

Titre: Développement d'un outil d'appariement client-fournisseur dans un
Title: contexte de transformation numérique

Auteur: Hugo Savi
Author:

Date: 2022

Type: Mémoire ou thèse / Dissertation or Thesis

Référence: Savi, H. (2022). Développement d'un outil d'appariement client-fournisseur dans
Citation: un contexte de transformation numérique [Mémoire de maîtrise, Polytechnique
Montréal]. PolyPublie. <https://publications.polymtl.ca/10173/>

 **Document en libre accès dans PolyPublie**
Open Access document in PolyPublie

URL de PolyPublie: <https://publications.polymtl.ca/10173/>
PolyPublie URL:

**Directeurs de
recherche:** Robert Pellerin, & Christophe Danjou
Advisors:

Programme: Maîtrise recherche en génie industriel
Program:

POLYTECHNIQUE MONTRÉAL

affiliée à l'Université de Montréal

**Développement d'un outil d'appariement client-fournisseur dans un contexte
de transformation numérique**

HUGO SAVI

Département de mathématiques et de génie industriel

Mémoire présenté en vue de l'obtention du diplôme de *Maîtrise ès sciences appliquées*

Génie industriel

Février 2022

POLYTECHNIQUE MONTRÉAL

affiliée à l'Université de Montréal

Ce mémoire intitulé :

Développement d'un outil d'appariement client-fournisseur dans un contexte de transformation numérique

présenté par Hugo SAVI

en vue de l'obtention du diplôme de *Maîtrise ès sciences appliquées*

a été dûment accepté par le jury d'examen constitué de :

Camélia DADOUCHI, présidente

Robert PELLERIN, membre et directeur de recherche

Christophe DANJOU, membre et codirecteur de recherche

Taha ARBAOUI, membre

REMERCIEMENTS

Je tiens à remercier mes directeurs de recherche, M. Robert PELLERIN et M. Christophe DAN-JOU, professeurs à Polytechnique Montréal, pour l'encadrement et le précieux support qu'ils m'ont apporté tout au long du projet d'étude dans des conditions de travail parfois compliquées par la pandémie mondiale.

Je tiens à remercier Mme Amira BOUTOUCHENT, M. Mehdi DRISSI et tous les membres de leur équipe pour l'intégration et la confiance qu'ils m'ont accordée au sein de l'entreprise partenaire de ce projet de recherche. Ils m'ont été d'une grande aide et d'un grand support tout au long de ce projet.

Pour terminer, je remercie ma famille, en particulier mes parents, ainsi que mes amis, qui malgré la situation mondiale et la grande distance pour certains d'entre eux ont su me donner un énorme support moral et des encouragements.

RÉSUMÉ

Le contexte actuel de transformation numérique et d'industrie 4.0 pousse l'intégralité des entreprises à se mettre à jour au niveau de la numérisation de leurs processus ou de l'automatisation de leurs procédés. Cependant, ce milieu est très complexe et les entreprises qui désirent s'y lancer nécessitent un guide pour s'y retrouver et s'orienter vers les personnes les plus aptes à répondre au mieux à leurs besoins. Notre partenaire industriel pour ce projet de recherche, la start-up montréalaise BRIDGR, est présente dans ce type d'accompagnement, mais fait face à une problématique d'efficacité de leur processus d'appariement entre leurs clients et le ou les fournisseurs répondant le mieux à leurs exigences. Nous avons cherché à répondre à cette problématique en proposant un outil d'appariement client/fournisseur à partir d'un modèle de système de recommandation basé sur le contenu appliqué au domaine de la transformation numérique.

Nos recherches, basées sur la littérature scientifique et sur des observations réalisées chez notre partenaire, nous ont permis de mieux cerner les différentes caractéristiques « fournisseur » essentielles à mettre en avant lors de la recherche de fournisseurs, mais également de comprendre les critères de réussite d'un système de recommandation et les différentes perspectives d'évaluation possibles.

De plus, lors de ce projet, BRIDGR a exprimé des exigences de construction pour l'outil d'appariement final que nous avons dû prendre en compte lors de notre développement de celui-ci comme la présence de leur expertise avant de transmettre les résultats au client.

La méthodologie DRM (Design Research Methodology) a été employée afin de répondre à nos besoins spécifiques. En effet, cette méthodologie de recherche mixte lie des études empiriques à des approches scientifiques dans le but d'identifier les éléments clés à mettre en place pour la réussite de notre outil d'appariement.

Cette proposition est donc un outil d'appariement entre une requête client et une liste de fournisseurs faisant partie de la base de données du partenaire. Cet outil d'appariement se décompose en deux couches, comme requis par le partenaire, afin de pondérer l'impact des différentes caractéristiques « fournisseur » sur le résultat final.

La première couche relève d'un filtrage de tous les fournisseurs présents dans la base de données par rapport à des critères exclusifs que sont les expertises du domaine de la transformation numérique auxquelles répond le fournisseur, les langues parlées ou encore le type de service/produit fourni. Cette première étape fournit une liste composée de tous les fournisseurs ayant passé le filtre. Cette liste plus courte est ensuite passée dans la seconde étape qui va l'ordonner. Ce tri est réalisé en fonction de la similarité entre les fournisseurs restants et la requête émise par le client par rapport à d'autres critères comme le secteur d'activité principal ou encore le nombre d'années d'expérience. À cette étape, nous avons fourni huit algorithmes différents pour traiter toutes les variantes de calcul de similarité pertinentes identifiées dans la littérature. Pour le bon fonctionnement de notre méthode, nous avons dû modifier le format des requêtes client et également de la base de données fournisseur. Nous avons également classifié les différentes expertises présentes dans le spectre de la transformation numérique.

Nous avons appliqué la méthode aux fournisseurs actuels de la base de données de BRIDGR pour des requêtes fictives, car aucune ancienne requête n'était utilisable. De plus, nous n'avons pas encore implémenté l'outil en situation réelle et nous avons uniquement fait des tests à l'interne. De ce fait, nous n'avons pas de résultats concernant la précision ou la satisfaction client, mais le critère de rapidité d'exécution autour duquel la problématique de notre partenaire était basée est résolu. Nous avons également identifié pour la suite des pistes de réflexion afin d'améliorer notre outil et ainsi de prendre en compte encore plus de caractéristiques pertinentes.

ABSTRACT

The current context of digital transformation and industry 4.0 is pushing all companies to update and automate their processes. However, this environment is very complex and companies that want to get started need a guide to find their way around and orient themselves towards the people best able to meet their needs. Our industrial partner for this research project, the Montréal-based start-up BRIDGR, is involved in coaching but is currently facing a problem of efficiency in their process of matching their clients with the supplier(s) that best meet their requirements. We sought to address this issue by proposing a customer/supplier matching tool based on a content-based recommendation system model applied to the digital transformation domain.

Our research, based both on scientific literature and on observations made at our partner's, allowed us to better identify the various essential supplier characteristics to be put forward during the search for a supplier, but also to understand the success criteria of a recommendation system and the various possible evaluation perspectives. Furthermore, during this project, BRIDGR expressed construction requirements for the final matching tool that we had to consider during our development of the tool such as the presence of their expertise in a second step before transmitting the results to the client.

The DRM (Design Research Methodology) was used to meet our specific needs. Indeed, this research methodology is mixed and links empirical studies to scientific approaches in order to identify the key elements to be put in place for the success of our matching tool.

This proposal is therefore a matching tool between a customer query and a list of suppliers from the partner's supplier database. This matching tool is decomposed into two layers as required by the partner to weight the impact of different supplier characteristics on the final result. The first layer is a filtering of all suppliers in the database against exclusive criteria such as the supplier's expertise in the digital transformation domain, the languages spoken, or the type of service/product provided. This first step provides a list of all suppliers that have passed the filter. This shorter list is then passed to the second step which will sort it. This sorting is done according to the similarity between the remaining suppliers and the request made by the customer against other criteria such as the main sector of activity or the number of years of experience. At this stage we provided eight different algorithms to handle all the relevant similarity calculation variants identified in the

literature. For the proper functioning of our method, we had to modify the format of the client queries and of the supplier database. We also classified the different expertises present in the digital transformation spectrum.

We applied the method to the current suppliers in the BRIDGR database for dummy queries because no old queries were usable. In addition, we have not yet implemented the tool in real life and have only tested it internally. As a result, we do not have any results concerning accuracy or customer satisfaction, but the speed of execution criterion around which our partner's problem was based has been solved. We have also identified avenues of reflection for the future to improve our tool and thus consider even more relevant characteristics.

TABLE DES MATIÈRES

REMERCIEMENTS	III
RÉSUMÉ.....	IV
TABLE DES MATIÈRES	VIII
LISTE DES TABLEAUX.....	XI
LISTE DES FIGURES.....	XII
LISTE DES SIGLES ET ABRÉVIATIONS	XIII
LISTE DES ANNEXES.....	XIV
CHAPITRE 1 INTRODUCTION	1
CHAPITRE 2 REVUE DE LITTÉRATURE.....	3
2.1 Les méthodes de recommandation et la production participative	3
2.1.1 Terminologie	3
2.1.2 Point commun : le profil utilisateur.....	4
2.2 Stratégie de recherche et analyse de l'état de l'art	5
2.2.1 Stratégie de recherche.....	5
2.2.2 Revue des travaux de la littérature sur les systèmes de recommandation.....	8
2.3 Revue critique.....	25
2.4 Conclusion.....	28
CHAPITRE 3 MÉTHODOLOGIE DE RECHERCHE	29
3.1 Objectifs de recherche	29
3.2 Méthodologie de recherche	30
3.2.1 Revue de littérature	31
3.2.2 Première étude descriptive	32
3.2.3 Développement de l'outil d'appariement	33

3.2.4 Évaluation des performances de l'outil d'appariement.....	33
3.3 Conclusion.....	33
CHAPITRE 4 ÉTUDE DESCRIPTIVE	34
4.1 Mise en contexte.....	34
4.2 Processus d'appariement client/fournisseur	35
4.3 Critères de réussite de l'outil.....	37
4.4 Conclusion.....	41
CHAPITRE 5 CONSTRUCTION DE L'OUTIL D'APPARIEMENT	43
5.1 Partie intrants (expertises et critères exclusifs, inclusifs, etc.)	43
5.1.1 Caractéristiques des fournisseurs et des requêtes clients	43
5.1.2 Expertises et mots-clés.	44
5.1.3 Plages des caractéristiques	46
5.1.4 Caractéristiques exclusives, inclusives et de collecte de données.....	47
5.2 Filtre	49
5.3 Tri.....	51
5.3.1 Caractéristiques générales	52
5.3.2 Prise en compte des années d'expérience.....	53
5.3.3 Prise en compte des mots-clés.....	54
5.3.4 Conclusion.....	57
5.4 Conclusion.....	58
CHAPITRE 6 ÉVALUATION DES PERFORMANCES DE L'OUTIL	59
6.1 Processus global d'évaluation des performances de l'outil	59
6.2 Résultats et analyses	63
6.2.1 Tests de la somme des rangs de Wilcoxon.....	63

6.2.2 Tests de scalabilité et de rapidité.....	69
6.3 Conclusion.....	73
CHAPITRE 7 CONCLUSION ET RECOMMANDATIONS.....	74
7.1 Conclusion.....	74
7.2 Perspectives d'amélioration	76
RÉFÉRENCES.....	79
ANNEXES	84

LISTE DES TABLEAUX

Tableau 2.1 : Résultats retenus par notre recherche bibliographique	8
Tableau 2.2 : Tableau synthétique des éléments d'un système de recommandation.....	27
Tableau 4.1 : Critères industriels de l'outil d'appariement et méthodes d'évaluation.....	38
Tableau 4.2 : Liste complète des critères d'évaluation de l'outil.....	41
Tableau 5.1 : Récapitulatif des caractéristiques des profils fournisseur et des requêtes client	45
Tableau 6.1 : Résultats du test de la somme des rangs de Wilcoxon de la première campagne	63
Tableau 6.2 : Résultats des tests de la somme des rangs de Wilcoxon pour la deuxième campagne	65
Tableau 6.3 : Résultats des tests de la somme des rangs de Wilcoxon pour la troisième campagne	66
Tableau 6.4 : Résultats des tests de la somme des rangs de Wilcoxon pour la quatrième campagne	67
Tableau A.1 : Tableau récapitulatif de la recherche bibliographique.....	84
Tableau D.1 : Liste des secteurs d'activités.....	89

LISTE DES FIGURES

Figure 2.1 : Processus de recherche bibliographique systématique	7
Figure 2.2 : Phases principales dans une structure de système de recommandation (adapté de Grida et al.,2020).....	9
Figure 2.3 : Structure générale d'un marché à microtâches (adaptée de Schnitzer 2019).....	15
Figure 2.4 : Recommandation de tâches dans un marché à microtâches (adapté de Schnitzer 2019)	16
Figure 2.5 : Algorithme TOP-K-W (extrait de Safran & Che 2015)	17
Figure 2.6 : Base hiérarchique de connaissances sémantique (adapté de Li et al., 2002).....	22
Figure 2.7 : Système de flux d'étude (adapté de Jatnika et al., 2019)	23
Figure 2.8 : Illustration du MD-algorithm et du MDB-tree (adapté de Sun et al. 2013)	24
Figure 3.1 : Phases de la méthodologie DRM liées aux phases du projet de recherche	31
Figure 4.1 : Symboles du formalisme BPMN (tiré de la norme ISO/CEI 19510 : 2013).....	35
Figure 4.2 : Diagramme BPMN du processus.....	36
Figure 5.1 : Logigramme de fonctionnement de l'algorithme	49
Figure 6.1 : Logigramme du protocole d'évaluation	60
Figure 6.2 : Résultats de la première campagne sur la scalabilité.....	70
Figure 6.3 : Résultats de la deuxième campagne sur la scalabilité	70
Figure 6.4 : Résultats de la troisième campagne sur la scalabilité	71
Figure 6.5 : Résultats des tests réalisés "à vide"	72
Figure 7.1 : Logigramme du scénario proposé d'évaluation de la satisfaction client.....	77
Figure B.1 : Liste des expertises	85
Figure C.1 : Liste des glossaires par expertise	86

LISTE DES SIGLES ET ABRÉVIATIONS

CSAT Customer Satisfaction Score ou Score de Satisfaction du Client

TF Term Frequency ou Fréquence du Terme

SR Système de Recommandation

LISTE DES ANNEXES

Annexe A : Tableau récapitulatif de la recherche bibliographique.....	84
Annexe B : Liste des expertises	85
Annexe C : Liste des glossaires par expertise	86
Annexe D : Liste des secteurs d'activité.....	89

CHAPITRE 1 INTRODUCTION

Depuis son apparition en 2011 lors du Forum Mondial de l'Industrie à Hanovre, le terme Industrie 4.0 est omniprésent dans le domaine du génie industriel, et ce, sous plusieurs formes : usine intelligente (smart factory), fabrication digitale (digital manufacturing) ou encore usine du futur. Ce terme désignant une nouvelle forme d'industrie se décline en fonction du secteur industriel concerné ou encore du pays. Les piliers principaux de cette industrie nouvelle sont l'interconnexion entre les appareils, les composants et les personnes, la transparence des informations, l'assistance technique et enfin la prise de décisions décentralisées. Les piliers sont eux-mêmes soutenus par le développement et l'utilisation de nouvelles technologies comme l'Internet des Objets (IoT) ou encore de nouveaux procédés de fabrication comme la fabrication additive. Cette révolution numérique, ou quatrième révolution industrielle pousse les dirigeants d'entreprises et gestionnaires à revoir leurs modèles d'affaires et à considérer l'intégration de nouvelles technologies dans leurs processus afin de demeurer compétitifs.

Cette transformation numérique amène les différentes entreprises du secteur à identifier tous les nouveaux acteurs de ce milieu en constante évolution. Cela comprend par exemple les fournisseurs de nouvelles technologies ou encore les experts et autres consultants spécialisés pouvant les accompagner dans leur processus de transformation. L'exploration de ce milieu ultra-compétitif en constante et rapide évolution représente un défi très important, notamment pour les petites ou moyennes entreprises. Elles n'ont en effet pas le luxe de consacrer autant de ressources que les grandes entreprises.

L'accompagnement des petites ou moyennes entreprises en matière de transformation numérique est donc primordial. Cet accompagnement peut prendre différentes formes, telles que le suivi de l'évolution de la maturité numérique du client ou encore l'établissement d'une feuille de route de projets. L'un des éléments primordiaux reste le fait que les projets établis ou recommandés ne peuvent être réalisés de manière autonome. Ils nécessitent l'implication d'un ou plusieurs fournisseurs spécialisés. La recommandation de fournisseurs ou de formateurs spécialisés dans une technologie ou une formation très précise est donc une étape charnière pour le bon déroulement du processus.

Les clients étant dans la plupart des cas des néophytes du milieu, leurs requêtes peuvent être floues ou peu précises. De plus, les positions de néophytes pour les entreprises impliquent également une maîtrise approximative, voire inexistante, du vocabulaire spécifique du milieu. C'est pour cela qu'il est nécessaire d'attribuer une grande importance à la compréhension précise de ces requêtes.

La réalisation de ce mémoire s'est faite en collaboration avec un partenaire industriel spécialisé dans le domaine de la transformation numérique et de l'industrie 4.0. Ce partenaire faisait face aux problématiques évoquées ci-dessus lors de leur appariement manuel de fournisseurs avec leurs clients en fonction des critères mis en avant. Leur méthode actuelle est cependant très chronophage en raison de la taille de leur base de données de fournisseurs.

Devant cette problématique industrielle, ce mémoire vise à développer un système d'appariement de fournisseurs de solutions 4.0. Le système d'appariement doit offrir un découpage du milieu de la transformation numérique et de l'industrie en termes d'expertises diverses et variées qui peuvent traiter de domaines extrêmement variés tels que l'infonuagique ou encore l'essor des technologies propres. Ce découpage permet d'identifier plus facilement le domaine d'expertise demandé par le client et les spécifications supplémentaires sont quant à elles traitées sémantiquement par l'existence de labels. Ce système peut également être relié aux autres outils du partenaire afin d'associer par exemple des recommandations de projets à des fournisseurs directement.

Ce mémoire s'organise de la façon suivante. Tout d'abord, une revue de littérature autour des termes de notre problématique est réalisée au chapitre 2. Ensuite, le chapitre 3 vient préciser la méthodologie de recherche mise en place tout au long de la réalisation de ce projet. Ensuite, les chapitres 4 et 5 traitent de la construction de l'outil d'appariement requête/fournisseur qui sera ensuite évalué dans le chapitre 6. Finalement, nous terminerons ce mémoire par une conclusion et des recommandations.

CHAPITRE 2 REVUE DE LITTÉRATURE

Ce chapitre permet de prendre connaissance des progrès récents liés à notre sujet. Le chapitre débute avec une présentation des méthodes et domaines importants liés aux différentes méthodes de recommandation avant de poursuivre avec notre stratégie de recherche permettant d'identifier les travaux scientifiques pertinents de ces domaines. Les articles résultant de cette stratégie de recherche sont ensuite présentés. Le chapitre se termine par une revue critique qui souligne les limitations et manques des travaux retenus afin de préciser les objectifs de recherche spécifiques de ce projet.

2.1 Les méthodes de recommandation et la production participative

Cette première partie a pour but d'introduire brièvement les méthodes de recommandation classiques connues ainsi que de traiter du concept de production participative afin de mettre en avant les différents points communs existant, mais également de déterminer des mots-clés pouvant nous servir pour notre stratégie de recherche.

2.1.1 Terminologie

Le terme « recommandation » possède plusieurs sens. L'Académie française dans sa 9^e édition, l'actuelle, en donne six différentes dont la majorité s'éloigne de notre domaine. En général, une recommandation consiste en un « *avis que l'on donne à autrui pour l'engager à faire ou ne pas faire quelque chose* »¹.

La recommandation est donc un échange entre deux personnes, l'une possédant un savoir, étant experte dans un domaine quelconque, et l'autre cherchant à acquérir ce savoir, des informations sur ce fameux domaine. La communication et le partage de l'information semblent être les fondements de la recommandation. Dans notre contexte actuel d'abondance de données et d'informations diverses et variées sur tous les domaines existants, il est souvent très difficile pour les néophytes d'un domaine de s'y retrouver. Des techniques informatiques ont donc été proposées afin de faciliter et d'automatiser l'accès aux informations. La plus répandue est la méthode des

¹ Définitions successives du mot « recommandation » sur le Site Web de l'Académie française : <https://www.dictionnaire-academie.fr/article/A9R0923>

systèmes de recommandation qui repose sur une forme particulière de filtrage qui vise à présenter les éléments d'information susceptibles d'intéresser l'utilisateur de ce système comme cela est défini par Adomavicius et Tuzhilin (2005). Cela peut fonctionner, entre autres, à partir de caractéristiques de référence (par exemple le type de film pour un système de recommandation dans le domaine du cinéma) grâce auxquelles le système arrive à prédire l'avis qu'aurait l'utilisateur sur tel ou tel item du catalogue puis recommande ceux qui possèdent l'avis prédit le plus favorable.

La production participative (ou crowdsourcing en anglais) aussi connue sous le nom d'externalisation ouverte est un concept qui est défini comme étant un « mode de réalisation d'un projet ou d'un produit faisant appel aux contributions d'un grand nombre de personnes »². Le journaliste américain Jeff Howe, qui est à l'origine en 2006 du terme crowdsourcing, en donne une explication simplifiée dans Howe (2006) : « des gens comme vous et moi qui utilisent leur temps libre à créer du contenu, à résoudre des problèmes... ». Il existe donc des plateformes mettant en relation des internautes désireux de résoudre des problèmes contre rémunération avec des entreprises qui elles ont un problème sur les bras, mais pas le temps pour le résoudre par exemple. Certes, certains problèmes ne nécessitent aucune expertise spécialisée afin d'être résolue, mais d'autres en nécessitent. Certaines plateformes servent donc à mettre en relation des entreprises faisant face à un problème nécessitant une expertise particulière qui n'est pas présente au sein de l'entreprise avec un travailleur externe qui possède cette expertise. Un exemple de ce type de plateforme spécialisée est Amazon Mechanical Turk. Ce type de marché est appelé marché à microtâches. L'existence de ces plateformes implique également l'existence d'une méthode d'association entre les travailleurs et les tâches proposées, et ce, en fonction des expertises de ces travailleurs.

2.1.2 Point commun : le profil utilisateur

Dans les deux situations explorées brièvement plus haut, le but de la recommandation est de faire en sorte qu'elle soit le plus possible centrée sur les attentes de l'utilisateur. Il est en effet peu utile de recommander quelque chose qui est très loin des préférences de l'utilisateur. Afin de saisir ces

² Article sur la production participative sur le Site Web Bank Innovation Competence Center : <http://www.baicc.news/la-production-participative-crowdsourcing-la-puissance-de-la-foule/>

fameuses attentes, il est donc nécessaire de construire un profil utilisateur qui regroupe les caractéristiques préférées par l'utilisateur en fonction des données récoltées par le système dans les items qui lui seront recommandés. Ces items doivent également posséder un ensemble de caractéristiques afin de les distinguer et recommander le plus adéquat. Ce point commun implique donc un profil utilisateur, un profil item, des caractéristiques communes sur ces profils, mais également des méthodes afin de comparer ces caractéristiques.

Si nous considérons les différentes caractéristiques comme des vecteurs, il existe plusieurs manières de calculer la similarité entre des vecteurs et donc de déterminer quels items correspondent le mieux aux requêtes clients.

2.2 Stratégie de recherche et analyse de l'état de l'art

2.2.1 Stratégie de recherche

Notre revue de littérature est effectuée sur les bases de données Scopus ainsi que sur Inspec et Compendex, mais également sur Proquest afin de ne pas passer à côté de thèses pouvant se révéler intéressantes. Cette revue nous a permis d'identifier les articles traitant de systèmes de recommandation recommandant des experts, mais également des articles traitant de la recommandation de tâches dans un contexte de crowdsourcing ou bien encore de méthode de comparaison de textes. À la suite de l'analyse effectuée précédemment ainsi que de plusieurs recherches initiales nous avons pu effectuer notre recherche avec les mots-clés suivants : **((("recommendation system" OR "recommender system") AND ("expert" OR "expertise") AND ("supplier" OR "service" OR consultan*)) OR (("task assignment" OR "task recommendation" OR "recommendation algorithms ") AND ("crowdsourcing" OR "micro-task market")) OR ((similar* OR "comparaison" OR "similarity-based") AND ("words" OR "mots" OR phrase* OR text* OR "text based") AND (analys* OR seman* OR syntax*))).**

Ayant effectué une recherche bibliographique systématique, le choix et la justification des mots-clés sont primordiaux dans le processus de recherche. Le premier groupe de mots-clés dans la recherche **((("recommendation system" OR "recommender system") AND ("expert" OR "expertise") AND ("supplier" OR "service" OR consultan*))** correspond à la recherche traitant des

systèmes de recommandations en tant que tels. Le premier bloc correspond aux différentes appellations des systèmes de recommandation rencontrées lors des recherches initiales. Le deuxième bloc traite de la présence d'experts ou de la recommandation par des experts. Le dernier bloc met en avant la recherche de recommandation concernant des fournisseurs, consultants ou des fournisseurs de services.

Le deuxième groupe de cette recherche ((`"task assignment"` OR `"task recommendation"` OR `"recommendation algorithms"`) AND (`"crowdsourcing"` OR `"micro-task market"`)) aborde de son côté l'aspect des algorithmes et des modèles de production participative et notamment l'aspect d'assignation des tâches qui nous intéresse le plus. En effet, la problématique du partenaire industriel est d'associer des tâches ou des projets à des fournisseurs en fonction des caractéristiques demandées pour la tâche et des compétences du fournisseur.

Le troisième et dernier groupe de cette recherche ((`similar*` OR `"comparaison"` OR `"similarity-based"`) AND (`"words"` OR `"mots"` OR `phrase*` OR `text*` OR `"text based"`) AND (`analys*` OR `seman*` OR `syntax*`)) traite enfin de la comparaison et de la recherche de similarité entre différents mots ou groupes de mots ou textes quelconques. Les différents types d'analyse à savoir sémantique et syntaxique sont également présents afin de ne pas passer à côté d'un modèle qui ne traiterait que l'une des deux solutions.

Le tableau en Annexe A fait un récapitulatif des trois groupes de recherche et de leurs concepts.

Les recherches préliminaires nous ont aussi permis de limiter nos recherches aux documents traitant au moins d'informatique (Computer Science, Computer Technology) ou bien de technologies d'information (Information Technology). Nous limitons également notre recherche à des papiers en anglais ou en français. Ces limitations permettent de supprimer une partie des textes qui ne seront pas intéressants pour nous, car traitant de domaines très éloignés du cas d'étude. Cependant, certains articles peuvent passer ses critères de sélection sans répondre aux besoins spécifiques de cette revue de littérature. La lecture attentive de chaque article retenu peut ainsi mener à l'exclusion de certains papiers dans un deuxième temps. La figure 2.1 suivante présente ainsi le processus de recherche bibliographique systématique qui a été mis en place afin de déterminer les papiers les plus pertinents pour notre cas d'étude. Le tableau 2.1 suivant présente les résultats obtenus en appliquant ce processus.

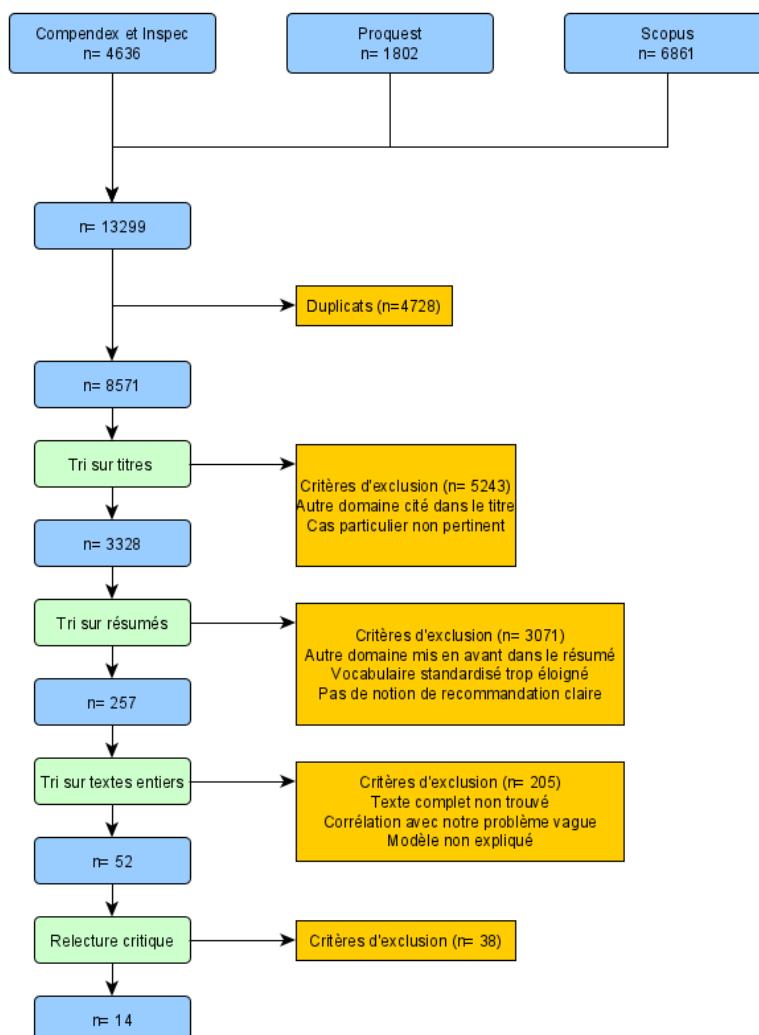


Figure 2.1 : Processus de recherche bibliographique systématique

Tableau 2.1 : Résultats retenus par notre recherche bibliographique

	Articles	Auteurs
1	A structured framework for building recommender system	Grida et al. (2020)
2	Building an expert recommender chatbot	Cerezo et al. (2019)
3	Preselection of documents for personalized recommendations of job postings based on word embeddings	Schnitzer et al. (2019)
4	Task Recommendation in Crowdsourcing Platforms	Schnitzer (2019)
5	Task Recommendation in Crowdsourcing Systems	Yuen et al. (2012)
6	Task Recommendation in Crowdsourcing Systems	Yuen (2016)
7	Real-time recommendation algorithms for crowdsourcing systems	Safran et Che (2015).
8	Comparaison de textes : quelques approches...	Negre (2013)
9	Comparaison de trois mesures de similarité utilisées en documentation automatique et analyse textuelle.	Lelu, (2002)
10	An Approach for Measuring Semantic Similarity between Words Using Multiple Information Sources	Li et al. (2002)
11	Word2Vec Model Analysis for Semantic Similarities in English Words	Jatnika et al. (2019).
12	Privacy-preserving Multi-keyword Text Search in the Cloud Supporting Similarity-based Ranking	Sun et al. (2013)

2.2.2 Revue des travaux de la littérature sur les systèmes de recommandation

2.2.2.1 Définition des systèmes de recommandation

La première question à se poser est qu'est-ce qu'un système de recommandation ? Les systèmes de recommandation sont une forme spécifique de filtrage de l'information visant à présenter les

éléments d'information (films, musique, livres, nouvelles, images, pages Web, etc.) qui sont susceptibles d'intéresser l'utilisateur. Un système de recommandation est basé sur le profil d'un utilisateur à certaines caractéristiques de référence, et cherche à prédire l'avis que donnerait un utilisateur.³ Ce profil est dynamique dans le temps et s'actualise en fonction des données disponibles. Ces caractéristiques dépendent du type de système de recommandation utilisé, mais nous reviendrons sur les différents types des SRs plus tard. Les domaines d'utilisation des systèmes de recommandation sont variés, mais gravitent la plupart du temps autour du commerce électronique. Les utilisations les plus parlantes sont par exemple la recommandation d'objets similaires à celui ou ceux que nous venons d'acheter ou souvent achetés avec ces derniers sur les sites marchands ou encore la recommandation de films ou de séries similaires à ceux que nous avons l'habitude de regarder sur des plateformes de vidéos à la demande.

Dans (Grida et al., 2020), les auteurs passent en revue les différents points importants à étudier pour la mise en place d'un système de recommandation. Pour cela, ils traitent de chacun des types de systèmes de recommandation un par un en énumérant leur fonctionnement ainsi que leurs différents avantages et inconvénients spécifiques. En plus, ce texte décrit également une structure classique afin de construire un système de recommandation qui est décrite dans la figure 2.2 suivante.

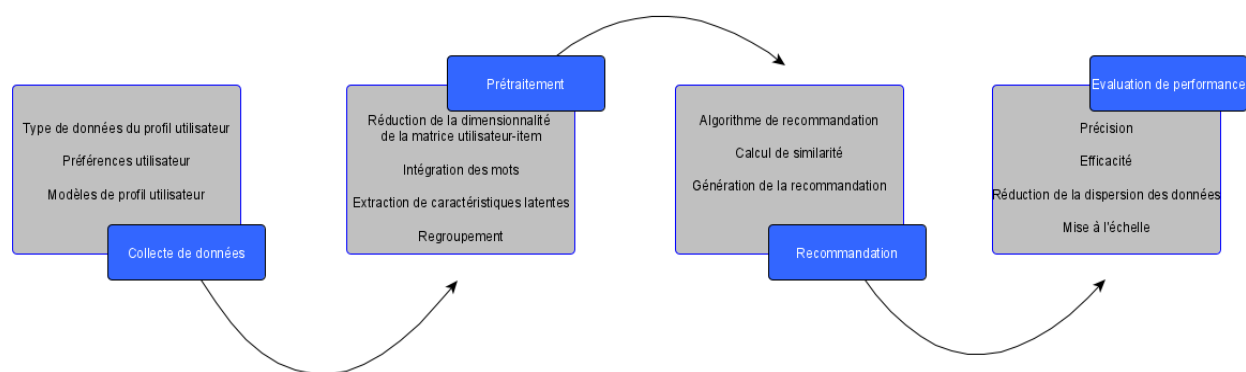


Figure 2.2 : Phases principales dans une structure de système de recommandation (adapté de Grida et al.,2020)

³ Page Wikipédia du concept de *Système de recommandation*. Modifiée pour la dernière fois le 9 janvier 2021.

Selon (Grida et al., 2020), une structure de recommandation classique se sépare en quatre phases principales chacune se séparant en différentes sous-phases successives. La première phase est celle de la collecte de données qui consiste à recueillir les préférences de l'utilisateur. Pour cela, il existe deux méthodes différentes, l'explicite et l'implicite. La première, l'explicite, utilise les données collectées explicitement donc à travers de notations réalisées par l'utilisateur ou bien de likes ou dislikes attribués à certains items. Certains systèmes de recommandation interactifs peuvent même recueillir des retours utilisateurs en temps réel. Cette méthode permet d'avoir des recommandations très précises, mais requiert un minimum d'implication pour l'utilisateur. La seconde, l'implicite, consiste à suivre le comportement des utilisateurs comme les pages ouvertes ou encore sur les boutons cliqués. Le point fort de cette méthode est le fait que l'utilisateur n'a rien à faire à proprement parler; le simple fait d'être sur le site en question collectera des données pour son profil utilisateur. Cependant, cette approche souffre du problème de démarrage à froid dans lequel le système peine à donner des recommandations pertinentes dans le cas où peu de données ont été récoltées. Ce problème intervient généralement lors du démarrage d'un nouveau système de recommandation et que les données récoltées sont trop dispersées pour assurer la précision des méthodes de prédiction de l'utilité.

Une fois les données collectées il est nécessaire de les stocker quelque part, et c'est là que selon (Grida et al.,2020) intervient le profil utilisateur. Ils nous présentent également les différents types de profils utilisateurs présents dans la littérature ainsi que leurs avantages et inconvénients propres. Les six types de profil présentés sont les suivants :

- Modèle de profil par mots-clés;
- Modèle de profils par vecteurs;
- Modèle de profil sémantique;
- Modèle de profil basé sur l'incertitude;
- Modèle de profil probabiliste; et
- Modèle de profil par histogramme.

Toujours dans (Grida et al.,2020), la seconde phase présentée est celle du prétraitement. Elle peut être présente sous plusieurs formes différentes en fonction du système de recommandation. L'étape

de prétraitement est bénéfique pour analyser le contenu des articles afin d'extraire des caractéristiques comme les concepts ou les mots-clés en utilisant des algorithmes de recherche d'information.

La troisième phase présente dans la structure de (Grida et al., 2020) est la phase de recommandation qui contient elle-même trois étapes. La première est l'algorithme de recommandation en tant que tel qui peut être basé sur le filtrage collaboratif ou bien basé sur le contenu. Il existe également des méthodes hybrides utilisant différents éléments de chacune de ces deux méthodes. D'après Adomavicius et Tuzhilin (2005), les systèmes de recommandation basés sur le filtrage collaboratif essaient de prédire l'utilité des items pour un utilisateur particulier en prenant en compte les items déjà notés par d'autres utilisateurs qui sont considérés comme semblables à notre utilisateur cible. Les systèmes basés sur le contenu prédisent quant à eux l'utilité d'un item pour un utilisateur grâce à l'utilité assignée par cet utilisateur à d'autres items semblables à l'item cible.

La deuxième étape est celle de calcul de similarité qui peut être effectué de multiple manière en fonction du type de similarité recherché (Corrélation de Pearson⁴, Similarité basée sur le Cosinus⁵).

Enfin, la dernière phase de la structure d'un système de recommandation selon (Grida et al., 2020) est celle d'évaluation des performances. Plusieurs critères de performance sont mis en avant, tels que la précision, l'efficacité, la possibilité de changer d'échelle ou encore le traitement de la dispersion des données. Ces critères sont ceux décrits comme les plus importants dans ce papier et s'intéressent à la différence entre ce qui est recommandé par le système et ce qui était attendu par l'utilisateur, mais également aux méthodes utilisées par le système afin de pallier le problème de démarrage à froid évoqué précédemment.

2.2.2.1.1 Exemples de systèmes de recommandation présentant des similarités avec notre contexte

⁴ Cours sur la corrélation de Pearson par l'Université de Liège en Belgique : http://www.biostat.ulg.ac.be/pages/Site_r/corr_pearson.html

⁵ Explications simples sur le fonctionnement de la similarité en cosinus : <https://ichi.pro/fr/comprendre-la-similarite-cosinus-et-son-application-278464028813674>

Dans notre recherche bibliographique, deux articles se sont retrouvés en avant, le chatbot recommandant des experts en programmation (Cerezo et al., 2019) ainsi qu'un système de recommandation recommandant des annonces pour des postes en fonction de mots-clés utilisés par les deux partis (Schnitzer et al., 2019).

Dans (Cerezo et al., 2019), le produit mis en avant est un chatbot, c'est-à-dire une intelligence artificielle avec qui nous pouvons discuter afin de lui demander qui sont les experts dans un domaine spécifique. Cette intelligence artificielle est sur la plateforme Discord⁶ et recommande des experts dans certains packages, méthodes ou cours concernant un langage de programmation à code ouvert (Pharo). Le fonctionnement de ce chatbot est séparé en trois éléments indépendants qui sont :

- Le module de classification des phrases de l'utilisateur qui reconnaît et classe les requêtes utilisateurs de manière empirique dans différentes catégories. La requête est classifiée à l'aide d'un algorithme basé sur la fréquence des termes (TF). Cette fréquence est ensuite comparée à celles contenues dans une base de données de fréquences qui correspondent à telle ou telle catégorie de requêtes;
- L'identification des concepts clés qui comme son nom l'indique identifie le concept clé de la requête utilisateur. Dans ce module l'algorithme utilisé est celui appelé IDF (inverse document frequency); et
- La recommandation d'expertise qui repose sur deux approches différentes qui sont tout d'abord le minage d'expertise puis l'utilisation du profil des experts. Ces deux approches permettent d'identifier les meilleurs experts pour la requête. Dans ce cadre la méthode utilisée est basée sur le contenu.

Le second exemple qui semble pertinent est proposé dans (Schnitzer et al., 2019). Cet article traite d'un SR qui recommande des annonces de postes à pourvoir en comparant des listes de tags. En effet, il y a un grand nombre d'annonces de postes qui sont postées chaque jour sur Internet. Il est donc nécessaire d'aider les demandeurs d'emploi à trouver des offres de postes pertinentes. Ce modèle de système de recommandation fonctionne en deux étapes :

⁶ Plateforme de messagerie instantanée notamment utilisée dans les domaines des jeux en ligne ou du codage

- La première étape est la création de profils utilisateurs grâce à la présélection d'annonce de travail nous correspondant. Cette présélection est basée sur le principe de groupe et repose sur une collecte de données exposées à l'utilisateur. Le nombre de clics des différents utilisateurs sert de base pour le tri des différentes annonces. Les différentes annonces de travail sont ensuite positionnées dans un espace vectoriel complexe en fonction de leurs caractéristiques telles que le niveau d'études requis ou encore le domaine concerné. Les groupes réalisés auparavant sont ensuite transposés dans cet espace vectoriel et donc l'intégralité des annonces est positionnée; et
- Une fois ce profil créé, le système de recommandation formule ses recommandations de deux façons différentes : la première basée sur un algorithme des k-voisins les plus proches (k-nearest neighbors ou KNN) qui recommande les annonces en fonction de leur proximité avec la moyenne des autres annonces auxquelles l'utilisateur a déjà postulé. La seconde approche du SR est quant à elle basée sur la création d'un hyperplan dans cet espace vectoriel basé sur l'appréciation des annonces antérieures. Les coordonnées de cet hyperplan sont basées sur les différents retours que l'utilisateur concerné a donnés sur d'autres annonces antérieures. Une fois cet hyperplan créé, les annonces ayant le score le plus élevé par rapport à cet hyperplan sont ensuite recommandées à l'utilisateur.

2.2.2.2 Production participative et marchés à microtâches

La seconde étape primordiale de notre découverte des résultats de cette recherche se concentre sur l'étude des concepts de production participative et de marché à microtâches. Comme expliqué plus haut, la production participative repose sur une implication d'acteurs extérieurs à l'entreprise ou organisme faisant la requête initiale. En exemple de grands projets de production participative connus nous pouvons citer l'encyclopédie participative Wikipédia qui propose à ses utilisateurs de partager leurs connaissances sur des sujets extrêmement variés en étant soumis à juste quelques règles ou encore le projet *Galaxy Zoo*⁷ qui invite les internautes à collaborer afin de classifier des galaxies. Les marchés à microtâches sont quant à eux une manière de mettre en relation des entreprises en recherche d'expertises spécialisées, mais de manière ponctuelle. Dans cette situation, il

⁷ Projet Galaxy Zoo : <https://www.zooniverse.org/projects/zookeeper/galaxy-zoo/>

n'est pas rentable pour l'entreprise d'embaucher un expert de ce domaine si c'est pour ne pas utiliser ses compétences de manière récurrente. Ce genre de marché associe ce genre d'entreprises à des travailleurs possédant les compétences requises et disponibles pour réaliser une tâche à la fois.

2.2.2.2.1 Sélection des tâches et recommandation dans des marchés à microtâches

Dans (Yuen et al., 2012), les auteurs proposent deux éléments intéressants concernant les marchés à microtâches. Tout d'abord, ils réalisent une étude qui s'intéresse à la sélection des tâches par les travailleurs et enfin, ils proposent une structure de recommandation de tâches. Afin de réaliser leur étude, les auteurs se servent de données collectées d'Amazon Mechanical Turk, qui est une plateforme de production participative lancée par Amazon en 2005. Cette étude s'intéresse aux données traitant la similarité des différentes tâches effectuées par le même travailleur. Elle montre qu'en général les travailleurs préfèrent réaliser des tâches qui sont semblables à des tâches qu'ils ont déjà réalisées. De plus, il semblerait que si jamais leur travail précédent n'avait pas été accepté par l'entreprise proposant la tâche, les travailleurs ont tendance à ne plus refaire de tâches de cette catégorie. Enfin, la dernière donnée importante à relever est que les critères de sélection les plus importants pour les travailleurs sont le montant de la rémunération et le type de tâche à réaliser.

La seconde partie du papier traite, comme dit plus haut, l'établissement d'une structure de recommandation de tâche dans ce contexte de marché à microtâches. Afin d'établir cette structure, il est nécessaire de mettre en relation les tâches et les travailleurs. Une matrice de recommandation de tâches basée sur les préférences tâches-travailleurs est donc réalisée. Dans cette matrice, la préférence d'un travailleur envers une tâche est notée sur une échelle d'entier de 1 à 5 avec, 1 le travailleur ne parcourt pas les informations relatives à la tâche et 5, il fournit un travailleur qui est accepté par le demandeur.

Le problème étudié dans cette seconde partie est donc le remplissage de cette matrice. Ce remplissage se fait par factorisation puis décomposition. Les préférences en termes de tâches sont ensuite modélisées à l'aide d'une distribution gaussienne sphérique ainsi que d'une interférence bayésienne. Une fois cette matrice complétée, il est donc facile d'extraire les possibilités qu'un travailleur réalise chaque tâche et donc de sélectionner les tâches présentant la plus haute probabilité.

Dans la poursuite de ses travaux Yuen (2016), présente plusieurs éléments en rapport avec la recommandation de tâches dans un contexte de production participative. Il y complète son travail

de 2012 en introduisant une nouvelle dimension, celle des catégories de tâches. En appliquant la même méthode que précédemment, l'auteur obtient une manière de classer les tâches par travailleur et vice-versa. L'auteur teste également les performances de son modèle. Pour cela, il compare un jeu de données provenant d'Amazon Mechanical Turk avec une simple factorisation des différentes matrices de son modèle en s'intéressant à l'erreur absolue moyenne (MAE) ainsi qu'à l'écart quadratique moyen (RMSE).

Dans Schnitzer (2019), l'auteur présente le domaine de la recommandation des tâches sur des plateformes de production participative, ainsi que les fondements sur lesquels est basée son étude. La structure d'un marché à microtâches est illustrée dans la figure 2.3 ci-dessous. Dans les marchés de microtâches, l'organisation de production participative est appelée le demandeur de la microtâche (requester) et le contributeur est appelé le travailleur (crowd worker). La plateforme fournit le marché de microtâches, qui publie les microtâches et gère l'affectation des travailleurs aux microtâches. Le demandeur propose une microtâche sur le marché des microtâches et donc demande sa complétion, qui est assurée par un travailleur. Ainsi, le travailleur fournit ses compétences et connaissances et reçoit une rémunération, qui peut également être gérée par le marché des microtâches.

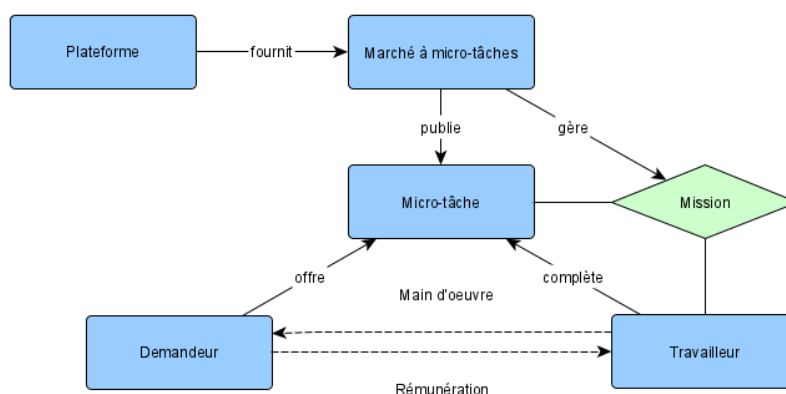


Figure 2.3 : Structure générale d'un marché à microtâches (adaptée de Schnitzer 2019)

Les marchés de microtâches présentent le problème de la nature dynamique des marchés de tâches, où les tâches deviennent indisponibles après leur première interaction (achèvement). Lorsque les demandeurs proposent leurs microtâches sur le marché des microtâches, comme le montre la figure 2.4 ci-dessous, ils fournissent une description textuelle de la tâche, ainsi que d'autres critères tels que le paiement proposé, le temps de travail prévu et le nombre d'heures de travail, etc. Le marché

des microtâches peut exploiter l'historique des affectations du travailleur, analyser les descriptions de tâches respectives, trouver des similitudes entre les tâches et les recommander au travailleur. Par conséquent, du point de vue conceptuel de la recommandation de tâche présentée dans la figure 2.4, le marché des microtâches est en mesure de recommander des missions au travailleur sur la base des similitudes entre les microtâches.

Dans ses travaux qui traitent de la recommandation de tâches en général, l'auteur traite de la similarité entre les différentes microtâches afin de recommander. Tout d'abord, les microtâches sont classifiées en catégories, car d'après l'auteur ces dernières aident le travailleur à faire un choix. Par la suite, chaque tâche possède des attributs qui sont répartis entre quatre catégories :

- Attributs factuels (Paie, employeur, pays, temps alloué, etc.);
- Attributs de contenu (Vocabulaire utilisé dans la description de la tâche mesuré par TF-IDF ou par n-grams);
- Attributs de structure (Style d'écriture utilisé, Nombre de mots, Nombre de signes de ponctuation, etc.); et
- Attributs sémantiques (Sujets, sentiments extraits à l'aide de logiciels de sémantiques et de bibliothèques externes (SentiWordNet)).

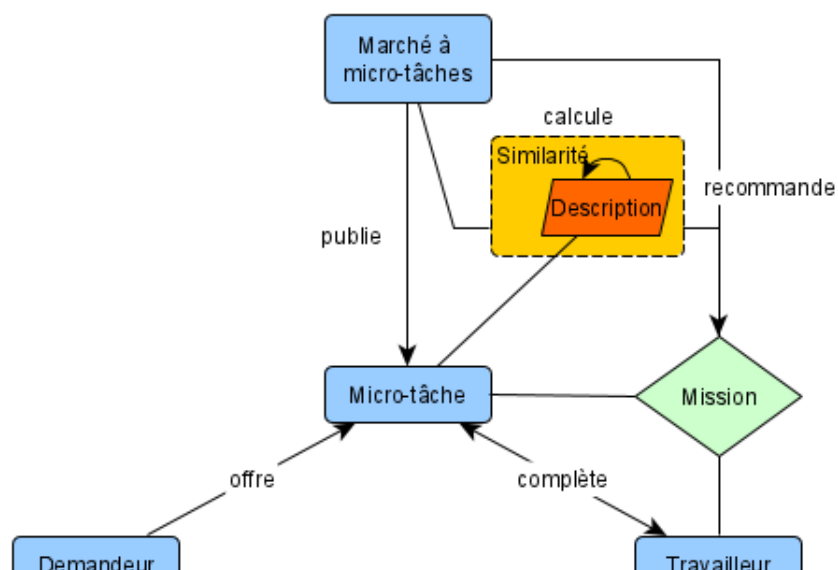


Figure 2.4 : Recommandation de tâches dans un marché à microtâches (adapté de Schnitzer 2019)

Dans l'optique d'évaluer ces différentes approches de classification, un corpus catégorisé à la main est nécessaire. Cette classification est ensuite comparée à la classification réalisée par différents algorithmes de classification que sont Random Forest, K Nearest Neighbors (IBk), Bayes naïf, la classification basée sur des règles (JRip), l'arbre de décision (J48) ainsi qu'une machine à vecteur de support (SMO). Une fois cette évaluation réalisée, la classification la plus efficace est celle réalisée sur les attributs de contenu mesuré par TF-IDF.

2.2.2.2.2 Algorithmes de recommandation en temps réel

Safran & Che (2015) présentent des algorithmes de recommandation en temps réel dans un contexte de systèmes de production participative. Le premier a pour but de déterminer les k meilleures tâches pour un travailleur donné (TOP-K-T) tandis que le second a pour but de déterminer les k meilleurs travailleurs disponibles pour une tâche donnée (TOP-K-W). La figure 2.5 présente l'algorithme brut TOP-K-W.

Input: $C \rightarrow W$: the category-worker data structure.
TaskCat: list of the weighted categories of given task.
k: the number of workers to be recommended

Output: *L*: list of top-*k* workers for recommendation.

```

1 Initialize output array L
2 passtonext  $\leftarrow 0$ 
3 for i  $\leftarrow 1$  to TaskCat.size do
4   Ctarget  $\leftarrow$  getCategory(TaskCat[i])
5   Cweight  $\leftarrow$  getWeight(TaskCat[i])
6   numselected  $\leftarrow$  cascaderound(Cweight  $\times$  k)
7   if numselected + passtonext > Ctarget  $\rightarrow$  W.size then
8     passtonext  $\leftarrow$  passtonext + numselected - Ctarget
       $\rightarrow$  W.size
9     L  $\leftarrow$  all workers in Ctarget  $\rightarrow$  W (not already in L)
10  else
11    L  $\leftarrow$  first (numselected + passtonext) workers from
      Ctarget  $\rightarrow$  W (not already in L)
12  end
13 end

```

Figure 2.5 : Algorithme TOP-K-W (extrait de Safran & Che 2015)

Les deux algorithmes sont basés sur plusieurs éléments primordiaux à traiter. Tout d'abord il y a l'utilisation de catégories. Ces catégories servent à catégoriser les tâches, mais également les travailleurs en fonction de leurs compétences. Une fois cette catégorisation faite, cela accélère le travail de tri et d'attribution des travailleurs à la tâche, car l'algorithme commence par traiter les

travailleurs étant les meilleurs dans la catégorie principale de la tâche en question. Il existe également un lien entre les catégories et les travailleurs. Il s'agit en effet d'un score par couple travailleur/catégorie qui prend en compte le taux d'acceptation de la catégorie par le travailleur, le score de préférence de la catégorie pour le travailleur et enfin la similarité en cosinus entre le profil du travailleur et la description de la catégorie. Chaque fois que le travailleur remplit une tâche, tous ses scores de catégories sont mis à jour de tels que le taux d'acceptation et le score de préférence de catégorie reflètent toutes les tâches effectuées jusqu'à présent. Dans les deux algorithmes, il y a un arrondi en cascade qui intervient au moment de choisir le nombre de tâches ou de travailleurs qui proviennent de telle ou telle catégorie en fonction des catégories de la tâche ou du travailleur initial. Le fait de séparer ainsi permet de ne pas présenter k fois le même profil de travailleur à un demandeur de tâche ou k fois le même type de tâche à un travailleur en recherche.

2.2.2.3 Comparaison de vecteurs

Cette dernière section traite des mesures de similarité entre vecteurs dans le but de pouvoir comparer et de déterminer quels couples de vecteurs sont les plus similaires. Pour cela nous allons nous intéresser au travail réalisé par Elsa Negre pour un cahier du LAMSADE (Laboratoire d'Analyse et de Modélisation des Systèmes d'Aide à la DEcision) traitant des méthodes de comparaison de textes. En premier, Negre (2013) présente les deux grandes approches de comparaison entre vecteurs à savoir l'approche syntaxique et l'approche sémantique. L'approche syntaxique repose la comparaison de vecteurs sur les chaînes de caractère. La comparaison repose donc uniquement sur la forme des vecteurs. En reprenant l'exemple utilisé par l'auteur, « voiture » et « voiturier » sont très proches alors que « voiture » et « automobile » ne le sont pas du tout. Dans cette approche, il est nécessaire de transformer chaque document en un vecteur. Une fois cette transformation réalisée, il existe un grand nombre de mesures de similarité syntaxique existantes. L'auteur liste les mesures suivantes :

- La similarité Cosinus qui calcule le cosinus de l'angle entre les représentations vectorielles des documents à comparer. (Baeza-Yates et Ribeiro-Neto, 1999) et (Huang, 2008) ;
- Le coefficient de corrélation de Pearson qui calcule le cosinus de l'angle entre les représentations centrées réduites des documents vectorisés. (Huang, 2008) ;

- La distance euclidienne qui calcule la similarité comme la distance entre les représentations vectorielles des documents ramenées à un seul point. (Huang, 2008);
- Le coefficient de Jaccard qui est le rapport entre la cardinalité (le nombre d'éléments) de l'intersection des ensembles considérés (les représentations vectorielles des documents) et la cardinalité de l'union de ces ensembles. (Jaccard, 1901) et (Huang,2008);
- La distance de Levenshtein qui calcule la similarité sous forme de chaîne de caractères des documents. La distance est ici le nombre minimal d'opérations pour transformer une chaîne de caractères en l'autre. Les opérations sont par exemple le remplacement d'un caractère ou bien l'ajout ou la suppression d'un caractère. (Levenshtein, 1966) et (Navarro, 2001); et
- L'indice de Dice qui calcule la similarité en se basant sur le nombre de termes communs aux deux documents. (Baake, 2006).

Pour sa part, la similarité sémantique est un concept selon lequel deux documents se voient attribuer une similarité selon la ressemblance de leur signification ou contenu. Il existe trois types d'approches en termes de similarité sémantique. La première est l'approche vectorielle qui consiste à déterminer le sens d'un mot en se servant des autres mots utilisés dans des phrases. Dans cette approche, les mots sont considérés comme des vecteurs ayant un sens et la similarité sémantique est ensuite calculée à l'aide de méthodes de similarité vectorielle. Les deux méthodes les plus utilisées dans cette approche sont celles des vecteurs sémantiques, qui se sert des mots voisins pour déterminer le sens d'un mot et celle du bi-clustering (ou co-clustering).

La deuxième approche est l'approche topologique qui s'appuie sur un réseau sémantique des mots. Dans cette approche, la similarité des mots est déterminée grâce à leurs positions relatives dans la hiérarchie de la base de connaissances. Cette base de connaissances est un arbre où chaque nœud est un concept (tous les mots correspondant à ce concept y sont) et où tous les nœuds enfants de ce nœud sont les concepts hyponymes de notre concept et où tous les nœuds parents sont les concepts hyperonymes de notre concept. Dans cette approche, il existe en grand nombre de méthodes d'évaluation de la similarité dont un certain nombre sont présentées par l'auteur :

- Edge-based : estimation de la distance entre les nœuds correspondants aux concepts à comparer. Plus la distance est courte plus ils sont similaires;

- Leacock et Chodorow : Même méthode que précédemment, mais en prenant en compte la profondeur de l'arbre. En effet plus les arcs entre les concepts sont bas dans l'arbre plus ces derniers sont similaires. Par exemple « voiture » et « voiture de sport » sont très similaires;
- Wu et Palmer : métrique qui mesure la profondeur de deux concepts donnés dans la taxonomie WordNet et la profondeur de leur plus bas ancêtre commun et les combine pour obtenir un score de similarité;
- Resnik : Retourne le contenu de l'information du plus bas ancêtre commun des deux concepts;
- Lin : Normalisation de la méthode de Resnik ci-dessus; et
- Jiang et Conrath : Méthode basée également sur celle de Resnik, mais qui l'inverse.

La troisième approche est l'approche statistique ou basée sur un corpus. Dans cette approche, où chaque méthode est différente, il est notamment possible de compter le nombre de co-occurrences d'un couple de termes dans un corpus afin de déterminer leur similarité. Contrairement aux deux autres approches, cette dernière ne nécessite pas une compréhension du vocabulaire du domaine ou du contexte étudié.

Lelu (2002) présente une comparaison de trois mesures de similarité qui sont utilisées en documentation automatique et en analyse textuelle sur la base de quatre éléments :

- Découpage des unités textuelles;
- Choix de codage des textes ;
- Choix des opérations offertes aux usagers; et
- Transformation des vecteurs-données.

Les trois mesures de similarités concernées par la comparaison sont le cosinus de Salton, la distance du khi-deux ainsi que le cosinus dans l'espace distributionnel. Des comparaisons empiriques sont ensuite effectuées. Le corpus est préparé et la requête (pas trop large ni trop précise) est effectuée.

La comparaison est ensuite effectuée à l'aide de courbes rappel/précision. Le rappel correspond au nombre de documents pertinents retrouvés par rapport au nombre de documents retrouvés tandis que la précision correspond au nombre de documents pertinents retrouvés par rapport au nombre de documents pertinents existants. Le document se conclut en mettant en avant qu'au moins pour une stratégie d'indexation automatique le cosinus distributionnel donne de meilleurs résultats que le cosinus « TF-IDF » de Salton qui est lui-même meilleur que la distance du khi-deux.

Dans (Li et al., 2002), une méthode permettant de mesurer la similarité sémantique entre les mots en utilisant plusieurs sources d'information est présentée. Cette méthode repose sur une base de connaissance sémantique hiérarchique comme traités précédemment dans Nègre (2013). La figure 2.6 ci-dessous présente cette base. La proposition en termes de similarité qui est faite est que la similarité entre deux mots soit une fonction de plusieurs variables. Ces variables étant la longueur du chemin entre les deux, la profondeur du sous-arbre contenant les deux mots ainsi que la densité sémantique locale des deux mots. Les auteurs essaient ensuite d'établir les contributions de chacune des variables dans le calcul final de la similarité. En ce qui concerne la longueur du chemin entre deux mots, plus cette distance tend vers 0 plus la similarité entre les mots est censée tendre vers 1. La profondeur influe le calcul de similarité, car les concepts plus hauts dans l'arbre sont plus vagues tandis que ceux plus bas sont beaucoup plus précis. Deux mots dont le dernier ancêtre commun est haut dans l'arbre devraient donc avoir une similarité moindre que deux mots dont le dernier ancêtre commun est plus bas. La densité sémantique locale est un paramètre plus dur à considérer selon les auteurs. En effet, il n'est possible à obtenir que grâce à des filets sémantiques.

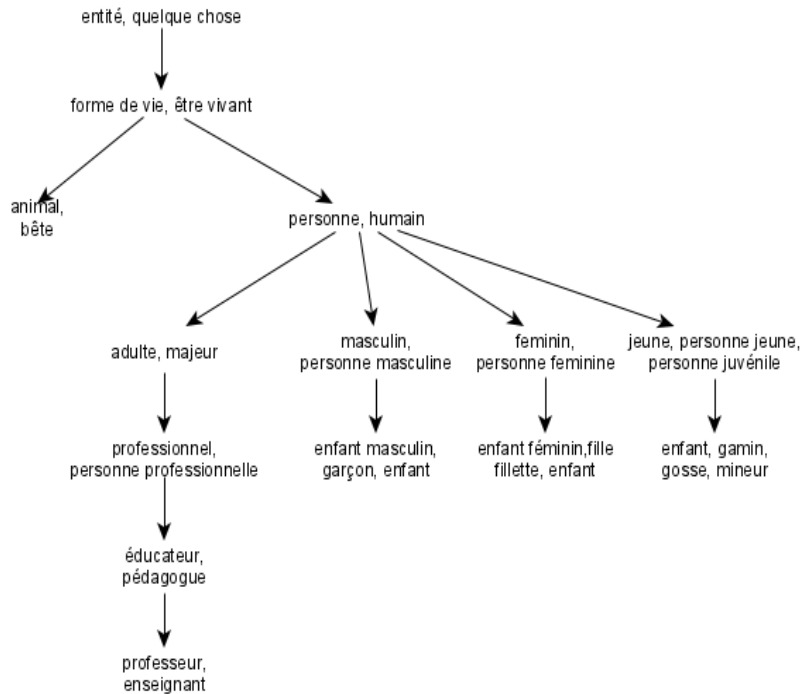


Figure 2.6 : Base hiérarchique de connaissances sémantique
(adapté de Li et al., 2002)

Afin de pouvoir déterminer l'efficacité de cette méthode combinant à la fois longueur, profondeur et contenu de l'information, il est nécessaire de la comparer au bon sens humain tout d'abord puis à d'autres stratégies de mesure de similarité sémantiques. Dans cette expérience dix mesures de similarité sont comparées et c'est finalement un produit de formes non linéaires de la longueur du chemin le plus court ainsi que de la profondeur qui obtient le meilleur résultat en termes de coefficient de similarité avec les méthodes humaines. Les auteurs nous proposent donc une formule pour mesurer la similarité qui est la suivante :

$$S(w1, w2) = e^{-\alpha l} * \frac{e^{\beta h} - e^{-\beta h}}{e^{\beta h} + e^{-\beta h}} \quad (2.1)$$

Avec $\alpha > 0$ et $\beta > 0$ qui sont les paramètres décrivant la contribution respectivement de la longueur du chemin le plus court et de la profondeur de l'arbre. En se basant sur les jeux de données de test, les auteurs ont réussi à déterminer les valeurs optimales pour les paramètres qui sont les suivantes $\alpha = 0.2$ et $\beta = 0.6$.

Dans (Jatnika et al.,2019), les auteurs présentent la représentation en Word2Vec pour une analyse des similarités sémantiques dans des mots anglais. La représentation Word2Vec est un modèle qui sert à transformer les mots en vecteurs et ainsi à appliquer dessus les méthodes de similarité vectorielle comme notamment celle de la similarité en cosinus. Les différentes étapes de l'étude réalisée par les auteurs sont décrites dans la figure 2.7 ci-dessous. Dans cette expérience, les deux paramètres de Word2Vec dont l'influence est mesurée sont la taille de la fenêtre et le nombre de dimensions. La similarité est calculée à l'aide de la formule de la similarité en Cosinus.

$$similarity = \cos \theta = \frac{\vec{x} \cdot \vec{y}}{||\vec{x}|| ||\vec{y}||} \quad (2.2)$$

Les résultats de cette première équation sont ensuite recalculés à l'aide de l'équation de corrélation

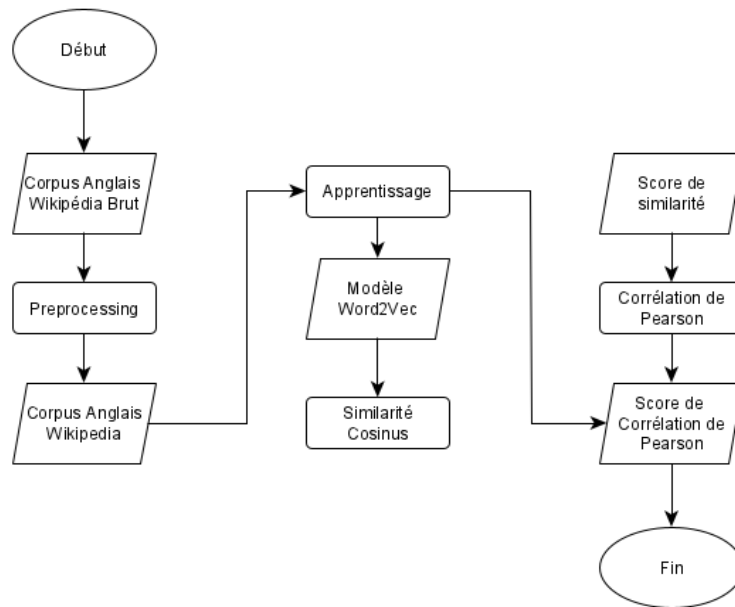


Figure 2.7 : Système de flux d'étude (adapté de Jatnika et al., 2019)

de Pearson afin de déterminer la corrélation entre les résultats obtenus et ceux déterminés par les jeux de test. Les deux études menées en parallèle chacune sur un jeu de tests différent ont mené toutes les deux aux mêmes résultats.

Sun et al. (2013) proposent également un schéma de recherche à plusieurs mots-clés avec en point d'orgue l'assurance de la conservation du secret de la recherche. Ce schéma supporte également un classement basé sur la similarité. Plusieurs éléments mis en avant dans ce papier sont remarquables. Le modèle est composé de trois entités, l'utilisateur de données, le possesseur de données et le serveur infonuagique.

Le modèle met en avant une efficacité de la recherche et celle-ci passe par non pas une recherche linéaire, mais par un index basé sur un arbre. L'algorithme utilisé est appelé MD-Algorithm et la représentation en arbre utilisée est appelée MDB-Tree. Dans cette représentation, chaque niveau de l'arbre représente un attribut de la base de données et chaque élément de la base de données se voit attribuer une valeur d'attribut pour cet attribut. Tous les éléments possédant la même valeur pour un attribut sont regroupés en un nœud fils qui sert ensuite de base à un niveau inférieur et à un autre attribut. La figure 2.8 nous montre un exemple d'arbre avec 3 attributs et 11 éléments dans la base de données.

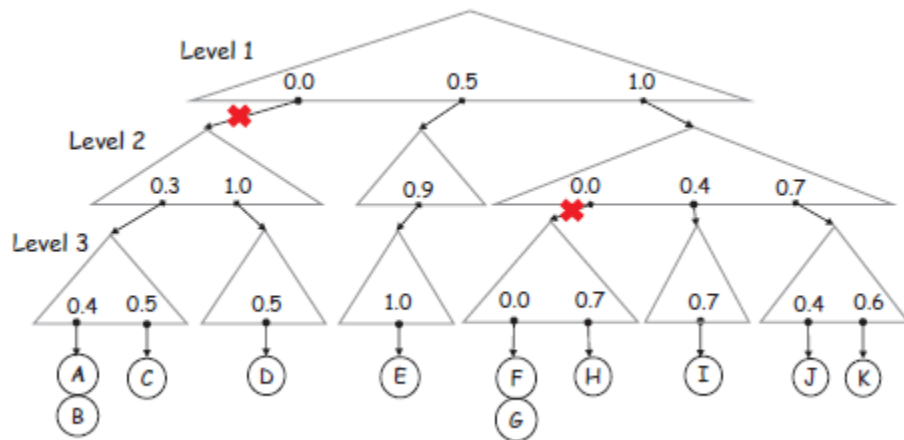


Figure 2.8 : Illustration du MD-algorithm et du MDB-tree (adapté de Sun et al. 2013)

L'un des paramètres de recherche les plus importants est le seuil de prédiction pour chaque attribut qui détermine le niveau maximal qu'il est possible d'atteindre pour cet attribut. Le critère de sélection des différents éléments est le suivant, si la prédiction du score pour l'élément considéré est inférieure ou égale aux scores des top-k éléments déjà traités le procédé de recherche revient au nœud parent et sinon il continue au niveau du nœud enfant. Ce procédé de recherche dans l'arbre est ainsi répété jusqu'à ce que les k éléments présentant les meilleurs scores de similarité sont déterminés. Afin de calculer la similarité entre deux éléments, les auteurs utilisent une fonction de

similarité en cosinus. La réduction du seuil de prédiction permet de réduire le temps nécessaire à parcourir l'arbre, mais si jamais le seuil est fixé trop bas la précision commence à baisser drastiquement. En ce qui concerne la position du mot-clé le plus important en premier dans la et donc que cet attribut soit testé en premier dans l'arbre réduit grandement le nombre de nœuds qui sont parcourus par l'algorithme sans pour autant en réduire la précision.

2.3 Revue critique

L'état de l'art des travaux traitant des sujets des systèmes de recommandation et des algorithmes d'assignement des tâches dans des systèmes de production participative montre qu'il existe déjà un grand nombre de méthodes et d'approches connues et utilisées. Cependant, aucune ne traite de recommandation fournisseur dans le cadre de projets de transformation numérique.

Nous avons vu dans le début de cette revue de littérature qu'il existait deux grandes catégories de systèmes de recommandation, à savoir les SR basés sur le contenu et ceux basés sur le filtrage collaboratif. Si nous nous référons aux cas de SR que nous avons pu étudier au cours de cette revue de littérature nous nous rendons compte que les cas proches de notre contexte d'étude sont résolus en développant un système de recommandation basé sur le contenu. Cette observation nous guide vers ce type de système de recommandation plutôt que vers ceux basés sur le filtrage collaboratif.

Nous pouvons également mettre en avant que la grande majorité des algorithmes de production participative traitent du sens du travailleur vers la tâche. C'est-à-dire que ces algorithmes prennent un travailleur et une liste de tâches et renvoient les meilleures tâches pour cet utilisateur. Ce sens est donc l'inverse de celui de notre problématique. De plus il est nécessaire de mettre ici en avant que les tâches et les projets ne sont pas les mêmes choses. Des tâches composent un projet, mais toutes les tâches d'un projet ne sont pas forcément du même type et donc ne requièrent pas les mêmes expertises.

Enfin, les approches syntaxiques reposent sur la similarité de forme des différents vecteurs qu'elles comparent et non pas sur le sens de ces vecteurs. Cela peut convenir pour des vecteurs dont le sens n'est pas important comme des vecteurs de nombres, par exemple, mais cela semble difficile à utiliser dans le cas de vecteurs où le sens est primordial comme pour des vecteurs de phrases ou de mots-clés.

En utilisant les éléments d'un système de recommandation selon l'architecture décrite dans (Grida et al., 2020), le tableau 2.2 permet de caractériser les modèles issus de la littérature. Comme nous pouvons le voir dans le tableau les cinq derniers papiers étudiés, ceux traitant plus particulièrement de similarité syntaxique ou sémantique se concentrent uniquement sur cet aspect de similarité et de comparaison et traitent peu ou pas du tout des autres aspects importants selon (Grida et al., 2020). D'après ce tableau il n'existe pas de système de recommandation présenté utilisant tous les éléments mis en avant par Grida (2020) et spécialisé dans le domaine de la transformation numérique. Il est nécessaire de se spécialiser dans le domaine de la transformation numérique, car c'est un domaine qui change très vite et souvent, mais qui nécessite néanmoins une précision importante. Le dernier point que nous voulions aborder ici est donc que nous allons nous servir des différents éléments rencontrés au cours de nos lectures pour développer un outil répondant aux attentes de notre partenaire industriel. Cet outil ne pourra cependant pas être un système de recommandation. En effet, l'un des points primordiaux d'un système de recommandation est la présence d'un profil utilisateur dynamique. Or, dans notre cas comme nous désirons recommander des fournisseurs pour réaliser un projet, qui est quelque chose d'éphémère c'est-à-dire qu'il disparaîtra une fois rempli, nous sommes dans l'impossibilité de créer ce profil utilisateur dynamique. C'est pour cela que nous allons développer un système d'appariement.

Tableau 2.2 : Tableau synthétique des éléments d'un système de recommandation

Textes	Collecte de données			Prétraitement				Recommandation			Évaluation des performances			
	Type de données	Préférences utilisateur	Motifs de profils	Réduction des mots	Intégration des mots	Extraction des caractéristiques	Regroupement	Algorithme de recommandation	Calcul de similarité	Génération de la recommandation	Précision	Efficacité	Réduction de la dispersion	Mise à l'échelle
Grida et al. (2020)	X	X	X	X	X	X	X	X	X	X	X	X	X	X
Cerezo et al. (2019)	X					X	X	X		X		X		
Schnitzer et al. (2019)		X	X				X	X	X		X			
Schnitzer (2019)	X					X	X		X		X			X
Yuen et al. (2012)	X	X	X	X		X		X	X					
Yuen (2016)	X	X	X	X		X		X	X		X	X	X	X
Safran et Che (2015)	X	X		X		X	X	X	X	X		X		
Negre (2013)	X								X		X	X		
Leblu (2002)	X								X		X	X		
Li et al. (2002)	X								X		X	X		
Jatnika et al. (2019)	X								X		X	X		
Sun et al. (2013)	X								X		X	X		

2.4 Conclusion

Cette revue de littérature a permis de constater que l'assignation de fournisseurs à des tâches nécessitant des compétences particulières est un sujet qui possède de multiples aspects de résolution possibles. Cette revue a également permis de confirmer qu'il n'existe aucun exemple de ce type de système dans un contexte de fournisseurs de technologies ou de services liés à la transformation numérique. La contribution de ce mémoire sera donc de réunir tous les éléments nécessaires à la construction d'un système de recommandation basé sur des éléments inspirés de ceux découverts lors de cette revue bibliographique tout en les adaptant à un contexte de transformation numérique. Les adaptations nécessaires seront généralement orientées vers l'utilisation d'un système d'appariement et non de recommandation. La méthodologie employée est décrite dans le prochain chapitre.

CHAPITRE 3 MÉTHODOLOGIE DE RECHERCHE

Ce chapitre vise à préciser nos objectifs de recherche ainsi que la méthodologie de recherche employée pour répondre à la problématique du partenaire industriel. Nous présentons ensuite l'approche méthodologique utilisée.

3.1 Objectifs de recherche

Nous avons établi dans les chapitres précédents qu'il n'existe aucun système de recommandation ou d'appariement pour fournisseur dans le domaine de la transformation numérique. En effet, certains projets ou articles présentent des éléments pouvant répondre partiellement à la problématique comme (Safran et Che, 2015) qui présentent un algorithme de recommandation en temps réel des k meilleurs travailleurs disponibles pour effectuer une tâche donnée dans un contexte de marché à microtâches. Cependant, aucun article ne nous présente un système de recommandation présentant tous les éléments principaux mis en avant par (Grida et al., 2020) dans son architecture de système de recommandation.

Notre objectif principal consiste donc à développer un système d'appariement qui permet d'associer de potentiels clients avec les fournisseurs correspondant le mieux aux réquisitions préliminaires énoncées dans une requête client.

De cet objectif principal ainsi que de l'architecture découlent quatre sous objectifs traitant chacun une partie de la construction du système d'appariement :

Objectif 1 : *Identifier les éléments principaux d'un système d'appariement ainsi que les méthodes de calcul de similarité.*

Objectif 2 : *Déterminer les données d'entrée ainsi que les critères de sélection des fournisseurs dans un cas d'étude.*

Objectif 3 : *Développer un algorithme de filtre et tri des fournisseurs prenant en compte les critères établis précédemment.*

Objectif 4 : *Évaluer la qualité du système d'appariement complet en l'appliquant au cas d'étude.*

3.2 Méthodologie de recherche

Cette section détaille la méthodologie proposée afin d’atteindre chacun des sous-objectifs énoncés plus haut. La méthode proposée doit s’intégrer au processus en place chez notre partenaire industriel afin de comprendre son fonctionnement, de proposer une solution et enfin d’évaluer l’évolution du fonctionnement une fois la solution mise en place. Cette contrainte nous permettra en revanche d’accéder à une liste de fournisseurs préqualifiés. L’intégration avait également pour but d’accéder à d’anciennes requêtes déjà traitées par le partenaire industriel de manière manuelle afin de les comparer aux résultats que nous obtiendrons à l’aide de notre système. Nous avons opté ici pour une méthodologie mixte connue sous nom de *Design Research Methodology*, qui est à la fois empirique avec des phases descriptives, mais également expérimentales avec des phases prescriptives.

Cette méthode, développée par Blessing et Chakrabati (2002) possède une structure rigoureuse permettant de mener à bien des projets de recherche appliquée. De plus, cette méthode permet de valider les résultats sans avoir recours à une implantation complète. Dans notre cas, aucun modèle d’appariement dédié au contexte de transformation numérique et d’Industrie 4.0 n’existe et nous devons ainsi analyser le fonctionnement existant, proposer des solutions aux problématiques majeures et intégrer cela dans un outil d’appariement validé. La méthodologie DRM est composée de quatre phases distinctes pouvant se répéter qui sont les suivantes:

- La formulation des critères de réussite ;
- L’étude descriptive I ;
- L’étude prescriptive ; et
- L’étude descriptive II

Nous pouvons lier ces étapes globales avec les différentes étapes de notre projet de recherche, comme le montre la figure 3.1 suivante :

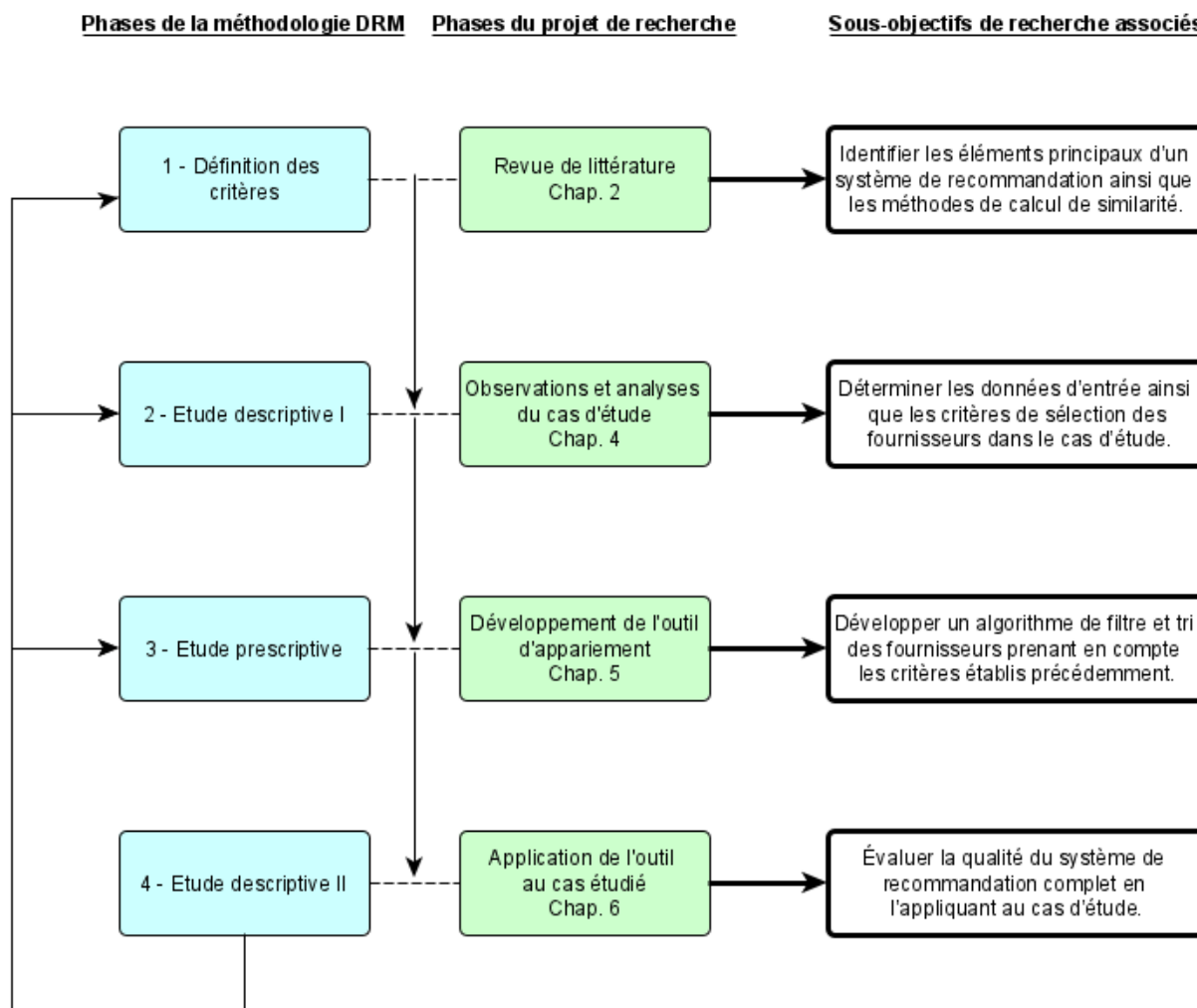


Figure 3.1 : Phases de la méthodologie DRM liées aux phases du projet de recherche

3.2.1 Revue de littérature

Une première revue de littérature concentrée à la fois sur les systèmes de recommandation, mais également sur les marchés à microtâches ainsi que sur les méthodes de comparaison entre des textes nous a permis, d'une part de valider le besoin et la pertinence de notre projet dans un cadre de transformation numérique, mais également d'autre part d'identifier certaines des publications scientifiques pertinentes pour l'atteinte de nos objectifs.

La revue critique de la littérature a été effectuée. L'article de (Grida et al., 2020) nous a notamment permis d'identifier les éléments primordiaux d'un système d'appariement qui concernent la

collecte de données, le prétraitement des données, la formulation d'une recommandation et enfin l'évaluation des performances dudit système.

3.2.2 Première étude descriptive

Cette première étude descriptive a pour but de définir les attentes précises du partenaire ainsi que ses besoins pour finalement améliorer la compréhension globale du projet pour l'intégralité de ses acteurs. Une analyse des contraintes du partenaire sera réalisée lors de cette étape et la finalité recherchée ici est l'identification des entrants et des critères de sélection des fournisseurs pour notre problème et par la même occasion répondre au deuxième sous-objectif explicité plus haut.

3.2.2.1 Définition du cas d'étude

Lors de cette première étude descriptive, nous analyserons les besoins et les contraintes d'une entreprise désirant améliorer et automatiser son processus d'appariement client fournisseur afin d'accélérer ce même processus et donc de traiter plus de requêtes client. L'entreprise concernée est une start-up montréalaise nommée BRIDGR, spécialiste dans l'accompagnement de la transformation numérique. Dans l'optique d'aider le plus d'entreprises possible à passer le cap de la transformation numérique, l'un des outils proposés par la start-up est une place du marché (marketplace) virtuelle sur laquelle les entreprises voulant être aidées remplissent une requête à laquelle un membre de BRIDGR va répondre. Cette réponse prend la forme d'une liste de fournisseurs qualifiés pour la requête et présents sur la base de données de BRIDGR. Cependant, le problème de cette méthode est qu'elle est extrêmement chronophage, car réalisée à la main.

3.2.2.2 Collecte d'information chez le partenaire

Comme décrit plus tôt il existe déjà un processus d'appariement mis en place chez le partenaire, mais il doit être amélioré. Dans ce but, il est nécessaire pour nous d'identifier les entrants actuels et ceux devant être ajoutés. Il sera également nécessaire d'identifier les critères présents dans les requêtes clients ainsi que de les classer en termes d'importance.

Ce travail de collecte d'information a été effectué en explorant les documents et les bases de données de l'entreprise ainsi que lors de réunion avec les responsables recherche et développement du partenaire.

3.2.3 Développement de l'outil d'appariement

La phase de développement de l'outil sera réalisée sous la forme d'une étude normative dont la première étape est de combiner les informations recueillies lors de l'étape précédente, à savoir l'étude descriptive, avec les méthodes et les descriptions de systèmes de recommandation critiquées lors de la revue de littérature et notamment les travaux de (Grida et al., 2020) dans le but de construire un système d'appariement permettant d'automatiser le processus d'appariement.

3.2.4 Évaluation des performances de l'outil d'appariement

Une fois l'outil développé nous l'appliquerons au cas du partenaire BRIDGR. Cette seconde étude descriptive se déroulera comme une série de tests dont nous analyserons les résultats afin d'évaluer notre outil par rapport à différents critères établis par le partenaire industriel et déterminés à l'aide de la littérature. Nous réaliserons ainsi des tests traitant de la rapidité d'exécution de notre outil ou encore de la précision des résultats qu'il retourne. Pour ce faire, à cause du manque d'anciens cas de BRIDGR utilisables et donc de référence, nous devons utiliser un générateur de requête. Ce générateur nous permettra d'avoir autant de requêtes au bon format afin de tester l'intégralité des éléments que nous aurons jugés nécessaires.

3.3 Conclusion

Nous appliquerons donc, pour mener notre étude, la méthodologie de recherche DRM mise au point par Blessing et Chakrabati (2002). Cette méthodologie de recherche nous permettra de mener notre étude au travers d'un cas concret. À l'aide des informations recueillies lors de la revue de littérature et de la première étude descriptive, nous développons un outil d'appariement entre requêtes clients et fournisseurs qualifiés. Nous présenterons lors du prochain chapitre la première étude descriptive du cas d'étude afin d'améliorer la compréhension du contexte ainsi que des besoins du partenaire.

CHAPITRE 4 ÉTUDE DESCRIPTIVE

La première étude descriptive constitue la seconde étape de la méthodologie de recherche DRM, telle qu'exposée au chapitre précédent. Une mise en contexte à propos de la transformation numérique et de l'entreprise partenaire BRIDGR est d'abord proposée dans ce chapitre, suivie d'une description du cas d'étude. Les contraintes et les besoins du partenaire seront identifiés lors de la phase d'observation de ce chapitre à travers une description du fonctionnement actuel.

4.1 Mise en contexte

Nous vivons dans une époque d'évolutions constantes en ce qui concerne les technologies utilisées dans le monde du travail. Cette évolution touche notamment le monde manufacturier. Différents événements récents, comme les différentes crises économiques ou encore la pandémie mondiale, ont poussé les entreprises de différents secteurs à revoir leur manière de travailler afin de répondre aux nouvelles exigences. Cependant, ce changement de modèle ou de méthodes de travail ne se fait pas en un claquement de doigts. Une préparation adéquate ainsi qu'un suivi constant sont nécessaires pour assurer une transition efficace qui remplit ses objectifs. Les différents aspects dans lesquels des progrès sont réalisés chaque année sont tellement larges que personne ne peut s'estimer spécialiste de l'intégralité de ces secteurs. Le spectre d'expertises s'étend des technologies de fabrication additive aux techniques de marketing digital et plus encore. Une grande majorité des entreprises ont donc besoin d'aide afin de développer leurs connaissances à propos de ces expertises, et ce quel que soit leur secteur d'activités. Afin de réaliser des projets d'implantation ou d'intégration de nouvelles technologies, de formation à l'utilisation de nouvelles techniques de management ou bien encore de conseil sur telle ou telle stratégie, il est nécessaire de s'entourer des acteurs les plus qualifiés. Des entreprises se sont donc spécialisées dans l'accompagnement d'autres entreprises lors de la recherche de ces fameux fournisseurs les plus qualifiés.

BRIDGR, notre partenaire industriel pour ce projet de recherche, s'inscrit dans ce contexte d'accompagnement en transformation numérique décrit plus haut. Cette start-up montréalaise fondée en 2016 utilise plusieurs outils, dont un questionnaire d'auto-évaluation du niveau de maturité numérique des entreprises. Leur produit qui nous intéresse ici est leur place du marché (marketplace) qui sert à mettre en relation des clients formulant une requête et un ou plusieurs fournisseurs répondant un maximum à ces critères présents sur la base de données de BRIDGR. La

principale problématique à laquelle BRIDGR a dû faire face lors de la réalisation de ces appariements dans le passé est le temps. En effet, comme cet appariement était jusqu'à présent fait à la main et que le nombre de fournisseurs potentiels est très élevé, il était extrêmement chronophage. Le temps requis ralentit notamment la réalisation des autres activités de l'entreprise. La volonté d'automatiser ce procédé vise à augmenter la rapidité des appariements et donc leur nombre, mais également de débloquer les autres activités de l'entreprise.

4.2 Processus d'appariement client/fournisseur

Cette section relate les différentes observations effectuées chez le partenaire lors de notre collecte d'information sur le terrain. Le processus d'appariement entre requêtes client et profils de fournisseurs lié à la place du marché proposée par BRIDGR est formalisé.

Afin de cartographier le processus mis en place, le formalisme BPMN (Business Process Modeling Notation) est ici utilisé. Ce formalisme normalisé (ISO/CEI 19510 : 2013) qui permet de représenter les activités, les acteurs ainsi que les flux physiques et les flux d'information de la chaîne d'approvisionnement de notre cas d'étude. Les symboles du formalisme BPMN sont présentés dans la figure 4.1 et le diagramme présentant le processus dans son intégralité dans la figure 4.2.

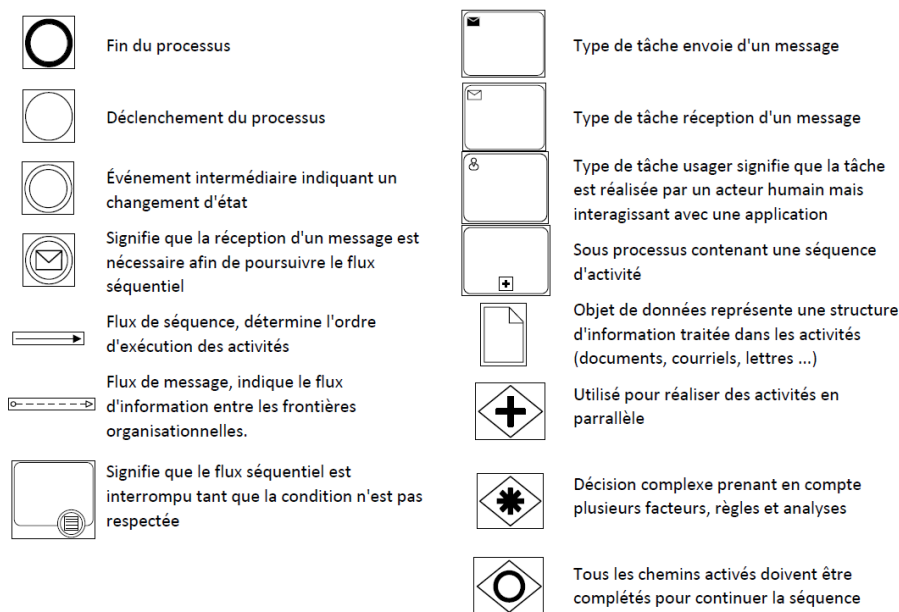


Figure 4.1 : Symboles du formalisme BPMN (tiré de la norme ISO/CEI 19510 : 2013)

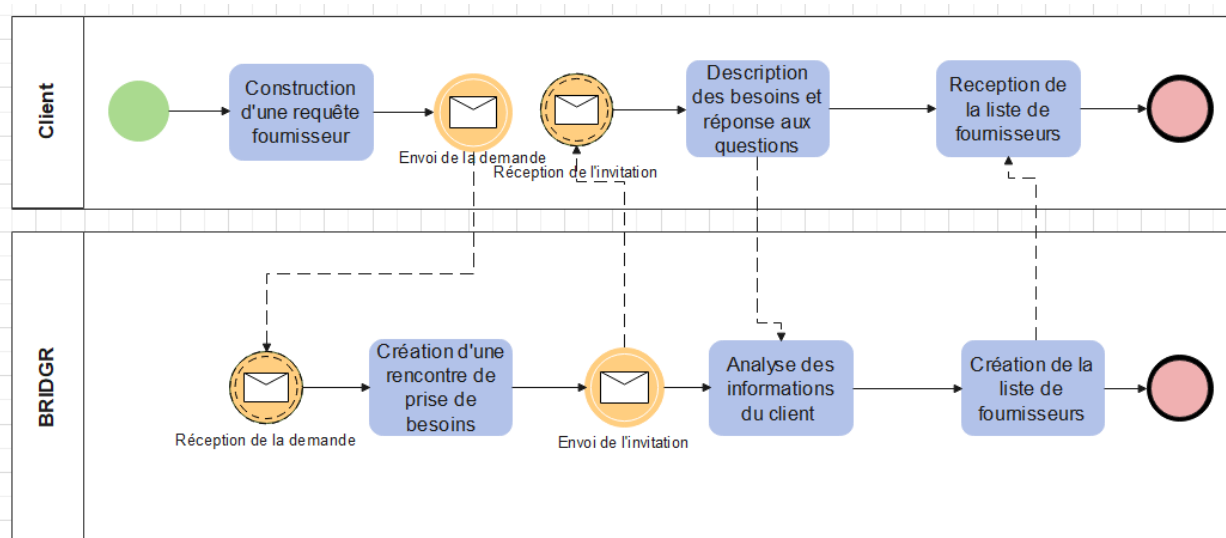


Figure 4.2 : Diagramme BPMN du processus

En général, le processus débute par un contact entre BRIDGR et le client établi par le client au moyen d'un courriel ou d'un appel téléphonique. Dans ce premier contact, le client précise dans les grandes lignes sa requête en donnant quelques critères primordiaux pour lui. Ensuite, une rencontre (en personne ou par visioconférence) est prévue afin d'effectuer de manière complète la prise de besoins du client. Lors de cette rencontre, le client énonce ses besoins en explicitant à la fois la problématique à laquelle il fait face, mais également ses contraintes spécifiques. Une fois cette première étape remplie, les experts de BRIDGR posent des questions couvrant des aspects supplémentaires que le client n'a pas forcément pensé à aborder lors de sa présentation initiale comme une restriction géographique ou bien une langue parlée par le fournisseur.

Une fois les informations nécessaires recueillies, les experts BRIDGR procèdent ensuite à la recherche des fournisseurs correspondant le mieux aux critères requis par le client. La base de données de fournisseur de BRIDGR est donc fouillée à la main par les experts qui connaissent, en plus des caractéristiques classiques du profil des fournisseurs contenues dans la base de données, des informations non classifiées formellement comme des façons de travailler. Ces informations supplémentaires, acquises par l'expérience et l'échange avec les fournisseurs, influencent le choix des experts BRIDGR quant aux fournisseurs à recommander à tel client. Une liste de recommandations de fournisseurs ainsi que leurs informations de contact est ensuite proposée au client qui peut les contacter à sa guise.

4.3 Critères de réussite de l'outil

L'analyse des processus de BRIDGR s'est effectuée à partir de juillet 2020 (en distanciel), puis dans leurs locaux depuis octobre 2020 sous la forme de stages de recherche. Durant cette intégration, nous avons pu mener ensemble une vingtaine de réunions avec les cofondateurs ainsi que le responsable du pôle Recherche et Développement à propos de leurs attentes en ce qui concerne l'outil d'appariement. Tout d'abord, nous nous sommes mis d'accord sur les questions que nous allions poser dans les requêtes des clients, et donc sur les critères pour les fournisseurs qui en découleraient. Les critères en premier sont le secteur d'activité, le nombre d'années d'expertise, les préférences en termes de situation géographique (fournisseur situé sur le même continent ou bien dans le même pays, etc.) ou encore le type d'expertise proposé par le fournisseur, à savoir s'il propose des formations, des produits, des conseils, l'intégration de solutions, etc.

Étant donné le spectre énorme d'expertises faisant partie du domaine de la transformation numérique, il nous fallait également déterminer une classification les regroupant toutes. Une fois cette classification terminée, il sera possible de l'intégrer, à la fois aux requêtes client, mais également lors de la qualification des fournisseurs. L'un des critères qui revenait le plus souvent lors des entrevues de prise de besoins était la volonté d'être apparié avec des fournisseurs possédant un minimum d'expérience. Nous nous sommes donc mis d'accord sur le fait qu'il fallait une question relative au nombre d'années d'expérience du fournisseur dans la ou les expertises désirées.

De plus lors de ces entrevues avec les responsables, une notion de différents niveaux de critères pour les fournisseurs a été établie. En effet, pour des expertises ne comprenant pas beaucoup de fournisseurs, il serait dommage de passer à côté d'un fournisseur, car il ne correspond pas à 100% à la requête client, mais plus à 95%.

Le dernier point abordé de manière primordiale par l'équipe est celui sur la présence de l'expertise humaine de BRIDGR qui devait confirmer ou infirmer les propositions faites par l'outil en amont de l'envoi de cette liste de recommandations au client. Un point supplémentaire serait une forme d'apprentissage pour l'outil qui enregistrerait les modifications apportées par l'équipe de BRIDGR à sa proposition et qui en tiendrait compte lors de ses prochaines, notamment si ces modifications se répètent.

L'un des critères qui ne fut pas directement établi lors de ces entrevues, mais plus compris de manière tacite est celui de la rapidité d'exécution. Étant donné que la problématique majeure de l'ancien processus était sa chronophagie, l'amélioration dans cet aspect est primordiale.

Le tableau 4.1 suivant présente ainsi un récapitulatif des critères complets désirés par BRIDGR. Ces critères sont également accompagnés de pistes évoquées lors de réunions sur les façons de les évaluer une fois l'outil développé.

Tableau 4.1 : Critères industriels de l'outil d'appariement et méthodes d'évaluation

Critères
Rapidité
Flexibilité de l'algorithme par rapport à la priorisation des caractéristiques fournisseur
Présence de l'expertise BRIDGR
Scalabilité de l'algorithme
Apprentissage de la solution
Format d'architecture reproductible
Précision de l'outil

Certains de ces critères industriels mis en avant par le partenaire industriel sont cependant plus des exigences de construction que des critères d'évaluation. En effet, certains des « critères » de réussite du projet sont à intégrer lors de la phase de développement de l'outil et sont donc, par construction, validés lors de la phase d'évaluation finale du modèle. Nous allons donc supprimer ces exigences de construction du tableau final des critères d'évaluation afin d'y laisser uniquement les critères que nous allons évaluer dans le chapitre 6.

Ces exigences de construction sont les suivantes :

- Flexibilité de l'algorithme par rapport à la priorisation des caractéristiques fournisseur;
- Présence de l'expertise BRIDGR sur les résultats de l'algorithme; et

- Format d'architecture reproductible.

Les autres critères traitent pour la scalabilité de l'algorithme la réaction de l'algorithme face à l'évolution de nombre de requêtes soumises simultanément ou bien face à l'évolution du nombre de fournisseurs. La notion d'apprentissage traite de l'apprentissage de l'algorithme afin qu'il prenne en compte les changements réalisés par les experts BRIDGR ou bien les retours effectués par les clients dans ses listes de recommandations futures. Enfin, la notion de précision traite de la différence entre ce qui est proposé par notre système et ce qui est attendu par le client.

Afin d'avoir un aperçu complet et d'ensuite sélectionner les meilleurs critères d'évaluation possible, nous avons également reconduit un état de l'art à propos des critères d'évaluation des systèmes de recommandation.

D'après Drachsler et al. (2015), il existe plusieurs types de critères selon lesquels les systèmes de recommandations peuvent être évalués. Il est possible d'évaluer les performances du système de recommandation et il est également possible de mesurer les effets centrés sur l'utilisateur. La précision ou encore le rappel du système traitent des performances du système, mais nécessitent une estimation de la vérité, c'est-à-dire une référence à laquelle comparer. D'après Zaier (2010), il existe également une métrique qui réunit ces deux dernières afin d'établir à la fois la précision et le rappel du système. Cette nouvelle métrique s'appelle la mesure F1.

D'après McNee et al. (2006), bien que les fonctions d'erreur et les métriques d'aide à la décision soient importantes pour mesurer les performances des systèmes de recommandation, il existe des aspects de ces derniers qui ne sont pas couverts. En utilisant Miller (2003) et Adomavicius et Tuzhilin (2005), nous pouvons ainsi mettre en avant des métriques qui traitent plus des effets centrés sur l'utilisateur. Ces métriques moins utilisées sont :

- La couverture décrite dans (Miller et al., 2004) correspond au pourcentage des items pour lesquels le système a pu fournir des recommandations;
- Le taux d'apprentissage traité dans (Herlocker et al., 2004) qui évalue la qualité des recommandations en fonction de la quantité de données d'apprentissage disponibles;
- La nouveauté et l'« heureux hasard » qui sont rassemblés sous le terme « serendipity » en anglais sont traités dans Terveen et Hill (2001). Ces métriques ont pour but d'inciter les

systèmes de recommandation à ne pas recommander des items évidents aux clients. Elles favorisent donc les recommandations non évidentes en pénalisant les évidentes;

- La satisfaction des utilisateurs qui est mise en avant dans les travaux de Cosley et al. (2003) comme étant fonction du nombre d'erreurs produites par le système;
- La similarité moyenne introduite par Miller (2003) et qui traite de la similarité moyenne du groupe utilisé afin de réaliser le calcul de la recommandation; et
- Les ressources utilisées qui correspondent indirectement à la complexité de l'algorithme. Ces ressources peuvent prendre plusieurs formes comme le temps de réponse ou encore la consommation de ressources système.

En combinant les deux sources d'informations disponibles à ce moment-là, d'un côté celles extraites de la littérature (théorique) et de l'autre celles fournies par BRIDGR (industrielles), nous pouvons désormais établir une liste complète des critères d'évaluation de notre outil d'appariement. De nouveaux critères possibles déterminés à l'aide de la littérature, ceux qui nous intéressent le plus, sont ceux liés aux performances du système. Dans la liste ci-dessus, la couverture semble intéressante, notamment pour déterminer quels types de fournisseurs il faudrait ajouter à la base de données de fournisseur de BRIDGR, mais ne peut pas vraiment être utilisée. Dans le cadre d'un système de recommandation, cela répond à la question de combien, parmi les fournisseurs de la base de données, sont fréquemment recommandés. Ce n'est pas applicable dans notre cas, car nous n'avons pas de notion de prédiction de l'utilité comme dans un système de recommandation. Le taux d'apprentissage correspond complètement à la métrique industrielle mise en avant par l'équipe pour traiter de l'apprentissage de la solution concernant l'expertise supplémentaire ajoutée par les experts BRIDGR. La notion de nouveauté semble ne pas correspondre à notre contexte de travail, car elle est très liée au ressenti instantané face à la recommandation. De plus, lorsqu'on recherche un fournisseur afin de réaliser une tâche spécifique, le but est de trouver quelqu'un correspondant parfaitement, la nouveauté ou l'inventivité de la recommandation n'entre pas en compte. Ce n'est pas comme lorsque l'on cherche un film à regarder. Dans ce cas-là, une notion de surprise dans la recommandation peut intervenir et être, pour certains utilisateurs, un point positif supplémentaire. La satisfaction des utilisateurs est une bonne métrique à utiliser ici, mais pas uniquement sur l'utilisation de l'outil d'appariement, mais sur le processus d'appariement complet. Cela pourrait prendre la forme d'un sondage simple proposé aux clients ayant utilisé ce processus avec

BRIDGR : « Est-ce que les fournisseurs recommandés correspondaient à ce que vous recherchez ? ». Le dernier point ici est celui des ressources utilisées. Dans notre cas, le but est de développer un outil d'appariement entre requêtes client et profils fournisseurs. Pour cela, nous allons utiliser des méthodes simples de programmation afin de comparer différents algorithmes pour déterminer lequel est le plus efficace. Comme un grand nombre de tests différents devront être réalisés afin de tester les différentes métriques définies ci-dessus, le temps d'exécution de notre programme se doit d'être court et peu coûteux en ressources système. Cette métrique est donc importante et sera mesurée en premier (notamment la rapidité), mais ne pourra pas réellement être déterminante pour le choix entre plusieurs algorithmes à moins que la différence ne soit vraiment abyssale.

Tableau 4.2 : Liste complète des critères d'évaluation de l'outil

Critères	Méthodes d'évaluation
Rapidité	Chronométrage de la solution
Scalabilité de l'algorithme	Tests d'influence des facteurs
Satisfaction des utilisateurs	Mesure du CSAT ⁸
Précision de l'outil	Tests de Wilcoxon ⁹

Le tableau 4.2 présente une liste complète des critères d'évaluation combinant des critères venus de la littérature et des critères industriels issus d'entrevues avec BRIDGR.

4.4 Conclusion

Cette première étude descriptive nous a permis de mieux appréhender les différents enjeux et contraintes auxquels font face les acteurs du milieu de la transformation numérique, et notamment celles qui à l'instar de notre partenaire ont pour but d'aider les autres lors de cette transformation. Lors de cette étude, nous avons donc mis en avant les besoins en termes de fonctionnalités présentes

⁸ Le CSAT est un indicateur qui mesure la satisfaction globale du client de manière basique. Méthode découverte dans : <https://blog.smart-tribune.com/fr/mesurer-satisfaction-client> (consulté le 07/10/2021)

⁹ Le Test de Wilcoxon est un test statistique qui permet de réfuter une hypothèse selon laquelle deux méthodes de calcul sont identiques.

dans l'outil d'appariement. Ces besoins comprennent une réduction du temps pour effectuer le processus d'appariement, une liste de critères (expertise, expérience par expertise, secteur d'activité principal, localisation géographique, langues couramment parlées, type d'expertise recherché) à inclure dans l'outil afin d'identifier les meilleurs fournisseurs, une précision importante associée à l'expertise BRIDGR dans l'optique d'améliorer constamment cette précision à travers un processus d'apprentissage. Ainsi, à la fois les critères et les données d'entrées nécessaires pour notre outil d'appariement ont été identifiés, ce qui constitue notre deuxième sous objectif de recherche. Le prochain chapitre présente les différents aspects de l'outil d'appariement ainsi que les méthodes et outils utilisés afin de les développer.

CHAPITRE 5 CONSTRUCTION DE L'OUTIL D'APPARIEMENT

Ce chapitre présente de façon détaillée le développement de notre outil d'appariement entre requêtes client et profils fournisseurs dans un contexte de transformation numérique. Cet outil, basé sur l'architecture d'un système de recommandation, a pour but d'identifier les fournisseurs les plus pertinents pour répondre à une demande client formulée selon un format réfléchi en fonction de critères plus ou moins importants. Cet outil doit aussi améliorer les performances globales du cas d'étude décrit dans le chapitre précédent. Nous allons tout d'abord, en partie 5.1, nous intéresser aux différents intrants que nous allons utiliser dans notre outil. Nous traiterons ici le format des requêtes client, celui des profils fournisseur ou encore les différents critères identifiés influençant le choix d'un fournisseur. Dans un second temps nous nous intéresserons aux différentes couches de notre outil avec, en partie 5.2, la couche qui relève d'un filtrage des fournisseurs et en partie 5.3 celle qui traite du tri de ces fournisseurs filtrés.

5.1 Partie intrants (expertises et critères exclusifs, inclusifs, etc.)

Le premier aspect qu'il est nécessaire de considérer est l'aspect des intrants de notre système d'appariement, c'est-à-dire les informations nécessaires au bon fonctionnement de ce dernier. L'établissement de ces intrants repose sur les informations recueillies chez le partenaire industriel ainsi que sur la littérature à travers des exemples de systèmes de recommandation fonctionnels.

5.1.1 Caractéristiques des fournisseurs et des requêtes clients

Lors de plusieurs entrevues réalisées chez le partenaire industriel, nous avons identifié les critères qui correspondent aux caractéristiques que BRIDGR considérerait comme les plus cohérentes dans un contexte de recherche de fournisseurs en termes de transformation numérique. Cette mise en avant de critères à prendre en compte de critères afin de recommander des fournisseurs nous oriente vers un système de recommandation basé sur le contenu. Afin de conserver une vision globale du projet et de ne pas nous restreindre uniquement aux informations fournies par BRIDGR pour déterminer les meilleures caractéristiques fournisseur, nous avons décidé de nous intéresser aux caractéristiques les plus présentes dans ce genre de système de recommandation dans la littérature.

Dans Pazzani et Billsus (2007), les auteurs discutent de la personnalisation pour les utilisateurs en évoquant plusieurs caractéristiques pouvant nous intéresser. Ces caractéristiques sont le genre

favori de films, le nom des équipes de sports favorites ou encore les sections préférées sur un site de nouvelles. Dans son article sur le site Internet Interstices, Negre (2018) nous propose des exemples concernant un hypothétique système de recommandation basé sur le contenu à propos de livres. Dans cet exemple, les caractéristiques de livre qui sont mises en avant sont le genre du livre, le/la ou les auteurs, le titre, le nombre de pages ou encore le prix moyen. Il y a également la présence de mots-clés qui servent à venir préciser les catégories sans pour autant en rajouter énormément. Une partie des critères mis en avant dans la littérature sont transposables à notre contexte et correspondent directement à des caractéristiques identifiées lors des entrevues avec le partenaire. Ces caractéristiques sont la mise en place de catégories de fournisseurs en fonction de leurs expertises. Un second type de catégories pour les fournisseurs est leur secteur d'activité principal. En nous servant des autres exemples mis en avant dans la littérature, nous pouvons identifier d'autres caractéristiques importantes. Le budget, qui est une caractéristique importante mise en avant par BRIDGR, correspond quant à lui à la caractéristique prix moyen qui est présente dans l'exemple. Les autres exemples de caractéristiques présents dans la littérature considérée sont impossibles à transposer à notre contexte. En effet, le nom d'un auteur ou encore le nombre de pages n'ont pas d'équivalents évidents en ce qui concerne des fournisseurs. Une liste provisoire des caractéristiques présentes dans le profil fournisseur et donc également dans les requêtes remplies par les clients est présentée dans le tableau 5.1 suivant.

Comme nous pouvons le remarquer, certaines caractéristiques qui étaient mises en avant par le partenaire dans leur processus d'appariement tel que le budget du client ou encore la date de début souhaitée pour le projet ne possèdent pas d'équivalent à positionner dans le profil fournisseur. Nous savons qu'il serait possible de les associer en créant des variables intermédiaires cependant

5.1.2 Expertises et mots-clés.

La première étape dans l'établissement de ces caractéristiques de profils fournisseurs et de requêtes client est l'élaboration de ces catégories décrivant l'ensemble du spectre des possibilités de l'industrie 4.0 sans pour autant tomber dans la précision à outrance qui pourrait nuire aux résultats obtenus par notre outil. Afin de déterminer les expertises comprises dans l'industrie 4.0, nous nous sommes intéressés à plusieurs rapports publiés récemment. Les référentiels utilisés sont *Le guide des technologies de l'industrie du futur* publié par l'Alliance Industrie du Futur en Mars 2018, le rapport *Industrie du futur, prêts, partez !* publié par l'institut Montaigne en Septembre 2018 ainsi

Tableau 5.1 : Récapitulatif des caractéristiques des profils fournisseur et des requêtes client

Profils fournisseur	Requêtes client
Expertises fournisseur	Expertises requête
Type d'activité fournisseur	Type d'activité requête
Langues principales	Langue principale
Nombre d'années d'expérience fournisseur pour chaque expertise	Nombre d'années d'expérience requises pour chaque expertise de la requête
Secteur d'activité principal fournisseur	Secteur d'activité requête
Situation géographique du fournisseur	Restriction géographique de la requête

que le *Guide pratique de l'usine du futur* publié par Usine du futur en octobre 2015. Ces trois référentiels se composent de fiches consécutives traitant de groupes de technologies ou de méthodes courantes dans un contexte d'industrie 4.0. En listant les titres de ces fiches puis en les regroupant par ressemblance, nous avons identifié une liste de 84 expertises regroupées dans 11 catégories d'expertises qui sont décrites dans l'Annexe B.

Une fois cette liste d'expertises établie, il nous faut reconsidérer la présence de mots-clés dans le but de préciser les catégories, telles que mentionnées par Negre (2018). En effet, nos expertises couvrent le spectre des possibilités de la transformation numérique dans son intégralité sans pour autant rentrer complètement dans les détails. En exemple facilement compréhensible, nous pouvons citer celui de la fabrication additive. Il existe de nos jours, selon le dossier publié dans Usine Nouvelle (2019), environ sept méthodes différentes de fabrication additive qui dépendent du matériau pris en compte (bois, plastique ...) ou bien de la méthode utilisée (stéréolithographie, par frittage laser ou encore extrusion de la matière). Ces méthodes sont toutes complètement différentes et pourtant, dans notre classification de la transformation numérique, elles se trouvent toutes dans l'expertise « Fabrication additive ». Nous allons donc mettre en place des mots-clés pour chaque expertise afin d'augmenter la précision sans ajouter de nouvelles expertises. Dans le but d'emmagasinier un maximum de mots-clés pour chacune des expertises déjà définies, nous avons décidé de rechercher des glossaires spécifiques pour chacune de ces expertises. Afin de varier les sources, nous avons recherché plusieurs glossaires différents et réalisé des regroupements entre les

glossaires afin d’obtenir des mots-clés pour toutes les expertises couvrant le plus d’éléments possible de cette expertise. Certains glossaires sont communs à plusieurs expertises, car ils sont plus généraux et couvrent des catégories d’expertises plus que des expertises elles-mêmes. La liste des glossaires pour chaque expertise est présente en Annexe C.

5.1.3 Plages des caractéristiques

Nous avons défini les différentes caractéristiques ainsi que les différents choix possibles pour les expertises. Il nous faut donc définir les plages de choix pour chacune des caractéristiques, c’est-à-dire les différentes possibilités entre lesquelles les clients pourront choisir pour chacune des caractéristiques au moment de remplir leur requête. En ce qui concerne les expertises, nous avons déjà défini précédemment qu’il est possible de choisir entre 86 choix. Les types d’activité disponibles ont été définis par consensus lors d’entrevues avec l’équipe de BRIDGR. En effet, lors de leurs différents appariements réalisés dans le passé, les clients demandent quasi exclusivement parmi quatre types de services. Ces derniers sont la recherche de produits (« Produit »), la recherche de conseil (« Conseil »), celle de formations à propos de méthodes de management par exemple (« Formation ») ou finalement celle de l’intégration de systèmes ou de logiciels commerciaux à des systèmes déjà en place (« Intégration »). En général, plusieurs de ces expertises sont liées, par exemple si une entreprise vend un produit spécifique, elle peut également s’occuper de la formation des usagers et de l’intégration de leurs solutions. Dans ce cas, le fournisseur choisira « Produit » afin de couvrir toutes ces possibilités. Pour les langues, nous avons recherché les 30 langues les plus parlées dans le monde. D’après le tableau¹⁰ trouvé, ces 30 premières langues correspondaient (en 2017) à plus de 4,5 millions de locuteurs à travers le monde. Nous aurions également pu permettre uniquement les langues déjà présentes dans la base de données fournisseur, mais nous avons opté pour l’optique s’orientant vers une expansion de la base de données fournisseurs.

En ce qui concerne les autres caractéristiques, le nombre d’années d’expérience n’a pas vraiment de plage à définir à part le fait d’être un nombre entier positif, pour les secteurs d’activités nous avons utilisé la liste de 146 secteurs d’activités déjà utilisés par BRIDGR dans leurs autres outils. Cette liste est présentée dans l’Annexe D. La dernière caractéristique traitée est celle de la

¹⁰ Tableau des 80 langues les plus parlées dans le monde (consulté le 20/08/2021) :

https://www.axl.cefan.ulaval.ca/Langues/2vital_expansion_tablo1.htm

restriction géographique requise par le client. Comme la situation géographique des fournisseurs et des clients est déjà connue, c'est sur la distance entre les deux que doit s'appliquer la caractéristique. Nous avons donc défini des paliers de distance qui peuvent décrire des restrictions utiles lors de la recherche de fournisseurs. Ces paliers sont :

- Pas de restrictions géographiques ;
- Le fournisseur est dans la même ville (- de 50 km) ;
- Le fournisseur est dans la même région (- de 200 km) ;
- Le fournisseur est dans le même pays ; et
- Le fournisseur est sur le même continent.

5.1.4 Caractéristiques exclusives, inclusives et de collecte de données

Comme évoqué par Benarouet (2018), l'un des problèmes rencontrés par un système de recommandation basé sur le contenu est la surspécialisation. Cette situation correspond à un trop grand nombre de caractéristiques précises formulées dans la requête tels que, malgré la grande taille de la base de données de fournisseur, aucun ou très peu ne correspondent parfaitement. Cette situation résulte donc sur un outil d'appariement qui ne répondrait pas à son but principal. Afin de pallier ce problème, nous allons nous intéresser à un critère évoqué par l'équipe de BRIDGR. Le fait de mettre plusieurs niveaux d'importance pour les critères de sélection des fournisseurs peut, soit se mettre en place en modifiant des coefficients, soit en décrivant plusieurs types de critères. Les caractéristiques de requêtes clients mises en avant dans le tableau 5.1 correspondent à celles qui possèdent un équivalent dans le profil fournisseur. Cependant il y en avait d'autres, présentes dans le processus initial d'appariement, à savoir la date de début souhaitée ou encore une estimation du budget. Ces informations sont importantes afin de bien sélectionner le fournisseur, mais comme elles ne correspondent pas à des informations dans la base de données du côté fournisseur, il est impossible de s'en servir dans la partie algorithme.

Ces considérations font apparaître ici trois catégories de caractéristiques. Une première catégorie qui déterminerait les fournisseurs possibles pour répondre à la requête client, donc ceux qui répondent aux caractéristiques les plus importantes. Une deuxième catégorie qui sert à classer les fournisseurs sélectionnés par la première catégorie à l'aide de caractéristiques importantes, mais

pas indispensables (ce qui permettrait d'éviter la surspécialisation de notre outil). Et enfin, la troisième catégorie qui correspond aux caractéristiques que le partenaire ne souhaitait pas utiliser pour le moment, mais s'inscrivant dans l'aspect expertise spécifique ajoutée par BRIDGR après le passage dans l'algorithme. La répartition des différentes caractéristiques dans ces catégories a été réalisée par consensus avec l'équipe de BRIDGR lors de plusieurs réunions.

En parcourant le tableau 5.1, les caractéristiques qui nous semblent les plus importantes à mettre dans la première catégorie sont les expertises, le type d'activité et les langues parlées. La première caractéristique correspond aux catégories de connaissances et/ou de savoir-faire recherchées par le client. Dans cette optique, il semble logique de ne pas proposer de fournisseurs qui ne connaissent rien dans le domaine demandé par le client. Le type d'activité proposé est également une caractéristique exclusive, car si le client recherche absolument une formation sur un sujet précis, il semble incohérent de lui proposer des fournisseurs spécialisés dans l'intégration de solutions par exemple. En ce qui concerne les langues, même si de nos jours une grande partie de la planète parle anglais, il est difficile de communiquer avec un fournisseur sans partager au moins un langage.

Les caractéristiques des requêtes client que l'on peut placer dans la deuxième catégorie, celle des caractéristiques inclusives, sont le nombre d'années d'expérience, le secteur d'activité principal et la situation géographique du fournisseur parmi celles présentées dans le tableau 5.1. À cette liste, nous pouvons également ajouter le nouveau niveau de précision déterminé précédemment, les mots-clés utilisés pour préciser les expertises. Les années d'expérience requises sont importantes, mais il serait dommage de passer à côté d'un fournisseur possédant toutes les bonnes expertises seulement parce qu'il possède huit années d'expérience alors que le client a mis dix dans sa requête. C'est la même chose en ce qui concerne les secteurs d'activités et la situation géographique. C'est normal d'avoir des préférences et ces caractéristiques se doivent d'entrer dans le processus de sélection. Les mots-clés servent, quant à eux, à préciser les expertises et donc à augmenter la précision de notre outil.

Le dernier groupe de caractéristiques, celles servant à accumuler de la donnée, est composé, comme décrit précédemment, de la date de début souhaitée ainsi que de l'estimation du budget. Cette donnée servira à augmenter l'expertise fournie par BRIDGR après la réalisation de l'appariement par l'algorithme. BRIDGR se servira également de ces données afin de recueillir un maximum d'informations liées à leurs différents produits présents.

5.2 Filtre

La première partie de notre outil de filtre considéré est celle concernant les caractéristiques exclusives. Ce type de caractéristiques, qui contient les expertises recherchées par le client, le type d'activité recherché par le client ainsi que la langue qu'il faut que le fournisseur parle afin de parler avec le client, regroupe les caractéristiques qui excluent les fournisseurs n'y répondant pas complètement. Il nous faut donc un filtre qui parcourt tous les éléments de la base de données des fournisseurs afin d'écarter ceux qui ne correspondent pas. Pour réaliser cet algorithme de filtrage plutôt simple, nous avons utilisé le langage de programmation Python sur l'environnement de développement intégré IDLE. Nous avons fait ces choix, car nous avons déjà réalisé des projets passés sur ces outils, et nous avons donc une certaine expérience dans leur utilisation. La figure 5.1 suivante présente un logigramme expliquant le fonctionnement global de l'algorithme que nous avons construit puis nous reviendrons plus en détail sur ces différents éléments dans cette partie et la prochaine.

Les intrants de cet algorithme seront donc :

- La liste des requêtes clients sous forme de tableau dans un fichier en .csv ; et
- La base de données des fournisseurs sous forme de tableau dans un fichier en .csv.

La sortie de cet algorithme sera

- Une liste réduite des fournisseurs répondant aux caractéristiques exclusives présentes dans la requête client considérée.

Nos premiers tests avec des conditions complètement exclusives se sont révélés très intéressants. En effet, en ne permettant qu'une langue pour chaque fournisseur et en forçant le fait que les fournisseurs doivent avoir exactement toutes les expertises demandées par le client nous avons peu ou pas du tout de résultats pour nos requêtes de test. Comme chaque requête client peut contenir plusieurs expertises, il y a un risque de surspécialisation (Benarouet, 2018) ce qui réduit le nombre de résultats passant l'étape du filtre et donc le nombre de résultats donnés finalement au client. En effet, à chaque expertise demandée en plus par le client il faut que les clients possédants déjà toutes les expertises précédentes possèdent également cette dernière. Comme nous ne pouvons pas aisément faire l'hypothèse que l'intégralité des clients qui rempliront les requêtes soit des spécialistes du domaine de la transformation numérique ou de l'industrie 4.0, nous devons prendre

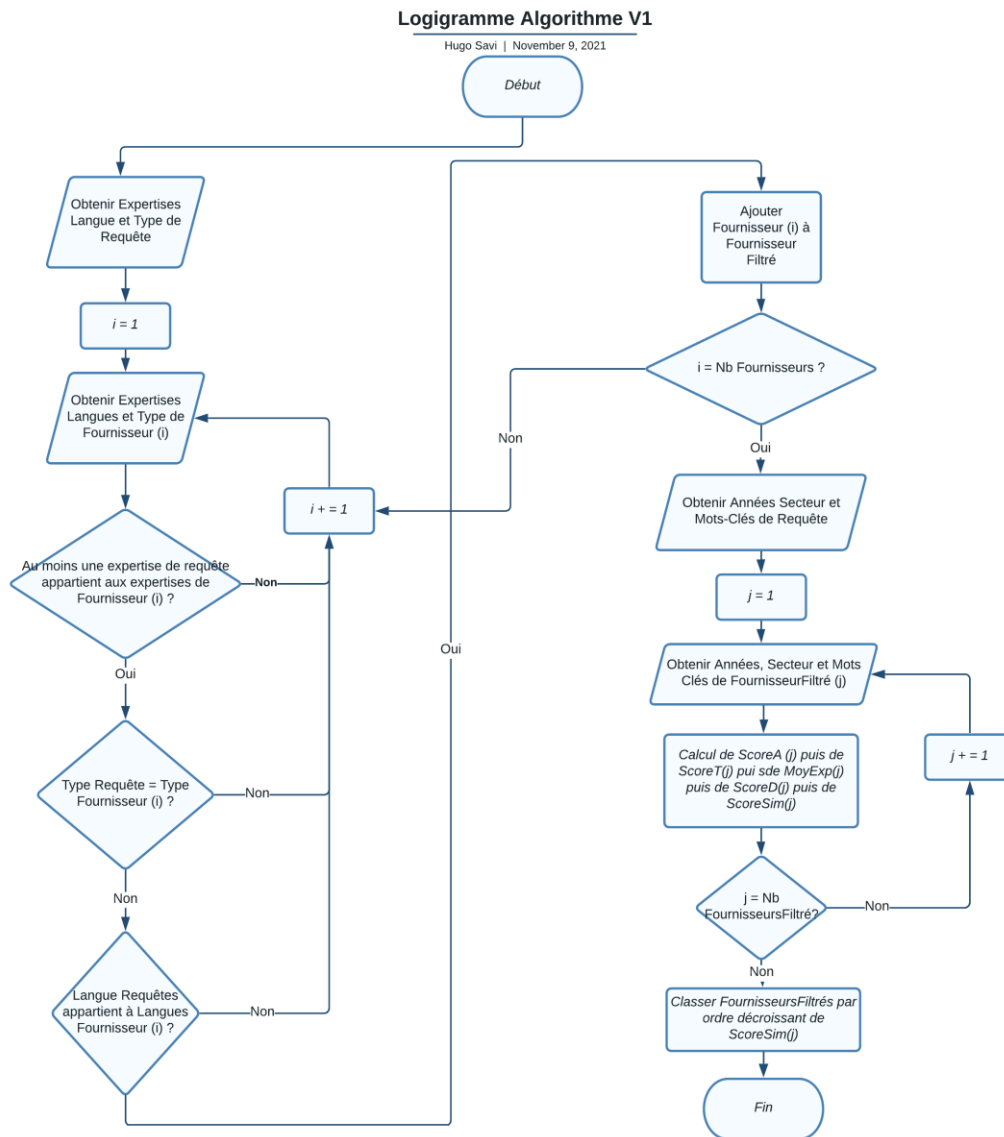


Figure 5.1 : Logigramme de fonctionnement de l'algorithme

en compte la possibilité qu'ils associent deux expertises (ou plus) dans une requête sans que celles-ci aient de vrais liens entre elles. Dans ce cas-ci, il est fort probable qu'aucun fournisseur de la base de données, aussi riche soit-elle, ne puisse correspondre à ces caractéristiques.

Notre solution pour pallier ce problème est d'autoriser le fait de correspondre uniquement à une des expertises clients afin de passer l'étape du filtre. Cette décision est basée sur le fait que dans un projet il est, dans la grande majorité des cas, question de plusieurs tâches qui peuvent généralement être toutes différentes. Il faudra cependant prendre en compte cet aspect de nombre

d'expertises auxquelles le fournisseur correspond dans la deuxième étape de l'outil qui ordonnera les fournisseurs ayant passé le filtre. En ce qui concerne les langues, nous avons décidé de laisser jusqu'à trois langues possibles pour les fournisseurs. Cette valeur peut tout à fait être modifiée, nous avons choisi cette limitation de manière arbitraire. Nous nous attendons à ce qu'une grande partie des fournisseurs mette au moins l'anglais qui est considéré comme la langue des affaires¹¹. En prenant exemple sur le Québec, nous pouvons facilement estimer qu'une bonne partie des fournisseurs québécois veuille également mettre le français en langage parlé. En prenant en compte ces considérations, nous avons décidé de laisser le choix aux fournisseurs pour le nombre de langues à positionner entre une et trois. Cela permet de traiter les cas où deux ou trois langues sont importantes dans une région ou un pays sans pour autant forcer les fournisseurs à faire un choix qui pourrait les faire passer à côté de certains projets.

Maintenant que nous avons traité la première étape, celle du filtrage, nous allons nous intéresser à l'autre type de caractéristiques importantes utilisables dans notre cas, les caractéristiques inclusives.

5.3 Tri

Une fois la liste de fournisseurs disponibles écumée de ceux ne répondant pas aux caractéristiques exclusives identifiées, il est temps d'ordonner ceux qui restent en fonction de leur respect (ou non) des caractéristiques inclusives. Pour rappel, ce type de caractéristiques contient le nombre d'années d'expérience pour chacune des expertises, des mots-clés pour chacune des expertises, le secteur d'activité principal du client ainsi que les possibles restrictions géographiques entre le client et les fournisseurs.

5.3.1 Caractéristiques générales

Tout d'abord, nous allons nous intéresser à ce dernier point. Nous n'avons malheureusement pas eu le temps de nous consacrer à cet aspect lors du développement des versions initiales de l'outil. En effet, la prise en compte de cet aspect nécessite l'utilisation d'une API extérieure à Python

¹¹ Site Internet Amplexor consulté le 15/09/2021 : <https://blog.amplexor.com/fr/les-10-langues-les-plus-demandées-dans-le-monde-des-affaires>

classique afin de gérer les données de géolocalisation. Nous avons décidé de laisser cet aspect de côté pour l'instant, mais il reste une piste pour une amélioration future de l'outil pour BRIDGR.

Les autres caractéristiques traitées ici le sont de manières différentes. En effet, le but de cette seconde étape qui, elle aussi sera réalisée en langage Python sur l'environnement de développement intégré IDLE, est de trier les fournisseurs ayant passé l'étape du filtre. Tout d'abord, nous avons décidé de classer les caractéristiques en deux types, les bonus et les malus. Les caractéristiques bonus qui contenaient, initialement, le secteur d'activité et les mots-clés correspondent aux caractéristiques dont la présence apporte un bonus au profil du fournisseur, mais dont l'absence n'est pas préjudiciable pour ce dernier. Les caractéristiques malus qui contenaient, à la base, uniquement les années d'expérience, sont des caractéristiques dont la présence est « normale » pour un profil fournisseur, mais dont l'absence pénalise le profil. De plus les années d'expérience ainsi que les mots-clés diffèrent du secteur d'activité, car il n'y a qu'un seul secteur d'activité par requête client, mais il y a autant de nombre d'années d'expérience requis ou de liste de mots-clés qu'il y a d'expertise présentes dans la requête client. Comme un grand nombre d'algorithmes de tri, notre méthode pour ordonner les fournisseurs ayant passé le filtre est de calculer un score de similarité qui est fonction des différentes caractéristiques, à la fois de la requête client et profil fournisseur, puis d'ordonner les fournisseurs en fonction du score obtenu.

En prenant en compte les éléments que nous avons traités avant, notamment la dépendance à l'expertise, la formule de calcul du score $S_{f,r}$ pour un fournisseur donné et pour une requête donnée est la suivante :

$$S_{f,r} = \frac{\sum_{k=1}^{n_{e,r}} S_m(m_{k,f}, m_{k,r}) + S_a(a_{k,f}, a_{k,r})}{n_{e,r}} + S_d(d_f, d_r) \quad (5.1)$$

Avec :

- $n_{e,r}$ le nombre d'expertises différentes présentes dans la requête;
- S_m la fonction traitant l'aspect mots-clés par expertise qui est fonction de $m_{k,f}$ et $m_{k,r}$
qui représentent respectivement les mots clés fournisseur de la $k^{\text{ième}}$ expertise de la requête et les mots clés requête de la $k^{\text{ième}}$ expertise de la requête;

- S_a la fonction traitant l'aspect années d'expérience par expertise qui est fonction de $a_{k,f}$ et $a_{k,r}$ qui représentent respectivement les années d'expérience fournisseur de la k ème expertise de la requête et les années d'expérience requête de la k ème expertise de la requête; et
- S_d la fonction traitant l'aspect secteur d'activité qui est fonction de d_f et d_r qui représentent respectivement le secteur d'activité du fournisseur et le secteur d'activité de la requête.

Nous allons désormais traiter les différentes fonctions de caractéristiques en commençant par celle du secteur d'activité. Comme décrit précédemment, cette caractéristique est un bonus, donc elle ne rapportera des points au fournisseur que si celui-ci possède le même secteur d'activité que celui précisé dans la requête. La fonction S_d est donc de la forme :

$$S_d(d_f, d_r) = \begin{cases} \alpha & \text{si } d_f = d_r \\ 0 & \text{sinon} \end{cases} \quad (5.2)$$

avec la valeur du coefficient α qui n'est pas encore déterminée, mais qui dépend de l'importance que l'on désire donner au secteur d'activité dans le classement

5.3.2 Prise en compte des années d'expérience

Pour prendre en compte les années d'expérience, plusieurs alternatives sont possibles. Une version où cette caractéristique intervient en tant que malus, mais également une version où elle intervient en tant que malus si le nombre d'années requis n'est pas atteint et en tant que bonus s'il est dépassé (dans une moindre mesure). Cette seconde possibilité est apparue lors de réunions avec l'équipe dans le but de pouvoir différencier deux fournisseurs possédant toutes les autres caractéristiques de manière similaire, mais qui ont une légère différence en termes d'années d'expérience, afin de favoriser le plus expérimenté.

Les premiers essais que nous avons réalisés sur cet aspect prenaient en compte la différence entre les années requises par la requête et les années d'expérience du fournisseur de manière brute dans le score. Cependant, cela apportait trop de disparités entre les différents fournisseurs. Par exemple, un fournisseur qui possédait le bon nombre d'années d'expérience pour une expertise et aucun mot-clé passait devant un fournisseur qui, sur le papier paraissait plus qualifié, mais à qui il manquait une ou deux années d'expérience par rapport au nombre requis. De plus, ce système n'était pas

borné en minimum et donc plus le nombre d'années requises augmentait plus le risque qu'un fournisseur obtienne un score très bas augmentait également. Nous avons donc opté pour une version bornée en nous intéressant au ratio entre le nombre d'années d'expérience du fournisseur et le nombre d'années d'expérience requises par la requête client.

Pour la première version (uniquement malus) la fonction S_a est de la forme :

$$S_a(a_{k,f}, a_{k,r}) = \min\left(\frac{a_{k,f}}{a_{k,r}}, 1\right) \quad (5.3)$$

Et pour la seconde version (qui prend en compte l'aspect malus, mais récompense les années d'expérience supplémentaires), la fonction S_a est de la forme :

$$S_a(a_{k,f}, a_{k,r}) = \begin{cases} \min\left(\frac{2a_{k,f}}{3a_{k,r}}, \frac{2}{3}\right), & \text{si } a_{k,f} \leq a_{k,r} \\ \min\left(\frac{a_{k,f}}{3a_{k,r}} + \frac{1}{3}, 1\right), & \text{sinon} \end{cases}$$

Nous avons codé ces deux versions différentes de la fonction S_a et les avons implémentées dans un algorithme global. Les différents résultats des tests réalisés sur les différentes versions de cet algorithme sont présentés dans le chapitre suivant.

5.3.3 Prise en compte des mots-clés

Afin de développer l'aspect de prise en compte des mots-clés, nous nous sommes intéressés de nouveau aux notions découvertes dans la dernière partie de la revue de littérature. En effet, afin de comparer plusieurs listes des mots-clés, nous allons utiliser des méthodes de calcul de similarité. Comme décrit dans Negre (2013), il existe différents types de méthodes de calcul de similarité, des méthodes syntaxiques (basées uniquement sur l'écriture des mots et non sur le sens) et des méthodes sémantiques (basées sur le sens des mots).

5.3.3.1 Méthodes syntaxiques

Les premières que nous avons décidé d'implanter, car plus simples à mettre en place, sont les méthodes syntaxiques. Voici la liste comme décrite dans Negre (2013) :

- La similarité Cosinus ;
- Le coefficient de corrélation de Pearson ;

- La distance euclidienne ;
- Le coefficient de Jaccard ;
- La distance de Levenshtein ; et
- L'indice de Dice.

Afin de déterminer lesquelles mettre en place, nous allons étudier leurs spécificités. Les plus grosses différences viennent du fait que les trois premières méthodes de calcul de similarité sont des méthodes qui s'appliquent uniquement sur des vecteurs. Afin de les appliquer, nous devons donc vectoriser nos mots-clés. Comme l'ordre des coordonnées dans les vecteurs est extrêmement impactant dans l'utilisation des ces trois premières méthodes et que nous ne pouvons pas forcer un ordre, que ce soit chez le client ou le fournisseur, nous devons nous rabattre sur les trois autres méthodes qui se basent sur des cardinaux d'ensemble et qui sont donc, par définition, indépendantes de l'ordre des éléments.

Une autre métrique de calcul de similarité est décrite dans Buscaldi et al. (2019). Il s'agit de la distance de Bray Curtis qui traite uniquement de l'abondance d'une espèce dans un échantillon. Comme le domaine d'application est extrêmement éloigné de notre domaine, nous avons décidé d'écarter cette méthode de calcul de similarité, car elle est trop spécifique aux domaines de l'écologie et de la biologie afin d'évaluer la dissimilarité entre deux échantillons en termes d'abondance de taxons présents.

À cette liste, nous allons ajouter une autre méthode qui compte le nombre de mots-clés en commun entre la requête client et le profil fournisseur.

De la même manière que pour les années d'expérience précédemment, nous allons développer plusieurs méthodes et donc plusieurs versions différentes de l'algorithme que nous allons tester dans le chapitre suivant. Encore de la même manière que précédemment, étant donné que le coefficient de Jaccard et l'indice de Dice sont des ratios bornés entre 0 et 1, nous avons adapté les deux autres méthodes pour qu'elles aussi soient bornées entre 0 et 1 et que la comparaison entre les quatre méthodes ait du sens.

Pour la première version (coefficient de Jaccard) la fonction S_m est de la forme:

$$S_m(m_{k,f}, m_{k,r}) = \frac{|m_{k,r} \cap m_{k,f}|}{|m_{k,f} \cup m_{k,r}|} \quad (5.5)$$

On rappelle ici que dans ce cas là $m_{k,f}$ et $m_{k,r}$ représentent les listes de mots-clés du profil fournisseur et de la requête client.

La deuxième version (indice de Dice) de la fonction S_m est de la forme:

$$S_m(m_{k,f}, m_{k,r}) = \frac{2 * |m_{k,r} \cap m_{k,f}|}{|m_{k,f}| + |m_{k,r}|} \quad (5.6)$$

La troisième version (distance de Levenshtein) de la fonction S_m est de la forme:

$$S_m(m_{k,f}, m_{k,r}) = 1 - \frac{\text{Levenshtein}(m_{k,r}, m_{k,f})}{\max(|m_{k,r}|, |m_{k,f}|)} \quad (5.7)$$

Nous rappelons ici que la distance de Levenshtein correspond au nombre minimal d'opérations simples pour transformer une chaîne de caractères en l'autre. Ici, nous avons adapté la méthode en utilisant le ratio entre la distance de Levenshtein et la longueur maximale des deux chaînes de caractères, qui est également la valeur maximale que peut atteindre la distance de Levenshtein. Ce cas serait atteint s'il est nécessaire de changer absolument tous les caractères d'une chaîne pour obtenir l'autre. De plus, nous avons ajouté le 1 auquel est soustrait le fameux ratio, car la distance de Levenshtein est une distance c'est-à-dire que plus la valeur augmente plus les chaînes sont éloignées.

La dernière version qui augmente du score lorsque les mots-clés sont en commun qui ne correspond pas à une méthode de calcul de similarité est la suivante :

$$S_m(m_{k,f}, m_{k,r}) = \frac{|m_{k,r} \cap m_{k,f}|}{|m_{k,r}|} \quad (5.8)$$

5.3.3.2 Méthodes sémantiques

Comme décrit dans Negre (2013), il existe en plus des méthodes syntaxiques, des méthodes sémantiques. Ces méthodes sont basées sur le sens des mots et nécessitent donc une référence en termes de sens des mots qui est fixée en amont. Parmi ces méthodes, il existe trois grandes catégories de méthodes. La première catégorie comprend les méthodes vectorielles, exactement les mêmes que celles présentées et écartées précédemment, mais avec une vectorisation basée, cette

fois-ci, sur le sens des mots. La deuxième catégorie comprend les méthodes hiérarchiques dans lesquelles la notion sémantique intervient en tant que distance entre les mots dans une base hiérarchique des connaissances sémantiques. La troisième catégorie regroupe les différentes méthodes d'analyses basées sur des corpus.

Dans le cadre de notre cas d'application, nous avons choisi de nous orienter méthodes hiérarchiques. Il nous faut donc construire des bases de connaissances sémantiques pour chacune des expertises. Nous nous sommes concertés, car cela prendrait beaucoup de temps de créer ces bases de connaissances pour toutes les expertises. Nous avons donc décidé de n'en créer que pour certaines expertises bien définies qui sont les plus présentes dans la base de données fournisseur. Derrière ce choix vient le fait que si elles sont plus présentes dans la base de données fournisseur, les requêtes qui les demandent auront plus de réponses et nous aurons un plus gros échantillon de test. Tout le protocole utilisé pour réaliser l'intégralité des tests est expliqué au début du chapitre suivant.

5.3.4 Conclusion

Une fois tous les scores calculés dans les différentes versions de l'algorithme, ils sont injectés dans la formule décrite à la page 63 et le score de similarité entre la requête client et le profil fournisseur est calculé. D'après les différentes versions du score, il est compris entre 0 et 2,5. Le score nul est théoriquement impossible à atteindre, car si un fournisseur passe l'étape du filtre cela signifie qu'il possède au moins une expertise en commun avec la requête et donc au moins les points ramenés par les années d'expérience. Une fois les scores de similarité de tous les fournisseurs calculés, l'algorithme termine de trier en ordonnant du fournisseur ayant obtenu le meilleur score à celui ayant obtenu le moins bon. Une liste de triples (numéro du fournisseur dans la base de données, nom du fournisseur, score obtenu) est ensuite fournie à l'utilisateur de l'outil.

5.4 Conclusion

Dans ce chapitre nous venons d'expliquer tous les principes qui nous ont menés au développement de notre outil d'appariement tel qu'il est actuellement. Ce chapitre reste cependant très théorique et rempli d'hypothèses. Le chapitre suivant servira donc à confirmer ou à infirmer toutes ces dernières.

CHAPITRE 6 ÉVALUATION DES PERFORMANCES DE L'OUTIL

Dans ce chapitre nous nous intéresserons au processus d'évaluation des différents critères de réussite déterminés dans le chapitre 4. Tout d'abord, nous nous concentrerons sur le processus d'évaluation dans sa globalité en nous intéressant aux étapes préliminaires nécessaires ainsi qu'au test réalisé. Dans un second temps, nous nous concentrerons sur les résultats obtenus et sur ce qu'ils signifient dans le cas de notre outil.

6.1 Processus global d'évaluation des performances de l'outil

La figure 6.1 ci-dessous présente les différentes étapes qui composent notre protocole d'évaluation des critères de réussite de notre outil. La satisfaction des utilisateurs, bien que présente dans le tableau récapitulatif 4.2, n'est pas évaluée dans ce logigramme. En effet, il s'agit d'un critère de réussite important, mais afin de pouvoir l'évaluer, nous nécessitons des retours utilisateurs comme le nom du critère l'indique. Notre outil reste à un stade de prototype pour l'instant et nous n'avons pas actuellement de retours utilisateurs. Cependant, nous avons tout de même réfléchi à comment évaluer les retours d'utilisateurs et à les prendre en compte et nous nous concentrerons sur ce point dans un second temps.

Nous allons désormais parcourir le logigramme afin d'expliquer chacune des étapes. Tout d'abord, il y a une étape d'adaptation de la base de données fournisseur. Cela consiste en un changement de format pour la base de données fournisseur afin que celle-ci corresponde au format choisi pour le bon fonctionnement de l'algorithme. En réalité, nous n'avons pas eu beaucoup de changements à faire à part séparer certaines colonnes comme les expertises en plusieurs colonnes pour plus de clarté dans notre code.

La deuxième étape est beaucoup plus importante et nous a demandé beaucoup plus de temps. Un des gros problèmes auquel nous avons fait face dans ce projet est le manque de référence. En effet, les anciens appariements réalisés par BRIDGR étaient, pour la plupart, inutilisables ou, a minima, dans un format beaucoup trop éloigné et incomplet. Nous avons cependant besoin de requêtes clients complètes afin de tester les différents aspects de notre outil. Pour cela, nous avons créé un programme Python de générateur de requêtes. Ce programme permet de créer des fichiers .csv d'ensemble de requêtes. Tous les éléments composant ces requêtes sont sélectionnés au hasard

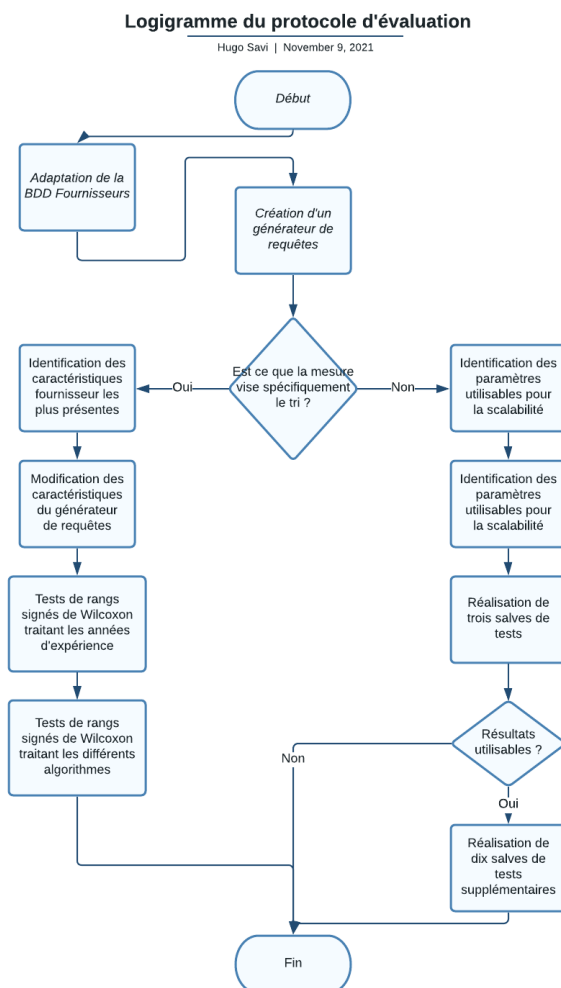


Figure 6.1 : Logigramme du protocole d'évaluation

parmi les possibilités. Chaque requête demande entre une et quatre expertises distinctes et pour chacune des expertises demandées, il y a entre un et cinq mots-clés.

La séparation dans le logigramme entre les évaluations visant le tri (le côté précision et la différence entre les différentes versions de l'algorithme) et celles traitant de l'algorithme dans son intégralité, est liée aux échantillons de requêtes que nous allons utiliser pour tester. Dans la partie où le tri est la fonctionnalité que nous avons décidé d'évaluer, il est nécessaire pour nous d'avoir de la matière à trier. Pour cela, nous devons faire en sorte de créer des requêtes pour lesquelles notre outil, avec la base de données fournisseur de BRIDGR adaptée seulement en format, allait fournir un nombre important de résultats. En effet, en générant un fichier contenant cent requêtes complètement aléatoires, sans restreindre leurs possibilités, nous obtenons seulement 5% des requêtes qui ont des

résultats et 1% qui a plus de trois résultats. Pour pallier ce problème, sans avoir à remplir la base de données de faux fournisseurs couvrant plus de domaines d'expertises que ceux présents actuellement, nous avons décidé de restreindre les possibilités de l'aléatoire dans le générateur. Pour ce faire nous allons mettre en avant les éléments les caractéristiques exclusives les plus répandues chez nos fournisseurs et restreindre la plage de sélection de ces caractéristiques à ces éléments. Afin de déterminer quelles étaient les caractéristiques recherchées pour un fournisseur qui avaient le plus de chance d'apparaître, nous avons analysé les caractéristiques présentes chez les fournisseurs actuels de la base de données fournisseur de BRIDGR. Une fois cette analyse réalisée, nous avons supprimé artificiellement des possibilités d'expertises dans le générateur toutes celles qui n'étaient présentes chez aucun fournisseur. Nous avons également restreint les langues aux quatre possibilités les plus présentes chez les fournisseurs actuels à savoir « Anglais », « Français », « Espagnol », et « Italien ».

Une fois ces modifications apportées au générateur de requêtes, nous avons pu créer 500 requêtes dans ces nouvelles conditions, puis mettre de côté celles qui obtenaient plus de six résultats. Nous avons choisi le seuil de six arbitrairement, afin de prendre en compte le fait que pour évaluer un tri et l'influence de tel ou tel algorithme sur un classement ou un score il faille avoir des résultats à trier. Nous avons ainsi réalisé plusieurs tests de la somme des rangs de Wilcoxon sur les versions de l'algorithme deux à deux. Ce test permet de prouver ou non qu'une hypothèse de base soit fausse. En général, l'hypothèse appelée hypothèse nulle est de la forme, « il n'y a aucune différence entre telle ou telle version ». Nous avons réalisé ces tests pour prouver la différence des quatre versions de calcul de score de similarité pour les mots-clés et pour les deux versions de calcul de score de similarité pour les années d'expérience. Une fois ces tests réalisés sur des requêtes ayant obtenu six résultats ou plus, nous avons voulu tester l'importance du seuil dans les résultats obtenus et nous avons réalisé à nouveau l'expérience pour dix résultats ou plus et pour quinze résultats ou plus. Tous les résultats de ces études seront montrés et analysés dans la partie 6.2 suivante.

En parallèle à ces études qui reposent sur le nombre de résultats passant l'étape du filtre, nous nous sommes intéressés aux paramètres de rapidité et de scalabilité que nous avons couplés. En étudiant

plusieurs sources¹² concernant les études de scalabilité dans les algorithmes, nous nous sommes rendu compte que les caractéristiques pouvant être mesurées sont les suivantes :

- Temps de réponse qui correspond à la rapidité ;
- Débit, qui correspond au nombre de requêtes qui sont process par unité de temps ;
- Utilisation du processeur de la machine réalisant les tests ;
- Utilisation de la mémoire de la machine réalisant les tests ; et
- Utilisation du réseau.

Dans notre cas, comme les tests se font en réseau local sur un ordinateur personnel, qui ne sera pas forcément celui sur lequel les appariements seront réalisés dans le futur, les trois derniers points ne correspondent pas. De plus, le débit n'est pas réellement une caractéristique à prendre en compte ici, car le nombre d'éléments traités par unité de temps dépend surtout du nombre d'éléments en entrée. Nous avons donc décidé de nous concentrer uniquement sur le temps de réponse. Nous avons décidé d'étudier l'influence du nombre de requêtes sur le temps de réponse, mais également du nombre d'expertises demandées par requêtes. Afin de faire varier ces caractéristiques l'une après l'autre, nous avons forcé des caractéristiques comme le nombre d'expertises par requêtes dans notre générateur de requêtes. Le logigramme présente une boucle traitant de la pertinence des résultats, car les premiers tests que nous avons réalisés avec des programmes annexes qui tournaient pendant que l'algorithme calculait fournissaient des résultats complètement différents pour des conditions semblant être identiques. Pour remédier à cela, nous avons réalisé à nouveau les mêmes tests en consacrant l'ordinateur uniquement à ce test et en vidant la mémoire à chaque itération.

Maintenant que les différents tests réalisés ont été présentés, nous pouvons passer à la présentation des résultats.

¹² Sources regroupées dans l'article sur le site internet Software Testing Help consulté le 7 octobre 2021 : <https://www.softwaretestinghelp.com/what-is-scalability-testing/>

6.2 Résultats et analyses

6.2.1 Tests de la somme des rangs de Wilcoxon

Tableau 6.1 : Résultats du test de la somme des rangs de Wilcoxon de la première campagne

Requêtes contenant 6 résultats ou plus pour les méthodes de calcul des années d'expérience							
Utilisation des scores de similarité				Utilisation des rangs du classement			
Fonctionnel	Jaccard	Dice	Levenshtein	F	J	D	L
Oui	Oui	Oui	Oui	Non	Non	Non	Non

Comme précisé plus tôt nous n'avions pas à notre disposition de requêtes client déjà traitées par BRIDGR et dans le format adéquat à l'utilisation par notre outil. Nous avons donc dû créer de toutes pièces des requêtes. Comme ne nous pouvions pas alors réellement tester la précision de notre outil nous avons décidé d'aider les futurs tests de précision en déterminant s'il existait des ressemblances entre nos différentes méthodes utilisées afin de calculer la similarité entre un fournisseur et une requête client.

6.2.1.1 Années d'expérience

Le test de la somme des rangs de Wilcoxon (ou test de Wilcoxon-Mann-Whitney) est un test non paramétrique qui compare deux mesures d'une variable quantitative effectuées sur les mêmes sujets. Dans notre cas, ils ont été réalisés sur deux types de variables distinctes dans un certain nombre d'itérations distinctes. Tout d'abord, ils ont été utilisés afin de prouver ou non la différence entre les deux méthodes de calcul de similarité concernant les années d'expérience. En effet, comme décrit dans la partie 5.1.3.2, nous avons décidé de mettre en place deux méthodes différentes ne prenant pas en compte exactement les mêmes éléments. Nous avons réalisé ces tests de la somme des rangs de Wilcoxon une fois sur les scores obtenus par des fournisseurs en réponse à des requêtes et une fois sur le rang qu'ont obtenu ces mêmes fournisseurs dans le classement fourni par l'outil d'appariement. Nous avons effectué ces tests de comparaison des méthodes V1 et V2 de calcul de la similarité pour le calcul des années d'expérience pour chacune des quatre versions de méthode de calcul de similarité des mots-clés (Fonctionnelle, Jaccard, Dice et

Levenshtein). Comme décrit dans la partie précédente, nous avons commencé par effectuer ces tests sur des requêtes recevant six résultats ou plus lors de leur passage dans l'outil d'appariement.

Le tableau 6.1 présente un récapitulatif des résultats obtenus en effectuant les huit tests de la somme des rangs de Wilcoxon. Les résultats sont présents dans la dernière ligne tandis que les trois premières servent à présenter les paramètres de chaque test. Nous rappelons ici que le résultat du test de la somme de rangs de Wilcoxon est la possibilité de rejeter ou non une hypothèse initiale. L'hypothèse initiale étant dans tous les cas de la même forme, par exemple pour le deuxième cas ici : « Il n'existe pas de différence entre les deux méthodes de calcul du score de similarité des années d'expérience qui impacte les scores obtenus par des fournisseurs répondant aux caractéristiques de requêtes recevant six résultats ou plus en utilisant la méthode du coefficient de Jaccard pour calculer le score de similarité des mots-clés ». Ces expertises sont longues et complexes, mais en réalité elles exposent juste tous les paramètres pris en compte afin de bien contrôler les jeux de données et de résultats comparés pour obtenir des résultats analysables.

La présence du mot « Oui » dans la case correspondant à l'hypothèse évoquée plus haut, à savoir la deuxième case de la quatrième ligne, revient à dire qu'il est possible de rejeter cette hypothèse. Cela revient à dire qu'il y a effectivement une différence notable entre les scores obtenus en utilisant les deux méthodes pour les années d'expérience. Nous avons réalisé ces tests pour chacune des quatre versions de méthode de calcul de la similarité pour les mots-clés afin de nous affranchir de leurs possibles influences que nous avons traitées plus tard.

Comme montré dans le tableau 6.1, les quatre tests réalisés en tenant compte des scores ont donné un résultat positif en ce qui concerne le rejet des hypothèses de base ce qui nous indique que les deux méthodes de calcul induisent une différence notable dans les scores obtenus par les fournisseurs et ce quelle que soit la méthode de calcul employée pour les mots-clés. A contrario, en ce qui concerne les tests réalisés avec les rangs obtenus par ces mêmes fournisseurs et non plus leurs scores bruts le résultat est complètement inversé. En effet, les résultats obtenus ne nous permettent pas de réfuter (ni de valider par ailleurs, subtilité du test de la somme des rangs de Wilcoxon) les hypothèses de base. Nous ne savons donc pas si l'utilisation de telle ou telle méthode pour les années d'expérience influe les rangs finaux obtenus.

Nous avons réitéré cette étude deux fois en ne prenant respectivement que des requêtes ayant reçu plus de dix et plus de quinze résultats (donc de fournisseurs passant l'étape du filtre). Les résultats sont consignés dans le tableau 6.2.

Tableau 6.2 : Résultats des tests de la somme des rangs de Wilcoxon pour la deuxième campagne

Requêtes contenant 10 résultats ou plus pour les méthodes de calcul des années d'expérience							
Utilisation des scores de similarité				Utilisation des rangs du classement			
F	J	D	L	F	J	D	L
Oui	Oui	Oui	Oui	Non	Non	Non	Non
Requêtes contenant 15 résultats ou plus pour les méthodes de calcul des années d'expérience							
Utilisation des scores de similarité				Utilisation des rangs du classement			
F	J	D	L	F	J	D	L
Oui	Oui	Oui	Oui	Non	Non	Non	Non

Comme nous pouvons le voir, les résultats sont les mêmes quel que soit le nombre minimal de résultats à scorer ou à trier. Ce seuil semble donc ne pas avoir d'importance par rapport aux résultats obtenus par l'outil d'appariement et nous pouvons donc justifier que les deux méthodes de calcul du score de similarité traitant des années d'expérience sont sensiblement différentes en termes de scores obtenus.

6.2.1.2 Similarité syntaxique des mots-clés

Nous avons également réalisé des tests de la somme des rangs de Wilcoxon pour comparer et possiblement différencier les quatre différentes méthodes de calcul de score de similarité traitant des mots-clés. De la même manière que précédemment, nous avons effectué ces tests à la fois sur les scores obtenus et sur les rangs obtenus. Pour plus de précision, nous ajoutons ici que ces tests ont été réalisés sur les mêmes jeux de données et de résultats que les tests traités précédemment qui concernaient les années d'expérience. Dans le cas des mots-clés, nous avons donc réalisé les comparaisons entre les méthodes deux à deux. Nous avons commencé en comparant notre méthode

fonctionnelle avec celle basée sur le coefficient de Jaccard puis nous l'avons comparée à celle basée sur l'indice de Dice et cela jusqu'à avoir les six comparaisons deux à deux traitées. De plus, pour s'affranchir de l'influence de la méthode de calcul employée pour traiter la similarité des années d'expérience nous avons effectué chacune de ces comparaisons deux fois, une fois par méthode de calcul du score de similarité utilisée.

Les résultats pour les requêtes comportant six résultats ou plus sont présentés dans le tableau 6.3 suivant :

Tableau 6.3 : Résultats des tests de la somme des rangs de Wilcoxon pour la troisième campagne

Requêtes contenant 6 résultats ou plus pour les méthodes de calcul des mots-clés (scores)											
FxJ		FxD		FxL		JxD		JxL		DxL	
V1	V2	V1	V2	V1	V2	V1	V2	V1	V2	V1	V2
Oui	Oui	Non	Non	Oui	Oui	Non	Non	Non	Non	Oui	Oui
Requêtes contenant 6 résultats ou plus pour les méthodes de calcul des mots-clés (rangs)											
FxJ		FxD		FxL		JxD		JxL		DxL	
V1	V2	V1	V2	V1	V2	V1	V2	V1	V2	V1	V2
Non	Non	Non	Non	Non	Non	Non	Non	Non	Non	Non	Non

De la même manière que précédemment, les « Oui » ou « Non » dans la quatrième et la huitième ligne correspondent à la possibilité de rejeter ou non l'hypothèse de base selon laquelle les deux méthodes comparées ne présentent pas de différence dans la configuration choisie. Les résultats obtenus montrent qu'en ce qui concerne les rangs le résultat est le même que pour les années, impossible de rejeter ou de confirmer l'hypothèse de base selon laquelle le changement de méthode entraîne un changement notable dans le classement.

En revanche, en ce qui concerne les scores obtenus, nous remarquons que pour certains couples de méthodes la différence peut être prouvée alors que pour d'autres non. De plus, les résultats sont les mêmes pour chaque couple de méthodes, quelle que soit la méthode de calcul pour les années

Tableau 6.4 : Résultats des tests de la somme des rangs de Wilcoxon pour la quatrième campagne

Requêtes contenant 10 résultats ou plus pour les méthodes de calcul des mots-clés (scores)											
FxJ		FxD		FxL		JxD		JxL		DxL	
V1	V2	V1	V2	V1	V2	V1	V2	V1	V2	V1	V2
Oui	Oui	Non	Non	Oui	Oui	Non	Non	Non	Non	Oui	Oui
Requêtes contenant 10 résultats ou plus pour les méthodes de calcul des mots-clés (rangs)											
FxJ		FxD		FxL		JxD		JxL		DxL	
V1	V2	V1	V2	V1	V2	V1	V2	V1	V2	V1	V2
Non	Non	Non	Non	Non	Non	Non	Non	Non	Non	Non	Non
Requêtes contenant 15 résultats ou plus pour les méthodes de calcul des mots-clés (scores)											
FxJ		FxD		FxL		JxD		JxL		DxL	
V1	V2	V1	V2	V1	V2	V1	V2	V1	V2	V1	V2
Oui	Oui	Non	Non	Oui	Oui	Non	Non	Oui	Oui	Oui	Oui
Requêtes contenant 15 résultats ou plus pour les méthodes de calcul des mots-clés (rangs)											
FxJ		FxD		FxL		JxD		JxL		DxL	
V1	V2	V1	V2	V1	V2	V1	V2	V1	V2	V1	V2
Non	Non	Non	Non	Non	Non	Non	Non	Non	Non	Non	Non

d'expérience employée. Nous pouvons donc prouver par les statistiques que les méthodes F et J soit fonctionnelle et basée sur le coefficient de Jaccard sont sensiblement différentes pour les scores obtenus par les fournisseurs. Nous pouvons faire de même pour la méthode fonctionnelle et celle

basée sur la distance d'édition de Levenshtein ainsi que pour celle sur l'indice de Dice et celle sur la distance de Levenshtein.

Comme pour la prise en compte des années d'expérience, nous avons réitéré ces tests pour des requêtes recevant plus de dix et plus de quinze résultats (encore une fois les mêmes que pour les années). Les résultats sont consignés dans le tableau 6.4 suivant :

Comme nous pouvons le voir, les résultats obtenus pour des requêtes de dix résultats ou plus, sont les mêmes que ceux obtenus pour six ou plus, à savoir pas de possibilité de rejeter les hypothèses de non-influence pour les rangs et les mêmes « couples » de méthodes différentes pour la prise en compte des scores. Cependant, en ce qui concerne les résultats obtenus pour des requêtes obtenant quinze résultats ou plus, nous remarquons une différence en ce qui concerne les résultats pour les scores. Les résultats pour les rangs sont les mêmes donc aucune information supplémentaire, mais nous remarquons que dans ce cas, avec un jeu de données et de résultats qui ne contient que des requêtes ayant reçu quinze résultats ou plus, nous obtenons une preuve statistique selon laquelle il y a une différence importante entre les scores obtenus par la méthode centrée sur le coefficient de Jaccard et celle basée sur la distance d'édition de Levenshtein.

6.2.1.3 Conclusion

Ces nombreux tests de la somme des rangs de Wilcoxon nous ont donc permis de montrer qu'il existe des différences entre certaines méthodes dans certains contextes, mais ils ne nous permettent pas de déterminer laquelle ou lesquelles de ces méthodes sont les meilleures. En effet, pour cela il nous faudrait une référence à laquelle comparer nos résultats. Nous n'avons pas de telle référence, car nous utilisons des requêtes artificielles afin de réaliser ces tests, et que par conséquent, elles ne sont pas passées par le processus d'appariement actuel de BRIDGR afin d'obtenir un résultat que l'on pourrait qualifier de « vrai » ou d'« optimal » dans notre contexte. De plus, les résultats obtenus dépendent beaucoup des jeux de requêtes utilisés et même si un grand nombre de requêtes ont été utilisées il sera probablement nécessaire de réitérer les tests avec d'autres jeux de données. Notamment en ce qui concerne les tests pour des requêtes à quinze résultats ou plus.

Cependant, les tests concernant les différentes méthodes de calcul pour les mots-clés en utilisant les scores et non pas les rangs peuvent nous mettre sur la piste d'éléments intéressants. Ces tests nous ont permis de prouver la différence entre les méthodes de calcul et donc de créer des « groupes » de ces méthodes. Une fois l'étape du prototype passée, il est temps de faire utiliser

l'outil à des vrais utilisateurs ces « groupes » de méthodes nous permettront de prédire les retours possibles des utilisateurs en fonction des méthodes utilisées et donc, au fil du temps, de déterminer quelle est la méthode qui semble convenir le mieux aux utilisateurs.

En ce qui concerne la considération des années d'expérience, il sera possible de réaliser la même méthode en traitant les retours utilisateurs en fonction de la méthode utilisée. Toutes ces considérations à propos de l'utilisation du retour utilisateur et de l'apprentissage que l'on pourra en faire seront traitées dans la partie 6.3.

6.2.2 Tests de scalabilité et de rapidité

Dans cette partie, nous allons traiter des tests mis en place pour mesurer la rapidité de l'outil et plus particulière de l'algorithme, mais également la scalabilité de cet outil en ce qui concerne le temps de réponse. Nous réalisons ces tests afin de s'assurer de la scalabilité, mais également de valider formellement la réponse à la problématique initiale.

6.2.2.1 Tests

Afin de mesurer le temps de réponse de notre outil, nous avons ajouté à chacun de nos algorithmes traitants chacun d'une combinaison de méthodes, une pour les années et une autre pour les mots-clés, un chronomètre qui nous affiche le temps total de calcul et d'affichage des résultats. Nous allons donc pouvoir mesurer la rapidité et essayer d'établir de manière empirique l'influence de deux facteurs qui sont le nombre de requêtes et le nombre d'expertises par requêtes. Nous avons donc créé huit jeux de requêtes donc les caractéristiques sont les suivantes :

- Jeu 20 : 100 requêtes avec une expertise par requête;
- Jeu 21 : 100 requêtes avec deux expertises par requête;
- Jeu 22 : 100 requêtes avec trois expertises par requête;
- Jeu 23 : 100 requêtes avec quatre expertises par requête;
- Jeu 30 : 100 requêtes avec une expertise par requête;
- Jeu 31 : 1000 requêtes avec une expertise par requête;
- Jeu 32 : 10000 requêtes avec une expertise par requête; et

- Jeu 33 : 100000 requêtes avec une expertise par requête.

Les jeux 20 et 30 sont différents même s'ils possèdent les mêmes caractéristiques.

Nous avons mesuré le temps que chacune des huit versions différentes de l'algorithme mettait pour calculer et afficher les résultats pour ces huit jeux de données et ce trois fois afin d'estimer si les mesures avaient du sens. En effet, si après trois fois dans les mêmes conditions, les résultats sont sensiblement différents, nous pourrions nous dire que nous ne pourrions pas réellement déterminer l'influence de chacun de paramètres sur le temps de calcul et d'affichage.

Les figures 6.2, 6.3 et 6.4 suivantes regroupent les résultats des trois campagnes de tests réalisés :

	2.7.2	2.7.1	2.8.2	2.8.1	2.9.2	2.9.1	2.10.2	2.10.1
20 x 1	1,85 s	1,88 s	1,94 s	1,72 s	1,75 s	3,1 s	1,77 s	1,79 s
21 x 1	2,04 s	1,64 s	1,91 s	2,06 s	1,98 s	3,39 s	2,5 s	1,97 s
22 x 1	2,15 s	2,17 s	2,1 s	2,27 s	2,29 s	2,81 s	3,26 s	2,6 s
23 x 1	1,48 s	1,59 s	1,73 s	2,41 s	2,34 s	2,47 s	2,06 s	2,15 s
30 x 1	1,89 s	1,82 s	1,61 s	2,12 s	2,15 s	2,47 s	2,61 s	1,9 s
31 x 1	14,88 s	15,52 s	14,95 s	19,56 s	20,35 s	26,99 s	24,05 s	18,68 s
32 x 1	118,89 s	116,05 s	109 s	160,38 s	198,97 s	232,03 s	193,15 s	156,58 s
33 x 1	1025,15 s	1058,54 s	1042,41 s	1645,39 s	1982 s	2111,76 s	1737,92 s	2367 s

Figure 6.2 : Résultats de la première campagne sur la scalabilité

	2.7.2	2.7.1	2.8.2	2.8.1	2.9.2	2.9.1	2.10.2	2.10.1
20 x 1	1,94 s	1,53 s	1,66 s	2,17 s	1,54 s	2,71 s	1,47 s	1,38 s
21 x 1	2,13 s	1,77 s	2,07 s	1,82 s	2,04 s	1,58 s	3,57 s	2,06 s
22 x 1	2,11 s	2,55 s	2,14 s	2,18 s	1,82 s	1,89 s	2,68 s	2,55 s
23 x 1	2,13 s	1,90 s	2,78 s	1,54 s	2,15 s	2,58 s	3,15 s	2,69 s
30 x 1	1,76 s	2,13 s	1,87 s	2,24 s	1,91 s	1,94 s	1,56 s	1,78 s
31 x 1	16,5 s	17,39 s	18,24 s	16,42 s	14,2 s	23,32 s	18,63 s	22,37 s
32 x 1	147,55 s	146 s	123,83 s	135,32 s	134,75 s	171,3 s	165,89 s	162,67 s
33 x 1	1442 s	1397,8 s	1226,37 s	1301,6 s	1239,48 s	1108,65 s	1546,53 s	1370,07 s

Figure 6.3 : Résultats de la deuxième campagne sur la scalabilité

	2.7.2	2.7.1	2.8.2	2.8.1	2.9.2	2.9.1	2.10.2	2.10.1
20 x 1	1,82 s	1,90 s	1,36 s	1,52 s	1,99 s	1,32 s	1,43 s	1,26 s
21 x 1	1,69 s	2,24 s	1,46 s	1,64 s	2,66 s	1,53 s	1,8 s	1,45 s
22 x 1	1,94 s	2,48 s	1,63 s	2,16 s	2,86 s	1,96 s	2,03 s	2,67 s
23 x 1	1,55 s	2,23 s	1,49 s	1,58 s	2,62 s	1,92 s	2,35 s	1,85 s
30 x 1	2,54 s	2,14 s	1,88 s	1,89 s	2,37 s	1,49 s	1,78 s	1,41 s
31 x 1	15 s	29 s	17,16 s	15,50 s	22,82 s	14,57 s	15,27 s	11,21 s
32 x 1	142,96 s	117,06 s	137,15 s	230,6 s	154,83 s	156,22 s	122,37 s	113,63 s
33 x 1	1293,81 s	1205,80 s	1250,28 s	1869,91 s	1221,08 s	1217,76 s	1124,69 s	1049,28 s

Figure 6.4 : Résultats de la troisième campagne sur la scalabilité

Dans ces tableaux, les colonnes représentent les différentes versions de l'algorithme utilisées avec les 2.7.x correspondant aux versions initiales sans base pour les mots-clés, les 2.8.x aux versions centrées sur le coefficient de Jaccard, les 2.9.x celles centrées sur l'indice de Dice et les 2.10.x celles utilisant la distance d'édition de Levenshtein. Le dernier chiffre correspond à la version de la méthode de calcul pour les années d'expérience avec 2.x.1 qui correspond à la première méthode explicitée dans la partie 5 et le 2.x.2 à la seconde méthode.

Les lignes correspondent aux jeux de données utilisés avec le jeu de requêtes en premier et la base de données fournisseur au cas où dans de futurs tests il faille changer de base de données fournisseur.

En comparant les différents résultats pour chacune des combinaisons jeux de données, version de l'algorithme, pour les trois campagnes de données, nous nous rendons compte que l'écart entre les différents temps de réponse est extrêmement grand. Par exemple, pour les 2.10.1 et le jeu de données à 100000 requêtes, nous passons de 2367s à 1050s.

Cette remarque nous permet de montrer que les différents temps de réponse ne semblent pas réellement dépendre des différentes versions, mais plus d'éléments que nous ne pouvons pas contrôler réellement.

Nous avons donc effectué une seconde vague de test avec uniquement Python ouvert sur la machine en plus de nettoyer la RAM entre chaque passage. Les résultats obtenus lors de cette seconde série de tests servant uniquement à valider ou non un environnement de tests d'ouvert sont présentés dans la figure 6.5 suivante. Pour effectuer ce test, nous avons réalisé dix fois, pour les versions 2.7.1 et 2.7.2, les mesures sur le jeu de requêtes 33, donc celui qui prenait le plus de temps et donc celui pour lequel les écarts de temps de réponse, s'il y en a, seront les plus faciles à déceler.

Tests à vide	
FV1	FV2
855,95 s	865,37 s
866,72 s	862,82 s
863,34 s	1117,57 s
809 s	1116,96 s
936,99 s	1132,17 s
862,3 s	857,26 s
819,2 s	868,12 s
819,1 s	837 s
958,46 s	1153,66 s
820,26 s	842,13 s

Figure 6.5 : Résultats des tests réalisés "à vide"

Nous pouvons ainsi nous rendre compte que malgré l'unique présence du programme Python en fonctionnement sur la machine les résultats ne sont pas uniformes. Nous remarquons par exemple pour la version FV2, correspondant à l'utilisation de notre première méthode pour les mots-clés et de la seconde pour les années, deux groupes de temps de réponse sont observables avec un autour de 860 s et un autre autour de 1120 s.

6.2.2.2 Conclusion

Ces tests nous ont permis d'identifier la présence d'un élément non mesurable pour nous intervenant sur la mesure du temps de réponse. En effet, même en allouant les ressources machine uniquement à ce test, les résultats fluctuent d'une itération à l'autre. Il est donc impossible pour nous, dans nos moyens actuels, de réaliser des tests en isolant l'influence de chacun de paramètres influant sur l'évolution du temps de réponse.

L'autre but de ces tests était d'évaluer la rapidité de l'outil et de la comparer au travail manuel qui était réalisé auparavant par les experts BRIDGR. Nous avons, au maximum, trois secondes pour

une centaine de requêtes, alors qu'auparavant, le travail pour une seule requête pouvait prendre jusqu'à plusieurs heures. Le critère de rapidité de l'outil est donc validé.

6.3 Conclusion

Cette partie nous a permis d'explicitier à la fois les processus mis en place pour évaluer et les difficultés auxquelles nous avons dû faire face pour évaluer notre outil. Ainsi, un grand nombre de critères industriels ou issus de la littérature ne peuvent pas réellement être évalués par manque de références auxquelles comparer ou bien de premiers retours clients.

CHAPITRE 7 CONCLUSION ET RECOMMANDATIONS

7.1 Conclusion

Le but initial pour ce projet était de développer un outil de recommandation permettant, en fonction des caractéristiques « fournisseur » requises par le client pour réaliser son projet, de recommander une liste de fournisseur classée par similarité du profil avec le profil « optimal » établi dans un contexte de transformation numérique. Ce projet a été réalisé grâce à notre partenaire industriel, la start-up montréalaise, BRIDGR. Il s'inspire de leur processus d'appariement entre une requête client et un profil réalisé manuellement actuellement et avait pour but d'accélérer et d'améliorer ce dernier.

Afin de répondre à la problématique de notre partenaire industriel, nous avons développé un outil d'appariement prenant en entrée une ou plusieurs requêtes clients ainsi que la base de données des fournisseurs de BRIDGR. Chaque requête ainsi que chaque profil fournisseur contiennent des éléments qui se répondent servant à identifier quels profils correspondent le mieux à chacune des requêtes. Ces critères ou caractéristiques sont variés et vont du secteur d'activité principal du fournisseur aux expertises recherchées par le client et donc proposées par le/les fournisseur(s).

L'apport de notre outil est donc de permettre un appariement entre client et fournisseur dans un domaine précis qui est celui de la transformation numérique. Il s'accompagne d'un début de cartographie du spectre des expertises disponibles dans ce domaine vaste. Cette cartographie se compose actuellement de regroupements d'expertises, des expertises à part entière et d'un début de liste de mots-clés pour chacune des expertises dans le but de préciser les spécialités existantes.

Notre outil est également composé de différents algorithmes prenant part dans l'étape du tri de notre outil. Ces différents algorithmes prennent en compte les méthodes identifiées expérimentalement et à l'aide de la littérature pour mesurer la similarité entre les requêtes client et les profils fournisseur.

Une fois notre outil construit en répondant aux exigences de constructions émises par le partenaire industriel nous avons pu l'évaluer sur les différents critères déterminés dans la littérature et au cours de réunions avec notre partenaire. Nous n'avons pas réellement pu obtenir des résultats concernant tous les critères par manque de référence, mais nous avons comparé les différents

algorithmes afin de les réunir pour faciliter les comparaisons et le choix d'un algorithme spécifique dans le futur. Nous avons également identifié qu'il nous manque un critère influant sur l'évolution du temps de réponse à identifier. Mais nous avons validé le critère d'accélération du temps pour obtenir un résultat du processus d'appariement de notre partenaire.

En effet, le critère de rapidité d'exécution est respecté, car le programme met beaucoup moins de temps à traiter une requête que ce que ne le faisait le processus initial. Les exigences de construction qu'étaient la présence de l'expertise BRIDGR ou encore la flexibilité au niveau de l'importance des différents critères sont également respectées pour que l'outil ressemble le plus possible à ce que voulait le partenaire.

Cependant, certains points sont entrés dans les discussions à un certain moment, mais n'ont pu être traités par manque de temps. Le premier d'entre eux concerne le calcul de la similarité au niveau des mots-clés. Dans les différentes versions actuelles, la similarité prise en compte est uniquement syntaxique or dans le domaine de la transformation numérique il existe un grand nombre de mots-clés différents pour traiter de concepts, au final, assez proches. Une prise en compte de la similarité sémantique basée sur une banque de connaissances hiérarchique qui nous permettrait de relier différents mots en concept. Cela nous permettrait d'évaluer la distance dans cette banque de connaissances entre les différents concepts demandés par le client dans sa requête et ceux présents dans les profils fournisseurs. Cela nous permettrait également d'enrichir notre banque de mots-clés pour chacune des expertises à l'aide des nouveaux mots-clés entrés par les clients dans leurs requêtes au lieu de les restreindre comme aujourd'hui à une liste prédéfinie de mots-clés pour chaque expertise. Cette notion d'apprentissage permettrait également d'améliorer encore la précision de l'outil comme évoqué à l'aide de retours clients. Une piste pour une utilisation optimale de l'outil par le partenaire est également d'élargir un maximum le spectre de fournisseurs disponibles sur leur plateforme, afin, dans un premier temps d'emmagasiner encore plus de mots-clés, mais surtout de satisfaire encore plus de clients. Cela passerait sûrement par la mise en place d'un processus de qualification des fournisseurs reprenant la même structure d'expertises et de mots clés que celle utilisée dans les requêtes clients.

Le principal point méritant une recommandation par rapport à ce projet est le manque de référence afin de véritablement évaluer la précision des différents algorithmes proposés. Ce manque de référence provenant de l'impossibilité d'utiliser les anciens projets du partenaire à cause de leur

trop grande différence avec le format actuel. Cette observation peut nous indiquer que dans un autre projet de ce type il serait intéressant d’aborder cette question du format des données d’entrées dès le début afin de trouver un compromis entre les formats nouveaux, mais difficilement testables par manque de référence et les formats originaux testables, mais manquant de précision ou de discernement entre les profils.

7.2 Perspectives d’amélioration

Comme décrit dans la partie 6.1, nous allons ici nous intéresser au scénario que nous aurions pu mettre en place afin de mesurer et de prendre en compte la satisfaction client dans notre outil. Le but ici est d’émettre des hypothèses sur les méthodes que nous aurions pu utiliser et de créer une entrée vers la partie discussions présente dans la suite.

La figure 7.1 suivante présente un logigramme de ce qui aurait pu être notre protocole de mesure de la satisfaction client. Après l’avoir expliqué plus en détail nous reviendrons également sur des possibilités de méthodes de prise en compte de ce retour et ainsi d’introduire une notion d’apprentissage pour notre outil.

Dans ce logigramme représentant un processus de récolte de la satisfaction client, l’indicateur de satisfaction CSAT (signifiant Customer Satisfaction Score) est choisi. Il existe différents indicateurs pouvant traiter de la satisfaction client celui-ci est le plus simple et correspond mieux à la notion de satisfaction voulant être mise en avant ici, à savoir le retour immédiat du public par rapport à un service qui lui est fourni. Des formulaires de satisfaction client contenant la question « Êtes-vous satisfait de la liste des fournisseurs que BRIDGR vous a fournie ? ». Afin de répondre à cette question, les utilisateurs auront cinq choix possibles :

- Pas du tout satisfait (niveau 1) ;
- Pas satisfait (niveau 2) ;
- Moyennement satisfait (niveau 3) ;
- Satisfait (niveau 4) ; et
- Extrêmement satisfait (niveau 5).

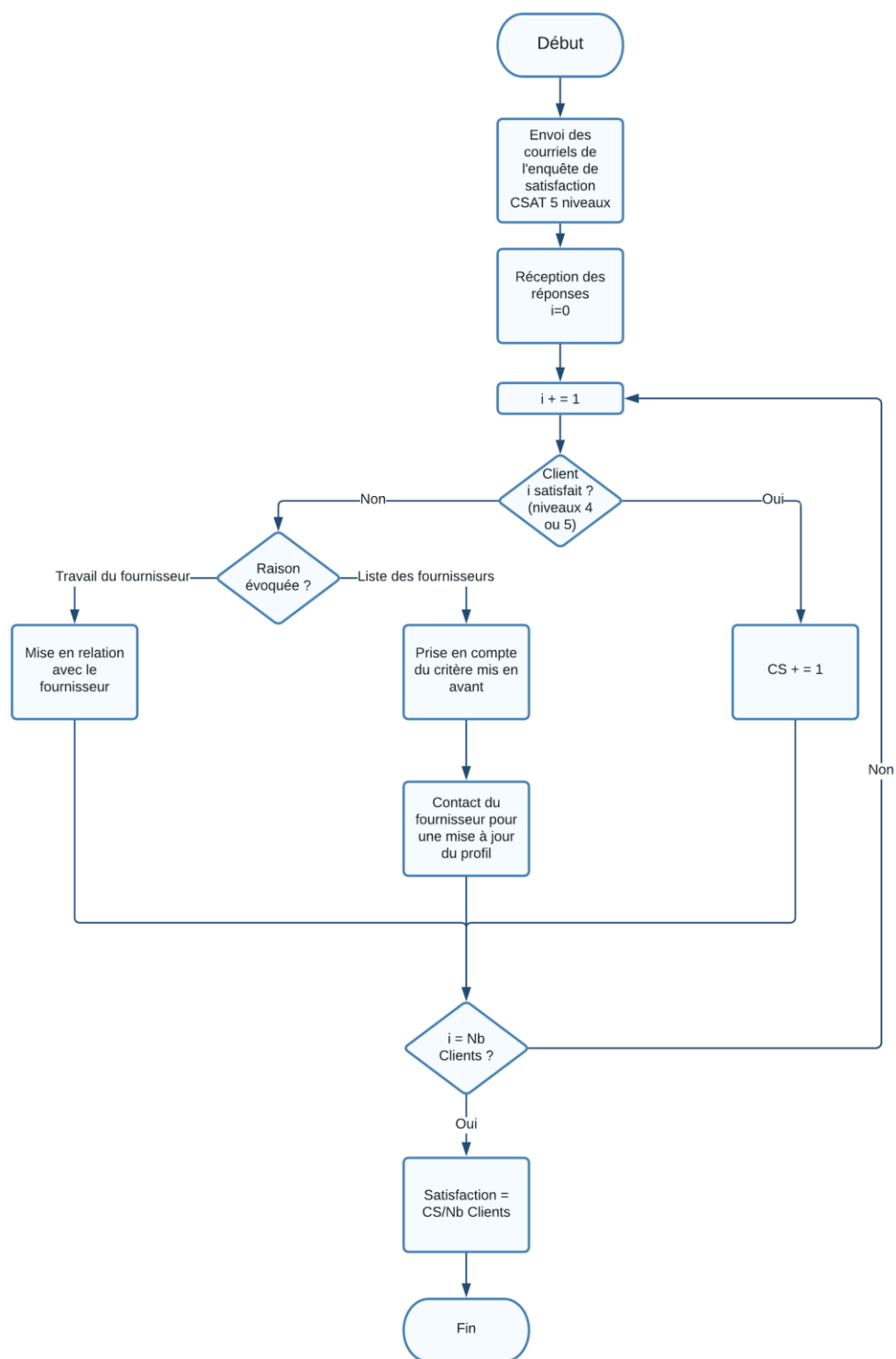


Figure 7.1 : Logigramme du scénario proposé d'évaluation de la satisfaction client

Une fois la satisfaction des clients déterminée, il est également nécessaire, si insatisfaction il y a, de déterminer son origine afin d'éviter de réitérer les mêmes vecteurs d'insatisfaction. Après réflexion avec l'équipe de BRIDGR, nous avons conclu qu'il y avait deux types d'insatisfaction possibles après utilisation de l'outil d'appariement. Le premier réside en une insatisfaction envers le travail fourni par le fournisseur choisi. Ce problème ne réside pas dans un élément que BRIDGR peut corriger, il est nécessaire que le client communique directement avec son fournisseur. Cependant, si ce genre de mécontentement se répète il est possible que BRIDGR considère une autre qualification du fournisseur et/ou une suppression de la base de données. Le second type de source de mécontentement repose dans la liste fournie par BRIDGR et l'outil d'appariement au client. Si jamais le client n'est pas satisfait de la liste fournie, cela repose sûrement sur un ou plusieurs critères de sélection. BRIDGR peut ainsi effectuer un travail afin d'affiner la sélection au niveau de ces critères dans le but, encore une fois d'améliorer l'outil et ainsi d'augmenter la satisfaction globale. Dans ce processus, les clients satisfaits ou extrêmement satisfaits (niveaux 4 ou 5) sont comptabilisés et ce nombre est ensuite comparé au nombre total de clients (nombre connu) afin d'établir le taux de satisfaction.

Cette ébauche de processus reste une ébauche et possède de nombreuses améliorations possibles notamment au niveau de l'automatisation de ce processus. En effet, le but de l'outil était d'accélérer en automatisant le processus d'appariement. De plus, en fonction du moment auquel les questionnaires de satisfaction sont envoyés la bijection des causes de l'insatisfaction peut être mise de côté et nous pouvons nous intéresser aux retours concernant le véritable résultat de l'outil, la liste de fournisseurs. Pour ce faire, nous avons pensé à une liste de résultats « interactive ». Cette liste inclurait des boutons « like » ou « dislike » qui permettrait au client de noter chacun des fournisseurs qui lui seraient proposés. Le fait d'inclure ce genre de note permettrait peut-être à l'outil de se rapprocher d'un système de recommandation à filtrage collaboratif comme évoqué dans la revue de littérature du chapitre 2. Si jamais le fait de « dislike » un ou plusieurs fournisseurs pose un problème nous pourrions uniquement ajouter la possibilité au client de nous indiquer quel fournisseur il a finalement choisi parmi la liste fournie. Le fait de sélectionner le troisième fournisseur, par exemple, nous communique ainsi des informations sur les deux premiers qui, malgré leur meilleur classement selon notre méthode de calcul, ne furent pas finalement choisis par le client.

RÉFÉRENCES

Adomavicius, G. and Tuzhilin, A., (2005) *Toward the next generation of recommender systems: a survey of the state-of-the-art and possible extensions*, Transactions on Knowledge and Data Engineering, vol. 17, pp. 734-749.

Baake, M., Grimm, U., and Giegerich, R. (2006). *Surprises in approximating levenshtein distances*. Journal of Theoretical Biology, pages 279-282

Baeza-Yates, R. A. and Ribeiro-Neto, B. (1999). *Modern Information Retrieval*. Addison-Wesley Longman Publishing Co., Inc., Boston, MA, USA.

Benouaret, I., (2017) *Un système de recommandation contextuel et composite pour la visite personnalisée de sites culturels*. Autre [cs.OH]. Université de Technologie de Compiègne. Français. ffNNT : 2017COMP2332ff. fftel-01767997f

Bisson, G. and Hussain, S. F. (2008). *Chi-sim : A new similarity measure for the coclustering task*. In ICMLA, pages 211-217.

Blei, D. M., Ng, A. Y., and Jordan, M. I. (2003). *Latent dirichlet allocation*. J. Mach. Learn. Res., 3 :9931022.

Blessing, L., Chakrabati, A., (2002) *DRM: A Design Research Methodology* Communication présentée à Proceedings of Les Sciences de la Conception.

Cosley, D., Lam, S. K., Albert, I., Konstan, J. A., and Riedl, J., (2003) *Is Seeing Believing? How Recommender Interfaces Affect Users' Opinions*, in CHI Lett. vol. 5, A. Press, Ed., 2003, pp. 585-592.

Deerwester, S., Dumais, S. T., Furnas, G. W., Landauer, T. K., and Harshman, R. (1990). *Indexing by latent semantic analysis*. JOURNAL OF THE AMERICAN SOCIETY FOR INFORMATION SCIENCE, 41(6) :391-407

Drachsler, H., Verbert, K., Santos, O. C., et Manouselis, M., (2015) *Panorama of Recommender Systems to Support Learning*, en, in Recommender Systems Handbook, F. Ricci, L. Rokach et B. Shapira, éd., Boston, MA : Springer US, 2015, p. 421-451, isbn : 9781489976369 9781489976376.

doi : 10.1007/978-1-4899- 7637-6_12. adresse : http://link.springer.com/10.1007/978-1-4899-7637-6_12

Gabrilovich, E. and Markovitch, S. (2007). *Computing semantic relatedness using wikipedia-based explicit semantic analysis*. In In Proceedings of the 20th International Joint Conference on Artificial Intelligence, pages 1606-1611.

Guide des technologies de l'industrie du futur par l'Alliance Industrie du Futur, Edition Mars 2018, http://www.industrie-dufutur.org/content/uploads/2018/03/Guide-des-Technologies_2018_V3.pdf

Guide pratique de l'usine du futur par la Fédération des Industries de Mécanique et l'Alliance Industrie du Futur, Edition Octobre 2015, <https://www.pfa-auto.fr/wp-content/uploads/2016/03/Guide-pratique-Usine-Automobile-du-Futur.pdf>

Huang, A. (2008). *Similarity Measures for Text Document Clustering*. In Holland, J., Nicholas, A., and Brignoli, D., editors, New Zealand Computer Science Research Student Conference, pages 49-56.

Herlocker, L. J., Konstan, A. J., Terveen, G. L., and Riedl, T. J., (2004) *Evaluating collaborative filtering recommender systems*, ACM Trans. Inf. Syst., vol. 22, pp. 5-53.

Hofmann, T. (1999). *Probabilistic latent semantic indexing*. In Proceedings of the 22nd annual international ACM SIGIR conference on Research and development in information retrieval, SIGIR '99, pages 5057, New York, NY, USA. ACM.

Howe, J. (2006) *The Rise of Crowdsourcing*. Wired. Tiré de : <https://www.wired.com/2006/06/crowds/> [consulté le 15 avril 2021].

Jaccard, P. (1901). *Étude comparative de la distribution florale dans une portion des alpes et des jura*. Bulletin de la Société Vaudoise des Sciences Naturelles, 37 :547579.

Jatnika, D., Bijaksana, M. A., Suryani, A. A., (2019). *Word2Vec Model Analysis for Semantic Similarities in English Words*. In 4th International Conference on Computer Science and Computational Intelligence 2019 (ICCSCI), 12-13 September 2019

Jiang, J. J. and Conrath, D. W. (1997). *Semantic similarity based on corpus statistics and lexical taxonomy*. CoRR, cmp-lg/9709008.

Leacock, C. and Chodorow, M. (1998). *Combining local context and wordnet similarity for word sense identification*. In Fellbaum, C., editor, MIT Press, pages 265-283, Cambridge, Massachusetts.

Lelu, A. (2002) *Comparaison de trois mesures de similarité utilisées en documentation automatique et analyse textuelle*. JADT 2002 : 6es Journées internationales d'Analyse statistique des Données Textuelles

Levenshtein, V. (1966). *Binary codes capable of correcting deletions, insertions, and reversals*. Soviet Physics Doklady, 10(8) :707-710.

Li, Y., Bandar, Z. A., McLean, D., (2002). *An Approach for Measuring Semantic Similarity between Words Using Multiple Information Sources*. In IEEE transactions on knowledge and data engineering, Vol. 15, No. 4, July/August 2003

Lin, D. (1998). *An information-theoretic definition of similarity*. In In Proceedings of the 15th International Conference on Machine Learning, pages 296-304. Morgan Kaufmann.

McNee, S. M., Riedl, J., and Konstan, J. A., (2006) *Being accurate is not enough: how accuracy metrics have hurt recommender systems*, CHI '06: CHI '06 extended abstracts on Human factors in computing systems, pp. 1097-1101, 2006.

Miller, N. B., (2003) *Toward a personal recommender system*, University of Minnesota, 2003, p. 185.

Miller, B. N., Konstan, J. A., and Riedl, J. T., (2004) *PocketLens: Toward a personal recommender system*, ACM Trans. Inf. Syst., vol. 22, pp. 437-476, 2004.

Navarro, G. (2001). *A guided tour to approximate string matching*. ACM Comput. Surv., 33(1) :31-88

Negre, E. (2013) *Comparaison de textes : quelques approches....*, Cahier du LAMSADE vol. 338, Avril 2013, <https://hal.archives-ouvertes.fr/hal-00874280>

Page Wikipédia du concept de *Système de recommandation*. Modifiée pour la dernière fois le 9 janvier 2021 [consultée le 15 Avril 2021].

Pazzani M.J., Billsus D. (2007) *Content-Based Recommendation Systems*. In: Brusilovsky P., Kobsa A., Nejdl W. (eds) *The Adaptive Web. Lecture Notes in Computer Science*, vol 4321. Springer, Berlin, Heidelberg. https://doi.org/10.1007/978-3-540-72079-9_10

Plateforme de crowdsourcing Amazon Mechanical Turk <https://www.mturk.com>

Porter, M. (1980). *An algorithm for suffix stripping*. *Program*, 14(3) :130_137.

Projet Galaxy Zoo : <https://www.zooniverse.org/projects/zookeeper/galaxy-zoo/> [consulté le 16 avril 2021].

Rapport *Industrie du futur, prêts, partez !* par l'Institut Montaigne publié en septembre 2018, <https://www.institutmontaigne.org/publications/industrie-du-futur-prets-partez>

Resnik, P. (1995). *Using information content to evaluate semantic similarity in a taxonomy*. In *IJCAI 1995*, pages 448453.

Safran, M., Che, D. (2015). *Real-time recommendation algorithms for crowdsourcing systems*. Southern Illinois University, King Saud University. <http://dx.doi.org/10.1016/j.aci.2016.01.001>

Salakhutdinov, R. and Mnih, A. (2007). *Probabilistic matrix factorization*. In *NIPS '07: Proceedings of the Twenty-First Annual Conference on Neural Information Processing Systems*. Curran Associates, Inc.

Schnitzer, S. (2019) *Task Recommendation in Crowdsourcing Platforms* <http://tuprints.ulb.tu-darmstadt.de/8549>

Sun, W., Wang, B., Cao, N., Li, M., Lou, W., Hou, Y.T., Li, H., (2013) *Privacy-preserving Multi-keyword Text Search in the Cloud Supporting Similarity-based Ranking* ASIA CCS'13, May 8–10, 2013, Hangzhou, China.

Terveen, G. L. and Hill, W., (2001) *Beyond Recommender Systems: Helping People Help Each Other*, in *Human-Computer Interaction in the New Millenium*, J. M. Carroll, Ed. Reading, MA, USA., 2001.

Wu, Z. and Palmer, M. (1994). *Verbs semantics and lexical selection*. In *Proceedings of the 32nd annual meeting on Association for Computational Linguistics*, ACL '94, pages 133138, Stroudsburg, PA, USA. Association for Computational Linguistics.

Yuen, M-C., King, I., Leung, K-S., (2012) *Task Recommendation in Crowdsourcing Systems* CrowdKDD'12, August 12, 2012, Beijing, China.

ANNEXE A : TABLEAU RÉCAPITULATIF DE LA RECHERCHE
BIBLIOGRAPHIQUE

Tableau A.1 : Tableau récapitulatif de la recherche bibliographique

Premier groupe de recherche	Systèmes de recommandation	Experts	Fournisseurs ou consultants
	« recommendation system » « recommender system »	« expert » « expertise »	« supplier » « service » consultan*
Deuxième groupe de recherche	Assignation des tâches et algorithmes		Marché à microtâches et associés
	« task assignment » « task recommendation » « recommendation algorithms »		« crowdsourcing » « micro-task market »
Troisième groupe de recherche	Similarité et comparaison	Mots et phrases	Analyses sémantiques et syntaxiques
	similar* « comparaison » « similarity-based »	« words » « mots » phrase* text* « text based »	analys* seman* syntax*

ANNEXE B : LISTE DES EXPERTISES

Commerce, gestion et affaires	Internet industriel et production de données	Production additive, matériaux et procédés innovants	Autres
Comptabilité	Capteurs autonomes et communicants	Composants électroniques	IA et science des données
Finance	Internet industriel des objets/ IIOT	Matériaux et alliages métalliques innovants	Outils facilitant le travail à distance
Ventes	Réseaux industriels	Assemblage innovant	
Marketing	Systèmes de télécommunication	Automatisation et robotique des processus	
Gestion et stratégie des entreprises	Composants/Objets intelligents	Composites gros volumes	Cloud et Cybersécurité
Gestion des connaissances et de la documentation	Analyse de données (Big Data Analytics)	Fabrication additive	Cybersécurité
Gestion de projet	Gestion et prise de décision en temps réel	Fonctionnalisation de surface	Surveillance et contrôle à distance
Export et internationalisation		Formage et usinage innovants ou optimisés	Systèmes de contrôle numérique
E-commerce	Outils de simulation avancée	Matériaux intelligents et adaptatifs	Cloud
Gestion Qualité Hygiène Sécurité Environnement	Réalité virtuelle	Matériaux non métalliques et fluides bio-sourcés	Blockchain
Management Agile	Conception de produits et de machines	Micro-fabrication	Entreposage et gestion des données
Gestion du risque	Simulation	Processus Near Net Shape	
Intelligence économique	Réalité Augmentée	Technologies de soudage hautes performances	Aspect écologique
Développement d'affaires	Jumeau Numérique	Contrôle Non Destructif innovant	Gestion du cycle de vie
Investissement	Conception basée sur un modèle		Ecodesign
Marketing digital		Robots collaboratifs et machines intelligentes	Durabilité
Conseil en matière de transformation des entreprises	Production/Logistique		Logistique inversée
Gestion de produit	Gestion de l'inventaire	Mise à niveau des machines	Bâtiments écologiques
	Digitalisation de la chaîne d'approvisionnement	Fiabilité des systèmes mécatroniques	Gestion et efficacité énergétique
Ressources Humaines	Gestion de la relation client/fournisseur	Machines intelligentes	Responsabilité sociétale des entreprises
Technologies centrées sur l'homme	Qualité 4.0	Machines programmables innovantes ou optimisées	
Formation en action	Gestion et supervision des opérations	Applications robotiques collaboratives	CleanTech
		Cobots et exosquelettes	
	Maintenance		
Conception et ingénierie de logiciels	Systèmes d'information d'entreprise		
Développement Web	Fabrication agile et flexible		
Développement mobile	Gestion logistique et transport		
Conception de logiciels et infographie	Contrôle et audit des processus		
Développement de logiciels			
Intégration des systèmes et des logiciels			
Intelligence d'affaires			
Systèmes embarqués			

Figure B.1 : Liste des expertises

ANNEXE C : LISTE DES GLOSSAIRES PAR EXPERTISE

Fabrication additive	Hubs.com	Sculpteo.com	Rigid Ink	mwa3D.com	Compétitivité	dyssca.org	accounting.com	daysimple.com
Contrôle Non Destructif innovant	Nde-nd.org	asnt.org	infrastructureps.com	holowayrnt.com	France	investopedia.com	nationwide.com	
Technologies de soudage hautes performances	weldingtech.net	weldguru.com	highspeedtraining.co.uk		Ventes	saleshacker.com	learn.g2.com	
Matériaux et alliages métalliques innovants	aimm.com	materialiseeng training.co.uk			Marketing	precisionmarketinggroup.com	malchimp.com	marketingterms.com
Processus Near Net Shape	metaltek.com	twi-global.com			Gestion et stratégie des entreprises	thebalancecareers.com	intralocus.com	
Fonctionnalisation de surface	steticonind.co in	coatingguyz.com			Gestion des connaissances et de la documentation	krill.com	knowledgepoint.com.au	
Assemblage innovant	arebasolutions. com	thinkautomation fr.com			Gestion de projet	smartsheet.com	projectmanager.com	
Automatisation et robotique des processus	uipath.com				Export et internationalisation	eximguru.com	legacy.trade.gov	
Composites gros volumes	compositestextile sl.com	compositestextile ions.com			E-commerce	outletboxdesign.com	ecommerceguide.com	tychesoftware.com
Composants électroniques	westflorida.com ponents.com	digikey.com			Gestion Qualité Hygiène Sécurité Environnement	isobva.de	community.intellex.com	
Formage et usinage innovants ou optimisés	wisconsinmetal tech.com				Management Agile	agilealliance.org	scrum.org	
Matériaux intelligents et adaptatifs	arebasolutions. com				Gestion du risque	quantivate.com	trackernetworks.com	
Matériaux non métalliques et fluides bio-sourcés	biohassdeconsul tancy.com				Intelligence économique	madfirmasd.org	encyclopedia.com	insightsoftware.com
Micro-fabrication	wikipedia.org				Développement d'affaires	ambusinvestdevelopment.co.uk	businessballs.com	
					Investissement	moniesseuse.ca	jpmorgan.com	investorjunkie.com
					Marketing digital	marketingterms.com	lyfemarketing.com	
					Conseil en matière de transformation des entreprises	managementconsulted.com	linnolt.com	
					Gestion de produit	productplan.com	altalis	

Figure C.1 : Liste des glossaires par expertise

Technologies centrées sur l'homme	encyclopedia2.thefre edictionary.com	copnizant.com			Capteurs autonomes et communicants	nl.com	omega.co.uk		
Formation en action	twinkl.com				Internet industriel des objets/ IOT	belimo.com	dzone.com	blog.temboo.com	
Développement Web	careerfoundry.com	wearelagram. com			Réseaux industriels	rammeter.com	la.omron.com		
Développement mobile	thinkapps.com	willowtreasaps. com	answethub.com		Systèmes de télécommunication	techtvrb.org	lamcommun.c om	itservices.usc.edu	
Conception de logiciels et infographie	xenimgage.com	icons8.com			Composants/Objets intelligents	qblis.com	arm.com		
Développement de logiciels	tda.gov	whodevtable.co m			Analyse de données (Big Data Analytics)	biodata-madesim ple.com	biodataanalyt ics.com	statistics.com	
Intégration des systèmes et des logiciels	sourcecmaking.com	icoreolutions. com			Gestion et prise de décision en temps réel	gartner.com	decision-makin g-solutions.co m	home.ubalt.edu	
Intelligence d'affaires	silense.com	insightssoftware .com	solutionsreview.com	item.com	Gestion de l'inventaire	stitchlabs.com	clearinventory .com		
Systèmes embarqués	bargroup.com	loftize.com	glossartech.com		Numérisation de la chaîne d'approvisionnement	supplychaingame changer.com	blindsight.com		
					Gestion de la relation client/fournisseur	mbaskool.com	zycus.com		
					Qualité 4.0	asq.org	qccentral.com		
Réalité virtuelle	ket-studios.com	vrjossary.org	arpost.co		Gestion et supervision des opérations	vocabulary.com	informtl.com		
Conception de produits et de machines	design1st.com	ukdesign.cc	zallaco.com		Maintenance	dpsl.com	ldcon.com	de-spinetech.com	enworkorder.co m
Simulation	acm-signal.miskz.org				Systèmes d'information d'entreprise	quizlet.com	opengroup.org		
Réalité Augmentée	wikitude.com	arpost.co	ar.rocks		Fabrication agile et flexible	bulip.co	gartner.com		
Jumeau Numérique	digitalwincconsortiu m.org	brighttalk.com			Gestion logistique et transport	inboundlogistics.4 gm	frontapp.com	zipline logistics.com	ehol.com
Conception basée sur un modèle	mburteomobd.com	moadsafe.com			Contrôle et audit des processus	quantivate.com	process-control s.com	emerson.com	dummies.com

Figure C.1 : Liste des glossaires par expertise (suite)

Mise à niveau des machines	merriam-webster.com			Cybersecurté	globalknowledge.com	csrc.nist.gov	
Fiabilité des systèmes mécaniques	flashcardmachine.com	rampf-group.com		Surveillance et contrôle à distance	wikibooks.org		
Machines intelligentes	infineon.com			Systèmes de contrôle numérique	wikibooks.org		
Machines programmables innovantes ou optimisées	infineon.com	learn.q2.com		Cloud	solutionsreview.com	cloudwatchhub.eu	
Applications robotiques collaboratives	cobotrends.com			Blockchain	connect.comptia.org	objectcomputing.com	
Cobots et exosquelettes	blog.robotiq.com			Entreposage et gestion des données	oracle.com	cloudmoyo.com	
IA et science des données	dataquest.io	wikipedia.org	appen.com	Gestion du cycle de vie	productlifecycle-management.com	3d-printing-expert.com	
				Ecodesign	cleantechnica.com		
				Durabilité	earthshiftglobal.com	iqd.com	
				Logistique inversée	informit.com		
				Bâtiments écologiques	bpcgreenbuilders.com	greenspec.co.uk	
				Gestion et efficacité énergétique	nga.com	dexma.com	energy.gov
Outils facilitant le travail à distance	wrike.com	proofhub.com		Responsabilité sociétale des entreprises	satelllinstitute.org		
				CleanTech	cleantechnica.com	stats.oecd.org	

Figure C.1 : Liste des glossaires par expertise (suite et fin)

ANNEXE D : LISTE DES SECTEURS D'ACTIVITÉ

Tableau D.1 : Liste des secteurs d'activité

Administration gouvernementale	Construction navale	Industrie composants électriques et électroniques	Mines/métaux	Santé et forme
Agriculture	De gros	Industrie de l'édition	Musées/Institutions	Semi-conducteurs
Animation	Des sports	Industrie de la recherche	Médecine douce	Services aux consommateurs
Apprentissage en ligne	Diversissement/Production de films	Industrie militaire	Médias de diffusion	Services d'information
Architecture/Planification	Dotation/Recrutement	Information technologique IT	Médicaments	Services des installations
Arts et orfèvrage	Défense/Espace	Installations/services de loisirs	Nanotechnologie	Services environnementaux
Articles de luxe/Bijoux	Développement de programme	Institutions religieuses	Nourriture/Boisson	Services financiers
Articles de sport	Electronique grand public	Jeux d'argent/Casino	Organisation civique et sociale	Services individuels/familiaux
Arts performants	Emballages/Conteneurs	Jeux d'ordinateur	Organisation politique	Services juridiques
Assurances	Enseignement primaire/secondaire	Journaux/Journalisme	Papier/produits forestiers	Services événementiels
Automobile	Enseignement supérieur	Judiciaire	Philanthropie	Soin de la santé mentale
Autre industrie	Entreposage	L'automatisation industrielle	Plastiques	Supermarchés
Aviation/aérospatiale	Externalisation/Délocalisation	La musique	Pratique du droit/lobbing d'avocats	Sécurité informatique et réseaux
Banque	Fabrication de chemins de fer	La photographie	Pratique médicale	Sécurité Enquête
Banque d'investissement/Venture	Films	La production de médias	Production alimentaire	Textile
Banque/Hypothèque	Forces de l'ordre	La sécurité publique	Produits chimiques	Traduction/Localisation
Beaux-Arts	Formation professionnelle	Laitier	Produits de beauté	Transport
Bibliothèque	Fournitures/équipements de bureau	Les affaires internationales	Publication en ligne	Télécommunications
Biens de consommation	Gestion de l'investissement	Livraison de colis/fret	Pétrole/Energie/Solaire/Energie verte	Utilitaires
Biotechnologies/Technologies vertes	Gestion de l'éducation	Logiciels informatiques et ingénierie	Pêche/rie	Versé/Céramique/Béton
Bureau exécutif	Groupes de réflexion	Logistique/Achats	Relations gouvernementales	Vins/Spiritueux
Bureau législatif	Génie civil	Loisirs/Voyage	Relations publiques/RP	Vétérinaire
Capital risque/VC	Génie mécanique ou industriel	Machinerie	Renouvelables/Environnement	Vêtements/Mode
Collecte de fonds	Hospitalité	Marché des capitaux	Ressources humaines/RH	Internet
Commerce de détail	Hôpital/Soins de santé	Maritime	Restaurants	Je tabac
Commerce international/Développement	Immobilier commercial	Marketing/Publicité/Vente	Règlement extrajudiciaires des différends	À but non lucratif/bénévoles
Compagnies aériennes/Aviation	Immobilier/hypothèque	Matériaux de construction	Rédaction/édition	Événement en l'air
Conception	Importeur/Exporter	Matériel informatique	Réseaux informatiques	Équipement médical
Conception graphique/conception Web	Impression	Meubles	Sans fil	Étude de marché
Conseil en gestion				
Construction				