

L'EXTRACTION TERMINOLOGIQUE OUTILLÉE À L'ÉPREUVE DE L'EXPÉRIMENTATION : PROPOSITION D'EXERCICE AVEC DES ÉTUDIANTS EN TRADUCTION

Experimenting with terminology extraction tools: suggestions for an exercise with translation students

Rudy LOOCK *UMR « Savoirs, Textes, Langage » -CNRS,
Université de Lille, France*

Résumé

Cet article revient sur un travail d'extraction terminologique effectué avec des étudiants en seconde année de master de traduction spécialisée, qui ont été placés en situation d'expérimentation afin d'exploiter une même base de données linguistiques sur une thématique spécialisée à l'aide de quatre outils différents : un logiciel de Traduction assistée par ordinateur (TAO), un concordancier, un outil d'extraction en ligne, et un *Large Language Model* (LLM) ou grand modèle de langage en français. Les étudiants ont alors dû fournir une analyse de ces résultats, l'objectif final étant la compilation d'un glossaire monolingue. Cet article propose un retour d'expérience sur ce travail qui vise à développer chez les étudiants un regard critique vis-à-vis des résultats générés par les outils, soit une compétence humaine et particulièrement nécessaire dans un contexte technologique en évolution constante.

Mots-clés : Terminologie – Extraction terminologique – Outil – Formation – Expérimentation

Abstract

This article reports on a terminology extraction project conducted with second-year master's students enrolled in a specialized translation program. Students engaged in terminological extraction tasks based on a specialized linguistic corpus, with four different tools: a computer-assisted translation (CAT) software, a concordancer, an online terminology extraction platform, and a large language model (LLM). Students were subsequently required to perform a comparative analysis of the outputs generated by each tool, with the ultimate objective of compiling a monolingual glossary. The article

offers a reflective evaluation of the project, which aimed to develop students' evaluative judgment on the results generated by the tools — a human competence increasingly essential in the evolving landscape of language technologies.

Keywords: Terminology – Terminological extraction – Tool – Training – Experimentation

INTRODUCTION

Les débats actuels autour de l'intelligence artificielle (IA), dans le monde de la traduction comme ailleurs, se cristallisent souvent autour de la place de l'humain vis-à-vis de ces nouvelles innovations technologiques : l'IA représente-t-elle une menace ? Faut-il intégrer ces outils ou posent-ils trop de problèmes, qu'ils soient techniques, éthiques, ou écologiques ? Comment faire en sorte que l'humain reste au centre de toutes ces innovations et qu'elles lui soient véritablement utiles et profitables ? S'agissant spécifiquement de l'IA générative, chaque jour voit apparaître de nouveaux outils ou de nouvelles fonctionnalités pour les outils existants. Il peut être très facile pour tout un chacun, les étudiants en particulier, non seulement de s'y perdre, mais de consacrer du temps à développer des compétences qui risquent de finir par devenir obsolètes.

Dans les formations en traduction, où l'on a le plus souvent intégré la traduction automatique neuronale (TAN) depuis longtemps maintenant, cela pose la question de savoir à quoi nous formons les étudiants quand il s'agit d'outils. Que privilégier ? La maîtrise des outils les plus fréquemment utilisés sur le marché pour améliorer l'employabilité ? L'accomplissement de tâches avec des outils différents pour en maîtriser un maximum ? Une approche globale avec des compétences transférables d'un outil à l'autre ? Comment garantir la maîtrise technique, mais aussi l'analyse nécessaire des résultats, avec suffisamment de recul et sans perdre les compétences en traduction qui représentent l'essentiel ? Face à une offre pléthorique et à une évolution constante, il nous a semblé intéressant de prendre du recul et de nous attarder sur le développement de compétences bien humaines : savoir choisir l'outil le plus pertinent pour effectuer une tâche donnée et analyser de façon critique les résultats générés.

C'est dans ce cadre que nous proposons dans cet article de revenir sur un travail effectué avec des étudiants en seconde année de master de traduction spécialisée, travail qui les place dans une situation d'expérimentation vis-à-vis d'outils différents pour accomplir une seule et même tâche, à savoir une extraction terminologique, à l'aide (i) d'un logiciel de Traduction assistée par ordinateur (TAO), (ii) d'un concordancier, (iii) d'un outil d'extraction en ligne, et (iv) d'un *Large Language Model* (LLM) ou grand modèle de langage en français.

1. POINTS DE DÉPART

1.1. Importance de la terminologie : compétences et enjeux

Les compétences terminologiques font partie des compétences à développer chez les traducteurs et traductrices en formation. Comme le rappellent Grossmann & Lavault (2000, p. 155), à partir du moment où l'on traduit, « [c]omme M. Jourdain faisait de la prose, tout le monde fait de la terminologie sans le savoir ». Le référentiel de compétences mis en place par le réseau

European Master's in Translation (EMT) de la Commission européenne (EMT, 2022 pour la dernière version en date) mentionne explicitement la terminologie dans 2 compétences des 36 qui y sont listées (compétences 4 et 14), tandis que d'autres y sont indirectement associées, comme la recherche documentaire. Nombreuses sont donc les formations universitaires qui proposent un ou des enseignement(s) de terminologie, où les outils informatiques occupent une place de choix, qu'il s'agisse de logiciels de bureautique, d'outils intégrés dans les logiciels de Traduction assistée par ordinateur (TAO), ou encore d'outils de corpus (voir section suivante).

Les enjeux relatifs à l'enseignement de la terminologie sont multiples. Afin de compiler un glossaire, les étudiants doivent en effet apprendre à définir et à repérer ce qu'est un terme pertinent, trouver un équivalent correspondant dans une autre langue dans le cas d'un travail bilingue, et intégrer des informations qui seront utiles au moment de la traduction (p. ex. définition, exemple en contexte, informations complémentaires relatives à l'usage du terme). Passer de l'« exercice de terminologie » à la création d'une ressource sous forme de glossaire exploitable au moment de la traduction représente un enjeu essentiel et n'est pas sans difficultés (Grossmann & Lavault, 2000, p. 166-169). Il peut alors être utile de faire le lien entre enseignement de terminologie et enseignement pratique de traduction en mettant en place un décroisement pédagogique (Loock & Moulard, 2025), qui permet aux étudiants d'exploiter une ressource créée dans le cadre d'un enseignement pour traduire un texte spécialisé dans un autre. Peut également être pertinent la mise en place de collaborations avec des donneurs d'ordre extérieurs, par exemple avec l'Organisation mondiale de la propriété intellectuelle (OMPI) (Frérot, 2018), ou encore la création de ressources en ligne, comme les plateformes ARTES (Aide à la rédaction de textes scientifiques, Pecman & Kübler, 2011) ou UCOTerm (Veroz-González *et al.*, 2019). Enfin, des exercices de simulation via des agences virtuelles de traduction ou *skills lab* (van Egdom *et al.*, 2020), où le glossaire doit servir de ressource à une équipe de traducteurs différente de celle qui aura effectué le travail terminologique est également un dispositif pédagogique intéressant.

Ce qu'ont en commun toutes ces approches, c'est de faire sortir les étudiants de l'exercice de terminologie et de les amener à développer une réflexion d'anticipation vers l'utilisation d'une ressource, soit une compétence pertinente dans le cadre de la professionnalisation des futurs diplômés.

1.2. Une approche outillée

La recherche terminologique et la compilation de glossaires sont aujourd'hui essentiellement outillées, avec pour objectif final la création de ressources exploitables au moment de la traduction et intégrables à l'environnement de travail. Au-delà de la maîtrise des concepts terminologiques, il est donc crucial pour les étudiants d'apprendre à utiliser des outils d'extraction terminologique, qui, comme tous les autres outils d'aide à la traduction, foisonnent à l'heure actuelle (voir Rothwell *et al.*, 2023 pour un inventaire global récent et

Rothwell *et al.*, 2025 pour un panorama des outils utilisés dans les formations universitaires du réseau EMT). Les outils informatiques sont ainsi aujourd'hui omniprésents dans la formation aux métiers de la traduction, en écho à la réalité du terrain, et la recherche terminologique est suffisamment concernée pour que l'on puisse parler aujourd'hui de « terminologie numérique » (Vezzani, 2022, p. 34).

Des choix doivent être faits quant à l'utilisation des outils en cours avec les étudiants, sur des critères pédagogiques mais aussi parfois financiers. Il est impossible de former à tous les outils existants, et à l'inverse se concentrer sur un seul fait courir le risque que les compétences acquises ne soient pas transférables à d'autres outils. On retrouve là l'opposition classique entre apprendre à utiliser un outil pour effectuer une tâche et apprendre à effectuer une tâche en utilisant un outil, ce qui n'est pas la même chose : dans le premier cas, l'accent est mis sur la maîtrise technique de l'outil afin d'obtenir un résultat tandis que dans le second on se concentre sur la tâche à réaliser, et l'outil n'est qu'un moyen pour y parvenir, avec concentration sur l'évaluation de la pertinence du résultat.

Parmi les outils existants, on trouve des outils intégrés dans les logiciels de TAO (Traduction assistée par ordinateur) mais aussi le tableur Excel de la suite Microsoft Office. Le recours à des bases de données linguistiques (corpus) à exploiter via des concordanciers hors ligne ou des plateformes en ligne est une autre possibilité, parfois perçue (à tort) comme une approche académique (voir p. ex. Bowker, 2011 ; Kübler, 2003 ; Looock, 2016). Enfin, les nouveaux outils d'intelligence artificielle générative qui ont fait leur apparition pour le grand public en 2022 peuvent également être exploités comme outils d'extraction terminologique. Selon la dernière enquête ELIS (*European Language Industry Survey*), les outils les plus utilisés pour la gestion terminologique sont Multiterm et Excel, mais on constate une montée en puissance des nouveaux outils d'IA générative (ELIS 2025, p. 34, 38) tandis que l'exploitation de corpus via un concordancier reste limitée. Nous passons en revue tour à tour ces trois types d'outils.

1.2.1. Les outils associés à la TAO

Le lien entre terminologie et Traduction assistée par ordinateur (TAO) est naturel, puisque dans ce type d'environnement régulièrement utilisé en traduction spécialisée, il importe d'identifier et de gérer les termes spécifiques d'un domaine pour assurer la cohérence et la précision des traductions. Au-delà des mémoires de traduction, plusieurs outils facilitent ce processus, notamment MultiTerm, Xbench et MemoQ LiveDocs. MultiTerm¹, développé par RWS, est un logiciel dédié à la gestion terminologique intégré dans Trados qui permet de créer, d'organiser et de consulter des bases de données terminologiques multilingues. Xbench², quant à lui, est

1. <https://www.trados.com/fr/product/multiterm/>

2. <https://www.xbench.net/>

un outil polyvalent conçu pour le contrôle qualité en traduction, qui offre des fonctionnalités d'extraction terminologique en analysant les segments traduits pour identifier les termes potentiellement incohérents ou incorrects. L'outil LiveDocs³ intégré au logiciel de TAO memoQ, que nous retiendrons ici pour notre expérimentation, permet également d'effectuer des extractions de termes pertinents à partir de documents existants. Un éditeur spécifique permet d'examiner les candidats termes retenus pour filtrage et validation avant l'intégration à une base terminologique.

En complément, les professionnels ont également souvent recours à l'outil Microsoft Excel pour des raisons de flexibilité : les glossaires compilés peuvent en effet être importés dans différents outils de TAO. Par exemple, avec MultiTerm Convert, il est possible de transformer des glossaires existants sous format Excel en bases de données compatibles avec MultiTerm, facilitant ainsi leur intégration dans les projets de traduction.

1.2.2. L'extraction par concordancier

Il est possible d'effectuer des extractions terminologiques à partir d'un ensemble de données linguistiques (corpus) à exploiter à l'aide d'un concordancier, que celui-ci soit hors ligne ou intégré dans une interface en ligne. À partir d'un ensemble de documents rassemblés pour la circonstance, sur une thématique bien définie et spécialisée (on parle alors de corpus DIY pour « *Do-It-Yourself* », voir p. ex. Sánchez-Gijón, 2009 ou Loock, 2016), il est possible de compiler une base de données dont l'exploitation permettra (i) d'affiner la compréhension du texte source, (ii) d'améliorer la qualité du texte cible en consultant les usages linguistiques au sein de textes rédigés par des experts d'un domaine, ou encore (iii) de compiler des glossaires à partir d'extractions terminologiques. Cette exploitation permet d'extraire et d'afficher toutes les occurrences d'un mot ou d'une expression dans un corpus de textes, d'observer leur fréquence et leur environnement linguistique, ou encore de dresser des listes de mots. Parmi les plus connus, nous pouvons citer AntConc, Wordsmith Tools ou LancsBox⁴. Il est également possible d'exploiter ses propres corpus via des sites internet comme le portail Sketch Engine,⁵ qui regroupe plus de 800 corpus pour une centaine de langues mais permet également de compiler ses propres corpus. Un outil associé est le site internet OneClick Terms⁶, un extracteur terminologique en ligne à partir de documents monolingues ou bilingues. Les outils de corpus que nous retiendrons ici sont AntConc (Anthony, 2023) et OneClick Terms.

3. <https://www.memoq.com/tools/livedocs/>

4. Ces concordanciers sont disponibles aux adresses suivantes :

<https://www.laurenceanthony.net/software/antconc/>

<http://corpora.lancs.ac.uk/lancsbox/>, www.lexically.net/wordsmith/.

5. <https://www.sketchengine.eu/>

6. <https://terms.sketchengine.eu/>

1.2.3. Les Large Language Models (LLM)

Récemment, avec l'avènement de nouveaux outils d'intelligence artificielle générative, à savoir les LLM (*Large Language Models*) avec ChatGPT comme figure de proue, un nouvel outil d'extraction terminologique semble avoir fait son apparition. Selon la dernière enquête ELIS (ELIS, 2025), les outils d'IA générative sont en effet régulièrement utilisés pour des tâches en lien avec la terminologie, mais essentiellement au sein des agences de traduction et des départements de langues : entre 30 % des départements et 40 % des agences utilisent l'IA pour l'extraction terminologique, à quoi il faut ajouter 60 % des agences et 10 % des départements pour la recherche terminologique. Il s'agit par exemple, à partir de *prompts* (instructions génératives) plus ou moins complexes, d'interroger le LLM pour qu'il génère une liste de termes à partir d'un corpus de textes, d'un site internet, ou encore d'un domaine spécialisé. Le LLM que nous retiendrons ici est ChatGPT-4, le travail ayant été mené fin 2024.

1.3. L'expérimentation comme objectif

À notre connaissance, la comparaison entre les résultats fournis par différents outils d'extraction terminologique a fait l'objet d'assez peu d'études ces dernières années, en particulier dans le cadre de la formation des futurs traducteurs (voir néanmoins Inkpen *et al.* 2016 ; Costa *et al.* 2014 ; Costa *et al.*, 2016 ; ou encore Muegge, 2023, qui propose une comparaison intéressante d'outils pour l'extraction terminologique, y compris ChatGPT). Or, dans une période marquée par la prolifération d'outils et les développements technologiques rapides, c'est justement cette comparaison qui est pertinente pour les étudiants, qui auront à faire leurs propres choix lorsqu'ils seront en activité. Il est également important qu'ils comprennent que tous les outils ne fournissent pas les mêmes résultats, et que le contrôle qualité, la prise de recul vis-à-vis de ceux-ci, représentent des étapes importantes si l'étudiant souhaite développer une approche critique et non une simple consommation « passive ». Il nous a donc semblé important de placer nos étudiants dans la position d'expérimentateurs afin qu'ils puissent s'exercer au développement de ce regard critique.

Le travail terminologique requiert dans un premier temps une sélection de données de travail, à partir desquelles les termes candidats et toutes les informations les concernant feront l'objet d'une sélection. Il s'agit alors de développer une *data literacy*, soit une « capacité à collecter, gérer, évaluer, et exploiter des données de façon critique » (Ridsdale *et al.*, 2015, p. 11, cité par Krüger & Hackenbuchner, 2022, p. 392, nous traduisons). Nous avons déjà ailleurs (Loock, 2025) avancé qu'une telle formation aux données était nécessaire pour la compilation et l'exploitation de bases de données linguistiques (corpus) ; à partir du moment où l'extraction terminologique se fait sur des données *in vivo*, il semble naturel d'affirmer que le développement de compétences terminologiques passe par une capacité à sélectionner et évaluer ces données. Ceci nous semble revêtir une importance particulière à

une époque où les textes qui alimentent le web sont de plus en plus générés ou traduits de façon automatique, et par conséquent de qualité hautement discutable (Thompson *et al.*, 2024 ; Bevendorff *et al.*, 2024). Dans un second temps vient le choix des outils, et là encore, l'esprit critique jouera un rôle important puisqu'il est nécessaire d'en sélectionner et d'évaluer les résultats obtenus. Qu'il s'agisse des données ou des outils, faire ses propres choix, par essence subjectifs, comprendre la variabilité des résultats que l'on peut obtenir et savoir les analyser, nous semblent nécessiter un recours à l'expérimentation par les étudiants eux-mêmes.

Nous proposons donc ici un retour d'expérience sur un travail effectué avec des étudiants inscrits en seconde année d'une formation en traduction spécialisée, en l'occurrence la formation de master « Traduction spécialisée multilingue » (TSM) de l'Université de Lille en France, au cours de l'année universitaire 2024-2025, et qui place les apprenants dans une telle situation d'expérimentation. Dans le cadre d'un cours consacré aux outils de corpus, nous avons demandé aux étudiants (n=23) d'effectuer une tâche d'extraction terminologique avec différents outils (n=4) à partir d'une même base de données linguistiques monolingues (français) compilée par leurs propres soins et portant sur une thématique spécialisée de leur choix. Il leur a été demandé d'observer et d'analyser les résultats fournis par les différents outils et de fournir un rapport détaillé critique sur la qualité des extractions obtenues, en prenant en compte les points forts et les faiblesses de chacune d'entre elles. L'objectif final reste la création d'une ressource terminologique, à savoir un glossaire monolingue. Nous présentons ce travail de façon détaillée dans la section suivante, avant d'en discuter les résultats dans la troisième section.

2. PRÉSENTATION DU TRAVAIL D'EXTRACTION TERMINOLOGIQUE

Dans le cadre d'un cours consacré aux outils de corpus qui vise à développer des compétences de compilation et d'exploitation de bases de données linguistiques comme outils d'aide à la traduction et de recherche en traductologie, l'évaluation finale a pris la forme d'un travail de recherche expérimentale en terminologie.⁷

Les consignes générales, données en octobre 2024, ont été les suivantes ; nous les détaillons ci-après :

1. Compiler un **corpus spécialisé monolingue en langue cible (FR)** sur la thématique de votre choix.
2. Opérer une **extraction terminologique** à l'aide des **4 outils** suivants :
(i) AntConc, (ii) OneClick Terms, (iii) memoQ, (iv) ChatGPT.
3. **Comparer et analyser de façon critique** les résultats obtenus.
4. Livrer un **glossaire monolingue de 25 termes**.

7. Ce travail a représenté 75 % de la note finale pour cet enseignement, les 25 % restants relevant d'un exercice de métacognition (réflexion par les étudiants sur leur propre apprentissage) comme décrit dans Loock (2025).

Les critères d'évaluation, annoncés en amont, ont été les suivants : (i) description et justification du corpus choisi : thématique, choix des textes, compilation (organisation, propriété des fichiers) ; réalisation des extractions terminologiques avec les différents outils ; analyse critique ; glossaire de 25 termes.

2.1. Choix de la thématique spécialisée

Il a été demandé aux étudiants de compiler dans un premier temps une base de données linguistiques en français (langue cible dans leur cas) sur une thématique suffisamment spécialisée pour qu'une recherche terminologique soit nécessaire et pertinente. Le nombre minimal de tokens attendu était de 10 000 mots, à partir d'au moins 5 sources différentes. Le Tableau 1 dresse la liste des thématiques retenues ainsi que le nombre de tokens des corpus respectifs.

Thématique choisie	Nombre de tokens
La torture moderne	55 377
Les conditions carcérales en France	20 646
L'ischémie-reperfusion	12 939
Le dressage équestre	17 260
La gestion des déchets électriques et électroniques	32 319
Les batteries à électrolytes solides pour véhicules électriques	11 895
Le peptide relié au gène de la calcitonine	11 478
Les matériaux biosourcés	13 800
L'oncologie pédiatrique	25 472
La faune endémique française	14 803
Le dépistage de la trisomie 21 pendant la grossesse	28 014
Les bioréacteurs à membrane	28 357
Le concept d'intelligence en psychologie	55 944
Les cardiopathies congénitales	16 524
Le diabète	43 882
Les memecoins	27 442
Les outils malveillants des APT (attaques persistantes avancées)	19 303
La chimie organique	18 205
La biomécanique sportive	13 083
L'aérodynamisme en Formule 1	84 266
Les trous noirs	14 346
Le diagnostic précoce de la maladie d'Alzheimer	106 384
Les perturbateurs endocriniens	20 784

Tableau 1. Liste des thématiques retenues par les étudiants et nombre de tokens

Les étudiants étaient libres de choisir leur thématique de travail, qui pouvait être en lien avec un projet de traduction qu'ils avaient à effectuer par ailleurs. On note des thématiques plus ou moins spécialisées, ce qui ne peut pas être sans incidence sur les résultats obtenus (voir discussion des résultats). Une réflexion sur la représentativité de leur corpus était attendue.

Ce type d'extraction à partir d'un corpus permet une observation *in vivo*, pour une recherche terminologique pragmatique (centrée sur l'usage), discursive (ancrée dans le texte), et cognitive (Gauthier, 2020). L'utilisateur a alors accès au contexte, puisque les termes identifiés apparaissent au sein de ce que l'on appelle des « segments riches en connaissances » (Planas *et al.*, 2014), ce qui pourra être utile au moment de renseigner certaines rubriques du glossaire demandé.

2.2. Outils utilisés

À partir de ce corpus DIY (*Do-It-Yourself*), il leur a été demandé d'opérer quatre extractions terminologiques distinctes, avec quatre outils différents et selon les consignes suivantes :

- **AntConC** : compiler votre liste de termes en appliquant une *stoplist*⁸ et le corpus de référence⁹ en français fourni. Fournir un tableau avec les résultats générés par l'outil WordList (50 premiers) ainsi que par l'outil n-gram pour les groupes de 2, 3 et 4 mots. Discuter la validité des résultats ainsi obtenus.
- **OneClick Terms** : compiler votre liste (*single-words* et *multi-words*). Fournir un tableau avec les résultats fournis par l'outil (50 premiers). Discuter la validité des résultats ainsi obtenus.
- **memoQ** : à l'aide de l'outil LiveDocs, opérer une extraction terminologique en appliquant la *stoplist* FR fournie par le logiciel. Fournir un tableau avec les résultats fournis par l'outil (50 premiers). Discuter la validité des résultats ainsi obtenus.
- **ChatGPT-4** : compiler votre liste à partir des prompts de votre choix (préciser les prompts utilisés, n'hésitez pas à expérimenter) et fournir un tableau avec la liste de 50 termes proposés par l'outil. Discuter la validité des résultats ainsi obtenus.

8. Une *stoplist* est une liste de termes à exclure lors d'une recherche afin que ceux-ci n'apparaissent pas dans les résultats. De façon typique, une *stoplist* contient de nombreux mots grammaticaux (articles, prépositions, certains adverbes...) qui ne sont pas pertinents pour la compilation de glossaires. La *stoplist* fournie aux étudiants pour le français a été téléchargée sur <https://countwordsfree.com/stopwords> ; pour le logiciel memoQ, en revanche, c'est la *stoplist* intégrée au logiciel qui a été utilisée.

9. Un corpus de référence est un corpus qui permet une comparaison avec le corpus à l'étude. Dans le cas présent, il s'est agi d'un corpus regroupant des textes de presse générale en français original, en l'occurrence le sous-corpus français du *TSM press corpus* (Loock 2019), qui contient plus de 2,5 millions de mots. Une telle comparaison avec des textes issus de la presse généraliste doit permettre de dégager des termes spécifiques au domaine.

Pour des raisons éthiques, et conformément à la *Charte des bonnes pratiques d'utilisation des intelligences artificielles génératives dans les travaux pédagogiques* mise en place par l'Université de Lille, les outils utilisés sont tous accessibles gratuitement pour les étudiants : si AntConc, OneClick Terms et ChatGPT-4 le sont pour tous les utilisateurs, le logiciel de TAO memoQ et son outil intégré LiveDocs le sont pour nos étudiants dans le cadre d'une collaboration avec la société qui développe le logiciel.

2.3. Livrables

Les étudiants étaient tenus de livrer un rapport analytique avec les résultats obtenus et un bilan des avantages et inconvénients de chaque outil, en analysant de façon critique les extractions générées (quel outil a fourni les informations les plus précises ? avec quelle efficacité et rapidité ? quid de leur accessibilité et de leur ergonomie ?). Au-delà du rapport, un glossaire monolingue était attendu, avec les 25 termes (composés d'un ou plusieurs mots) les plus pertinents pour la thématique choisie, et en renseignant les éléments traditionnels suivants : terme, catégorie grammaticale, définition, exemple en contexte, informations collocationnelles, informations complémentaires (variantes, remarques relatives à l'usage). Ce glossaire devait prendre la forme d'un tableau Excel à joindre au rapport. Par manque de place, nous n'analyserons pas ici les glossaires livrés, mais nous concentrerons sur l'analyse critique des extractions terminologiques. Le corpus de travail au format .db (*database*) était également attendu pour effectuer les vérifications nécessaires.

3. RÉSULTATS ET DISCUSSION

3.1. Utilisation des outils et extractions terminologiques

La compilation de la base de données linguistiques monolingues (corpus DIY) ne représentait pas un enjeu particulier, cette compétence ayant été développée en première année de master. S'agissant des outils, les étudiants ont rencontré assez peu de difficultés techniques pour effectuer leurs extractions terminologiques et ont été en mesure de fournir les listes de termes attendues. Là encore, les étudiants avaient été formés l'année précédente à l'utilisation de LiveDocs dans memoQ et du concordancier AntConc ; l'outil OneClick terms ne requerrait pas quant à lui de compétences spécifiques. L'enjeu de ce travail n'était de toute façon pas prioritairement technique. Les étudiants ont toutefois rapporté quelques difficultés lors de l'utilisation de memoQ et ont estimé qu'il n'était pas l'outil le plus ergonomique (voir ci-dessous) ; la qualité de la *stoplist*, notamment, a été fortement critiquée puisque son utilisation limite insuffisamment le bruit présent dans la liste de termes. Quelques soucis d'utilisation de ChatGPT au moment de soumettre le corpus en lien avec le *prompt* (instruction générative) ont été signalés (types de fichiers non pris en charge, ou fichiers trop nombreux), mais les étudiants ont su contourner la difficulté, en fusionnant ou en convertissant leurs fichiers.

On note toutefois deux soucis liés à une insuffisante maîtrise des outils que ce travail a mis au jour. Premièrement, si les étudiants ont bien intégré, en plus d'une *stoplist*, un corpus de référence (*reference corpus*) au sein du concordancier AntConc afin d'affiner la liste de termes générée en mettant en avant les termes spécifiques au corpus de travail (*target corpus*), la majorité des étudiants a retenu les résultats de l'onglet « *Word* » et non « *Keyword* », et a donc travaillé sur une liste de termes qui n'avait pas été affinée (ce qu'ils n'ont d'ailleurs pas manqué de faire remarquer mais sans comprendre l'origine du problème). Cette erreur assez généralisée indique que la manipulation n'a pas été comprise et que la formation à l'outil n'a pas atteint tous ses objectifs. Les Figures 1 et 2 illustrent la différence de résultats entre les deux listes générées pour le corpus relatif au dépistage de la trisomie 21 pendant la grossesse, onglet « *Word* » (Figure 1) et onglet « *Keyword* » (Figure 2) : dans le second cas, des termes comme *trimestre*, *taux*, *âge* ou *cas* n'ont pas été intégrés à la liste.

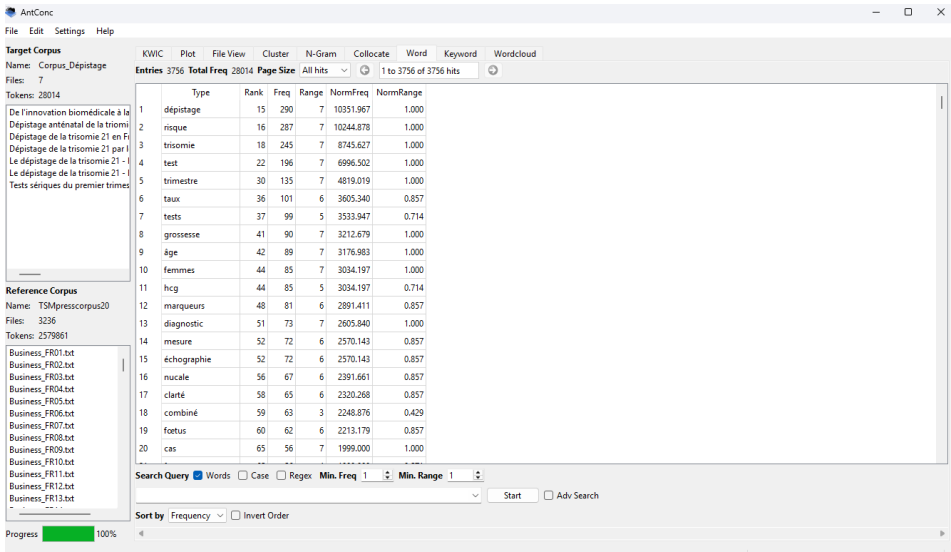


Figure 1. Extraction à l'aide du concordancier AntConc via l'onglet « *Word* » (sans prise en compte du corpus de référence)

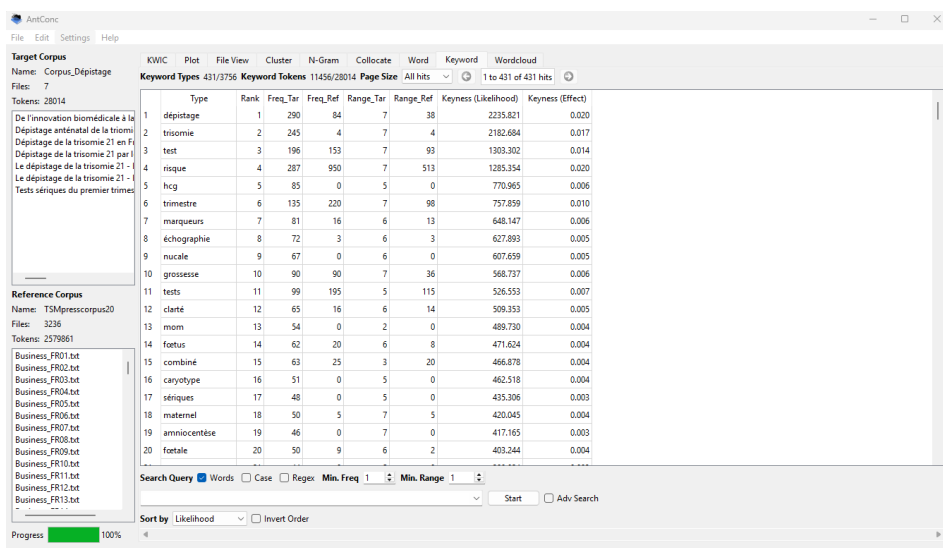


Figure 2. Extraction à l'aide du concordancier AntConc via l'onglet « Keyword » (avec prise en compte du corpus de référence)

Deuxièmement, et bien qu'il s'agisse là de l'outil le plus récent, on constate une expérimentation qui reste limitée avec l'outil ChatGPT. Rares sont en effet les étudiants qui ont testé différents types de *prompts* afin de voir si cela avait une influence sur les résultats ; bien souvent ils se sont contentés de la première liste générée par l'outil à partir de *prompts* basiques (*zero-prompts*) comme en (1) alors que des *prompts* plus élaborés (2), où l'on spécifie le rôle que doit endosser l'outil, la tâche précise, et le format de sortie, ou bien où l'on affine la requête grâce à une série de *prompts* supplémentaires ont permis de générer des résultats plus pertinents.

(1) Avec les 5 articles que je te mets ci-joint, peux-tu me faire une extraction terminologique de 50 mots qui pourront m'aider à créer un glossaire monolingue sur le diabète pour mon devoir, merci.

(2) Bonjour ChatGPT, aujourd'hui j'ai besoin de toi en tant que professionnel de la gestion terminologique. Tu dois créer une base terminologique sous forme de tableau Excel avec une colonne pour le mot et une colonne pour la fréquence du mot, la thématique est le peptide associé au gène de la calcitonine. Tu dois choisir les 50 termes les plus utilisés dans les 7 fichiers sources ci-joints dans le dossier « Corpus-CGRP.zip ». Attention, il ne faut pas utiliser les mots présents dans le document « stop_words_french.txt ».

3.2. Regard critique sur les résultats obtenus

Les étudiants ont généralement fait preuve d'un bon esprit critique au moment d'analyser les résultats générés par les différents outils : ils ont évalué la pertinence des termes retenus, la présence de bruit, le rôle de la *stoplist*, la

possibilité de générer des groupes de mots (n-grams), ou encore la facilité d'utilisation de l'outil (ergonomie). Des considérations professionnelles comme le coût ou l'accessibilité ont également été évoquées. Leurs conclusions sont au final assez homogènes. L'outil de TAO memoQ est considéré comme l'outil fournissant les moins bons résultats, avec beaucoup de termes non pertinents et une ergonomie améliorable. Par exemple, pour la thématique du peptide associé au gène de la calcitonine, l'outil a proposé des termes comme *IC à 95 %*, *mois*, *à long terme*, *2 à 8 crises*, soit des mots ou des séries de mots qui ne représentent pas des termes pertinents.

L'analyse pour les trois autres outils varie en fonction des étudiants, mais une majorité se dégage pour dire que OneClick Terms est le plus facile d'utilisation, bien que les listes de termes composés d'un seul mot soient de qualité très variable (les listes de termes composés de plusieurs mots donnent en revanche de meilleurs résultats). Ainsi, toujours pour la thématique du peptide associé au gène de la calcitonine, l'outil a proposé des termes candidats pertinents comme *rimégépant*, *prophylaxie*, *triptans*, *prophylactique* pour les termes simples et *récepteur d'entrée*, *arsenal thérapeutique*, *anticorps monoclonaux anti-cgrp* pour les termes complexes mais aussi des termes non spécifiques tels que *mois*, *traitement*, *jours*, *3 mois*, *nombre de jours*... Le concordancier AntConc est plutôt plébiscité pour ses résultats, mais critiqué pour sa prise en main parfois difficile et un format d'exportation peu exploitable pour une utilisation dans un cadre professionnel (simple liste au format .txt). Quant à ChatGPT, si les étudiants louent la rapidité de l'outil et sa capacité à générer des résultats facilement exploitables (moyennant des *prompts* pertinents, les résultats peuvent être générés sous forme de tableau où les termes sont classés par ordre alphabétique, voir Tableau 1 par exemple sur la thématique des conditions carcérales), ils s'interrogent néanmoins sur sa fiabilité. Des hallucinations ont en effet été repérées, comme la présence dans la liste générée de termes qui ne figurent *pas* dans le corpus de travail. Des réponses affichées par l'outil comme « Ce tableau repose sur une estimation basée sur les occurrences *probables* dans ton corpus » (nos italiques) ont pu semer le doute sur la précision des résultats, amenant les étudiants à faire preuve de méfiance sur les résultats fournis.

1	Surpopulation carcérale	26	Convention européenne des droits de l'homme
2	Conditions de détention indignes	27	Responsabilité de l'État
3	Traitements inhumains ou dégradants	28	Rapport sur les prisons
4	Règles pénitentiaires européennes	29	Conditions sanitaires dégradées
5	Droits fondamentaux	30	Manque d'intimité en cellule

6	Prisons françaises	31	Isolation prolongée
7	Administration pénitentiaire	32	Vétusté des locaux
8	Matelas au sol	33	Infrastructures pénitentiaires vétustes
9	Taux d'occupation carcérale	34	Stigmatisation des détenus
10	Prohibition des traitements inhumains	35	Mouvements de protestation carcérale
11	Capacité carcérale opérationnelle	36	Infestation parasitaire carcérale
12	Insalubrité carcérale	37	Violence en milieu carcéral
13	Contrôleur général des lieux de privation de liberté	38	Sécurité en détention
14	Hygiène pénitentiaire	39	Manque de personnel pénitentiaire
15	Conditions matérielles de détention	40	Rapport annuel du CGLPL
16	Comité européen de prévention de la torture	41	Recours contre les conditions de détention
17	Réinsertion sociale	42	Risques sanitaires en prison
18	Dignité humaine	43	Mauvaises pratiques carcérales
19	Normes minimales de détention	44	Détenus vulnérables
20	Rapport parlementaire	45	Protection juridique des détenus
21	Instruments de soft law	46	Principes de prévention des abus
22	Recours juridictionnels	47	Encellement individuel
23	Défaillances systémiques	48	Recommandations du CPT
24	Mécanismes de reconnaissance mutuelle	49	Alternatives à l'incarcération
25	Sanctions pénales	50	Réformes pénitentiaires

Tableau 1. Extraction terminologique à l'aide de ChatGPT à partir d'un corpus sur les conditions carcérales

L'analyse critique comparative a parfois été présentée de façon très synthétique sous forme de tableau comparatif en fonction de leurs avantages et limites, ou bien en fonction de critères (rapidité, ergonomie, précision des résultats), ce qui permet une lecture facile et rapide des résultats. Quelques étudiants ont même proposé un score afin de classer les outils. Le Tableau 2 liste les avantages et limites des quatre outils tels que fréquemment listés par les étudiants.

Outil	Avantages	Inconvénients
AntConc	- Outil personnalisable (<i>stoplist</i>, corpus de référence), qui permet d'affiner l'extraction - L'outil n-gram permet d'identifier des concepts clefs - Gratuité	- Bruit (termes pas toujours assez spécialisés, p. ex. <i>mois, cas, également, système</i>)¹⁰ - Absence de lemmatisation¹¹ - Nécessite une certaine prise en main - Exportation peu pratique
OneClick Terms	- Rapide et ergonomique - Gratuité - Listes et statistiques claires - L'outil <i>multi-words</i> permet d'identifier des concepts clefs	- Mots uniques pas toujours pertinents - Présence de beaucoup de noms propres inutiles
memoQ	- Liste unique pour les termes composés d'un ou plusieurs mots - Intégration dans le logiciel de TAO (lien direct avec le projet de traduction)	- Peu de résultats exploitables malgré la <i>stoplist</i> - Processus long et peu ergonomique (nécessite la création d'un projet) - Nécessite une licence

10. Voir ci-dessus : cette limite mentionnée est due à une mauvaise exploitation.

11. Bien que ceci soit possible, il n'avait pas été demandé aux étudiants de lemmatiser leur corpus, à savoir de permettre le regroupement des différentes formes fléchies d'un mot (p. ex. *important/importante/importants/importantes*), ceci étant incompatible avec l'utilisation d'une *stoplist*.

ChatGPT	- Intuitif - Rapide et puissant - Peut fournir des listes sous forme de tableaux	- Nécessité de s'y reprendre à plusieurs fois - Résultats variables en fonction du <i>prompt</i> - Incertitude vis-à-vis des résultats - Hallucinations avérées
---------	--	--

Tableau 2. Avantages et inconvénients des quatre outils testés par les étudiants.

Au final, les analyses proposées par les étudiants ont été pertinentes, même si les outils n'ont pas toujours été exploités de façon optimale, notamment pour exclure des termes trop génériques et constituer ainsi un glossaire pertinent. On constate d'ailleurs un manque de projection vers le glossaire, la création d'une ressource terminologique demeurant l'objectif final du travail : quasiment aucun étudiant ne mentionne l'exploitation des résultats pour renseigner le glossaire demandé.

CONCLUSION

L'objectif de cet article était de présenter un travail demandé à des étudiants en fin de parcours universitaire afin de les placer en situation d'expérimentation pour effectuer une tâche d'extraction terminologique à l'aide de différents outils et créer un glossaire. Il s'agissait de les amener à développer un esprit critique vis-à-vis des résultats générés. L'exercice a été plutôt réussi, même si des limites sont apparues, dans la maîtrise de l'outil comme dans la capacité des étudiants à affiner leur façon d'interroger les outils.

S'agissant d'une première expérience, l'exercice est sans aucun doute perfectible, et il ne s'agissait ici que de partager cette initiative qui nous a semblé répondre à certains des défis qui attendent les futurs diplômés en traduction :

- Évoluer au sein d'un environnement hautement informatisé avec de nombreux outils à disposition ;
- S'adapter aux évolutions en adoptant une démarche expérimentale ;
- Porter un regard critique sur des résultats générés automatiquement ;
- Savoir verbaliser ce regard critique selon des critères bien définis (dans le cadre d'un travail de consultance par exemple).

Dans une période où se pose de façon cruciale la question de ce qui doit être enseigné à de futurs spécialistes des métiers de la traduction dans un contexte de plus en plus informatisé et caractérisé par une forme d'incertitude autour de la façon dont l'activité traduisante se réalisera à l'avenir, il nous a paru important de placer les étudiants dans le rôle d'expérimentateurs. L'analyse critique est, et restera, une compétence avant tout humaine, et il est donc crucial de la développer en complément des compétences linguistiques, les deux étant naturellement liées : ce n'est que parce que l'on maîtrise suffisamment bien une langue de départ et une langue d'arrivée dans toutes

leurs subtilités que l'on peut analyser de façon experte les résultats fournis par des outils informatiques et décider si oui ou non ils sont pertinents et exploitables. Il s'agit alors d'adopter une forme de « co-intelligence » telle que définie par Ethan Mollick, spécialiste de l'impact des technologies d'intelligence artificielle, en particulier les grands modèles de langage comme ChatGPT, sur l'éducation et le travail. Selon lui, dans son ouvrage *Co-Intelligence : Living and Working with AI* publié en 2024, il importe de mettre en place une synergie entre l'intelligence humaine et les outils dans une dynamique de collaboration. Mollick invite ainsi à respecter quatre principes, parmi lesquels utiliser l'IA de façon systématique (« *Always invite AI to the table* ») pour identifier comment elle peut nous être utile, tout en restant au centre (« *Be the human in the loop* ») et en partant du principe que l'outil d'IA que l'on est en train d'utiliser est le pire disponible (« *Assume this is the worst AI you will ever use* »). Le travail proposé avec les étudiants respecte ces trois principes (nous n'avons pas mentionné le quatrième proposé par Mollick, « *Treat AI like a person* », qui mériterait une discussion approfondie qui n'est pas pertinente ici) puisqu'il invite les étudiants à expérimenter les outils à leur disposition tout en se plaçant aux commandes, dans la façon d'interroger l'outil, qui peut varier, et dans la validation ou non des résultats, et en ayant conscience de leur instabilité, du fait d'évolutions à venir.

Remerciements

Je tiens à remercier les étudiant(e)s de seconde année du master « Traduction spécialisée multilingue », en particulier Noëlie Collin, Joane Bourat, Étienne Delarche, et Carla Mariani, dont le travail est cité dans l'article. Merci également à Nathalie Moulard pour les échanges sur le sujet.

RÉFÉRENCES

- Anthony, L. (2023). AntConc (Version 4.2.4). Tokyo, Japon: Waseda University. Disponible sur <https://www.laurenceanthony.net/software>
- Bevendorff, J., Wiegmann, M., Potthast, M., et Stein, B. (2024). Is Google getting worse? A longitudinal investigation of SEO spam in search engines. Dans N. Goharian, Tonello, N., He, Y., Lipani, A., McDonald, G., Macdonald, C. & Ounis, I. (Éds.), *Advances in Information Retrieval, ECIR 2024, Lecture Notes in Computer Science*, 14610. Springer.
- Bowker, L. (2011). Off the record and on the fly: Examining the impact of corpora on terminographic practice in the context of translation. Dans A. Kruger, K. Wallmach & J. Munday (Éds.), *Corpus-based Translation Studies: Research and Applications* (p. 211-236). Bloomsbury.
- Costa, H., Pastor, G. C. & Muñoz, I. D. (2014). A comparative user evaluation of terminology management tools for interpreters. Dans *Proceedings of the 4th International Workshop on Computational Terminology*, (p. 68-76).
- Costa, H., Zaretskaya, A., Corpas Pastor, G. & Seghiri, M. (2016). Nine terminology extraction tools: Are they useful for translators? *Multilingual*, 27(3), 14-20.
- European Language Industry Survey 2025. Disponible sur <https://elis-survey.org>.
- EMT 2022. Référentiel de compétences 2022 du réseau European Master's in Translation (EMT) de la Commission européenne. Disponible sur https://commission.europa.eu/system/files/2023-01/emt_competence_fw_k_2022_fr.pdf.
- Frérot, C. (2018). Enseignement de la terminologie appliquée à une formation universitaire professionnalisante : illustration d'une collaboration avec l'Organisation Mondiale de la Propriété Intellectuelle. *Myriades*, 4, 35-52.
- Gauthier, L. (2020). Dictionnaires/glossaires vs corpus : ce que les corpus font (ou ne font pas) à la terminologie pour traducteurs. Séminaire « Penser la traduction », Université de Haute-Alsace, Mulhouse, France.
- Grossmann, F. & Lavault, E. (2000). Les enjeux professionnels et didactiques d'une formation en terminologie. *Lidil - Revue de linguistique et de didactique des langues*, 21, 155-177.
- Inkpen, Paribakht, T.S., Faez, F. & Amjadian, E. (2016). Term evaluator: a tool for terminology annotation and evaluation. *International Journal of Computational Linguistics and Applications*, 7(2), 145-165.
- Krüger, R. & Hackenbuchner, J. (2022). Outline of a didactic framework for combined data literacy and machine translation literacy teaching. *Current Trends in Translation Teaching and Learning E.*, 9, 375-432.
- Kübler, N. (2003). Corpora and LSP translation. Dans F. Zanettin, S. Bernardini & D. Stewart (Éds.), *Corpora in Translator Education* (p. 25-42). St Jerome.
- Loock, R. (2016). *La Traductologie de corpus*. Villeneuve d'Ascq, Presses Universitaires du Septentrion.

Loock, R. (2019). Parce que le « grammaticalement correct » ne suffit pas : le respect de l'usage grammatical en langue cible. Dans M. Berré, B. Costa, A. Kefer, C. Letawe, H. Reuter & G. Vanderbauwhede (Éds.), *La formation grammaticale du traducteur* (p.179-194). Presses universitaires du septentrion.

Loock, R. (2025). Enseigner la compilation et l'exploitation de corpus monolingues spécialisés pour la traduction : retour sur expériences et suggestions. *ILCEA*, 57.

Loock, R. & Moulard, N. (2025). 3 enseignements, 2 intervenants, 1 texte hyperspécialisé à traduire : présentation d'un travail terminologique effectué en opérant un décroisement pédagogique avec des étudiants en traduction. *Syn-thèses*, 16, 8-25

Mollick, E. (2024). *Co-intelligence: Living and working with AI*. Londres, WH Allen.

Muegge, U. (2023). *Terminology Extraction for Translation and Interpretation Made Easy*. Auto-publication.

Pecman, M. & Kübler, N. (2011). ARTES: an online lexical database for research and teaching in specialized translation and communication. Dans *Proceedings from International Workshop on Lexical Resources (WoLeR) 2011* (p. 87-93).

Planas E, Picton, A. & Josselin-Leray, A. (2014). Exploring the use and usefulness of KRCs in translation: Towards a protocol. Colloque "Terminology and Knowledge Engineering 2014".

Ridsdale, C., Rothwell, J., Smit, M., Ali-Hassan, H., Bliemel, M. Irvine, D., Kelley, D., Matwin, S. & Wuetherick, B. (2015). *Strategies and Best Practices for Data Literacy Education. Knowledge Synthesis Report*. Dalhousie University, Canada.

Rothwell A., Moorkens, J. Fernández-Parra, M., Drugan, J. & Austermuehl, F. (2023), *Translation Tools and Technologies* (1^{re} éd.), Londres, Routledge.

Rothwell, A., Moorkens, J. & Svoboda, T. (2025). Training in translation tools and technologies: Findings of the EMT survey 2023. Disponible sur <https://arxiv.org/abs/2503.22735>

Sánchez-Gijón, P. (2009). DIY Corpora in the Specialised Translation Course. Dans A. Beeby, P. Rodríguez-Inés & P. Sánchez-Gijón (Éds.), *Corpus Use and Translating: Corpus Use for Learning to Translate and Learning Corpus Use to Translate* (p. 109128), John Benjamins.

Thompson B., Preet Dhaliwal, M., Frisch, P., Domhan, T. & Federico, M. (2024). A shocking amount of the web is machine translated: Insights from multi-way parallelism. Dans L. W. Ku, A. Martins & V. Srikumar (Éds.), *Findings of the association for computational linguistics 2024* (p. 1763-1775).

van Egdom, G.-W., Konttinen, K., Vandepitte, S., Fernández-Parra, M., Loock, R. & Bindels, J. (2020). Empowering translators through entrepreneurship in simulated translation bureaus. *Hermes - Journal of Language and Communication Studies*, 60, 81-95.

Veroy-González, M. A., Díaz-Alarcón, S. & Marcos-Aldon, M. (2019). Proyecto de innovación docente: UCOTerm, sitio web para la difusión de recursos para la traducción científico-técnica. Dans J. J. Gázquez-Linares, M. del Mar Molero

Jurado, A. B. Barragán Martín, M. del Mar Simón Márquez, Á. Martos Martínez, J. G. Soriano Sánchez & N. F. Oropesa Ruiz (Éds.), *Innovación docente e investigación en arte y humanidades* (p. 375-389).

Vezzani, F. (2022). *Terminologie numérique : conception, représentation et gestion*. Peter Lang.