

LA TRADUCTION EST MORTE, VIVE LA TRADUCTION !

M. GUIDERE. Professeur des Universités (Paris 8). Directeur de recherches à l'INSERM

Institut National de la Santé et de la Recherche Médicale, France.

mathieu.guidere@inserm.fr

Résumé : Cet article explore la révolution de la traduction générative à l'ère de l'intelligence artificielle (IA), en analysant tant ses fondements théoriques que ses implications pratiques. Après avoir défini la traduction générative à l'ère de l'IA générative, l'auteur détaille comment ces technologies transforment le rôle du traducteur, désormais positionné comme expert en post-édition et gestion de projets hybrides. S'appuyant sur des exemples concrets, il met en lumière les atouts – tels que la gestion des contextes complexes et la production de rendus créatifs – ainsi que les limites, notamment en termes de gestion des nuances culturelles et de biais inhérents aux données d'entraînement.

L'article propose également une classification des modèles de traduction générative et illustre la redéfinition du concept d'équivalence, qui passe d'une correspondance absolue à une approche nuancée de la similarité sémantique et stylistique. Enfin, l'auteur interroge et repense les dichotomies classiques de la traductologie (théorie versus pratique, traduisible versus intraduisible, art versus science, auteur versus traducteur, etc.) et revisite les théories de la traduction – notamment l'approche interprétative du sens et la théorie du skopos, ouvrant ainsi de nouvelles perspectives pour une traductologie plus intelligente et plus attentive aux enjeux cognitifs de la traduction à l'heure de l'IA.

Mots-clés : Traductologie, traduction générative, intelligence artificielle, modèles de traduction, post-édition, équivalence, cognition, théorie interprétative du sens, théorie du skopos.

Abstract: This article explores the revolution of generative translation in the era of artificial intelligence (AI), analyzing both its theoretical foundations and practical implications. After defining generative translation within the context of generative AI, the author details how these technologies are transforming the role of the translator, who is now positioned as an expert in post-editing and the management of hybrid projects. Drawing on concrete examples, the article highlights the strengths—such as handling complex contexts and producing creative renderings—as well as the limitations, particularly regarding the management of cultural nuances and biases inherent in training data. The article also proposes a classification of generative translation models and illustrates the redefinition of the concept of equivalence, shifting from an absolute correspondence to a nuanced approach to semantic and stylistic similarity. Finally, the author questions and rethinks classic dichotomies in translation studies (theory versus practice, translatable versus untranslatable, art versus science, author versus translator, etc.) and revisits translation theories—particularly the interpretative theory of meaning and the skopos theory, thereby opening up new perspectives for a smarter translation studies discipline more attentive to the cognitive challenges of translation in the age of AI.

Keywords: Translation studies, generative translation, artificial intelligence, translation models, post-editing, equivalence, cognition, interpretative theory of meaning, skopos theory.

Introduction

L'avènement de l'intelligence artificielle (IA) a profondément bouleversé le domaine de la traduction, remettant en question les pratiques traditionnelles et ouvrant de nouvelles perspectives pour les professionnels du langage. Cette révolution technologique soulève une question cruciale : la traduction humaine est-elle vouée à disparaître ?

La question se pose car les systèmes de traduction automatique neuronale (NMT) ont atteint des niveaux de performance sans précédent, produisant des traductions de plus en plus fluides et précises. Ils offrent des avantages indéniables en termes de rapidité et de volume de traitement, permettant de traduire des quantités considérables de textes en un temps record. De plus, ces systèmes sont capables de maintenir une cohérence terminologique sur l'ensemble d'un corpus, aspect particulièrement important dans les domaines techniques et spécialisés.

Cette évolution rapide a conduit à une remise en question du rôle traditionnel du traducteur humain, suscitant des inquiétudes quant à l'avenir de la profession. Mais malgré ses progrès impressionnants, l'IA peine encore à saisir les nuances culturelles, les subtilités linguistiques et les contextes complexes qui sont essentiels à une traduction de haute qualité. La traduction ne se résume pas à un simple transfert linguistique, mais implique une compréhension des cultures source et cible, ainsi qu'une capacité d'adaptation et de créativité proprement humains (Guidère, 2016).

La transformation du rôle du traducteur

Plutôt que de disparaître, le rôle du traducteur humain évolue vers celui d'un expert en communication interculturelle et en gestion de projets linguistiques complexes. Les traducteurs sont de plus en plus amenés à travailler en collaboration avec les systèmes d'IA, en se concentrant sur la post-édition, la révision et l'optimisation des traductions automatiques. Cette nouvelle approche requiert des compétences spécifiques en évaluation de la qualité des traductions automatiques et en adaptation des textes aux spécificités culturelles et contextuelles du public cible.

Dans ce contexte, la spécialisation des traducteurs devient un atout majeur. Les domaines nécessitant une expertise pointue, tels que la traduction juridique ou médicale, continuent de valoriser l'intervention humaine. La capacité des traducteurs à comprendre et à transmettre les subtilités d'un texte dans ces domaines spécialisés reste inégalée par les systèmes d'IA.

Mais même dans ces domaines, l'avenir de la traduction semble se dessiner autour d'une collaboration étroite entre l'humain et la machine. Cette approche hybride permet de combiner les forces de l'IA (rapidité, cohérence, traitement de grands volumes) avec l'expertise humaine (sensibilité culturelle, créativité, jugement contextuel).

Les traducteurs humains deviennent de plus en plus des superviseurs et des post-éditeurs. Ils sont chargés de réviser et d'optimiser les traductions générées par l'IA, apportant leur expertise linguistique et culturelle pour affiner le résultat final. Cette synergie conduit à une augmentation de la productivité et de la qualité des traductions, tout en permettant aux traducteurs de se concentrer sur les aspects les plus complexes et créatifs de leur métier.

Mais pour s'adapter à ce nouveau rôle, les traducteurs doivent développer de nouvelles compétences, notamment : la maîtrise des outils de traduction assistée par ordinateur, la compréhension approfondie des systèmes de traduction automatique, la capacité à évaluer et améliorer les traductions générées par l'IA, ainsi que la veille technologique continue pour rester à jour avec les évolutions du domaine.

Pour ne pas disparaître, les traducteurs doivent également se positionner comme des experts dans des domaines spécifiques, tels que la traduction juridique, médicale ou technique. Cette spécialisation leur permet d'apporter une expertise que l'IA ne peut pas encore égaler dans ces domaines complexes et sensibles où la décision ne peut pas être laissée à la machine.

Dans les traductions non spécialisées, le rôle du traducteur doit évoluer vers celui d'un expert en communication interculturelle pour réaliser l'évaluation de la qualité des traductions, la post-édition et l'optimisation des textes générés par l'IA, ainsi que l'adaptation du contenu aux spécificités culturelles du public cible.

En outre, la collaboration homme-machine ouvre de nouvelles perspectives pour l'industrie de la traduction. En effet, elle permet le développement de méthodologies combinant traduction automatique neuronale et post-édition humaine, la création de glossaires thématiques et de mémoires de traduction pour améliorer les performances de l'IA, ainsi que la formation continue des traducteurs pour maîtriser les nouvelles technologies et outils de traduction générative.

Définition de la traduction générative

L'essor de l'IA générative constitue un tournant qui redéfinit les contours mêmes de la traduction. La traduction générative (TG) exploite des modèles d'IA capables de produire des textes traduits qui vont bien au-delà de la

simple substitution lexicale ou de la traduction statistique. C'est pourquoi nous proposons de définir et d'analyser ce nouveau type de traduction en nous appuyant sur les avancées technologiques et en illustrant notre propos par des exemples concrets issus des pratiques actuelles.

Qu'est-ce que l'IA générative ?

L'IA générative se caractérise par sa capacité à produire du contenu original à partir d'un ensemble de données d'entraînement, sans se limiter à une reproduction ou à une simple reformulation de textes existants. Contrairement aux systèmes traditionnels de traduction automatique (T.A ou T.A.O), qui s'appuyaient sur des règles préétablies, des corpus bilingues ou des méthodes statistiques, l'IA générative, grâce aux réseaux de neurones profonds et aux architectures de type transformeur, *apprend* des représentations linguistiques et contextuelles qui lui permettent de générer des phrases cohérentes et nuancées dans une langue cible.

Les modèles génératifs tels que GPT (Generative Pre-trained Transformer) fonctionnent en s'appuyant sur une pré-formation sur de vastes quantités de données textuelles. Ce pré-entraînement permet au modèle d'acquérir une connaissance étendue de la syntaxe, du vocabulaire et des structures discursives, qu'il peut ensuite mobiliser pour générer des contenus nouveaux. En traduction, cette capacité se traduit par la possibilité de générer des versions qui intègrent non seulement des correspondances linguistiques, mais aussi des éléments de cohérence contextuelle et culturelle.

Historiquement, la traduction automatique a connu plusieurs évolutions majeures. Dans un premier temps, les systèmes de traduction basés sur des règles (Rule-Based Machine Translation, RBMT) s'appuyaient sur des règles linguistiques codifiées manuellement. Ces systèmes présentaient l'avantage de produire des traductions souvent fidèles à la grammaire, mais ils étaient limités par la complexité d'élaborer un ensemble de règles exhaustives pour toutes les langues et contextes.

La révolution statistique, amorcée dans les années 2000, a permis d'atteindre un niveau de performance supérieur grâce à l'utilisation de modèles probabilistes et de corpus parallèles. Les systèmes de Traduction Automatique Statistique (Statistical Machine Translation, SMT) étaient capables d'exploiter d'énormes bases de données pour générer des traductions en se basant sur des cooccurrences de phrases et de segments de phrases.

La dernière évolution, celle de la Traduction Automatique Neuronale (Neural Machine Translation, NMT), a intégré les réseaux de neurones pour offrir des traductions plus fluides et cohérentes. Cependant, malgré leurs performances, ces systèmes restaient limités par leur capacité à traiter des contextes très larges ou des nuances subtiles de sens.

L'IA générative représente une nouvelle étape dans cette

évolution. Elle ne se contente plus de transférer des séquences de mots d'une langue à une autre, mais elle est capable de produire des traductions en « générant » du texte de manière créative et contextuellement pertinente. Ce mode opératoire, qui relève de la traduction générative, ouvre des perspectives inédites pour la traduction automatique.

Comment fonctionne la traduction générative ?

La traduction générative s'appuie sur des modèles d'IA capables de générer des phrases à partir d'un encodage contextuel profond. Le processus peut se décomposer en plusieurs étapes clés qui simulent le raisonnement humain :

1. Encodage du texte source : Le modèle commence par analyser le texte source pour en extraire des représentations contextuelles. Cette phase d'encodage permet de saisir non seulement les mots, mais également leurs relations syntaxiques et sémantiques.
2. Transformation contextuelle : Grâce aux mécanismes d'attention (attention mechanisms), le modèle identifie les éléments les plus pertinents du texte source pour chaque partie de la phrase traduite. Cette étape est essentielle pour garantir que le contexte, souvent étendu et complexe, soit correctement pris en compte.
3. Génération du texte cible : À partir des représentations encodées, le modèle génère la traduction. Contrairement aux approches traditionnelles, cette génération est effectuée de manière séquentielle, en choisissant chaque mot en fonction du contexte précédent et des connaissances acquises lors du pré-entraînement. Cela permet de produire des traductions qui s'adaptent finement aux nuances du texte source.
4. Post-traitement et ajustements : Dans certains cas, des mécanismes de post-édition automatisée ou des interventions humaines peuvent être appliqués pour corriger d'éventuelles incohérences ou pour adapter le style de la traduction à un contexte spécifique.

L'un des avantages majeurs de cette approche est sa capacité à générer des traductions qui semblent être créées « sur le moment », en s'appuyant sur une compréhension globale du texte. Par exemple, dans la traduction de textes littéraires ou de contenus marketing, où la fidélité au style et la sensibilité culturelle sont essentielles, la traduction générative peut produire des rendus qui s'éloignent des traductions littérales pour offrir une interprétation plus nuancée et adaptée.

Exemples concrets de traduction générative

Pour illustrer la portée de la traduction générative, examinons quelques exemples concrets issus de l'utilisation de modèles avancés d'IA.

Traduction de textes littéraires

Prenons un extrait de texte littéraire où la richesse stylistique

et la profondeur des émotions jouent un rôle déterminant. Une phrase telle que « La nuit enveloppait la ville d'un voile de mystère, et chaque ruelle semblait receler un secret oublié » demande non seulement une traduction fidèle aux mots, mais surtout une reproduction du ton poétique et de l'atmosphère mystérieuse.

Un système de traduction générative, entraîné sur de vastes corpus littéraires et doté d'un modèle de langage génératif, produit une traduction en anglais telle que : *"The night draped the city in a veil of mystery, with every alleyway seemingly harboring a forgotten secret."*

Cette traduction n'est pas seulement une correspondance des mots, mais une reconstruction du sentiment et du style de l'original. Le modèle, en intégrant le contexte global, parvient à adapter les métaphores et les images pour qu'elles résonnent avec le lecteur anglophone.

Traduction dans des contextes techniques et spécialisés

Dans des domaines spécialisés comme le droit ou la médecine (Guidère, 2020), la traduction doit impérativement préserver la précision terminologique et la cohérence des concepts. Par exemple, dans la traduction d'un article scientifique sur la biotechnologie, une phrase telle que « L'activation du gène CRISPR-Cas9 ouvre de nouvelles perspectives en thérapie génique » requiert une précision extrême.

Un système de traduction générative produit la traduction suivante : *"The activation of the CRISPR-Cas9 gene opens new horizons in gene therapy."*

Ici, la traduction générative démontre sa capacité à maintenir la cohérence terminologique et à adapter le registre scientifique tout en produisant un texte fluide et accessible. Ce type de traduction repose sur une base de données spécialisée et sur l'intégration de glossaires thématiques, souvent enrichis par la collaboration entre experts humains.

Traduction et adaptation interculturelle

Un autre exemple pertinent concerne la traduction de contenus à forte dimension interculturelle, tels que des publicités ou des messages institutionnels (Guidère, 2000, 2008). La traduction générative excelle dans la reformulation contextuelle, prenant en compte les différences culturelles pour adapter le contenu au public cible. Par exemple, un slogan publicitaire peut nécessiter des ajustements pour éviter des maladresses culturelles ou pour renforcer l'impact émotionnel dans une langue différente.

Supposons un slogan en français :

« La vie est belle, vivez-la pleinement »

Une traduction générative propose en anglais :

"Life is beautiful – live it to the fullest."

Le modèle ne se contente pas de traduire mot à mot, mais

prend en compte les idiomes et les constructions de la langue cible afin de proposer une version qui conserve l'esprit de l'original tout en s'adaptant aux normes culturelles et linguistiques de l'audience anglophone.

Les atouts de la traduction générative

La traduction générative se distingue par sa capacité à générer des traductions qui vont au-delà du simple transfert de mots, en intégrant des aspects créatifs et en s'adaptant aux spécificités du contexte. Cela permet d'obtenir des rendus qui respectent le style, le ton et l'intention du texte source.

Grâce aux architectures de transformeurs et aux mécanismes d'attention, ces systèmes sont capables de gérer des contextes étendus. Ils peuvent ainsi maintenir une cohérence narrative sur de longs passages et intégrer des informations dispersées dans le texte, ce qui constitue un avantage considérable par rapport aux systèmes de traduction plus traditionnels.

De plus, les modèles génératifs bénéficient d'une capacité d'apprentissage continu. En s'entraînant sur de nouveaux corpus et en intégrant des retours d'expérience issus de la post-édition humaine, ces systèmes peuvent améliorer leur précision et s'adapter aux évolutions linguistiques et culturelles.

Que ce soit dans la traduction littéraire, technique ou médicale, la traduction générative s'adapte à divers domaines d'application. Elle permet de traiter de grandes quantités de textes tout en garantissant une qualité capable de rivaliser avec celle des traductions humaines, en particulier lorsqu'elle est combinée avec une expertise de post-édition.

Les limites de la traduction générative

Bien que les modèles génératifs soient très performants pour saisir des contextes larges, ils peuvent encore rencontrer des difficultés dans la compréhension fine des sous-entendus culturels ou des références historiques spécifiques. Par exemple, des allusions littéraires ou des expressions idiomatiques très locales peuvent ne pas être entièrement correctement interprétées sans intervention humaine.

En outre, les modèles d'IA générative sont entraînés sur des données massives qui peuvent contenir des biais culturels ou linguistiques. Par conséquent, la traduction produite peut refléter ces biais, affectant la neutralité ou la pertinence du texte dans certaines cultures ou contextes spécialisés.

Même si les mécanismes d'attention améliorent la gestion du contexte, la traduction de documents extrêmement longs ou complexes peut entraîner des incohérences ou des répétitions. Cela nécessite souvent une révision pour garantir la fluidité et la cohérence du texte final. En effet, la post-édition demeure une étape essentielle pour corriger les erreurs, ajuster le style et garantir que les nuances culturelles et contextuelles soient correctement rendues.

Perspectives d'avenir de la traduction générative

La collaboration homme-machine est de plus en plus envisagée comme une solution synergique, combinant la rapidité et la capacité de traitement massif de l'IA avec la sensibilité culturelle et l'expertise humaine.

Une voie de recherche consiste à développer des systèmes hybrides qui intègrent à la fois l'apprentissage supervisé (basé sur des exemples corrigés par des humains) et l'apprentissage non supervisé, afin d'affiner la compréhension des contextes complexes et des subtilités linguistiques. Par exemple, des plateformes collaboratives pourraient permettre aux traducteurs de fournir des corrections qui seraient ensuite intégrées dans le modèle, créant ainsi un cycle d'amélioration continue.

Dans des domaines nécessitant une expertise pointue, comme la traduction juridique ou médicale (Guidère, 2000), l'intégration de bases de connaissances spécialisées (glossaires, bases de données terminologiques, etc.) est essentielle. L'avenir de la traduction générative pourrait inclure des modules dédiés qui s'enrichissent continuellement de ces ressources, permettant ainsi de produire des traductions non seulement précises, mais également contextuellement adaptées aux exigences des domaines spécialisés.

Une autre perspective réside dans la personnalisation des traductions en fonction des préférences stylistiques ou des besoins spécifiques des utilisateurs. Des modèles génératifs capables de s'adapter aux préférences d'un client ou d'un secteur d'activité pourraient offrir une grande valeur ajoutée. Par exemple, dans le domaine publicitaire, un système pourrait être configuré pour adopter un ton plus chaleureux ou plus formel selon la cible du message, tout en garantissant une cohérence globale.

Enfin, la transparence sur le rôle des IA dans la production des traductions, ainsi que la garantie d'un usage éthique et responsable, sont des enjeux d'avenir. Les institutions de traduction et les régulateurs devront collaborer pour définir des standards et des protocoles de validation afin de prévenir les dérives, notamment en termes de biais ou de perte de diversité linguistique.

Les modèles de traduction générative

Au cœur de la révolution IA se trouvent une multitude de modèles de langage génératifs, chacun apportant des caractéristiques particulières et des approches innovantes à la traduction. Nous explorons l'impact de ces modèles sur la traduction générative, en présentant leur fonctionnement, leurs avantages et leurs limites.

Classification des modèles de traduction générative

Les modèles de traduction générative peuvent être classés en plusieurs catégories en fonction de leur architecture et de leur mode de fonctionnement. Parmi les approches les plus courantes, nous distinguons :

- Les modèles séquentiels basés sur des réseaux de neurones récurrents (RNN) et leurs variantes ;
- Les modèles basés sur l'architecture Transformer ;
- Les modèles pré-entraînés et génératifs de type GPT (Generative Pre-trained Transformer) ;
- Les modèles hybrides et spécialisés pour la traduction d'un domaine en particulier (ex. Psychiatrie).

Chaque catégorie se distingue par la manière dont elle traite les données linguistiques et par la capacité à générer des traductions qui vont au-delà d'une simple correspondance lexicale.

Modèles séquentiels et réseaux de neurones récurrents (RNN)

Les premiers systèmes de traduction neuronale se sont appuyés sur des architectures séquentielles, notamment les réseaux de neurones récurrents (RNN) et leurs variantes, comme les Long Short-Term Memory (LSTM) et les Gated Recurrent Units (GRU). Ces modèles fonctionnent en traitant le texte source mot par mot, en maintenant une mémoire interne qui permet de capter le contexte dans une séquence linéaire :

- Encodage séquentiel : Le texte source est parcouru séquentiellement et chaque mot est transformé en une représentation vectorielle.
- Transmission de l'information : Grâce à leur mémoire interne, les RNN tentent de conserver l'information contextuelle sur l'ensemble de la séquence.
- Décodage séquentiel : Le modèle génère le texte cible en décodant ces représentations, un mot à la fois, en tenant compte du contexte accumulé.

Imaginons la phrase source « La météo annonce de la pluie pour demain ». Un système basé sur un RNN-LSTM va encoder cette phrase en une séquence de vecteurs et, au décodage, générer la traduction anglaise : « The weather forecast predicts rain for tomorrow ». Bien que ce modèle fonctionne parfaitement sur des phrases simples, il rencontre souvent des difficultés sur des textes longs, en raison de la limitation de sa capacité à maintenir une information contextuelle sur de longues séquences.

L'architecture Transformer

Introduit en 2017, le « Transformer » a révolutionné le traitement automatique du langage naturel (TALN / NLP) en s'appuyant sur des mécanismes d'attention (attention mechanisms) plutôt que sur des séquences récurrentes. Cette architecture permet de traiter l'ensemble d'un texte en parallèle, en évaluant les relations entre tous les éléments de la séquence simultanément. En traduction, le transformeur fonctionne de la manière suivante :

- Encodage global : Le modèle encode l'intégralité du texte source en générant des représentations vectorielles pour

chaque mot, enrichies par des informations sur la position et le contexte global grâce à des mécanismes d'attentivité multi-têtes.

- Mécanismes d'attentivité : Ces mécanismes permettent au modèle de pondérer l'importance de chaque mot par rapport aux autres, assurant ainsi une compréhension fine des relations contextuelles.
- Décodage contextuel : Le décodeur, également basé sur des mécanismes d'attentivité, génère le texte cible en se référant simultanément à l'ensemble des vecteurs encodés, ce qui permet de maintenir une cohérence sur des textes complexes et longs.

Prenons une phrase technique, par exemple : « L'activation de la protéine p53 induit une réponse cellulaire qui conduit à l'apoptose ». Un modèle de transformeur, grâce à sa capacité à saisir des dépendances longues et complexes, peut produire une traduction en anglais telle que « The activation of the p53 protein triggers a cellular response that leads to apoptosis », en veillant à conserver la précision scientifique et la cohérence du contexte.

Modèles pré-entraînés et génératifs (GPT et similaires)

Les modèles de la famille GPT (Generative Pre-trained Transformer) représentent une évolution des architectures « Transformer », combinant un pré-entraînement sur de vastes corpus textuels avec une capacité de génération puissante. Ces modèles sont non seulement capables de comprendre le langage, mais également de générer du texte de manière autonome. Leur fonctionnement en traduction peut être décrit comme suit :

- Pré-entraînement sur des corpus massifs : Le modèle est d'abord entraîné sur une énorme quantité de textes non étiquetés, lui permettant d'acquérir une compréhension globale de la langue, des structures grammaticales, et des subtilités contextuelles.
- Fine-tuning pour la traduction : Ensuite, il peut être ajusté sur des corpus bilingues ou spécialisés pour améliorer ses performances dans la traduction.
- Génération autonome : Lors de la traduction, le modèle génère le texte cible en anticipant le mot suivant dans une séquence, de manière à produire des traductions fluides et adaptées au contexte.

Un modèle GPT ajusté pour la traduction pourrait traiter une phrase littéraire telle que « Le crépuscule baignait l'horizon d'une lumière dorée, teintant le monde d'une douceur éphémère » pour produire cette traduction en anglais : *"The twilight bathed the horizon in a golden light, tinting the world with an ephemeral sweetness."*

Cette traduction témoigne de la capacité du modèle à reproduire non seulement le sens, mais aussi la tonalité poétique et l'atmosphère de l'original.

Modèles hybrides et spécialisés

Outre les modèles « génériques » de traduction générative, des approches hybrides émergent pour répondre à des besoins spécifiques. Ces modèles combinent souvent plusieurs architectures ou intègrent des composants additionnels pour renforcer certaines compétences, comme la terminologie spécialisée ou la gestion des contextes culturels. Ils fonctionnent ainsi :

- Intégration de modules spécialisés : Dans des domaines comme le médical (Guidère, 2024), des modules dédiés peuvent être intégrés afin de gérer les terminologies précises et les concepts spécifiques.
- Combinaison de méthodes : Certains systèmes hybrides combinent des modèles de type « Transformer » avec des techniques de recherche de phrase ou d'alignement de segments, afin d'améliorer la qualité des traductions sur des corpus techniques ou scientifiques (Guidère, 2023).
- Collaboration avec la post-édition humaine : Ces modèles sont souvent conçus pour fonctionner en tandem avec des experts humains qui valident ou corrigent les traductions, garantissant ainsi une qualité optimale.

Dans le domaine médical, un système hybride peut prendre en entrée la phrase française « L'administration de ce médicament nécessite une surveillance étroite pour éviter tout effet indésirable majeur » et, après traitement par un modèle spécialisé couplé à une base de données terminologique, générer la traduction en anglais suivante : *"The administration of this medication requires close monitoring to prevent any major adverse effects."*

La précision terminologique et la clarté du message sont renforcées par l'intégration d'un module dédié au vocabulaire médical, assurant une adéquation parfaite avec les attentes du domaine (Guidère & Jehel, 2025).

Exemples illustratifs de modèles génératifs en traduction

Pour mieux comprendre les différences entre ces modèles, il est instructif d'examiner quelques cas d'application concrets.

Exemple d'un modèle Transformer dans un contexte littéraire

Considérons un extrait littéraire français décrivant une scène bucolique :

« Au cœur de la forêt, la lumière filtrée par les feuillages créait un jeu d'ombres et de lumières, donnant à l'endroit une atmosphère presque féerique. »

Un modèle Transformer, grâce à son attentivité globale, est capable de générer une traduction en anglais telle que :

"In the heart of the forest, the dappled light filtering through

the foliage created a play of shadows and light, endowing the place with an almost magical atmosphere.”

La capacité du modèle à identifier les relations entre les différents éléments visuels et sensoriels du texte permet d’offrir une traduction qui transmet l’essence poétique de l’original.

Exemple d’un modèle GPT appliqué à la traduction technique

Prenons un document technique dans le domaine de l’informatique :

« L’implémentation d’algorithmes de deep learning repose sur des frameworks tels que TensorFlow ou PyTorch, qui offrent une flexibilité et une performance accrues pour le traitement des données. »

Un modèle GPT, pré-entraîné et ensuite fine-tuné pour des contenus techniques, peut produire la traduction suivante :

“The implementation of deep learning algorithms relies on frameworks like TensorFlow or PyTorch, which provide enhanced flexibility and performance for data processing.”

Ici, le modèle GPT montre sa capacité à reproduire fidèlement des termes techniques tout en assurant la fluidité et la clarté du texte.

Exemple de modèle hybride dans un contexte spécialisé

Dans le domaine juridique, la précision terminologique est primordiale. Une phrase telle que :

« Le contrat stipule que toute violation des clauses entraînera des pénalités financières significatives. »

Peut être traitée par un modèle hybride combinant un Transformer et un module spécialisé intégrant un glossaire juridique, donnant la traduction :

“The contract stipulates that any breach of the clauses will result in significant financial penalties.”

L’ajout du module spécialisé assure une correspondance exacte des termes juridiques, tout en bénéficiant de la fluidité générée par le Transformer.

Atouts et limites des différents modèles de traduction générative

Les modèles « Transformer » et « GPT », en particulier, excellent dans la gestion des dépendances à long terme grâce aux mécanismes d’attention. Cela leur permet de produire des traductions qui captent le contexte global et les nuances subtiles du texte source.

Les modèles pré-entraînés peuvent être ajustés à une variété de tâches linguistiques. En effectuant un « fine-tuning » sur des corpus spécifiques, il est possible d’adapter la traduction générative à des domaines spécialisés, qu’il s’agisse de textes littéraires, techniques ou publicitaires. Ces modèles intègrent une compréhension implicite du style

et du ton, produisant des traductions qui paraissent naturelles et cohérentes.

De plus, la possibilité d’intégrer des retours issus de la post-édition humaine dans des systèmes hybrides ou pré-entraînés offre une voie d’amélioration continue. Cela permet aux modèles de s’adapter aux évolutions du langage et aux exigences spécifiques de chaque domaine.

Mais tous les modèles génératifs, notamment ceux pré-entraînés sur de larges corpus, peuvent reproduire des biais présents dans les données d’entraînement. Malgré leur capacité à saisir le contexte global, ils peuvent parfois commettre des erreurs sur des points culturels ou contextuels très fins, nécessitant une supervision humaine.

Par ailleurs, les architectures de type « Transformer » ou « GPT » nécessitent des ressources informatiques importantes pour leur entraînement et leur déploiement. Cette exigence technique peut constituer un obstacle pour certaines applications ou pour des organisations aux moyens limités.

Même si l’architecture « Transformer » améliore la gestion du contexte par rapport aux RNN, la traduction de documents très longs ou d’un contenu particulièrement complexe peut entraîner des incohérences ou des pertes d’information, ce qui exige la mise en place de stratégies de segmentation et une post-édition soignée.

Enfin, les modèles génératifs peuvent avoir du mal à capter certaines nuances propres à des langues ou à des registres culturels spécifiques, en particulier dans des domaines très localisés ou avec des expressions idiomatiques non présentes dans les corpus d’entraînement.

Ainsi, l’avenir des modèles de traduction générative semble résider dans une convergence des approches et dans l’intégration de retours d’expérience. En effet, les systèmes actuels de traduction générative s’améliorent continuellement grâce aux corrections apportées par des traducteurs humains lors de la post-édition. Ce retour d’information est essentiel pour affiner les modèles et réduire les erreurs contextuelles. Des plateformes collaboratives, où les traducteurs peuvent annoter et corriger les sorties générées, permettent d’enrichir les bases d’apprentissage et d’adapter les modèles aux évolutions linguistiques.

La tendance est ainsi à la fusion des approches neuronales avec des méthodes symboliques ou basées sur des règles. Par exemple, pour la traduction de textes contenant des éléments très structurés (tels que des documents légaux ou des manuels techniques), l’intégration d’un système d’alignement sémantique ou d’extraction de règles spécifiques peut renforcer la qualité de la traduction générative. Cette hybridation permet de combiner la flexibilité des modèles neuronaux avec la rigueur des systèmes symboliques.

En outre, chaque domaine (juridique, médical, technique,

etc.) présente ses propres exigences en termes de terminologie et de style. Les modèles hybrides spécialisés, en intégrant des glossaires, des bases de connaissances et des modules spécifiques, offrent la possibilité d'obtenir des traductions sur mesure. Cette personnalisation est d'autant plus cruciale dans un contexte où la précision et la conformité aux standards sectoriels sont indispensables.

Enfin, au-delà des aspects techniques, l'utilisation des modèles de traduction générative soulève des questions éthiques, notamment en ce qui concerne les biais présents dans les données d'entraînement. La transparence quant aux sources et aux processus de génération devient alors primordiale pour garantir une traduction neutre et respectueuse des divers contextes culturels. La mise en place de protocoles de validation et de régulation est indispensable pour assurer un usage responsable de ces technologies.

Redéfinition du concept d'équivalence

La notion d'équivalence en traduction a longtemps constitué l'un des concepts centraux de la traductologie. Historiquement, elle a servi de critère pour mesurer la fidélité d'une traduction par rapport à un texte source, qu'elle soit définie en termes de fidélité formelle ou dynamique. Cependant, l'avènement de l'intelligence artificielle générative a radicalement transformé la pratique de la traduction, soulevant la nécessité de redéfinir ce concept.

Évolution du concept d'équivalence en traduction

Dans la théorie classique, deux grandes approches ont longtemps dominé la discussion sur l'équivalence.

D'une part, l'équivalence formelle se concentre sur la fidélité aux structures linguistiques, aux registres et à la syntaxe du texte source. Ce paradigme tend à privilégier une correspondance mot-à-mot ou phrase-à-phrase, souvent associée à une traduction littérale.

D'autre part, l'équivalence dynamique ou fonctionnelle vise à recréer l'effet de sens et la réponse émotionnelle chez le lecteur de la traduction, même si cela implique une reformulation du texte source. Ici, la priorité est donnée à la transmission de l'intention communicative plutôt qu'à la correspondance formelle.

Ces deux conceptions reposaient sur des critères relativement statiques, où la tâche du traducteur était de « reproduire » le contenu de manière qu'il soit perçu comme identique ou très similaire au texte d'origine.

Or, l'essor de l'IA générative, avec ses capacités de création de contenu contextuel et de reformulation dynamique, a remis en cause la rigidité des anciennes définitions. Plutôt que de se limiter à une simple correspondance formelle ou à une fonction communicative unique, la traduction générative offre une pluralité de rendus qui varient selon des facteurs tels que :

- La richesse sémantique du texte généré, qui peut contenir des nuances et des reformulations inédites.
- La créativité stylistique, qui permet de produire des traductions qui se distinguent par des choix lexicaux ou des constructions syntaxiques novatrices.
- La variation contextuelle, qui adapte le rendu en fonction du public cible ou du domaine d'application.

Dans ce nouveau paradigme, l'équivalence n'est plus un absolu, mais un rapport de similarité qui peut être mesuré à différents niveaux : de la similarité lexicale à la correspondance des intentions communicatives, en passant par la conservation de l'ambiance culturelle et stylistique. L'enjeu ne consiste plus à établir si une traduction est « équivalente » au sens traditionnel, mais à évaluer le degré et le type de similarité qu'elle révèle.

Redéfinition de l'équivalence en termes de similarités

Dans le cadre de la traduction générative, l'équivalence se décline en plusieurs dimensions :

- La similarité sémantique : Il s'agit de vérifier que la traduction porte le même sens, même si les mots ou les structures syntaxiques diffèrent. Par exemple, un modèle d'IA peut traduire une expression idiomatique en conservant son sens implicite plutôt que de la traduire littéralement.
- La similarité stylistique : Ici, l'accent est mis sur le maintien du ton, du registre et de l'ambiance du texte source. Une traduction d'un texte littéraire par IA générative peut choisir des tournures élégantes et poétiques pour recréer une atmosphère semblable à celle de l'original.
- La similarité pragmatique : Au-delà du sens et du style, cette dimension concerne l'effet communicatif et l'impact sur le lecteur. Une traduction doit reproduire, dans une certaine mesure, l'intention, l'humour ou l'ironie du texte source.

Ainsi, l'équivalence devient une étiquette permettant de qualifier le degré de ressemblance entre deux textes, en prenant en compte ces multiples facettes. Plutôt que d'établir une correspondance exacte, il s'agit de mesurer la qualité et la profondeur de la similitude.

Méthodes d'évaluation de la similarité traductique

Avec l'essor de l'IA générative, des outils d'évaluation plus sophistiqués ont été développés pour mesurer cette similarité multidimensionnelle. Parmi ces outils, on retrouve :

- Les évaluations humaines, où des experts en traduction analysent le texte traduit en termes de fidélité sémantique, d'adaptation stylistique et d'efficacité communicative.

- Les métriques automatiques telles que BLEU, METEOR ou encore BERT Score, qui quantifient la correspondance sémantique et lexicale entre l'original et la traduction.
- Les approches hybrides, combinant des algorithmes de traitement du langage naturel avec des analyses qualitatives pour fournir une évaluation plus nuancée.

Ces méthodologies permettent de dépasser une approche dichotomique (équivalent ou non) pour adopter une évaluation graduée de la similarité, tenant compte des divers aspects du texte.

Exemples illustrant la redéfinition de l'équivalence

Traduction de textes littéraires

Prenons l'exemple d'un extrait poétique français :

Texte source : « Les feuilles d'automne dansent sous le vent, portées par un souffle de nostalgie. »

Une traduction générative traditionnelle aurait pu chercher à reproduire une correspondance formelle, en donnant par exemple : *"The autumn leaves dance under the wind, carried by a breath of nostalgia."*

Cependant, une IA générative ajustée pour capter le style littéraire peut produire :

"Autumn leaves waltz in the breeze, carried away by a wistful sigh." (Traduction créative)

Dans ce cas, bien que les mots exacts diffèrent, la traduction conserve le sens, l'ambiance et même la musicalité du texte original. Ici, l'équivalence se mesure en termes de similarité sémantique et stylistique plutôt qu'en correspondance stricte des mots, démontrant la redéfinition du concept en tant que rapport de similarité.

Traduction de messages publicitaires

Dans le domaine publicitaire, la fidélité à l'original doit s'allier à l'adaptation culturelle. Par exemple, considérons un slogan français :

Texte source : « Vivez l'instant, partagez l'émotion. »

Une traduction générative qui vise à capter l'essence du message plutôt que la simple formulation pourrait proposer : *"Seize the moment, share the feeling."* (Traduction dynamique)

Ici, l'accent n'est pas mis sur la correspondance mot-à-mot, mais sur la reproduction de l'impact émotionnel et de l'invitation à l'action. L'équivalence est redéfinie comme le degré de similarité dans la manière dont l'émotion et le message publicitaire sont transmis, ce qui démontre la flexibilité de la notion dans un contexte d'IA générative.

Traduction technique et spécialisée

Prenons l'exemple suivant dans le domaine médical :

Texte source : « L'administration de ce traitement nécessite

une surveillance continue pour éviter les effets indésirables. »

Une traduction générative bien calibrée, avec une adaptation terminologique, pourrait donner :

Traduction spécialisée : *"The administration of this treatment requires continuous monitoring to prevent adverse effects."*

Dans ce cas, l'équivalence se mesure par la similarité sémantique et la précision des termes techniques. La traduction générative doit alors aligner non seulement le sens général, mais aussi les spécificités terminologiques du domaine. La redéfinition de l'équivalence implique ici une évaluation sur plusieurs niveaux : la conservation du contenu informatif, la précision des concepts et la clarté du message dans le contexte spécialisé.

Vers une approche nuancée de la traduction

La redéfinition de l'équivalence implique de repenser le rôle du traducteur. Plutôt que de viser une reproduction parfaite, le traducteur devient un évaluateur de similarités, chargé de mesurer le degré d'adaptation entre le texte source et le texte généré. Ce rôle hybride se matérialise notamment par :

- La post-édition, qui permet d'affiner le rendu généré par l'IA pour atteindre un niveau de similarité acceptable sur les plans sémantique, stylistique et pragmatique.
- La gestion des compromis, où le traducteur doit arbitrer entre fidélité au contenu original et adaptation aux attentes du public cible.

Ainsi, l'équivalence se transforme en une notion fluide, où le traducteur utilise son expertise pour ajuster le rapport de similarité en fonction des contextes et des exigences spécifiques.

Cette redéfinition de l'équivalence remet en question les métriques traditionnelles. Alors que des scores tels que BLEU ou METEOR fournissaient une mesure quantitative de la correspondance lexicale, ils se révèlent parfois insuffisants pour appréhender la profondeur des similarités sémantiques et stylistiques. Ainsi, de nouvelles approches combinant analyses quantitatives et évaluations qualitatives sont en cours de développement, intégrant des systèmes d'évaluation basés sur l'apprentissage profond (deep learning) pour capter les nuances du sens et du style.

Par exemple, un algorithme d'évaluation s'appuyant sur des modèles de représentation sémantique (tels que BERT) peut comparer l'original et la traduction en termes de vecteurs de sens. Cette méthode permet de quantifier la similarité sémantique de manière plus fine, tout en tenant compte des variations stylistiques. Le traducteur, de son côté, peut utiliser ces indicateurs pour ajuster la post-édition, contribuant ainsi à un processus de traduction itératif et collaboratif entre l'humain et la machine.

Cette redéfinition de l'équivalence soulève également des

questions éthiques et culturelles. D'un côté, il existe le risque que la traduction générative, en cherchant à optimiser un certain degré de similarité, puisse imposer des normes uniformes qui atténuent la diversité culturelle et linguistique. D'un autre côté, l'acceptation de multiples rendus équivalents pourrait favoriser une plus grande flexibilité, mais aussi une incertitude quant à la « vérité » du texte traduit.

Pour illustrer ce point, prenons le cas d'un texte à forte charge culturelle, tel qu'un proverbe ou une expression idiomatique. Une IA générative peut proposer plusieurs traductions qui se situent à des degrés de similarité différents par rapport à l'original. Dans ce contexte, l'évaluation de l'équivalence devra tenir compte de la pertinence culturelle, en s'appuyant sur des spécialistes capables de juger si la traduction conserve l'esprit et les références implicites du texte source.

Vers une nouvelle pratique de la traduction

L'IA générative ne remplace pas le traducteur humain, mais réinvente son rôle. L'évaluation de l'équivalence devient un processus collaboratif où la machine fournit des rendus variés, et l'expert humain choisit et ajuste ceux qui répondent le mieux aux exigences contextuelles et culturelles. Cette collaboration est d'autant plus importante lorsque l'on considère la multiplicité des niveaux de similarité à prendre en compte.

Par exemple, dans le cas de la traduction d'un discours politique (Guidère, 2015), l'IA peut générer plusieurs versions mettant en avant différents aspects : une version fidèle aux faits, une version accentuant le ton émotionnel, et une version adaptée aux sensibilités culturelles du public cible. Le traducteur doit alors analyser ces propositions et choisir la version qui parvient à équilibrer les exigences de similarité sémantique et de résonance pragmatique avec le discours original.

Dans ce contexte évolutif, la formation des traducteurs doit également s'adapter. La maîtrise des outils d'IA générative et la compréhension des nouvelles métriques de similarité deviennent indispensables. Les futurs traducteurs devront être formés à :

- Interpréter les scores de similarité fournis par des systèmes d'évaluation automatisés.
- Apprécier les nuances entre différents rendus génératifs et choisir celui qui offre la meilleure adéquation au contexte.
- Gérer les aspects éthiques liés à la traduction, notamment en ce qui concerne la préservation de la diversité culturelle.

Enfin, la redéfinition du concept d'équivalence appelle à la mise en place de standards et de protocoles permettant de quantifier la similarité entre l'original et la traduction. Ces standards, élaborés en collaboration entre chercheurs,

professionnels du langage et développeurs d'IA, pourront offrir des repères communs pour évaluer la qualité des traductions génératives. Un tel cadre de référence contribuera à la reconnaissance et à la valorisation des traductions qui, tout en étant différentes sur le plan formel, restent profondément fidèles au sens et à l'impact du texte source.

Redéfinition des dichotomies classiques

L'IA générative ne se contente pas d'améliorer la qualité des rendus en traduction, elle bouleverse également les cadres théoriques traditionnels en redéfinissant de nombreuses dichotomies classiques de la traductologie. Autrefois perçues comme des oppositions irréciliables, ces oppositions trouvent aujourd'hui de nouveaux équilibres grâce à la capacité de l'IA générative à produire des textes présentant un degré de similarité et de créativité inédit.

Théorie versus Pratique

Traditionnellement, la traductologie s'est structurée autour d'un débat entre théorie – qui cherche à élaborer des modèles, des concepts et des normes de traduction – et la pratique, c'est-à-dire l'activité effective du traducteur dans le monde professionnel. La théorie pouvait apparaître comme abstraite et parfois éloignée des contraintes concrètes de la traduction quotidienne (Guidère, 2016).

Les systèmes d'IA générative intègrent dès leur conception des principes théoriques issus de l'apprentissage profond (deep learning) et du traitement automatique du langage (TAL), tout en étant mis en œuvre dans des environnements opérationnels. Ainsi, la frontière entre théorie et pratique s'amincit. Par exemple, un modèle GPT qui produit une traduction littéraire intègre en temps réel des paramètres issus de recherches théoriques sur la représentation sémantique et stylistique, tout en étant déployé pour répondre à des besoins pratiques tels que la localisation de contenus publicitaires ou la traduction de textes littéraires.

Dans le processus de post-édition, le traducteur humain ajuste des sorties générées par l'IA. Cette activité relève à la fois d'une application pratique de l'outil (pratique) et d'une compréhension des mécanismes internes du modèle – une connaissance théorique qui oriente la correction. Ainsi, le traducteur devient à la fois praticien et théoricien, réconciliant les deux pôles.

Traduisible versus Intraduisible

Une notion longtemps débattue en traductologie est celle de l'intraduisible, c'est-à-dire d'un élément culturel ou linguistique qui résisterait à toute transposition dans une langue cible. L'idée d'un intraduisible pouvait justifier des pertes, des compromis ou même des omissions.

La capacité de l'IA générative à produire plusieurs versions d'un même texte permet aujourd'hui d'explorer différentes stratégies pour rendre ce qui était auparavant considéré

comme intraduisible. Au lieu d'admettre l'impossibilité de traduire certains éléments, l'IA propose des variations qui capturent la polysémie ou l'ambivalence d'un concept. La notion d'intraduisibilité devient ainsi un continuum de niveaux de similarité plutôt qu'une impossibilité catégorique.

Prenons l'expression française « joie de vivre ». Tandis qu'un traducteur classique pouvait hésiter entre différentes formules, un système génératif peut proposer successivement en anglais :

- “Zest for life”
- “The art of enjoying life”
- “Exuberance in living”

Chaque version capte une facette différente de la signification originale, attestant que le phénomène intraduisible se transforme en une série de rendus équivalents selon le contexte.

Art versus Science

La traduction a toujours oscillé entre une dimension artistique – qui valorise l'intuition, la créativité et l'expression personnelle – et une dimension scientifique – axée sur la rigueur, la méthodologie et l'objectivité. D'un côté, l'art traduit le flair du traducteur, de l'autre, la science impose des règles et des normes mesurables.

L'IA générative conjugue ces deux dimensions. Les algorithmes d'apprentissage profond reposent sur des principes mathématiques rigoureux (science), tout en générant des textes qui peuvent relever d'une grande créativité stylistique (art). En ce sens, le processus de traduction ne se réduit plus à une opposition, mais se conçoit comme une fusion où la rigueur des modèles statistiques rencontre la souplesse d'une création littéraire.

Lorsqu'un texte poétique est traduit par un modèle génératif, la structure rythmique et les choix lexicaux reflètent une analyse statistique fine des corpus de référence, tout en produisant un résultat qui peut être perçu comme une œuvre d'art littéraire. La traduction « scientifique » qui respecte des règles de syntaxe et de grammaire se mue alors en une interprétation artistique, où l'original et le rendu créent une nouvelle entité aux multiples niveaux d'appréciation.

Auteur versus Traducteur

La figure de l'auteur, source du texte original, a souvent été opposée à celle du traducteur, qui intervient comme intermédiaire en réinterprétant l'œuvre. La question se pose alors de savoir si le traducteur peut se positionner en co-créateur ou s'il ne fait que retranscrire l'intention de l'auteur.

Avec la traduction générative, cette dichotomie se complexifie. Les systèmes d'IA produisent des textes qui, bien que générés à partir d'un corpus, ne se contentent pas de copier mécaniquement l'intention de l'auteur. Ils offrent

une multiplicité de versions, chacune pouvant être considérée comme une nouvelle interprétation. Le traducteur humain, en post-éditant ces rendus, se trouve à la croisée des chemins entre la fidélité à l'auteur et la création d'un nouveau discours, brouillant la frontière traditionnelle.

Lors de la traduction d'un roman, l'IA peut générer plusieurs variantes d'un même passage. Le traducteur, en choisissant parmi ces propositions et en les adaptant, devient à la fois médiateur de l'œuvre originale et créateur d'une version renouvelée. Cette synergie montre que l'opposition « auteur versus traducteur » se transforme en un continuum collaboratif, où chacun apporte sa contribution à l'œuvre finale.

Original versus Copie

La relation entre l'original et la copie a toujours été au cœur du débat en traduction. La copie est souvent perçue comme une simple reproduction du texte source, avec le risque de perdre l'unicité et l'authenticité de l'original.

La traduction générative remet en cause cette vision binaire. Plutôt que de considérer la traduction comme une copie littérale, elle la conçoit comme une nouvelle création qui se situe dans une relation de similarité avec l'original. L'objectif n'est plus la reproduction exacte, mais la création d'un texte qui dialogue avec l'original par le biais d'un rapport de correspondance sémantique et stylistique.

Un poème traduit par un modèle génératif peut présenter des structures syntaxiques différentes de l'original, tout en capturant son essence, ses images et son rythme. Plutôt que de dire que le texte traduit est une « copie », il est préférable de parler d'un « écho » créatif qui répond à l'original sans en être une réplique exacte. Cela montre que l'original et la traduction coexistent dans une relation de complémentarité, plutôt que d'opposition.

Traduction versus Imitation

La distinction entre traduction et imitation se fonde sur l'idée que la traduction doit rester fidèle au texte source, alors que l'imitation relève d'une reprise libre, voire d'une transformation créative qui s'éloigne de l'original.

Les modèles génératifs produisent des textes qui oscillent entre fidélité et créativité. La traduction générative n'est pas une imitation dans le sens de la copie ou de la parodie, mais une réinterprétation qui vise à restituer le sens et l'intention de l'original tout en offrant une nouvelle forme. L'imitation, dans ce contexte, n'est plus un défaut mais une stratégie de variation qui permet d'explorer divers niveaux de similarité.

Dans la traduction d'un slogan publicitaire, le modèle génératif peut proposer une version qui s'éloigne légèrement de la formulation originale afin de mieux résonner avec le public cible. Cette démarche peut être vue non pas comme une imitation, mais comme une traduction créative qui cherche à capturer l'essence du message. Par exemple, le

slogan « Le goût de l'authenticité » peut devenir « Savor the genuine spirit », une transformation qui traduit la notion d'authenticité de manière adaptée sans être une simple imitation mot pour mot.

Sacré versus Profane

La dichotomie entre sacré et profane a longtemps structuré l'approche des textes religieux, mythologiques ou culturels, dans lesquels le sacré est considéré comme inviolable, et le profane comme accessible et commun.

L'IA générative, en proposant plusieurs rendus pour un même passage, permet de reconsidérer ces notions. Un même texte peut être traduit de façon à accentuer une dimension sacrée ou, au contraire, à le rendre plus accessible – tout dépend des choix éditoriaux et des paramètres du modèle. La traduction générative ouvre la voie à une pluralité d'interprétations, où le caractère sacré ou profane n'est plus imposé de manière binaire mais modulé en fonction du contexte culturel et de l'intention communicative.

Lors de la traduction d'un texte liturgique, une version pourrait conserver une terminologie élevée et un style solennel, tandis qu'une autre version, destinée à un public plus large, adopterait un ton plus contemporain. Ainsi, l'IA générative permet de naviguer entre le sacré et le profane sans imposer une hiérarchie rigide, mais en proposant des niveaux de traduction adaptés à différents usages.

Fidélité versus Liberté

Le débat entre fidélité – l'exactitude de la reproduction du contenu et du style – et liberté – la nécessité d'adapter, de transformer et de réinterpréter le texte – est l'un des axes centraux de la traductologie. Traditionnellement, certains traducteurs prônaient une fidélité absolue, alors que d'autres revendiquaient la liberté créative pour mieux faire vivre le texte dans la langue cible.

La traduction générative offre une approche graduée dans laquelle fidélité et liberté ne s'excluent pas mutuellement mais se complètent. Les modèles d'IA sont capables de calibrer le degré de liberté en fonction du contexte et des besoins du texte : ils peuvent proposer une version très proche du texte source, ou bien, au contraire, offrir une version plus libre qui met en valeur la tonalité et l'impact émotionnel. L'évaluation se porte ainsi sur un continuum de similarité où le traducteur décide du point d'équilibre optimal.

Pour un texte publicitaire, une version fidèle pourrait rendre la structure syntaxique du message original, tandis qu'une version libre, proposée par le modèle, pourrait réorganiser le discours pour mieux captiver un public anglophone. Le choix entre ces deux options permet de mesurer le rapport de fidélité/liberté souhaité, démontrant que la notion d'équivalence s'enrichit d'une dimension d'adaptation contextuelle.

Le Mot versus l'Idée

Historiquement, la traduction a souvent été abordée comme un transfert de mots d'une langue à l'autre. Cependant, de nombreux théoriciens ont souligné que l'essence d'un texte réside dans les idées, les concepts et les images qu'il véhicule, et non dans la simple correspondance lexicale.

L'IA générative privilégie une approche qui met en avant la transmission des idées et des significations profondes. Au lieu de se limiter à aligner mot pour mot, les systèmes modernes analysent le contexte global pour générer un rendu qui capture l'essence de l'idée. Cette redéfinition fait passer le débat de la simple traduction des mots à une traduction de la pensée.

Dans la traduction d'un aphorisme, le modèle génératif ne se contente pas de substituer des mots équivalents, mais reformule l'énoncé pour en restituer le sens global. Par exemple, pour un proverbe tel que « Mieux vaut prévenir que guérir », une traduction générative peut produire : *"It is wiser to prevent than to cure."*

Ce rendu n'est pas une correspondance mot-à-mot, mais une retranscription de l'idée centrale, démontrant la primauté du sens sur la forme.

La Lettre versus L'Esprit

Enfin, la dichotomie entre la lettre – le texte strictement écrit – et l'esprit – l'intention, le contexte culturel et la signification profonde – a toujours été au cœur du débat sur la traduction. Tandis que certains traducteurs s'attachent à une reproduction fidèle du texte, d'autres insistent sur le besoin de transmettre l'esprit du message.

La traduction générative permet d'appréhender simultanément la lettre et l'esprit du texte. En analysant le contexte global et en générant des propositions multiples, l'IA offre des rendus qui cherchent à équilibrer la précision formelle et la profondeur interprétative. L'approche devient alors une quête pour trouver le juste milieu entre le rendu littéral et la transmission de la dimension culturelle et émotionnelle.

Lorsqu'un texte religieux ou philosophique est traduit, une version générée par IA peut présenter une correspondance acceptable au niveau lexical tout en apportant des nuances qui révèlent l'esprit de l'œuvre. Ainsi, pour un extrait comme « Chercher la vérité au-delà des apparences », le modèle peut offrir une traduction telle que : *"Seek truth beyond mere appearances."*

Ce rendu, en plus de respecter la lettre, insuffle une dimension qui évoque l'intention philosophique du texte original, réconciliant ainsi les deux pôles de la dichotomie.

Ainsi, l'introduction de l'IA générative dans la traduction permet de repenser en profondeur des dichotomies jadis irréconciliables, en proposant des modèles hybrides et adaptatifs. La frontière entre théorie et pratique se dissout

grâce à une technologie qui s'appuie sur des fondements théoriques tout en étant opérationnelle.

Repenser les théories de la traduction

Au-delà de cette refondation conceptuelle, l'essor de l'IA générative et des systèmes d'apprentissage profond oblige les chercheurs à repenser les fondements théoriques de la traduction. Deux théories, en particulier, se trouvent au cœur de ce renouvellement conceptuel : la théorie interprétative du sens et la théorie du skopos. Chacune d'elles, jadis rigoureusement définie dans le débat traductologique, se voit désormais interroger par les capacités de l'IA à produire des textes riches en nuances, tout en intégrant des processus automatisés qui questionnent la nature même du sens et de l'intention communicative.

La théorie interprétative du sens et l'impact de l'IA

La théorie interprétative du sens, issue des travaux de Danica Seleskovitch (1921-2001), s'est longtemps appuyée sur l'idée que la traduction ne consiste pas uniquement à transposer des mots d'une langue à une autre, mais à interpréter le sens profond du texte original. Cette approche met en avant l'importance du contexte, des connotations et des intentions implicites qui traversent le discours. Le traducteur, en tant qu'interprète du message, se doit de saisir non seulement le contenu lexical, mais aussi la dimension pragmatique, culturelle et stylistique qui en fait la richesse.

Dans ce modèle, le sens n'est pas une donnée fixe et objective ; il est plutôt le résultat d'un processus interactif entre le texte, son contexte d'élaboration et le lecteur. Ainsi, la traduction est envisagée comme un acte créatif et interprétatif, où la fidélité au sens réside dans la capacité à reproduire l'effet global produit par l'original (Guidère, 2016).

L'avènement de l'IA générative a considérablement modifié le rapport à cette notion de sens. Grâce aux modèles de langage pré-entraînés (comme GPT et ses successeurs), il est désormais possible de générer des traductions qui vont bien au-delà d'une simple correspondance lexicale. Ces systèmes sont capables de capter des schémas (patterns) complexes dans les corpus textuels, d'identifier des relations sémantiques subtiles et même de proposer des reformulations qui reprennent la tonalité, le registre et l'atmosphère de l'original.

Par ailleurs, l'IA générative oblige à repenser le rôle du traducteur en tant qu'interprète du sens. Autrefois, la compréhension du texte reposait exclusivement sur l'expertise humaine et sur des compétences culturelles et linguistiques accumulées au fil du temps. Aujourd'hui, les modèles d'IA, entraînés sur des milliards de textes issus de diverses sources, disposent d'une « mémoire illimitée » et d'une compréhension contextuelle qui leur permettent d'identifier, par exemple, qu'un même mot peut porter plusieurs sens en fonction du contexte. Ainsi, l'IA peut produire des traductions qui respectent la polysémie et

l'ambiguïté inhérente au langage.

Prenons l'exemple d'un extrait littéraire qui joue sur les ambivalences du terme « lumière ». Dans un texte poétique, ce mot peut signifier à la fois la clarté physique et une dimension symbolique (l'illumination, la révélation). Un traducteur humain, confronté à ce choix, interprétera le texte en fonction de son expérience et de son ressenti. Un système d'IA générative, quant à lui, analysera des milliers de contextes dans lesquels « lumière » est utilisé et pourra proposer plusieurs versions de traduction, par exemple :

- *"The light shone, both revealing the path and igniting a spark of hope."*
- *"A radiant glow not only illuminated the surroundings but also awakened a sense of wonder."*

Ces propositions démontrent que l'IA ne se contente pas de transposer des mots, mais offre une palette de rendus qui captent diverses facettes du sens. La traduction devient ainsi un espace de variations où le sens est interprété de manière dynamique, suggérant que l'interprétation du sens n'est pas un acte univoque mais un processus de négociation entre les multiples potentialités du langage.

Vers une théorie du sens « assistée par IA »

L'intégration de l'IA dans le processus de traduction amène à envisager une nouvelle théorie interprétative du sens, dans laquelle l'humain et la machine collaborent pour extraire et restituer la signification d'un texte. Dans ce modèle hybride, l'IA sert d'outil analytique permettant de dégager des tendances et des patterns sémantiques, tandis que l'expertise humaine intervient pour affiner la traduction, adapter le rendu aux exigences culturelles et décider du registre approprié.

Le traducteur, dans ce nouveau paradigme, joue le rôle de superviseur et d'éditeur critique. Après que l'IA générative ait proposé un ensemble de traductions, l'humain évalue la fidélité du rendu au sens global du texte, corrige les éventuelles incohérences et ajuste le style. Ce processus itératif permet d'aboutir à une traduction qui respecte la richesse interprétative de l'original tout en bénéficiant de la capacité de l'IA à traiter de grandes quantités d'information et à détecter des nuances subtiles.

Cette approche hybride influence également la formation des traducteurs. La maîtrise des outils d'IA devient indispensable, tout comme la capacité à interpréter les scores de similarité sémantique générés par des algorithmes d'apprentissage profond. Les traducteurs du futur devront conjuguer leur savoir-faire traditionnel avec une compréhension technique des modèles de traitement du langage, afin de naviguer efficacement entre les rendus générés automatiquement et les exigences d'une traduction de qualité.

La théorie du skopos et l'influence de l'IA

La théorie du skopos, issue des travaux de Hans Vermeer dans les années 1970, repose sur l'idée que la traduction doit être guidée par la finalité (le skopos) de l'acte de traduction. Plutôt que de chercher une correspondance rigide avec le texte source, le traducteur doit avant tout répondre aux attentes et aux besoins du public cible. Cette approche fonctionnaliste met en exergue l'importance du contexte d'utilisation du texte traduit, la culture du public destinataire et les objectifs communicatifs spécifiques qui orientent le choix des stratégies traductionnelles.

Selon la théorie du skopos, le succès d'une traduction se mesure par son aptitude à remplir sa fonction dans le nouveau contexte, ce qui implique souvent des adaptations et des transformations qui s'éloignent d'une fidélité formelle stricte. La finalité de la traduction (informative, persuasive, esthétique, etc.) devient ainsi le critère déterminant de la stratégie à adopter.

L'intelligence artificielle générative, par sa capacité à produire plusieurs versions d'un même texte, offre de nouvelles perspectives pour satisfaire les exigences du skopos. En effet, les systèmes d'IA peuvent être configurés pour tenir compte de paramètres spécifiques liés au contexte d'usage, permettant ainsi d'ajuster le style, le ton et même la structure du texte traduit en fonction des besoins du public cible.

L'un des atouts majeurs de l'IA générative est la possibilité de personnaliser la traduction. En intégrant des métadonnées sur le public cible ou des indications sur l'usage prévu (par exemple, un texte publicitaire, un document technique ou un contenu littéraire), le modèle peut générer des traductions adaptées qui optimisent la réception du message. Cette flexibilité s'inscrit parfaitement dans la logique du skopos, où l'objectif final de la traduction prime sur une simple correspondance textuelle.

Considérons la traduction d'un slogan publicitaire. Dans un contexte où l'objectif est de captiver un public jeune et dynamique, l'IA générative peut proposer des rendus qui adoptent un ton informel et percutant, par exemple :

- Version A (plus classique) : *"Experience the thrill of every moment."*
- Version B (adaptée à un public jeune) : *"Feel the rush – live every moment to the max!"*

Ici, la décision entre ces versions dépend du skopos défini par le commanditaire de la traduction. Le modèle d'IA, en tenant compte des données sur le public cible, permet de générer des options qui respectent la finalité communicative du message. Le traducteur peut alors sélectionner ou ajuster la version la plus appropriée.

Là où autrefois le traducteur devait s'appuyer uniquement sur sa connaissance du public et sur son intuition pour définir le skopos, l'IA offre désormais des outils analytiques

capables d'identifier des tendances culturelles et comportementales à partir de vastes ensembles de données. Cela permet de calibrer les traductions de manière plus précise et de proposer des rendus qui répondent efficacement aux attentes des utilisateurs finaux.

Les modèles d'IA générative peuvent intégrer des indicateurs sur les préférences stylistiques et culturelles du public cible. Par exemple, dans le cas de la traduction de contenus web, les algorithmes peuvent analyser les retours des internautes (taux de clics, temps de lecture, interactions sur les réseaux sociaux) pour optimiser le rendu textuel. Ce processus d'auto-apprentissage rend la notion de skopos dynamique et évolutive, en phase avec les évolutions des usages et des technologies de communication (Guidère, 2008).

Un site e-commerce international peut utiliser une plateforme de traduction générative pour adapter ses descriptions de produits en fonction des marchés locaux. Pour un même produit, la description traduite en anglais pour le marché britannique pourra insister sur des aspects de qualité et d'héritage culturel, tandis que la version pour le marché américain pourra adopter un ton plus direct et dynamique. La décision sur le style final sera basée sur des analyses de données comportementales, permettant ainsi une adaptation fine au skopos défini par les objectifs commerciaux et culturels de l'entreprise.

Vers une théorie du skopos augmentée par l'IA

La redéfinition de la théorie du skopos à l'ère de l'IA implique une réévaluation des rôles respectifs de la finalité communicative et de la technologie. D'une part, l'IA générative offre des outils permettant de calibrer le rendu en fonction de paramètres mesurables, tels que le public cible, les tendances du marché et les attentes culturelles. D'autre part, elle ouvre la voie à une approche itérative, où la traduction est constamment affinée par des retours d'expérience et des analyses de performance.

Les systèmes d'IA, en s'appuyant sur des algorithmes d'apprentissage profond, peuvent anticiper les réactions du public et proposer des traductions qui maximisent l'impact du message. Ainsi, la théorie du skopos se voit enrichie par une dimension prédictive qui permet d'optimiser le choix des stratégies traductionnelles. Le traducteur ne se contente plus de transposer un message, mais participe à un processus itératif de validation et d'amélioration continue, soutenu par des indicateurs quantitatifs et qualitatifs.

Pour les praticiens, cette évolution signifie que la formation ne se limite plus à la maîtrise des techniques de traduction classiques. Les traducteurs doivent désormais comprendre comment interpréter et exploiter les données fournies par l'IA, collaborer avec des spécialistes des données et se familiariser avec les outils d'analyse statistique et comportementale. Ce nouveau paradigme renforce la dimension stratégique de la traduction, où l'objectif est de produire un texte qui non seulement respecte la finalité initiale, mais qui s'adapte de manière proactive aux attentes

du public cible.

Repositionnement des formations

L'IA ne fait pas disparaître la traduction mais elle déplace le centre de gravité du travail vers une collaboration humain machine où le traducteur devient davantage superviseur, post éditeur et gestionnaire de projets linguistiques hybrides. Une formation qui reste centrée sur la seule production « de zéro » risque donc de manquer le nouveau cœur de la chaîne de valeur.

Concrètement, les cursus peuvent se repositionner en organisant l'apprentissage autour de scénarios de production assistée, avec une progression qui va de l'usage raisonné des outils à la capacité d'en auditer les sorties, puis à la maîtrise de dispositifs qualité, délais, confidentialité et traçabilité adaptés à des flux de traduction partiellement automatisés.

Le métier évolue ainsi vers un profil d'expert en communication interculturelle, capable d'évaluer une sortie automatique, de la corriger, puis d'adapter le texte à un contexte, un public et un objectif. Cela s'accompagne d'une montée de la spécialisation, notamment dans les domaines sensibles ou à forte technicité, où l'intervention humaine reste recherchée. En parallèle, la pratique devient plus itérative, avec un aller-retour entre propositions générées et arbitrages humains, y compris à partir de plusieurs variantes produites par un même système, ce qui rapproche la traduction de formes de direction éditoriale et de pilotage d'options.

Dans cette perspective, la compétence linguistique reste nécessaire mais elle devient insuffisante si elle n'est pas accompagnée d'une capacité à repérer les écarts de sens fins, les faux amis pragmatiques, les pertes d'allusions culturelles et les effets de registre. L'IA demeure en difficulté sur des nuances culturelles, des références localisées et certains contextes complexes, ce qui fait de l'expertise culturelle et contextuelle un point de différenciation du professionnel.

Il y a d'abord une littératie outillée, qui ne se limite pas à « utiliser » un système, mais à comprendre ses familles de modèles et leurs comportements, y compris l'intérêt de solutions hybrides spécialisées et l'apport de ressources comme glossaires et mémoires de traduction pour mieux cadrer la production.

Il y a ensuite une compétence de post édition avancée, pensée comme une pratique de décision, où l'on choisit un rendu parmi plusieurs, puis on le fait converger vers un cahier des charges de sens, de style et d'effet.

Enfin, il y a une compétence éthique et qualité au sens large en raison des biais possibles issus des données d'entraînement et de la nécessité de transparence, de validation et de standards de contrôle.

Ainsi, un cursus peut se structurer autour d'ateliers où l'on

confronte des sorties de traduction générative à des objectifs différents, information, persuasion, conformité sectorielle, puis où l'on justifie les choix de réécriture avec des critères de similarité et d'adaptation. Les évaluations peuvent alors mesurer la capacité à détecter les erreurs à enjeu, la cohérence sur des textes longs, la gestion du style et du registre, et la pertinence culturelle, plutôt que la seule « correction » linguistique.

En parallèle, des modules de gestion de projet linguistique et de veille technologique deviennent structurants, car le traducteur doit rester à jour et maîtriser les outils, comprendre les systèmes, évaluer et améliorer les rendus.

Le traducteur garant de l'autonomie cognitive

La question de la dépendance cognitive devient un point central pour les formations. L'usage intensif des systèmes de traduction automatique et des modèles génératifs peut produire ce que l'on appelle une « dette cognitive ». Cela correspond à une situation où certaines opérations intellectuelles sont progressivement déléguées à la machine, ce qui peut affaiblir la capacité du professionnel à mobiliser seul ses compétences linguistiques, analytiques et interprétatives.

Dans le domaine de la traduction, plusieurs mécanismes cognitifs sont concernés. Le premier est l'effet d'ancrage : lorsque le traducteur reçoit une proposition produite par un système, cette proposition agit comme un point de référence implicite. Même si elle contient des erreurs, elle oriente la réflexion et limite l'exploration d'autres solutions. Ce phénomène est bien connu dans la recherche sur la post édition et il peut conduire à des corrections superficielles plutôt qu'à une réévaluation complète du sens.

Un second mécanisme est l'automatisation du jugement : plus un traducteur travaille avec un système performant, plus il peut développer une confiance implicite dans les sorties produites. Cela peut réduire la vigilance critique et entraîner une validation trop rapide de formulations incorrectes ou approximatives. Dans les textes techniques, juridiques ou médicaux, cette baisse de vigilance peut avoir des conséquences importantes.

Un troisième mécanisme est la réduction de l'effort de formulation : lorsque la première étape de génération est prise en charge par la machine, le traducteur peut perdre l'habitude de construire mentalement plusieurs hypothèses de traduction. Or cette capacité d'exploration linguistique est au cœur de l'expertise traductive. Si elle n'est plus exercée régulièrement, elle peut s'atrophier avec le temps.

Pour cette raison, les formations devraient intégrer une dimension de pédagogie cognitive. L'objectif n'est pas de limiter l'usage des outils, mais de former des professionnels capables d'en comprendre les effets sur leur propre processus de pensée.

Plusieurs pistes pédagogiques peuvent être envisagées :

Une première consiste à développer une réflexivité cognitive. Les étudiants doivent apprendre à analyser leur propre processus de traduction. Des exercices de verbalisation du raisonnement, des journaux de traduction ou des protocoles de pensée à voix haute peuvent aider à prendre conscience de l'influence des suggestions automatiques.

Une deuxième piste consiste à maintenir des espaces de traduction sans assistance. Dans un cursus dominé par les outils, il reste important de pratiquer régulièrement la traduction directe sans pré traduction automatique. Cela permet d'entretenir les compétences fondamentales de compréhension, de reformulation et de création linguistique.

Une troisième piste concerne la formation aux biais cognitifs. Les étudiants peuvent être sensibilisés à des phénomènes comme l'ancrage, la complaisance technologique, le biais d'autorité ou l'illusion de précision. Comprendre ces mécanismes permet de développer une attitude critique face aux sorties des systèmes.

Une quatrième piste est l'apprentissage de stratégies de désancrage. Il peut s'agir par exemple de masquer temporairement la proposition automatique, de produire d'abord une traduction personnelle avant de consulter la sortie de la machine, ou de comparer plusieurs systèmes pour éviter qu'une seule solution devienne la référence implicite.

Une cinquième piste concerne la diversification des activités cognitives. Les formations peuvent inclure des exercices d'analyse stylistique, de réécriture créative, d'adaptation culturelle ou de traduction inversée. Ces activités sollicitent des compétences qui ne sont pas directement reproduites par les systèmes et qui renforcent l'agilité linguistique.

Enfin, il devient pertinent d'enseigner ce que l'on pourrait appeler l'hygiène cognitive du traducteur. Cela renvoie à la gestion consciente du rapport aux outils. Le traducteur expert ne se contente pas d'utiliser un système. Il sait quand l'utiliser, quand s'en détacher, quand vérifier ailleurs et quand repartir du texte source pour reconstruire complètement la traduction.

Conclusion

L'essor de l'intelligence artificielle générative incite à repenser en profondeur les théories de la traduction. La théorie interprétative du sens, qui valorisait le rôle du traducteur en tant qu'interprète de la signification profonde d'un texte, se trouve enrichie par la capacité de l'IA à extraire et à restituer des nuances sémantiques multiples. L'IA offre une pluralité de rendus qui invitent à voir le sens comme un continuum de significations possibles, réinterprétées à la lumière des contextes culturels et stylistiques.

De son côté, la théorie du skopos, centrée sur la finalité communicative, voit son champ d'application élargi par des outils capables de personnaliser et d'optimiser le texte traduit pour un public cible précis. Les capacités analytiques

et prédictives de l'IA permettent de mesurer l'impact des traductions et d'ajuster le rendu de manière itérative, en conciliant les exigences commerciales, culturelles et stylistiques.

Ainsi, la révolution de l'IA appelle non seulement à une révision des concepts et des modèles théoriques traditionnels, mais aussi à une redéfinition du rôle du traducteur. Loin de se substituer à l'expertise humaine, l'IA apparaît comme un outil complémentaire qui permet de repenser le processus de traduction de manière plus dynamique, collaborative et adaptée aux exigences du monde globalisé. Dans ce contexte, le traducteur devient à la fois médiateur, éditeur et stratège, capable de naviguer entre la richesse interprétative du sens et l'optimisation des rendus en fonction des besoins spécifiques du public.

En somme, la traduction générative ouvre la voie à une nouvelle ère de la traductologie, dans laquelle les réflexions classiques se dissolvent pour laisser place à un continuum de similarités et de variations, reflétant la complexité des langues, des cultures et des intentions communicatives. Cette redéfinition offre non seulement des opportunités inédites pour améliorer la qualité des traductions, mais invite également à repenser le rôle du traducteur et les critères d'évaluation de la traduction, en mettant en avant la richesse du dialogue entre l'original et sa nouvelle vie dans une langue différente.

Dans ce contexte, la compétence clé du futur traducteur n'est pas seulement linguistique ou technologique ; elle devient métacognitive. Le professionnel doit être capable de surveiller son propre raisonnement, d'identifier les influences de la machine sur ses décisions et de maintenir un niveau d'autonomie intellectuelle suffisant pour garantir la qualité et la responsabilité de son travail.

Références

- Guidère, M. & Jehel, L. (2025). *Phénotypage du suicide : un nouvel horizon pour la prévention et le traitement. / Suicide Phenotyping : a New Horizon for Prevention and Treatment*. EMC Psychiatrie. Elsevier Masson. 37-500-A-25. DOI : 10.1016/S0246-1072(25)50043-3. <https://www.em-consulte.com/article/1720324>
- Guidère M. (2024). *Rethinking Language in Mental Health: Multilingual Approaches to Mental Health*. Montreal: Psynum. ISBN 979-8335186100.
- Guidère M. (2023). *The Language Within: Exploring Mental Health Through Predictive Linguistics*. Montreal. ISBN 979-8876838797.
- Guidère M. (2020). *La Traduction médicale à l'heure de la pandémie*. Paris : L'Harmattan. Collection Traductologie.
- Guidère M. (2016). *Introduction à la traductologie : penser la traduction, hier, aujourd'hui, demain*.

Bruxelles / Paris : De Boeck Supérieur, 3^e édition.

Guidère M. (2015) dir. *Traductologie et Géopolitique*.
Paris : L'Harmattan.

Guidère M. (2008). *La Communication multilingue :
traduction commerciale et institutionnelle*. Bruxelles /

Paris : De Boeck Supérieur. Collection Traducto.

Guidère M. (2000). *Publicité et Traduction*. Paris :
L'Harmattan.

SOMMAIRE

Introduction	28
La transformation du rôle du traducteur	28
Définition de la traduction générative	28
Qu'est-ce que l'IA générative ?	29
Comment fonctionne la traduction générative ?	29
Exemples concrets de traduction générative	29
Les atouts de la traduction générative	30
Les limites de la traduction générative	30
Perspectives d'avenir de la traduction générative	31
Les modèles de traduction générative	31
Classification des modèles de traduction générative	31
Exemples illustratifs de modèles génératifs en traduction	32
Atouts et limites des différents modèles de traduction générative	33
Redéfinition du concept d'équivalence	34
Évolution du concept d'équivalence en traduction	34
Exemples illustrant la redéfinition de l'équivalence	35
Vers une approche nuancée de la traduction	35
Vers une nouvelle pratique de la traduction	36
Redéfinition des dichotomies classiques	36
Théorie versus Pratique	36
Traduisible versus Intraduisible	36
Art versus Science	37
Auteur versus Traducteur	37
Original versus Copie	37
Traduction versus Imitation	37
Sacré versus Profane	38
Fidélité versus Liberté	38
Le Mot versus l'Idée	38
La Lettre versus L'Esprit	38
Repenser les théories de la traduction	39
La théorie interprétative du sens et l'impact de l'IA	39
La théorie du skopos et l'influence de l'IA	40
Repositionnement des formations	41
Le traducteur garant de l'autonomie cognitive	41
Conclusion	42
Références	42