments de manipulation par l'être humain est d'environ 54% [28]: c'est tout juste un peu mieux que de jouer au pile ou face pour décider si oui ou non l'information est juste.

Au regard de ces phénomènes de manipulations de l'information et des vulnérabilités de l'être humain face à celles-ci, on peut s'interroger sur l'utilisation d'outils informatiques pour se protéger des influences extérieures qui pourraient faire prendre des décisions qui ne vont pas être les plus favorables à nos objectifs. Une aide informatique est d'autant plus nécessaire que l'on est constamment exposés à une grande quantité d'information accessibles facilement. Un tel outil permettrait également de guider la décision face à des informations contradictoires, en plus de prévenir les influences inconscientes. Il existe des plateformes sur internet qui offrent des services qui vont dans ce sens.

I. LES DECODEURS DU *MONDE*

Cette un équipe de journalistes du groupe éditorial Le Monde qui procède à la récupération, la compilation, l'explication et la classification des rumeurs. En rectifiant les faits, ils expliquent au public si une rumeur est vraie ou non.

II. HOAXBUSTER

Ce site dispose d'un moteur de recherche pour les articles qu'ils peuvent eux-mêmes cataloguer comme rumeur, information, désinformation, mise en garde, etc. HOAXY

CHAPITRE 4 – L'INTELLIGENCE ARTIFICIELLE COMME OUTIL D'AIDE A LA DECISION

Les limites humaines présentées précédemment justifient la nécessité d'utiliser des méthodes de détection automatique et cela est d'autant plus difficile que la quantité de contenu disponible sur le Web est immense, rendant le traitement manuel impossible [29]. C'est pour cette raison que l'intelligence artificielle est présentée comme une solution potentielle en ce qui concerne sa partie « artificielle » avec l'utilisation d'algorithmes informatiques (traitement automatique du langage) et son caractère « intelligent » car elle effectue des processus cognitifs (prise de décision), généralement attribués aux humains [30]. Nous allons donc présenter par la suite différentes méthodes utilisées pour analyser des textes qui pourraient permettre ensuite de détecter les manipulations de l'information :

I. NLP : LINGUISTIQUE INFORMATIQUE

A différence de la linguistique, qui « est la science qui étudie et formalise les règles d'utilisation de la langue », la linguistique informatique (Natural Language Processing en anglais)² « est l'ensemble des méthodes de traitement du langage qui permettent de 'comprendre' automatiquement du texte pour en extraire du sens ». Le mot « comprendre » dans la définition précédente ne concerne pas la compréhension humaine du langage (au niveau lexical, syntaxique ou même pragmatique), mais la capacité de transformer le texte en une représentation formelle qui permet d'effectuer des calculs et, par exemple, de trouver des textes similaires, d'extraire des informations précises ou même d'identifier des schémas de manipulations de l'information [31]. La transformation est donc une première étape préalable à l'analyse : Quelques exemples de transformations de textes (compréhension de la machine) sont :

- Tokenisation : méthode permettant de décomposer un texte en plusieurs parties, telles que des phrases et des mots
- Suppression des caractères non alphanumériques : Suppression des numéros et symboles dont on constate généralement qu'ils sont dus à des erreurs dans l'extraction des informations
- Suppression des *stopwords* : Les *stopwords* sont des mots courants qui ne contribuent généralement pas au sens d'une phrase (par exemple, le, la, à, et, etc.). L'élimination de ces mots réduit donc la quantité de texte à analyser ultérieurement.

²Le terme utilisé en français pour la NLP est Traitement Automatique du langage Naturelle (TALN).

- Conversion des minuscules : Conversion de tous les caractères alphanumériques en minuscules pour éviter de répéter le même mot avec un autre caractère
- *Stemming* : Technique permettant de supprimer les affixes d'un mot, pour avoir juste la racine. Par exemple, la racine de manger est mange (car il contient manges, mangeons, mangent, etc.) .

En Python, il existe les bibliothèques [NLTK](#) et [spaCy](#) qui peuvent effectuer ces tâches ou nous pouvons également utiliser des techniques de programmation traditionnelles pour identifier les chaînes de caractères.

A. Analyse lexicale : l'approche classique

L'approche classique du NLP vise à :

1. Formaliser les règles linguistiques qui caractérisent la syntaxe d'une phrase.
2. Décrire les relations sémantiques entre les différents éléments qui composent la phrase

Dans ce travail, nous nous limitons au premier objectif, l'analyse de la syntaxe. Donc, pour représenter la syntaxe d'une phrase, nous allons travailler avec deux techniques :

- **Étiquetage morphosyntaxique (*part of speech*)**: Processus consistant à attribuer la classe de grammaire à chaque mot d'un texte. Dans ce processus, nous trouvons une limite qui est l'ambiguïté de certains mots comme *masque* qui peut être à la fois un nom et aussi la conjugaison du verbe *masquer*. Néanmoins, cette processus présente des performances jusqu'à 98 % de précision sur des articles de presse écrits en anglais ou en français. En revanche, la performance descend vers 90 % dans les cas de textes de réseaux sociaux, car ils présentent beaucoup plus de fautes d'orthographe et de grammaire [31].
- **Lemmatisation** : La langue française, comme d'autres langues, a des mécanismes de flexion. Autrement dit, à partir de la modification d'une base (ou lemme) d'un terme, on peut générer différentes formes. Par exemple pour un verbe conjugué (e.g. parlent), le lemme s'agit de l'infinitif *parler* ou à partir d'un nom au féminin pluriel (*serveuses*) sera le masculin singulier *serveur*.

B. Limites du NLP

Bien qu'un grand nombre d'applications de NLP soient observées (Illustration XXX), la plupart sont identifiées comme « difficilement utilisable » ou « faible intérêt d'un point de vue métier ».

Figure 4.1 – Quelques tâches usuelles de traitement automatique des langues.

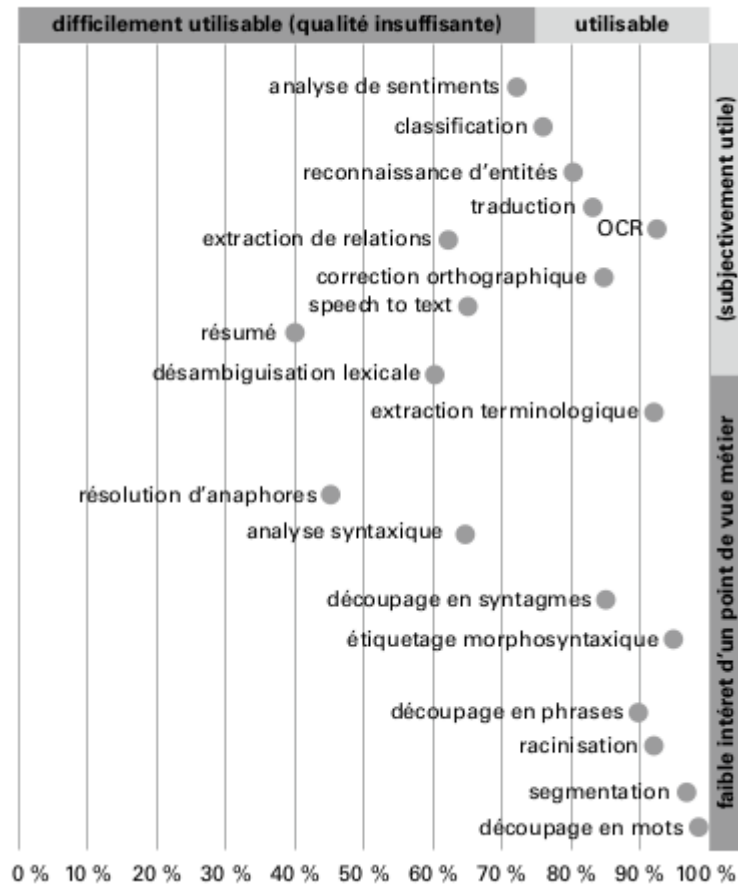


Figure 6 : Logo Grenoble IAE³

C'est principalement dû à ses limitations :

- Le langage humain est complexe en ce sens qu'il est très difficile de formuler des règles générales qui tiennent compte de son ambiguïté.
- Chaque langue a des caractéristiques spécifiques, donc chaque modèle linguistique est différent.

II. WEB MINING ET LA PRISE DE DECISION AUTOMATIQUE

Le web mining est l'utilisation de techniques informatiques pour découvrir et extraire automatiquement des informations de documents et de services sur le web [32]. Il fait partie du domaine de la *découverte de connaissances et extraction de données* (ou KDD - Knowledge Discovery and Data Mining en anglais) que comme le NLP, le but est de trouver des patrons utiles dans un

³ Pour insérer une légende sous l'image, dans l'onglet « Références », cliquer sur « Insérer une légende à la suite de « Figure 1 ». Cette fonctionnalité permet d'automatiser la numérotation des illustrations (figures, images, tableaux...) et de générer une table des illustrations à la fin de votre travail. Il est possible de modifier le style de légende.

ensemble de corpus de texte. Même si ce processus partage des techniques et des algorithmes de *machine learning*, la KDD diffère en ce qu'elle n'inclut pas des domaines tels que la robotique ou la vision par ordinateur [33]. La démarche automatique du web mining commence par une recherche de ressources, suivie d'un prétraitement d'informations, puis de la généralisation des tendances pour aboutir à l'analyse et la prise de décision.

A. Les étapes du Web Mining

i. Recherche de ressources web

À ce stade, l'objectif est de trouver les données à partir des sources disponibles sur le web telles que : les courriels, les bulletins d'information ou les textes des pages web au format HTML⁴. L'illustration XXX montre un exemple simplifié d'un document HTML, on observe que ce texte se distingue d'un texte normal par sa structure et sa syntaxe avec des balises (<html>, <head>, etc.).



Figure 7 : une page web vue (a) dans un navigateur web, (b) dans un éditeur de texte, (c) en tant que structure hiérarchique

Cette étape pourrait théoriquement être réalisée manuellement, mais dans le cas d'un grand nombre de pages web, elle n'est pas efficace. C'est pourquoi il existe des algorithmes informatiques qui effectuent le travail automatiquement, appelés algorithmes de *scrapping*. En Python, il existe deux bibliothèques largement utilisées pour effectuer cette tâche : [BeautifulSoup](#) y [Selenium](#).

ii. Sélection et pré-processing de données

À partir de ressources web extraites, il y a une sélection et prétraitement automatique du texte. Cette étape est caractérisée par une transformation des données initialement trouvées pour réduire le volume de texte et l'homogénéiser. Dans le cas des pages web, l'une des premières tâches serait de supprimer les balises <html> pour n'extraire que le texte qui nous intéresse. Nous avons évoqué ces techniques dans la section précédente (Partie 3.1 **Error! Reference source not found.**).

iii. Généralisation de tendances

⁴Pour connaître en détail le fonctionnement d'une page web, nous vous invitons à suivre ce cours.

À ce stade, on va chercher à identifier des tendances générales dans l'ensemble des corpus de texte à l'aide d'algorithmes. Deux approches peuvent être utilisées ce point :

1. Systèmes basés sur des règles (rule-based system) : approche classique du NLP, où nous créons un algorithme basé sur des règles linguistiques complexes pour extraire des patrons (évidemment identifiés auparavant).
2. Systèmes basés sur l'apprentissage automatique (*machine learning*)⁵ : Extraction des patrons à partir d'une analyse statistique basée sur des corpus de textes (quelques milliers d'exemples sont nécessaires). Voici quelques exemples de cette approche :
 - a. Classification: À partir d'une fonction objective, nous allons estimer la probabilité qu'une page web appartienne, ou non, à une catégorie donnée. Des exemples de catégories sont : page d'accueil, page contact, page erreur, etc.
 - b. Régression : Estimation de la valeur numérique à partir d'une fonction. Dans le cas de l'analyse des sentiments du texte, il peut s'agir de l'estimation de la valeur de la polarité⁶ d'une phrase.

iv. Analyse et prise de décision

Cette étape est menée par une validation et une interprétation humaine des résultats obtenus lors des étapes précédentes. Dans un contexte d'entreprise, cette étape est critique car les données sont placées dans un contexte spécifique avant un processus de prise de décision, qui se traduit généralement par une action automatique (sauvegarde de l'information, filtrage de l'information, rejet de l'information). L'application de la prise de décision automatique en le web mining se trouve par exemple dans l'indexation des pages web pour les moteurs de recherche, la suppression des publicités, et les systèmes de contrôle parental [34].

B. Le Web Mining et la détection des manipulations de l'information

Dans le processus de décision, il est impératif pour le décideur de pouvoir non seulement filtrer ces informations mais aussi de pouvoir distinguer celles qui sont trompeuses ou mensongères. En outre, parmi l'ensemble des informations du Web, le type d'information trouvé le plus souvent est textuelle [35], par conséquent, le filtrage des contenus indésirables du web mining semble adéquat pour accompagner le processus de décision dans un environnement des manipulations de l'information.

⁵La bibliothèque la plus utilisée à ce jour en Python est scikit-learn.

⁶La polarité fait référence à l'identification de l'orientation des sentiments (positif, neutre et négatif)

Au cours des deux dernières décennies, le développement des algorithmes linguistiques a permis, jusqu'à un certain niveau, la discrimination entre les informations factuelles et fausses dans les textes numériques [36]. En général, ces systèmes suivent les étapes du web mining de la manière suivante :

1. Recherche et pré-processing de l'information pour l'identification et l'extraction des indices de tromperie.
2. Découverte automatique de patrons généraux avec l'application des modèles de détection des tromperies établis à partir des indices identifiés précédemment.
3. Analyse des patrons pour procéder à la prise de décision de détection

Ce travail d'exploration se concentrera sur la première étape du processus, l'identification et l'extraction des indices.

i. Les indices de tromperie

Malgré l'absence d'un signe littéral de pour identifier les manipulations de l'information, les textes contiennent des indices spécifiques indiquant des différences évidentes entre les textes trompeurs et les textes véridiques. Quand nous utilisons des connaissances linguistiques pour identifier des indices, comme par exemple le sens des mots ou le rapport entre les syllabes et les mots, nous utilisons des indices linguistiques (*linguistic-based cues* en anglais).

Trouver des indices linguistiques par le biais d'une inspection manuelle est complexe, car cela est une tâche chronophage et les résultats peuvent varier d'une personne à l'autre. Donc, l'adoption des approches de NLP (analyse morphologique, syntaxique, lexicale ou sémantique) ainsi que des méthodes statistiques et de machine learning, ont permis d'extraire et d'analyser automatiquement des indices linguistiques à partir d'une grande quantité de corpus de texte⁷.

ii. Ressources en langue française

Dans la révision bibliographique, nous n'avons pas trouvé des bases de données sur la détection de indices linguistiques en français. Le Tableau 2 présente une liste des ressources trouvées en ligne sur détection de la tromperie (*deception*) à partir de différentes sources textuelles (avis, forums, critiques, etc.). Cela constate qu'il n'existe pas de ressources facilement accessibles en français, d'autre part, il en existent de nombreuses ressources en anglais et quelques uns en espagnol, en italien et en chinois.

⁷Ce domaine est connu sous le nom de Détection automatique des indices linguistiques (*Automatic Linguistics-Based Cues Detection* en anglais)

Tableau 2 : résumé des ressources disponibles pour la détection des tromperies .

Summary of available resources for deception detection, according to the source, size, domain, language and methodology to be obtained, where M=Manual, A=Automatic, SCE=sanctioned controlled experiment, UCE=unsanctioned controlled experiment and CR=Crowdfunding.

Work	Source	Size	Lang.	Method
Almela et al. (2012)	Opinions	600 statements	SP	M (SCE)
Chen and Chen (2014)	Web forum	632,234 messages	CH	M (SCE)
DePaulo et al. (2003)	Messages	120 samples	EN	M (UCE)
Fitzpatrick and Bachenko (2012)	–	35,090 words	EN	M
Fornaciari and Poesio (2012)	Hearing words	3015 utterances	IT	M
Fuller et al. (2009)	Written text	366 statements	EN	M
Jindal and Liu (2008)	Amazon rev.	5.8M reviews	EN	A (SCE)
Li et al. (2011)	Epinions reviews	6000 rev.	EN	M
Li et al. (2014)	Opinions	2924 opinions	EN	M
Lim et al. (2010)	Amazon rev.	50 reviews	EN	M
Litvinova et al. (2016)	Written texts	1800 documents	RU	M (SCE)
Mihalcea and Strapparava (2009)	Video and written	200 statements	EN	M (CR)
Newman et al. (2003)	Video and written	568 samples	SP	M (SCE)
Ott et al. (2011)	TripAdvisor rev.	800 opinions	EN	M (CR)
Pérez-Rosas and Mihalcea (2014)	Opinions	750 statements	EN, SP	M (CR)
Pérez-Rosas and Mihalcea (2015)	–	7168 sentences	EN	M (CR)
Rubin and Lukoianova (2015)	Texts	54 stories	EN	M
Rubin et al. (2015)	Transc. radio	144 news	EN	M
Wu et al. (2010)	TripAdvisor rev.	29,799 reviews	En	A

iii. Études liées à la détection d'indices

- Grammaire de la Fake News [37]: En fonction du format de l'information donnée, se pratiquent deux types d'écritures : « une grammaire traditionnelle possédant la cohérence de la narration pour les nouvelles longues répond une grammaire de l'énonciation pour les Tweets ou les Posts Facebook courts ». Cette étude présente une série de procédés grammaticaux comme :
 - Il y a la présentation d'une entité dont l'existence est reconnue (université, hôpital, etc.) , mais il ne s'agit pas forcément d'une institution unique (Université Grenoble-Alpes, CHU de Grenoble, etc.). De même, pour donner une légitimité à l'histoire, les docteurs, les chercheurs, etc. seront cités.
 - la création de personnage, avec description et nom.
- Détection des tromperies dans la communication par ordinateur [38]: Il s'agit d'un étude exploratoire centré sur l'identification des indices de tromperie et suggère que « la tromperie est modéré par des facteurs contextuels et les indices identifiés dans un type de contexte de tromperie peuvent ne pas convenir à d'autres contextes ». En outre, le travail utilise un modèle de classification avec 19 variables parmi lesquelles on trouve:
 - Verbes : Nombre d'occurrences de verbes
 - Longueur moyenne des phrases = nombre total de mots / nombre total de phrases
 - Longueur moyenne des mots = nombre total de caractères / nombre total de mots

- Pausalité = nombre total de signes de ponctuation / nombre total de phrases
- Emotivité = (nombre total d'adjectifs + nombre total d'adverbes) / (nombre total de noms + nombre total de verbes)

CHAPITRE 5 – PROBLEMATISATION

Au cours de la partie théorique, nous avons pu constater que le processus de prise de décision tant dans la vie personnelle que dans l'entreprise et la qualité de l'information (est-ce une information fidèle à la réalité ou transformé dans un but manipulateur ?) peut être affecté à cause des caractéristiques des plateformes numériques et à l'exploitation des vulnérabilités humaines. Ce phénomène est potentialisé par la quantité des informations à laquelle on est exposé dans nos sociétés actuelles. En réponse à cela, il existe actuellement certains outils de détection de manipulations de l'information, mais ils ne sont pas adaptés à l'entreprise car ils ne tiennent pas compte des limites importantes du processus de décision, en particulier le temps, la quantité et la simultanéité des informations. C'est dans cette conjoncture que des technologies telles que l'intelligence artificielle (en particulier le NLP et le machine learning) peuvent réduire ces limitations, permettant une diminution de la complexité du processus décisionnel dans l'entreprise.

Ce raisonnement nous amène à la question suivante :

Quelle solution informatique pour faciliter la prise de décision dans un environnement de manipulations de l'information ?

Ce document propose d'explorer l'hypothèse générale :

H0 : une solution utilisant le NLP et le machine learning facilite le processus de prise de décision en diagnostiquant les manipulations d'informations

PARTIE 2

-

PARTIE PRATIQUE

CHAPITRE 6 – STAGE CHEZ PROLEADS

Pour tester ces hypothèses je me suis basé sur mon expérience dans la mise en œuvre de la procédure du web mining au cours des missions du stage que j'ai effectué chez ProLeads.

I. PRESENTATION DE L'ENTREPRISE

ProLeads est une startup localisée à Saint-Martin d'Hères. Cette startup dont le démarrage de l'activité commerciale a commencé fin 2019, comprend actuellement trois associés. Son activité est centrée sur la réalisation d'un moteur de recherche d'entreprises B2B, au niveau français. Leur proposition de valeur est concentrée sur quatre points :

1. Recherches impartiales : il n'y a pas de financement par les entreprises référencées donc les résultats sponsorisés et la publicité sont inexistantes.
2. Libre choix du périmètre géographique : ciblage de un ou plusieurs territoires de France.
3. Souplesse des critères de recherche : libre choix des mots clés et pas de nomenclature imposée.
4. Confidentialité de vos données : Pas de commercialisation des données personnelles.

II. DESCRIPTION DES MISSIONS

Mes missions consistaient principalement à améliorer la qualité des résultats du moteur de recherche d'entreprises.

1. Mise en place d'une approche d'analyse linguistique permettant de traiter l'équivalence des formes lexicales des activités dérivant de la même racine présentes dans un corpus de textes décrivant l'activité d'entreprises. Par exemple "L'entreprise xxx vend des articles de sport" = "Nous vendons des articles de sport" = "Vente d'articles de Sport" = "Les articles de sport vendus par xxx ..."
2. Réalisation de modèles prédictifs pour classer les pages d'un site. L'indexation d'un site d'entreprise requiert des opérations de filtrage préalables afin d'éliminer les pages qui ne décrivent pas l'activité de l'entreprise (ex : pages de mentions légales, actualités, partenaires, etc).

III. MISE EN ŒUVRE DE MISSIONS

A. Mission 1 : analyse lexicale

Cette mission devait répondre à la problématique suivante :

Quand un utilisateur du moteur de recherche fait une requête, par exemple *usinage de pièces*, il s'attend à obtenir des résultats d'entreprises qui usinent des pièces, ce qui signifie qu'il s'attend à des résultats qui contiennent « *l'usinage de pièces...* » mais aussi « *Nous usinons des pièces ...* » ou « *Les pièces usinées...* ». C'est pourquoi le moteur de recherche devrait pouvoir traiter les équivalences lexicales et présenter à l'utilisateur les options plus pertinents liées à une forme lexicale.

Pour répondre à cette problématique et effectuer la mission, nous avons, dans un corpus de textes (plus de 3 millions de pages web) décrivant l'activité d'entreprises, utilisé des techniques de NLP avec Python pour traiter l'équivalence des formes lexicales qui suivent les structures syntaxiques ci-après :

- nom d'action + préposition + groupe nominaux
- forme verbale du nom d'action + préposition + groupe nominaux

Ce travail a été réalisé à partir d'une première liste de 174 noms d'action utilisée comme index pour extraire les structures. Une fois identifiés, ils sont exportés dans un fichier .tsv⁸ pour une manipulation ultérieure.

Cette action a donné les résultats⁹ suivants :

- 169 noms d'action ont été identifiés
- En moyenne, 269 formes différentes ont été identifiées pour chaque nom d'action. Les extrêmes sont *brevetage* sans structure identifié et *utilisation* avec 2661 structures identifiées

B. Mission 2 : classification des pages web

Cette deuxième mission s'attachait à traiter deux situations qui posent des problèmes pour l'indexation correcte des sites web d'entreprises car elles ajoutent des informations qui ne sont pas pertinentes pour la description de leurs activités.

- i. Au sein du site web d'une entreprise, on trouve différents types de pages qui, chacune, peut contenir du texte plus ou moins pertinents pour la description de l'activité de l'organisation. Par exemple, la page de contact, même si elle contient des informations sur les coordonnées de l'entreprise, la plupart du temps elle ne contient pas d'informations sur son activité.

⁸Format du document de texte représentant des données tabulaires sous forme de « valeurs séparées par des tabulations »

⁹Un exemple de document .tsv est présenté à l'annexe 8.2

- ii. De plus, pour des raisons techniques dans le processus de collecte d'informations, il est parfois possible que les données textuelles récupérées ne soient pas pertinentes, comme c'est le cas des pages d'erreur (404 not found), des pages en cours de maintenance, des pages piratées, etc.

Pour effectuer cette mission, nous avons suivi les étapes du Web Mining. Nous avons prélevé des échantillons du corpus initial (plus de 3 millions de pages web) en utilisant des algorithmes développés par ProLeads pour identifier deux types de pages : pages erreur et pages contact. Ensuite, nous avons sélectionné et prétraité le texte en utilisant les techniques NLP pour identifier des attributs pour générer un modèle prédictif de classification basé sur la régression logistique. Enfin, nous avons itéré les attributs du modèle en évaluant une fonction de perte dans chaque itération.

Résultats

- Le modèle de classification des pages d'erreur que nous avons développé montre une exactitude de 99.7 % avec une augmentation de 15 % (5 261 pages de plus) de pages erreur identifiées par rapport à l'approche précédente.
- Le modèle de classification des pages contact présente quant à lui une exactitude de 98 % avec une augmentation de 10% (10 927 pages de plus) de pages erreur identifiées par rapport à l'approche précédente.

PARTIE 3

-

SYNTHESE ET DISCUSSION

CHAPITRE 7 – EXPLORATION DE LA PRISE DE DECISION DANS UN CONTEXTE DE MANIPULATIONS DE L'INFORMATION

I. CHOIX DU SUJET

Afin d'explorer dans un autre périmètre la problématique, nous allons changer la source d'information du corpus de textes de la description de l'activité des entreprises à une type l'information destinée au grand public, les actualités numériques. Dans ce rapport nous avons examiné la façon dont les technologies d'information se prêtent aux manipulations de l'information et cela nous a incité à choisir les *fakenews* comme sujet à analyser. Le choix spécifique du COVID 19 a été fait en raison de sa nature complexe, étant donné la grande quantité d'informations générées sur ce sujet, associe a de nombreux cas de manipulations de l'information, mais aussi parce qu'il génère de nombreux débats en société et des questionnements en entreprise. Ainsi, si un manager souhaite prendre une décision relative à ce sujet, par exemple le port du masque par ses employés il sera confronté aux limites évoqués dans la revue de littérature. Ce qui est attendu de cette démarche est d'étudier l'adéquation de l'application des techniques de NLP et son lien avec la diminution de la complexité du processus de décision.

II. CHOIX DE LA METHODOLOGIE

Les méthodes employées pour réaliser les missions de stage se révèlent intéressantes pour traiter la problématique de la prise de décision dans un contexte de manipulations de l'information : En effet :

- Pour des raisons de marketing, les entreprises peuvent être amenées à présenter des informations qui ne sont pas équivalentes à l'ensemble des services qu'elles offrent en réalité. Autrement dit, l'action de transmettre une information pour créer une fausse impression ou conclusion.
- Parmi les limites identifiées dans le processus décisionnel, on trouve la capacité insuffisante à traiter toutes les informations disponibles (**Error! Reference source not found.**). Dans l'entreprise, comme dans la vie personnelle, nous sommes soumis à une surcharge d'informations qui, avec un objectif de manipulation des informations, peut exploiter cette vulnérabilité et altérer le processus de prise de décision.

La méthodologie expliquée dans la partie **Error! Reference source not found.** a été suivie, car elle utilise techniques du NLP avec la possibilité d'ajouter algorithmes de machine learning. Ainsi cette approche sera explorée pour la détection de manipulations de l'information. Le corpus de textes

CONCLUSION

À travers de la revue de la littérature, nous avons pu constater que le processus de prise de décision tant dans la vie personnelle que dans l'entreprise et l'authenticité de l'information peut être affecté à cause des caractéristiques des plateformes numériques et à l'exploitation des vulnérabilités le raisonnement limité des humaines. Ce phénomène est potentialisé par la quantité des informations à laquelle on est exposé dans nos sociétés actuelles.

De plus, les résultats de ce travail d'exploration nous permettent de conclure que l'hypothèse initiale (**H0**), d'une solution utilisant le NLP et le machine learning facilite le processus de prise de décision en diagnostiquant les manipulations d'informations, est non infirmée. Basé sur mon expérience dans la mise en œuvre de la procédure du web mining une méthodologie pour l'identification des manipulations d'informations a été proposée. Cette méthodologie s'est révélée prometteuse car elle combine la capacité d'analyser simultanément de grandes quantités de texte, cela permet de surpasser des capacités limitées de l'humaine. De plus, dans l'exercice exploratoire du COVID 19, ce processus a permis d'identifier un type d'indice pour détecter les manipulations d'informations : la répétition exacte des titres sur une même rumeur indique une propension à classer une rumeur comme une *fakenews*.

La complexité technique et conceptuelle en termes d'identification des manipulations d'informations a été une limite mise en évidence lors de la phase d'exploration. En outre, le manque de ressources, telles que des jeux de données, facilement accessibles en français, rend difficile l'identification automatique de ce phénomène.

Un travail envisagé à court terme serait d'annoter un ensemble de pages web de rumeurs afin de créer un jeux de données en langue française et tester l'approche d'apprentissage automatique supervisé avec des indices linguistiques préalablement définis.

BIBLIOGRAPHIE

1. Cornelius, I. (2002). Theorizing information for information science. *Annual review of information science and technology*, 36(1), 392-425.
2. Définitions : information - Dictionnaire de français Larousse. (s. d.). Consulté 13 août 2020, à l'adresse <https://www.larousse.fr/dictionnaires/francais/information/42993?q=information#42898>
3. INFORMARE : sens de ce mot latin dans le dictionnaire. (s. d.). Consulté 13 août 2020, à l'adresse <http://www.dicolatin.com/FR/LAK/0/INFORMARE/index.htm>
4. Web 2.0 : définition de Web 2.0 et synonymes de Web 2.0 (français). (s. d.). Consulté 28 juillet 2020, à l'adresse http://dictionnaire.sensagent.leparisien.fr/Web_2.0/fr-fr/#Un_terme_surtout_marketing
5. *Données, information, connaissance, savoir : un peu de théorie du management*. (s. d.).
6. Basics of information technology. (s. d.). Consulté 13 août 2020, à l'adresse <https://www.slideshare.net/alihassan1993/fit-I02-introductionit>
7. A Brief History Of Information. (s. d.). Consulté 13 août 2020, à l'adresse https://www.liquidinformation.org/information_history.html
8. Khanna, A., & Kaur, S. (2019). Evolution of Internet of Things (IoT) and its significant impact in the field of Precision Agriculture. *Computers and electronics in agriculture*, 157, 218-231.
9. Joule, R.-V., Beauvois, J.-L., & Deschamps, J. C. (1987). *Petit traité de manipulation à l'usage des honnêtes gens*. Presses universitaires de Grenoble Grenoble.
10. Guéguen, N. (2014). *Psychologie de la manipulation et de la soumission*. Dunod.
11. Garcés, R. (2017). Dictadura Militar Argentina: la estrategia de comunicación durante la Guerra de Malvinas.
12. Susser, D., Roessler, B., & Nissenbaum, H. (2019). Technology, autonomy, and manipulation. *Internet Policy Review*, 8(2).
13. J.-B. Jeangène Vilmer, A. Escorcía, M. Guillaume, J. H. (2018). Les Manipulations de l'information : un défi pour nos démocraties. *Centre d'analyse, de prévision et de stratégie (CAPS) du ministère de l'Europe et des Affaires étrangères et de l'Institut de recherche stratégique de l'École militaire (IRSEM) du ministère des Armées*.
14. post-truth adjective - Definition, pictures, pronunciation and usage notes | Oxford Advanced Learner's Dictionary at OxfordLearnersDictionaries.com. (s. d.). Consulté 14 août 2020, à l'adresse <https://www.oxfordlearnersdictionaries.com/definition/english/post-truth>
15. Baumard, P., Harbulot, C., Huyghe, F.-B., Lucas, D., Moinet, N., Prats, C., ... Valantin, J.-M. (s. d.). LA GUERRE COGNITIVE.
16. Jean-Marie, D. (1965). La propagande politique. *Paris, PUF*.
17. Facebook's Ability to Target « Insecure » Teens Could Prompt Backlash | WIRED. (s. d.). Consulté 16 août 2020, à l'adresse <https://www.wired.com/2017/05/welcome-next-phase-facebook-backlash/>
18. Posetti, J., & Matthews, A. (2018). A short guide to the history of fake news' and disinformation. *International Center for Journalists*, 7.
19. High score, low pay: why the gig economy loves gamification | Business | The Guardian. (s. d.). Consulté 19 août 2020, à l'adresse

<https://www.theguardian.com/business/2018/nov/20/high-score-low-pay-gamification-lyft-uber-drivers-ride-hailing-gig-economy>

20. Bérard, C. (2009). Le processus de décision dans les systèmes complexes: une analyse d'une intervention systémique. Université Paris Dauphine-Paris IX; Université du Québec à Montréal.
21. La teoría de la racionalidad limitada de Herbert Simon. (s. d.). Consulté 24 août 2020, à l'adresse <https://psicologiaymente.com/inteligencia/teoria-racionalidad-limitada-herbert-simon>
22. Fonseca, C. (2016). Toma de decisión:¿ Teoría racional o de racionalidad Limitada. *Kálathos*, 7(1), 1-13.
23. Susser, D., Roessler, B., & Nissenbaum, H. (2018). Online manipulation: Hidden influences in a digital world. *Georgetown Law Technology Review*, Forthcoming.
24. Chaouachi, S. G., & Rached, K. S. Ben. (s. d.). Les déterminants de la tromperie perçue dans la publicité.
25. Biros, D. P., George, J. F., & Zmud, R. W. (2002). Inducing sensitivity to deception in order to improve decision making performance: A field study. *MIS quarterly*, 119-144.
26. EU report on weedkiller safety copied text from Monsanto study | Environment | The Guardian. (s. d.). Consulté 26 août 2020, à l'adresse <https://www.theguardian.com/environment/2017/sep/15/eu-report-on-weedkiller-safety-copied-text-from-monsanto-study>
27. Pénét, P. (2019). Fake science et ignorance stratégique: retour sur les récentes controverses autour de l'austérité et du glyphosate. *Etudes de communication*, (2), 85-102.
28. George, J. F., Giordano, G., & Tilley, P. A. (2016). Website credibility and deceive credibility: Expanding prominence-Interpretation Theory. *Computers in Human Behavior*, 54, 83-93.
29. Zhou, L., Burgoon, J. K., Twitchell, D. P., Qin, T., & Nunamaker Jr, J. F. (2004). A comparison of classification methods for predicting deception in computer-mediated communication. *Journal of Management Information Systems*, 20(4), 139-166.
30. Moukrim, B. (2019). Intelligence artificielle en santé: espoirs et craintes des médecins généralistes.
31. Lemberger, P., & Chaumartin, F. R. (2020). *Le traitement automatique des Langues: Comprendre les textes grâce à l'intelligence artificielle*. Dunod.
32. Etzioni, O. (1996). The World-Wide Web: quagmire or gold mine? *Communications of the ACM*, 39(11), 65-68.
33. Provost, F., & Fawcett, T. (2013). *Data Science for Business: What you need to know about data mining and data-analytic thinking*. « O'Reilly Media, Inc. »
34. Hashemi, M. (2020). Web page classification: a survey of perspectives, gaps, and future directions. *Multimedia Tools and Applications*, 1-25.
35. Su, K.-Y., Su, J., Wiebe, J., & Li, H. (2009). Proceedings of the ACL-IJCNLP 2009 Conference Short Papers. In *Proceedings of the ACL-IJCNLP 2009 Conference Short Papers*.
36. Saquete, E., Tomas, D., Moreda, P., Martinez-Barco, P., & Palomar, M. (2020). Fighting post-truth using natural language processing: A review and open challenges. *Expert Systems with Applications*, 141, 112943.
37. MONTCLAIR, F. (s. d.). Grammaire de la Fake News. Deux modalités de l'écriture de la Fake news en média court et média long. *INFORMATION, COMMUNICATION ET HUMANITÉS NUMÉRIQUES. Enjeux et défis pour un enrichissement épistémologique*, 387.

38. Zhou, L., Twitchell, D. P., Qin, T., Burgoon, J. K., & Nunamaker, J. F. (2003). An exploratory study into deception detection in text-based computer-mediated communication. In *36th Annual Hawaii International Conference on System Sciences, 2003. Proceedings of the* (p. 10-pp). IEEE.
39. Fausses informations : les données du Décodex en 2017. (s. d.). Consulté 31 août 2020, à l'adresse https://www.lemonde.fr/le-blog-du-decodex/article/2017/12/19/fausses-informations-les-donnees-du-decodex-en-2017_5231605_5095029.html
40. Publicité: le choc, c'est chic - Libération. (s. d.). Consulté 14 août 2020, à l'adresse https://www.liberation.fr/france/2010/06/09/publicite-le-choc-c-est-chic_657800
41. Balmissé, G. (2008). Recherche d'information en entreprise: une question de gouvernance. *Usages, usagers et compétences informationnelles au 21e siècle. Hermès*, 71-95.
42. Smoke gets in your eyes: 20th century tobacco advertisements | National Museum of American History. (s. d.). Consulté 14 août 2020, à l'adresse <https://americanhistory.si.edu/blog/2014/03/smoke-gets-in-your-eyes-20th-century-tobacco-advertisements.html>
43. Le nudge marketing au service de l'écologie et de la société de demain. (s. d.). Consulté 14 août 2020, à l'adresse <http://gnitekram.fr/le-nudge-marketing-au-service-de-lecologie-et-de-la-societe-de-demain/>

TABLES DES FIGURES

FIGURE 1 : MODELE HIERARCHIQUE DE LA CONNAISSANCE[41].	12
FIGURE 3 : PERSPECTIVE HISTORIQUE DE L'INFORMATION [6].	12
FIGURE 4 : LES PERIODES DE L'ERE DE L'INTERNET [8].	13
FIGURE 5 : PUBLICITE EN FAVEUR DU TABAC AUX ÉTATS-UNIS QUI PRESENTE (A) DES AVIS D'EXPERTS ET (B) DES CHIFFRES PRECIS. [42].	14
FIGURE 6 : PUBLICITE AU SERVICE LA SECURITE ROUTIERE [40].	15
FIGURE 8 : LOGO GRENOBLE IAE	25
FIGURE 9 : UNE PAGE WEB VUE (A) DANS UN NAVIGATEUR WEB, (B) DANS UN EDITEUR DE TEXTE, (C) EN TANT QUE STRUCTURE HIERARCHIQUE	26
FIGURE 10 : NUAGE DE MOTS DU CORPUS DE TEXTE	38

TABLES DES MATIERES

CHAPITRE 1 – L'INFORMATION, UNE RESSOURCE FIABLE ?.....	11
I. Clarifications conceptuelles	11
A. Définition.....	11
B. Repères historiques.....	12
C. L'information, aujourd'hui	13
II. Les manipulations de l'information	14
A. Les manipulations à l'ère numérique	16
B. Conséquences sur la société et l'entreprise	17
CHAPITRE 2 – PRISE DE DECISION DANS UN ENVIRONNEMENT DE MANIPULATION D'INFORMATIONS.....	19
I. Prise de décision et rationalité	19
II. Influencer le processus de prise de décision	19
A. Facebook et les adolescents.....	20
B. Cambridge Analytica et les élections.....	20
C. Uber et le comportement des travailleurs	20
CHAPITRE 3 – DETECTER LES MANIPULATIONS DE L'INFORMATION.....	22
I. Les Décodeurs du <i>Monde</i>	22
II. Hoaxbuster.....	22
CHAPITRE 4 – L'INTELLIGENCE ARTIFICIELLE COMME OUTIL D'AIDE A LA DECISION	23
I. NLP : linguistique informatique	23
A. Analyse lexicale : l'approche classique	24
B. Limites du NLP	24
II. Web Mining et la prise de décision automatique	25
A. Les étapes du Web Mining	26
B. Le Web Mining et la détection des manipulations de l'information	27
CHAPITRE 5 – PROBLEMATISATION	31
CHAPITRE 6 – STAGE CHEZ PROLEADS	33
I. Présentation de l'entreprise	33
II. Description des missions	33
III. Mise en œuvre de missions	33
A. Mission 1 : analyse lexicale.....	33
B. Mission 2 : classification des pages web	34
CHAPITRE 7 – EXPLORATION DE LA PRISE DE DECISION DANS UN CONTEXTE DE MANIPULATIONS DE L'INFORMATION	37
I. Choix du sujet	37
II. Choix de la méthodologie	37

