



Recommandation d'objets d'apprentissage basée sur des objectifs d'apprentissage en utilisant les modèles de plongement de phrases

Molka Tounsi Dhouib, Catherine Faron, Oscar Rodríguez Rocha

► To cite this version:

Molka Tounsi Dhouib, Catherine Faron, Oscar Rodríguez Rocha. Recommandation d'objets d'apprentissage basée sur des objectifs d'apprentissage en utilisant les modèles de plongement de phrases. PFIA 2022 - Conférence Nationale sur les Applications Pratiques de l'Intelligence Artificielle 2022, Jun 2022, Saint-Étienne, France. hal-03662801v2

HAL Id: hal-03662801

<https://hal.science/hal-03662801v2>

Submitted on 23 May 2022

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

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



HAL Authorization

Recommandation d'objets d'apprentissage basée sur des objectifs d'apprentissage en utilisant les modèles de plongement de phrases

M. Tounsi Dhouib¹, C. Faron¹, O. Rodriguez Rocha²

¹ Université Côte d'Azur, Inria, CNRS, I3S, Sophia Antipolis, France

² Teach on Mars, France

{dhouib, faron}@i3s.unice.fr, oscar.rodriguez@teachonmars.com

Résumé

Avec la transformation numérique, l'adaptation et le développement des compétences sont devenus des facteurs majeurs pour améliorer les performances des collaborateurs et des entreprises. Comprendre les besoins des collaborateurs et les aider à atteindre leurs objectifs de développement de carrière est un véritable défi aujourd'hui. Dans ce travail, nous partageons notre expérience pour mettre en place un système de recommandation automatique permettant aux apprenants de trouver des objets d'apprentissage pertinents en fonction de leurs objectifs d'apprentissage. Cette tâche de mise en correspondance se base principalement sur la détermination de la similarité sémantique entre les objectifs et le contenu textuel des objets d'apprentissage. Nous avons évalué de manière comparative trois modèles pré-entraînés de plongement de phrases de l'état de l'art pour la tâche de la recommandation d'objet d'apprentissage. Les résultats des expérimentations montrent que l'utilisation de ces modèles de plongement de phrases dans le processus de recommandation est plus performante que le modèle BM25 d'Elasticsearch classiquement utilisé dans l'industrie.

Mots-clés

Apprentissage intelligent, Recommandation de cours de formation pour les collaborateurs, Systèmes de recommandation.

Abstract

With the digital transformation, adaptation and development of skills have become key factors in improving employee and company performance. Understanding the needs of employees and helping them achieve their career development goals is a real challenge today. In this work, we share our experience to implement an automatic recommendation system that allows learners to find relevant learning objects according to their learning goals. This matching task is mainly based on determining the semantic similarity between goals and the text of learning objects. We comparatively evaluated three state-of-the-art pre-trained sentence embeddings models for the learning object recommendation task. Experimental results show that using sentence embeddings models in the recommendation process

outperforms the Elasticsearch BM25 model generally used in industry.

Keywords

Intelligent learning, Recommender systems, Recommendation of training courses for collaborators.

1 Introduction

Aujourd'hui, nous sommes confrontés à un contexte de mutation des modes de travail. Les entreprises doivent répondre aux attentes et aux besoins des collaborateurs en matière de formation et de développement des compétences, en les aidant à choisir la formation qui convient à leur parcours, à leurs compétences actuelles, aux besoins du projet mais aussi à leurs objectifs d'apprentissage.

Dans le cadre d'un projet de recherche collaboratif entre l'équipe de recherche WIMMICS¹ et la société *Teach on Mars*², nous souhaitons mettre en place un système de recommandation d'objets d'apprentissage qui correspondent aux besoins et aux objectifs des collaborateurs d'une entreprise. *Teach on Mars* est spécialisée dans l'e-learning et l'apprentissage mobile et développe une plate-forme d'apprentissage spécialisée dans la formation continue au sein des entreprises. L'objectif de *Teach on Mars* est de relier les collaborateurs des entreprises à la formation et aux communautés qui sont essentielles pour améliorer leur travail et leur performance. Notre système de recommandation repose sur deux composants : (i) d'un côté, le profil de l'apprenant qui est construit à partir de l'ensemble des compétences déjà acquises et de l'ensemble des compétences qu'il souhaite apprendre, (ii) de l'autre côté, une base de données qui contient 760 objets d'apprentissage manuellement identifiés par un expert de *Teach on Mars*, principalement écrits en français ou en anglais.

Dans cet article, nous proposons une approche basée sur le calcul de la similarité sémantique entre les objectifs d'apprentissage fournis par l'apprenant et l'ensemble des contenus textuels des objets d'apprentissage de la base. Nous rapportons le résultat des expériences que nous avons menées pour répondre au cas d'utilisation de *Teach on Mars* : nous comparons les performances de trois modèles pré-entraînés

1. <https://team.inria.fr/wimmics/>

2. <https://www.teachonmars.com/fr/>

de plongement de phrases de l'état de l'art pour la tâche de la recommandation des objets d'apprentissage par rapport au modèle BM25³ qui est l'algorithme de similarité utilisé par Elasticsearch pour représenter la pertinence d'un document par rapport à la requête.

Nos principales questions de recherche sont : (i) Quel est le meilleur modèle que nous pouvons utiliser dans notre cas, pour construire des représentations vectorielles des objets et objectifs d'apprentissage afin d'améliorer la qualité des recommandations ? (ii) Quels sont les meilleurs paramètres et méta-données à utiliser pour améliorer la qualité de notre système de recommandation ? (iii) Les modèles de plongement peuvent-ils vraiment améliorer la qualité de notre système de recommandations ?

Le reste de cet article est organisé comme suit. La section 2 donne un aperçu de l'état de l'art sur les systèmes de recommandation dans le domaine de l'éducation. La section 3 détaille notre approche de recommandation pour répondre au cas d'usage de *Teach on Mars*. La section 4 rapporte et discute les résultats de nos expériences menées sur les données de *Teach on Mars*. La section 5 conclut et donne quelques orientations pour des travaux futurs.

2 État de l'art

Les systèmes de recommandation (RS) sont des systèmes de filtrage d'information ayant comme objectif d'aider les utilisateurs à trouver des contenus, des produits ou des services en se basant sur les préférences des autres utilisateurs [13, 3]. En se basant sur plusieurs études sur les RS [7, 6, 8, 18], nous pouvons distinguer quatre techniques de recommandation : (i) recommandation basée sur le contenu qui se base sur l'analyse d'un ensemble des attributs des articles précédemment aimés par les utilisateurs afin de recommander ceux dont le contenu est le plus similaire [21]. Plusieurs travaux dans le domaine de l'e-learning ont utilisé cette technique [11, 15]. (ii) recommandation basée sur le filtrage collaboratif qui se base principalement sur l'analyse de l'utilisateur pour recommander des éléments appréciés par des utilisateurs similaires [10]. Parmi les travaux dans le domaine de l'e-learning, nous pouvons citer [22, 19]. (iii) recommandation basée sur la connaissance qui utilise des ontologies pour suggérer des articles à l'utilisateur en fonction du contexte de l'utilisateur, du contexte de l'article et de leurs relations modélisées par l'ontologie. Parmi les travaux dans le domaine de l'e-learning, nous pouvons citer [1, 17]. (iv) recommandation hybride qui combine deux ou plusieurs techniques citées ci-dessus pour améliorer la recommandation. Parmi les travaux dans le domaine de l'e-learning, nous pouvons citer [2, 9].

Dans ce travail, nous adoptons une approche de recommandation basée sur le contenu et plus précisément sur la mesure de similarité textuelle sémantique (STS) entre le profil de l'utilisateur qui est représenté par un ensemble d'objectifs d'apprentissage et le contenu textuel des objets d'apprentissage. Une méthode de correspondance stricte des

mots ou des modèles TF-IDF entre les descriptions textuelles donnerait de mauvais résultats, car elle ne prendrait pas en compte les relations syntaxiques et sémantiques des mots telles que les synonymes ou les polysémies. Pour pallier cela, les techniques de plongement de mots (word embedding) ont été utilisées avec beaucoup de succès dans les tâches de STS. Le plongement de mots est une représentation distribuée des mots qui exploite la sémantique des mots en les faisant correspondre à des vecteurs de nombres réels. L'inconvénient de cette méthode est l'incapacité de ces modèles à prendre en compte le contexte des mots et à approfondir les relations entre les mots de la phrase [14, 12]. L'état de l'art sur les modèles de plongement de mots a récemment évolué vers ce que l'on appelle le plongement de mots contextuel. Les modèles d'apprentissage profond basés sur des architectures de transformateurs, tels que USE (Universal Sentence Encoder) [4], BERT (Bidirectional Encoder Representations from Transformers) [5] et SBERT (Sentence-BERT) [16], ont montré les meilleures performances sur les benchmarks de la tâche STS.

Dans ce travail, nous avons choisi d'évaluer le bénéfice de l'utilisation de ces modèles de plongement contextuel de l'état de l'art par rapport au modèle BM25 et cela pour plusieurs raisons : (i) ils ont démontré leur efficacité pour la tâche de STS, (ii) ils fournissent des modèles multilingues, et permettent ainsi la définition d'une approche générique de recommandation qui peut être étendue à plusieurs langues, (iii) ils n'ont pas besoin de passer par une étape d'extraction d'entités nommées car ils permettent de représenter sous forme de vecteurs des phrases entières avec leurs informations sémantiques.

3 Approche proposée

Dans notre scénario de recommandation, lorsqu'un nouvel objectif d'apprentissage est exprimé par l'apprenant, un ensemble d'objets d'apprentissage les plus pertinents pour cet objectif devrait être automatiquement suggéré. Dans notre approche, nous supposons qu'un objet d'apprentissage est pertinent pour un objectif d'apprentissage si ces deux derniers sont sémantiquement similaires.

La figure 1 présente l'approche que nous proposons. Elle comprend deux étapes principales : (i) la représentation vectorielle des objets et objectifs d'apprentissage, et (ii) le calcul de la similarité sémantique entre objets et objectifs d'apprentissage.

3.1 Représentation vectorielle des objectifs et objets d'apprentissage

Cette première étape permet de représenter chaque objet ou objectif d'apprentissage par un vecteur qui capture la sémantique de ses phrases de telle manière que l'objectif d'apprentissage et tous les objets d'apprentissage pertinents soient proches dans l'espace vectoriel.

3.1.1 Modèles de plongement de phrases

Nous avons commencé par étudier les différents modèles contextuels et multilingues pour générer ces représentations vectorielles :

3. <https://www.elastic.co/fr/blog/practical-bm25-part-1-how-shards-affect-relevance-scoring-in-elasticsearch>

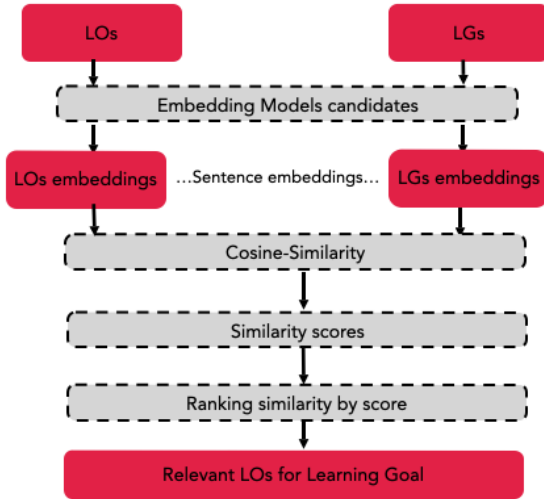


FIGURE 1 – Processus de recommandation des objets d’apprentissage.

Multilingual Universal Sentence Encoder (MUSE). MUSM⁴[20] convertit un texte de longueur variable en un vecteur de 512 dimensions. Il est destiné à être utilisé pour des tâches de classification de textes, clustering, recherche de similarités textuelles sémantiques, etc.

Sentence-BERT. Le modèle SBERT représente une modification de l’architecture du modèle BERT pré-entraîné qui utilise des structures de réseau siamois et triplet pour dériver des incorporations de phrases sémantiquement significatives qui peuvent être comparées en utilisant la cosinus-similarité. En effet, le principal inconvénient des modèles BERT est que la recherche de la paire la plus similaire est coûteuse en terme de temps de calcul [16].

SBERT fournit deux modes principaux pour la recherche sémantique : (i) La recherche sémantique symétrique où la requête et les textes du corpus ont la même longueur. Parmi les modèles de cette catégorie, on peut citer *Paraphrase*⁵, *Distiluse*⁶. (ii) La recherche sémantique asymétrique pour des requêtes courtes (c’est-à-dire une question ou un mot-clé) mais où les entrées dans le corpus sont plus longues. Parmi les modèles de cette catégorie, nous pouvons citer *MsMarco*.

Pour le traitement des objets d’apprentissage, si nous considérons uniquement leur titre, nous sommes dans le cas d’une recherche symétrique. Mais dans le cas où nous utilisons le contenu textuel de l’objet d’apprentissage, nous sommes plutôt dans le cas de la recherche asymétrique. L’inconvénient de la recherche asymétrique avec SBERT est qu’il n’existe pas de modèle multilingue pour générer les représentations vectorielles, ce qui ne répond pas à nos besoins par rapport aux données que nous avons, qui sont principalement en français et en anglais. Pour cette raison,

nous avons décidé d’utiliser les modèles de recherche symétriques, non seulement pour encoder le titre de l’objet d’apprentissage mais aussi pour encoder son contenu textuel. Nous avons identifié deux modèles multilingues à évaluer : (i) *Paraphrase* fait correspondre les phrases et les paragraphes à un espace vectoriel dense de 768 dimensions et peut être utilisé pour des tâches telles que le clustering ou la recherche sémantique. (ii) *Distiluse* fait correspondre les phrases et les paragraphes à un espace vectoriel dense de 512 dimensions et peut être utilisé pour des tâches telles que le clustering ou la recherche sémantique.

3.1.2 Méthode de calcul des représentations vectorielles des objets et objectifs d’apprentissage

Représentation vectorielle simple. La représentation vectorielle d’un objet d’apprentissage (ou d’un objectif d’apprentissage) x est un vecteur $V(x)$ qui représente le résultat direct obtenu à partir du modèle utilisé.

Représentation vectorielle moyenne des objectifs d’apprentissage en se basant sur les niveaux d’un thésaurus.

Nous considérons ici un thésaurus qui permet d’organiser l’ensemble des objectifs d’apprentissage selon différents niveaux hiérarchiques et dont les feuilles correspondent aux mots clés que nous utilisons pour annoter les contenus textuels. En utilisant un tel thésaurus, nous pouvons enrichir la représentation vectorielle simple d’un texte basée sur les seuls mots clés, avec les concepts des niveaux supérieurs dans le thésaurus.

La représentation vectorielle $V(lg)$ d’un objectif d’apprentissage lg est la moyenne des représentations vectorielles en considérant N niveaux du thésaurus :

$$V(lg) = \frac{1}{N} \sum_{i=1}^N V(c_i), \quad (1)$$

où $V(c_i)$ est la représentation vectorielle simple du concept c_i apparaissant dans lg et N est le nombre de niveaux considérés pour enrichir les représentations.

3.2 Similarité sémantique

Afin de pouvoir mesurer la similarité sémantique entre les représentations vectorielles V d’un objectif d’apprentissage LG et l’objet d’apprentissage LO , nous avons utilisé la métrique basée sur le cosinus.

$$sim(LO, LG) = \frac{V(LO) \cdot V(LG)}{\|V(LO)\| \cdot \|V(LG)\|}, \quad (2)$$

Un objet d’apprentissage lo est recommandé pour un objectif d’apprentissage lg si la valeur de $sim(lo, lg)$ est supérieure à un seuil donné.

4 Expérimentations

4.1 Données et protocole d’évaluation

Le jeu de données est composé de deux éléments principaux : les objectifs d’apprentissage (LGs) et les objets d’apprentissage (LOs).

4. <https://tfhub.dev/google/universal-sentence-encoder-multilingual/3>

5. <https://huggingface.co/sentence-transformers/paraphrase-multilingual-mpnet-base-v2>

6. <https://huggingface.co/sentence-transformers/distiluse-base-multilingual-cased-v1>

4.1.1 Objectifs d'apprentissage

Nous avons utilisé un référentiel interne à *Teach on Mars* qui a été défini par les experts de contenus de l'entreprise. Ce référentiel contient 166 concepts répartis en trois niveaux : (i) 5 catégories, (ii) 25 thématiques et (iii) 133 mots clés. Le tableau 1 montre le nombre de thématiques et de LOs par catégorie.

TABLE 1 – Référentiel interne des LGs

Catégories	Mots clés par catégorie	Thématiques par catégorie	Los par catégorie
Développement Personnel	33	6	194
Management and Leadership	36	6	181
Responsabilité Sociale d'entreprise	15	4	141
Business Performance	15	3	78
Innovation	39	6	190

4.1.2 Objets d'apprentissage

Un expert de *Teach on Mars* a identifié manuellement sur le Web (crawling) 1350 LOs. Il s'agit d'articles, de vidéos, ou de broadcasts. Chaque LO est associé manuellement par l'expert à une thématique unique du référentiel de *Teach on Mars*. Nous ne considérons dans cette étude que les LOs de type article car le texte est généralement plus long qu'une description de vidéo ou de broadcast. Nous avons obtenu 760 LOs dont 400 en langue française et 360 en anglais. La première étape du traitement consiste à récupérer le titre et le contenu textuel de chaque article.

Voici un exemple d'objet d'apprentissage :

Title: Learn to make decisions that last
take this quick test

Learning object text: [...] Everyday we
are faced with a multitude of decisions
that alter our lives in small or signi-
ficant ways. How you weigh up the pros
and cons of each decision and decide
which direction to take isn't always
easy [...]

4.1.3 Protocole d'évaluation

Afin de déterminer le meilleur modèle et la meilleure représentation vectorielle des LOs et LGs, nous avons défini 20 expérimentations dont les paramètres sont décrits dans la Table 2. Comme *baseline* nous avons utilisé le modèle BM25 d'Elasticsearch.

Afin de mesurer la correspondance entre les recommandations produites automatiquement et les recommandations produites manuellement par les experts, nous avons utilisé

les métriques de précision, rappel, F1 et "précision à N" en considérant les N LOs les mieux classés :

$$P@N = \frac{\text{EP parmi le top } N \text{ des ER}}{N}. \quad (3)$$

où EP représente les éléments pertinents et ER représente les éléments recommandés.

Les expériences ont été menées en utilisant la méthodologie de validation croisée 5 fois. La figure 2 présente la performance de notre système pour chaque paramètre testé en terme de précision, rappel, F1 score et precision@N.

TABLE 2 – Cadre expérimental

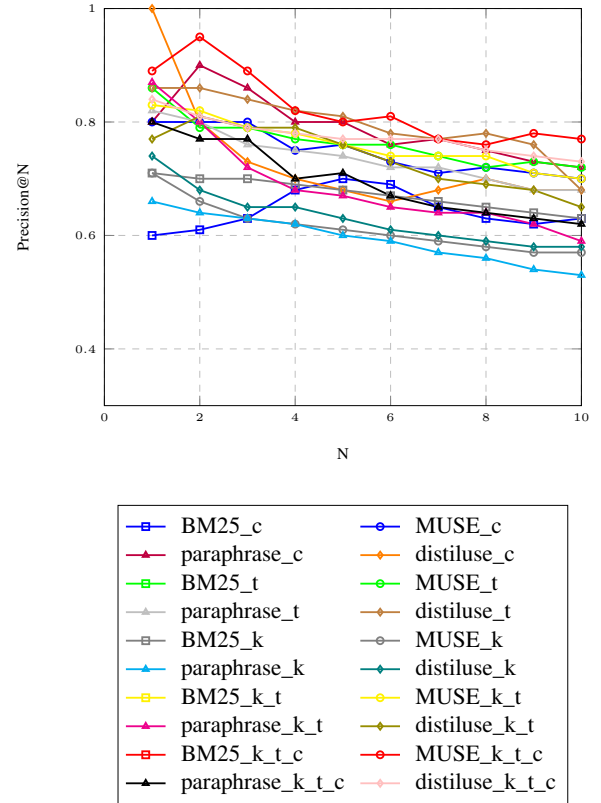
Expérimentations	Modèles	Paramètres
BM25_c	BM25	catégorie
BM25_t	BM25	thématique
BM25_k	BM25	mot clé
BM25_k_t	BM25	mot clé/ thématique
BM25_k_t_c	BM25	mot clé/ thématique catégorie
MUSE_c	MUSE	catégorie
MUSE_t	MUSE	thématique
MUSE_k	MUSE	mot clé
MUSE_k_t	MUSE	mot clé/ thématique
MUSE_k_t_c	MUSE	mot clé/ thématique/ catégorie
paraphrase_c	Paraphrase	catégorie
paraphrase_t	Paraphrase	thématique
paraphrase_k	Paraphrase	mot clé
paraphrase_k_t	Paraphrase	mot clé/ thématique
paraphrase_k_t_c	Paraphrase	mot clé/ thématique/ catégorie
distiluse_c	Distiluse	catégorie
distiluse_t	Distiluse	thématique
distiluse_k	Distiluse	mot clé
distiluse_k_t	Distiluse	mot clé/ thématique
distiluse_k_t_c	Distiluse	mot clé/ thématique/ catégorie

4.2 Résultat et discussion

Nous obtenons les meilleures performances du système en utilisant le niveau "Thématique" du référentiel interne de *Teach on Mars*, avec la meilleure valeur de précision en utilisant le modèle MUSE et la meilleure valeur de rappel en utilisant le modèle "Paraphrase". La meilleure valeur de F1 est obtenue en utilisant le modèle MUSE avec le niveau le plus précis du référentiel interne c.à.d "mot clé". En se basant sur la performance de notre système en terme de

Expérimentations	Seuil	P	R	F1
BM25_c	1.016	0.193	0.040	0.067
MUSE_c	0.130	0.350	0.370	0.360
paraphrase_c	0.240	0.310	0.280	0.290
distiluse_c	0.080	0.340	0.380	0.310
BM25_t	2.130	0.211	0.114	0.140
MUSE_t	0.190	0.610	0.450	0.520
paraphrase_t	0.270	0.400	0.610	0.480
distiluse_t	0.190	0.500	0.490	0.500
BM25_k	2.54	0.215	0.191	0.202
MUSE_k	0.220	0.520	0.590	0.550
paraphrase_k	0.360	0.390	0.61	0.480
distiluse_k	0.260	0.450	0.580	0.500
BM25_k_t	2.54	0.205	0.165	0.182
MUSE_k_t	0.130	0.370	0.570	0.470
paraphrase_k_t	0.270	0.330	0.590	0.430
distiluse_k_t	0.150	0.380	0.570	0.460
BM25_k_t_c	3.86	0.211	0.150	0.175
MUSE_k_t_c	0.080	0.360	0.560	0.440
paraphrase_k_t_c	0.300	0.380	0.510	0.440
distiluse_k_t_c	0.150	0.430	0.510	0.470

(a) Precision, Rappel et F1-scores.



(b) Précision@N

FIGURE 2 – Performances du système en fonction des différents paramètres expérimentés .

P@N, nous avons obtenu le meilleur résultat en rajoutant du contexte au troisième niveau du référentiel c.à.d “mot clé”, et ce en utilisant la moyenne des représentations vectorielles des trois niveaux du référentiel, c.à.d “mot clé”, “thématique” et “catégorie” avec le modèle MUSE (MUSE_t en vert avec un rond). Nous remarquons aussi que le modèle Distiluse a une meilleure valeur de cette mesure pour la P@1 (distiluse_c en orange) en utilisant le niveau le plus générique du référentiel (catégorie) et P@5 et P@8 en utilisant les thématiques (distiluse_t en marron).

Sans surprise, les modèles de plongement contextuel sont de loin plus performants que les modèles classiques qui se basent sur BM25 pour identifier les éléments pertinents. Le modèle MUSE permet une légère amélioration de la recommandation par rapport aux autres modèles de SBERT. De plus l’enrichissement contextuel des objectifs d’apprentissage en utilisant les relations hiérarchiques du référentiel interne de *Teach on Mars* ou d’un thésaurus d’une manière générale donne clairement de meilleurs résultats que l’utilisation de simples mots clés. L’introduction de la connaissance du domaine dans le processus de recommandation est bénéfique et permet d’améliorer les performances du système même en utilisant des modèles de plongement contextuel bien réputés.

5 Conclusion

Dans cet article, nous avons proposé un système de recommandation des objets d’apprentissage en fonction des objectifs de l’apprenant en nous basant sur le calcul de leur distance de similarité. Nous avons étudié en particulier la représentation vectorielle des descriptions, et évalué les performances du système en utilisant trois modèles différents de plongement de phrases, et en étudiant la meilleure façon pour générer ces représentations vectorielles en nous basant sur le référentiel interne de *Teach on Mars*. Nos expériences montrent que la précision de la recommandation des objets d’apprentissage en utilisant le modèle MUSE et avec l’enrichissement contextuel des objectifs d’apprentissage avec des relations hiérarchiques est de loin la meilleure configuration par rapport aux modèles classiques de recherche de correspondance comme BM25.

Les perspectives de ce travail sont tout d’abord de déterminer la meilleure représentation vectorielle des objets d’apprentissage. Les expériences montrent que la moyenne des représentations vectorielles des objectifs d’apprentissage permet d’obtenir la meilleure précision P@N. A court terme nous allons évaluer ce type de représentation pour les objets d’apprentissage. Deux autres perspectives intéressantes sont d’étudier comment présenter ces recommandations à l’apprenant et de définir des parcours d’apprentissage personnalisés, en prenant en compte les niveaux de difficulté des objets d’apprentissage et le niveau et l’histo-

rique d'apprentissage de l'utilisateur.

Références

- [1] Eiman Aciad and Farid Meziane. An adaptable and personalised e-learning system applied to computer science programmes design. *Education and Information Technologies*, 24(2) :1485–1509, 2019.
- [2] Soulef Benhamdi, Abdesselam Babouri, and Raja Chiky. Personalized recommender system for e-learning environment. *Education and Information Technologies*, 22(4) :1455–1477, 2017.
- [3] Jesús Bobadilla, Fernando Ortega, Antonio Hernando, and Abraham Gutiérrez. Recommender systems survey. *Knowledge-based systems*, 46 :109–132, 2013.
- [4] Daniel Cer, Yinfei Yang, Sheng yi Kong, Nan Hua, Nicole Limtiaco, Rhomni St. John, Noah Constant, Mario Guajardo-Cespedes, Steve Yuan, Chris Tar, Yun-Hsuan Sung, Brian Strope, and Ray Kurzweil. Universal sentence encoder, 2018.
- [5] Jacob Devlin, Ming-Wei Chang, Kenton Lee, and Kristina Toutanova. Bert : Pre-training of deep bi-directional transformers for language understanding. *arXiv preprint arXiv :1810.04805*, 2018.
- [6] Manqing Dong, Feng Yuan, Lina Yao, Xianzhi Wang, Xiwei Xu, and Liming Zhu. Trust in recommender systems : A deep learning perspective. *arXiv preprint arXiv :2004.03774*, 2020.
- [7] Qingyu Guo, Fuzhen Zhuang, Chuan Qin, Hengshu Zhu, Xing Xie, Hui Xiong, and Qing He. A survey on knowledge graph-based recommender systems. *arXiv preprint arXiv :2003.00911*, 2020.
- [8] Deepani B Guruge, Rajan Kadel, and Sharly J Halder. The state of the art in methodologies of course recommender systems—a review of recent research. *Data*, 6(2) :18, 2021.
- [9] Mushtaq Hussain, Wenhao Zhu, Wu Zhang, Syed Muhammad Raza Abidi, and Sadaqat Ali. Using machine learning to predict student difficulties from learning session data. *Artificial Intelligence Review*, 52(1) :381–407, 2019.
- [10] Shristi Shakya Khanal, PWC Prasad, Abeer Alsaadon, and Angelika Maag. A systematic review : machine learning based recommendation systems for e-learning. *Education and Information Technologies*, 25(4) :2635–2664, 2020.
- [11] Sucheta V Kolekar, Radhika M Pai, and Manohara Pai MM. Rule based adaptive user interface for adaptive e-learning system. *Education and Information Technologies*, 24(1) :613–641, 2019.
- [12] Goutam Majumder, Partha Pakray, Alexander Gelbukh, and David Pinto. Semantic textual similarity methods, tools, and applications : A survey. *Computación y Sistemas*, 20(4) :647–665, 2016.
- [13] Deuk Hee Park, Hyea Kyeong Kim, Il Young Choi, and Jae Kyeong Kim. A literature review and classification of recommender systems research. *Expert systems with applications*, 39(11) :10059–10072, 2012.
- [14] Elvys Linhares Pontes, Stéphane Huet, Andréa Carneiro Linhares, and Juan-Manuel Torres-Moreno. Predicting the semantic textual similarity with siamese cnn and lstm. *arXiv preprint arXiv :1810.10641*, 2018.
- [15] Mohammad Mustaneer Rahman and Nor Aniza Abdullah. A personalized group-based recommendation approach for web search in e-learning. *IEEE Access*, 6 :34166–34178, 2018.
- [16] Nils Reimers and Iryna Gurevych. Sentence-bert : Sentence embeddings using siamese bert-networks. *arXiv preprint arXiv :1908.10084*, 2019.
- [17] John K Tarus, Zhendong Niu, and Ghulam Mustafa. Knowledge-based recommendation : a review of ontology-based recommender systems for e-learning. *Artificial intelligence review*, 50(1) :21–48, 2018.
- [18] María Cora Urdaneta-Ponte, Amaia Mendez-Zorrilla, and Ibon Oleagordia-Ruiz. Recommendation systems for education : Systematic review. *Electronics*, 10(14) :1611, 2021.
- [19] Dianhui Wang and Ming Li. Stochastic configuration networks : Fundamentals and algorithms. *IEEE transactions on cybernetics*, 47(10) :3466–3479, 2017.
- [20] Yinfei Yang, Daniel Cer, Amin Ahmad, Mandy Guo, Jax Law, Noah Constant, Gustavo Hernandez Abrego, Steve Yuan, Chris Tar, Yun-Hsuan Sung, et al. Multilingual universal sentence encoder for semantic retrieval. *arXiv preprint arXiv :1907.04307*, 2019.
- [21] Philip S Yu. Data mining and personalization technologies. In *Proceedings. 6th International Conference on Advanced Systems for Advanced Applications*, pages 6–13. IEEE, 1999.
- [22] Yuwen Zhou, Changqin Huang, Qintai Hu, Jia Zhu, and Yong Tang. Personalized learning full-path recommendation model based on lstm neural networks. *Information Sciences*, 444 :135–152, 2018.