



REPUBLIQUE ALGERIENNE DEMOCRATIQUE ET POPULAIRE
MINISTRE DE L'ENSEIGNEMENT SUPERIEUR ET DE LA RECHERCHE SCIENTIFIQUE

UNIVERSITE IBN KHALDOUN - TIARET

MEMOIRE

Présenté à :

FACULTÉ DES MATHEMATIQUES ET D'INFORMATIQUE
DÉPARTEMENT D'INFORMATIQUE

Pour l'obtention du diplôme de :

MASTER

Spécialité : Intelligence Artificielle et Digitalisation

Par :

SENOUCI Ahmed
BOUAZA Mohamed

Sur le thème

Exploitation des Réseaux Sociaux pour l'Analyse des Sentiments à Base d'Aspects

Soutenu publiquement le 15/06/2025 à Tiaret devant le jury composé de :

Mr MEGHAZI Hadj Madani

MCB U.I.K.Tiaret

Président

Mme BOUBEKEUR Aicha

MAA U.I.K.Tiaret

Encadrant

Mr BERBER El Mehdi

MAA U.I.K.Tiaret

Examineur

2024-2025

Résumé:

Avec les avancées récentes de l'intelligence artificielle générative, les avis de clients en ligne sont devenus une source de connaissances incontournable pour diverses contributions et recherches scientifiques. L'analyse des sentiments, et en particulier l'analyse des sentiments à base aspects (ABSA), permet l'analyse fine des opinions en associant chaque sentiment à un aspect ou caractéristique spécifique d'un produit ou service quelconque.

Ce mémoire s'inscrit dans ce contexte et vise à concevoir et développer un outil d'ABSA appliquée aux avis des clients pour le domaine de la restauration. Contrairement aux approches d'analyses classiques, notre solution repose sur l'utilisation de larges modèles de langue (LLMs), exploitant des instructions d'incitation et de fine-tuning, pour adapter l'analyse au contexte des restaurants.

Les résultats expérimentaux montrent que le LLM, affiné et bien orienté, offre une analyse fine et contextualisée, tout en réduisant le besoin en annotation manuelle. Ce travail ouvre également des perspectives pour améliorer la robustesse, l'adaptabilité et l'application de cette approche à d'autres domaines.

Mots clefs: IA générative, LLM, ABSA, instruction d'incitation, finetuning.

Abstract:

With recent advances in generative artificial intelligence, online customer reviews have become an essential source of knowledge for various scientific contributions and research. Sentiment analysis, and in particular aspect-based sentiment analysis (ABSA), allows for detailed analysis of opinions by associating each sentiment with a specific aspect or characteristic of a given product or service.

This thesis fits into this context and aims to design and develop an ABSA tool applied to customer reviews for the restaurant industry. Unlike traditional analysis approaches, our solution relies on the use of large language models (LLMs), leveraging prompting and fine-tuning instructions to adapt the analysis to the restaurant context.

Experimental results show that the finetuned LLM and well-oriented offers detailed and contextualised analysis, while reducing the need for manual annotation. This work also opens up prospects for improving the robustness, adaptability and application of this approach to other fields.

Keywords: Generative AI, LLM, ABSA, prompting, fine-tuning.

ملخص :

مع التطورات الأخيرة في مجال الذكاء الاصطناعي التوليدي، أصبحت آراء العملاء عبر الإنترنت مصدرًا أساسيًا للمعلومات تُعتمد عليه في العديد من الأبحاث العلمية. يتيح تحليل المشاعر، ولا سيما التحليل القائم على الجوانب (ABSA)، فهماً دقيقاً للآراء من خلال ربط كل شعور بجانب أو خاصية محددة لمنتج أو خدمة معينة.

تهدف هذه الأطروحة إلى تصميم وتطوير أداة ABSA مخصصة لتحليل آراء العملاء في قطاع المطاعم. بخلاف المناهج التقليدية، يعتمد الحل المقترح على نماذج لغوية كبيرة (LLMs) مع توظيف تقنيات التعليم بالتوجيه والضبط الدقيق، بما يتيح تكييف التحليل مع خصوصية سياق المطاعم.

تُظهر النتائج التجريبية أن ال LLM المُحسن والمُوجّه جيّدًا توفر تحليلًا دقيقًا ومحدد السياق، مع تقليل الحاجة إلى الشرح اليدوي. يفتح هذا العمل أيضًا آفاقًا لتحسين متانة هذا النهج وقابليته للتكيف وتطبيقه على مجالات أخرى.

الكلمات المفتاحية: الذكاء الاصطناعي التوليدي، LLM، ABSA، التوليد المعزز/الموجه، تعليمات التحفيز، الضبط الدقيق.

Remerciements

Tout d'abord, nous tenons à exprimer notre profonde gratitude à notre directeur de projet, Mme. BOUBEKEUR Aicha, Professeur à l'université d'Ibn Khaldoun de Tiaret, Faculté des Mathématiques et d'Informatique, Département d'Informatique. Son accompagnement précieux, ses conseils avisés et sa disponibilité constante ont été des piliers essentiels tout au long de la réalisation de ce travail. Sa confiance et son soutien nous ont permis de surmonter les défis rencontrés et de mener à bien cette étude.

Nos remerciements vont aussi aux membres du jury qui ont accepté de consacrer leur temps et leur expertise à l'évaluation de ce travail

Nous tenons à exprimer notre reconnaissance envers tous nos collègues de travail et nos amis avec qui nous avons eu l'opportunité d'échanger. Leur soutien, leur bienveillance et les moments partagés ont rendu cette expérience d'autant plus enrichissante.

Un merci tout particulier à l'ensemble de l'équipe de la Faculté des Mathématiques et d'Informatique, Département d'Informatique de l'Université Ibn Khaldoun de Tiaret, pour l'environnement d'apprentissage stimulant et les ressources mises à disposition, qui ont été indispensables à la réalisation de ce projet.

Enfin, nous ne saurions terminer sans remercier nos familles respectives. Leur amour inconditionnel, leur patience, leurs encouragements constants et leur présence indéfectible ont été une source de motivation inépuisable. Ce travail leur est dédié.

Table des matières

	Page
Introduction générale	01
Chapitre 1 – L’analyse des sentiments	
1.1 Introduction	04
1.2 Analyse des sentiments – définitions	04
1.3 Chronologie de l’analyse des sentiments	05
1.4 Approches de l’analyse des sentiments	06
1.4.1 Approche lexicale	06
1.4.2 Approche d’apprentissage automatique traditionnel	07
1.4.3 Approche d’apprentissage automatique profond	07
1.4.4 Approche hybride	08
1.5 Niveaux d’analyse des sentiments	10
1.6 Applications de l’analyse des sentiments	11
1.7 Processus d’analyse des sentiments	12
1.8 Défis de l’analyse des sentiments	16
1.9 Conclusion	18
Chapitre 2 – Analyse des sentiments à base d’aspects	
2.1 Introduction	21
2.2 Analyse des sentiments à base d’aspects	21
2.3 Tâches ABSA	22
2.3.1 Tâches simples de l’ABSA (Single ABSA Tasks)	22
2.3.1.1 Extraction des termes d’aspect (ATE)	23
2.3.1.2 Catégorisation des termes d’aspect	23
2.3.1.3 Extraction de termes d’opinion	23
2.3.2 Tâches composées de l’ABSA	24
2.3.2.1 Extraction de paires aspect-opinion (AOPE)	24

Table des matières

2.3.2.2 Extraction de triplets aspect-opinion-sentiment (ASTE)	24
2.3.2.3 Prédiction de quadruplets aspect-opinion-sentiment (ASQP)	24
2.4 Evolution de l'ABSA	25
2.4.1 Modèles basés sur des transformateurs	28
2.5 Jeux de données de référence pour ABSA	30
2.6 Métriques de performance	32
2.7 Défis de l'ABSA	35
2.8 Large Modèles de Langue (LLM)	35
2.8.1 Principaux constituants des LLM	36
2.8.2 Fonctionnement des LLM	36
2.8.3 Apprentissage par transfert à partir des LLM	37
2.8.4 Méthodes de fine-tuning des LLM	38
2.8.5 Avantages et inconvénients de l'apprentissage par transfert	39
2.8.6 Domaines d'application des LLM	40
2.9 Relation entre LLM et apprentissage automatique	40
2.9.1 Facteurs clés de l'IA générative et des LLM	41
2.10 LLM dans l'analyse des sentiments	42
2.11 LLM dans l'apprentissage en contexte	42
2.11.1 Facteurs clés pour le choix des LLM	42
2.12 Conclusion	43

Chapitre 3 – Implémentation et résultats

3.1 Introduction	45
3.2 Formulation du problème ABSA	45
3.3 Choix techniques utilisés	46
3.4 Architecture globale du système ABSA	46
3.5 Structure globale du pipeline ABSA	47

Table des matières

3.6 Structure technique du pipeline ABSA	48
3.7 Jeu de données utilisé	50
3.8 Processus séquentiel détaillé de la solution ABSA	51
3.9 Résultats expérimentaux	53
3.9.1 Dataset après formatage	56
3.9.2 Présentation et évaluation des résultats	56
3.9.3 Perte d'entraînement	57
3.10 Conclusion	59
Conclusion générale	60
Bibliographie	61

Table des matières

	Page
Tableau 2.1 : Comparaisons de méthodes d'analyse des sentiments	08
Tableau 2.1 : Tâches simples et composées en ABSA	26
Tableau 2.2 : Jeux de données de référence	31
Tableau 3.1 : Choix techniques utilisées	47
Tableau 5 : Métriques de performances	58

Table des figures

	Page
Figure 1.1 : Différentes phases/niveaux de l'analyse des sentiments	10
Figure 1.2 : Étapes fondamentales de l'analyse des sentiments	13
Figure 2.1 : Quatre éléments de sentiment clés en ABSA	22
Figure 2.2 : Architecture du transformateur	28
Figure 2.3 : Matrice de confusion	32
Figure 3.1 : Architecture globale du système ABSA	47
Figure 3.3 : Dataset public NEUDM/absa-quad	52
Figure 3.4 : Processus de formatage des données	53
Figure 3.5 : Distribution des sentiments dans le jeu de données d'entraînement	55
Figure 3.6 : Tableau de distribution des catégories prédéfinies	55
Figure 3.7 : Répartition des sentiments par catégorie	56
Figure 3.8 : Evolution de la perte d'entraînement (Training Loss)	60

Table des abréviations

Abréviation Signification en français

AS	Analyse des sentiments
ABSA	Analyse des sentiments à base d'aspects
ATE	Extraction des termes d'aspects
OTE	Extraction des termes d'opinions
ATC	Classification des termes d'aspects
TOTE	Extraction des termes orientés vers la cible (<i>Target-Oriented Term Extraction</i>)
AOCE	Extraction contextuelle orientée vers les aspects
AOPE	Extraction des paires aspect-opinion
ASQP	Prédiction du quadruplet aspect-sentiment
LSTM	Mémoire à long terme (<i>Long Short-Term Memory</i>)
GRU	Unité récurrente simplifiée (<i>Gated Recurrent Unit</i>)
RNN	Réseau de neurones récurrent
CNN	Réseau de neurones convolutif
CBOW	Moyenne des mots en contexte (<i>Continuous Bag Of Words</i>)
SVM	Machine à vecteurs de support (<i>Support Vector Machine</i>)
GPT	Transformeur préentraîné génératif (<i>Generative Pre-trained Transformer</i>)
BERT	Représentations bidirectionnelles encodées à partir des transformeurs (<i>Bidirectional Encoder Representations from Transformers</i>)
NLP	Traitement automatique du langage naturel (<i>Natural Language Processing</i>)
DL	Apprentissage profond (<i>Deep Learning</i>)
LLM	Grand modèle de langage (<i>Large Language Model</i>)
ChatGPT	Agent conversationnel développé par OpenAI
OpenAI	Organisation développant les modèles GPT
GPU	Processeur graphique (<i>Graphics Processing Unit</i>)
TPU	Unité de traitement Tensoriel (<i>Tensor Processing Unit</i>)
PEFT	Ajustement efficace des paramètres (<i>Parameter-Efficient Fine-Tuning</i>)
TF-IDF	Fréquence de terme – Fréquence inverse de document (Term Frequency – Inverse Document Frequency)

Table des abréviations

BoW	Sac de mots (<i>Bag of Words</i>)
LoRA	Adaptation à faible rang (<i>Low-Rank Adaptation</i>)
RLHF	Apprentissage par renforcement à partir du retour humain (<i>Reinforcement Learning from Human Feedback</i>)
DAPT	Préentraînement supplémentaire sur des données spécifiques au domaine (<i>Domain-Adaptive Pre-Training</i>)

Introduction générale

A l'ère de l'IA actuelle, les plateformes interactives deviennent de plus en plus un espace incontournable d'expression et de communication pour leurs utilisateurs. Chaque jour, des millions d'utilisateurs partagent spontanément leurs opinions, critiques et recommandations à propos de produits ou de services, constituant ainsi une source riche et précieuse d'informations. Dans ce contexte, l'analyse des sentiments et plus spécifiquement l'analyse des sentiments à base d'aspects (*Aspect-Based Sentiment Analysis*, ou ABSA) s'impose comme une approche clé pour extraire, structurer et interpréter ce qui est partagé en ligne.

Contrairement à l'analyse des sentiments traditionnelle, qui se limite à déterminer le sentiment global d'un texte (positif, négatif ou neutre), l'ABSA permet de capturer des opinions plus fines et nuancées en associant les sentiments exprimés à des aspects ou caractéristiques spécifiques d'un produit ou service (par exemple, la qualité du service, le goût des plats ou l'ambiance dans un restaurant). Cette granularité fine offre une vision plus contextualisée des préférences et insatisfactions des clients, particulièrement utile pour des secteurs comme la restauration, où l'expérience client repose sur une multitude de critères.

Les approches classiques en ABSA reposent souvent sur des modèles d'analyse de sentiments simples pour l'extraction des aspects, la détection de sentiments, ou encore la reconnaissance des relations entre aspects et opinions. Ces approches exigent généralement un développement spécifique pour chaque composant, ainsi qu'un volume important de données annotées, ce qui limite leur portabilité et leur efficacité dans des domaines peu dotés en ressources.

Dans ce mémoire, nous adoptons une approche innovante en nous appuyant sur les larges modèles de langue, qui ont récemment démontré leurs performances remarquables dans diverses tâches de traitement du langage naturel. Plutôt que de concevoir des modèles spécifiques pour chaque sous-tâche de l'ABSA, nous exploitons les capacités avancées du modèle de langue, en les adaptant au contexte ciblé via des techniques d'apprentissage par transfert et de fine-tuning. Cette stratégie permet d'exploiter les connaissances linguistiques pré-acquises pour effectuer une analyse fine, contextualisée et plus facilement généralisable, même dans des contextes spécialisés comme celui de la restauration.

L'objectif principal de ce travail est ainsi de concevoir et d'évaluer une méthode ABSA appliquée aux avis reformatés à partir de datasets ABSA déjà existants, notamment ceux des concours SemEval (comme SemEval-2014 Task 4) dans le domaine de la restauration (restaurant reviews). Ces critiques ont été collectées de plateformes sociales de restaurants américains, principalement : « Yelp.com » et d'autres sources similaires de critiques gastronomiques. Ils sont réels, écrits par des utilisateurs, puis annotés manuellement pour créer des jeux de données de référence dans les tâches d'ABSA. En s'appuyant sur les LLM, nous cherchons à démontrer qu'il est possible d'améliorer la précision et la pertinence de l'analyse tout en réduisant les besoins en annotation manuelle et en développement de modèles spécifiques.

Ce mémoire est structuré comme suit : nous commencerons par une revue de la littérature sur l'analyse des sentiments, les approches ABSA et l'émergence des LLM dans ce domaine. Nous décrirons ensuite notre méthodologie, les données utilisées et les choix techniques effectués. Les résultats expérimentaux seront ensuite présentés et discutés, avant de conclure par une synthèse des apports de ce travail et des perspectives futures.

Chapitre 1

L'Analyse des sentiments

1.1 Introduction	04
1.2 Analyse des sentiments – définitions	04
1.3 Chronologie de l'analyse des sentiments	05
1.4 Approches de l'analyse des sentiments	06
1.4.1 Approche lexicale	06
1.4.2 Approche d'apprentissage automatique traditionnel	07
1.4.3 Approche d'apprentissage automatique profond	07
1.4.4 Approche hybride	08
1.5 Niveaux d'analyse des sentiments	10
1.6 Applications de l'analyse des sentiments	11
1.7 Processus d'analyse des sentiments	12
1.8 Défis de l'analyse des sentiments	16
1.9 Conclusion	18

1.1 Introduction

Dans l'ère du monde virtuel numérique, le volume important de contenus utilisateur générés par l'accroissement constant des plateformes sociales et des activités en ligne a ouvert la nouvelle voie pour les interactions et les échanges par le biais de moyen de communication non structuré (message, post, commentaire, chat, Blogs, critiques, ... etc.). Si le traitement est effectué manuellement, la gestion et l'analyse d'un tel volume de contenu généré prend beaucoup de temps, et l'une des méthodes les plus populaires pour traiter les sentiments d'un contenu non structuré est l'analyse des sentiments (AS). L'AS fournit un outil automatique, rapide et efficace pour identifier les opinions et les sentiments des évaluateurs. En outre, ce texte est la plus importante source d'information pour comprendre les sentiments du grand public pour une variété de domaines d'application tels que les produits, les restaurants, les films, les questions politiques, les plans éducatifs ou autres.

Ce chapitre explore non seulement les méthodes bien établies comme l'apprentissage automatique, l'apprentissage profond et les approches basées sur le lexique, mais aussi d'autres techniques d'AS comme les approches hybrides et basées sur les transformateurs avancés.

1.2. Analyse des sentiments – définitions

Dans la littérature scientifique [1], l'analyse des sentiments (AS) est désignée sous diverses dénominations, telles que l'extraction d'opinions (EO), l'étude de la subjectivité, l'analyse émotionnelle ou encore l'extraction des évaluations. Parmi ces termes, « analyse des sentiments » et « exploration d'opinions » sont les plus couramment employés. Selon [1], ces deux notions sont souvent perçues comme similaires et relèvent du même champ d'étude, qui peut être inclus dans le domaine plus large de l'analyse de la subjectivité. Les auteurs dans [[1] précise que l'analyse des sentiments constitue un domaine de recherche rattaché à l'exploration textuelle, définie comme le traitement algorithmique des opinions, des émotions et des aspects subjectifs contenus dans un texte. Il n'est donc pas surprenant que certains chercheurs aient exprimé une certaine confusion quant à la distinction entre « sentiment » et « opinion », débattant ainsi de la terminologie appropriée : doit-on parler d'analyse des sentiments ou d'extraction d'opinions ?

Selon le dictionnaire Merriam-Webster Collegiate [2] , un sentiment correspond à une attitude, une pensée ou un jugement influencé par une émotion, tandis qu'une opinion est décrite comme une vue, un jugement ou une appréciation formulée mentalement à propos d'un sujet particulier. La différence entre ces deux concepts est subtile, bien que chacun contienne des éléments de l'autre. En général, une opinion représente une position concrète adoptée par une personne vis-à-vis d'un sujet donné, tandis qu'un sentiment se rattache davantage à une émotion ressentie. Par exemple, la phrase « Je suis inquiet au sujet de la situation politique actuelle » traduit un sentiment, alors que « Je pense que la politique actuelle ne fonctionne pas bien » exprime une opinion. Si quelqu'un énonce la première phrase lors d'une discussion, on pourrait répondre : « Je partage ce sentiment. » En revanche, face à la deuxième phrase, on dirait plutôt : « Je suis d'accord » ou « Je ne suis pas d'accord. »

Cependant, les significations implicites des deux phrases sont étroitement liées. Le sentiment exprimé dans la première phrase est probablement motivé par l'opinion formulée dans la deuxième. Inversement, la première phrase suggère implicitement une opinion négative sur la politique, confirmée par la deuxième phrase. Bien que les opinions soient souvent accompagnées de sentiments, ce n'est pas toujours le cas. Par exemple, une phrase comme « Je pense qu'il gagnera les prochaines élections présidentielles » exprime une opinion sans nécessairement refléter une émotion particulière.

En résumé, l'analyse des sentiments désigne une méthode permettant de suivre et d'évaluer les avis des utilisateurs concernant un objet ou un sujet spécifique. Elle implique la mise en place d'un cadre structuré pour collecter et analyser les avis exprimés dans divers supports, tels que des articles de blog, des commentaires, des critiques ou des publications sur les réseaux sociaux.

1.3 Chronologie de l'analyse des sentiments

À l'origine, l'analyse des sentiments est apparue comme un domaine d'intérêt au milieu des années 1990, lorsque les chercheurs ont commencé à chercher des techniques pour traiter automatiquement les émotions et les sentiments dans un texte. Dans un premier temps, il y avait des approches basées sur des règles, visant à classifier le texte comme positif, négatif ou neutre à l'aide de recherches de mots-clés. Plus tard, au début des années 2000, des méthodes plus avancées sont apparues, telles que l'utilisation de machines à vecteurs de support (SVM) et

de modèles bayésiens naïfs dans l'analyse des sentiments, grâce à la prévalence de données étiquetées par caractéristiques, comme les critiques de films. Cela a marqué un changement vers le choix de caractéristiques issues de modèles capables d'apprendre à partir des données. La prochaine avancée majeure en termes de performance a commencé dans les années 2010 avec les techniques d'apprentissage profond, notamment les réseaux LSTM (Long Short-Term Memory), qui permettent une meilleure capture du contexte dans les données textuelles. En 2018, des modèles de type "Transformer", tels que BERT (Bidirectional Encoder Representations from Transformers) et GPT (Generative Pre-trained Transformer), ont vu le jour, redéfinissant considérablement l'analyse des sentiments, car ils surpassent en précision et en flexibilité d'application grâce à l'utilisation d'embeddings sensibles au contexte pré-entraîné. Depuis, le domaine s'est encore développé, notamment vers la détection multilingue des émotions et l'analyse des sentiments, en mettant l'accent sur l'équité dans l'IA. Cela illustre le progression conceptuelle d'un paradigme initial simple basé sur une approche à base des règles à un modèle plus sophistiqué basé sur l'analyse de grands ensembles de données, capturant ainsi la richesse du langage humain.

1.4 Approches de l'AS

En général, les approches d'analyse sentiments peuvent être divisées en quatre catégories : les approches basées sur le lexique, l'apprentissage automatique traditionnel et profond et les approches hybrides; où diverses contributions se sont allées chercher les meilleures méthodes pour accomplir la tâche avec plus de précision et moins de dépenses informatiques [3]:

1.4.1. Approche lexicale

Cette méthode nécessite l'utilisation d'un lexique de sentiments qui attribue des scores aux jetons collectés [4]. Il s'agit d'une technique non supervisée, et cette méthode repose sur la dépendance au domaine, car un même mot dans différents domaines peut avoir des significations différentes. Le problème peut être surmonté en adoptant une technique d'adaptation du dictionnaire. Les approches basées sur le lexique sont généralement divisées en deux méthodes : les méthodes basées sur le dictionnaire et les méthodes basées sur le corpus (dictionary-based and corpus-based methods) [5].

1.4.2. Approche d'apprentissage automatique traditionnel

Les méthodes d'apprentissage automatique séparent les ensembles de données en ensembles de données d'entraînement et en ensembles de données de test [06]. Les modèles d'apprentissage automatique peuvent apprendre de nombreuses informations connues à partir des ensembles de données d'entraînement. Les ensembles de données de test sont ensuite utilisés pour analyser l'effet des modèles. Les classificateurs d'apprentissage automatique conventionnels, tels que Naive Bayes (NB) [7], les machines à vecteurs de support (SVM) [8], l'entropie maximale (EM) [9], les arbres de décision (DT) [10], K-nearest neighbors (KNN) [11], la régression logistique (LR) [12] peuvent être utilisés pour la classification des sentiments, chacun présentant des avantages et des inconvénients. En apprentissage automatique, il existe deux principales approches pour analyser les sentiments : Apprentissage automatique non supervisé et Apprentissage automatique supervisé.

1.4.3 Approche d'apprentissage automatique profond

Les algorithmes d'apprentissage profond se sont imposés comme les précurseurs de l'analyse des sentiments, surpassant les méthodes traditionnelles. L'apprentissage profond représente un réseau de neurones artificiels (RNA) composé de trois couches ou plus, conçu pour gérer de vastes ensembles de données ainsi que leurs caractéristiques complexes telles que la non linéarité et les modèles complexes. Il effectue la transformation et l'extraction automatiques de caractéristiques, similaires aux processus d'apprentissage humain, en traversant plusieurs couches cachées. De nombreux modèles d'apprentissage profond privilégient les embeddings de mots comme caractéristiques d'entrée, qui peuvent être acquis à partir de données textuelles à l'aide de techniques telles que Word2Vec, la couche d'intégration ou les vecteurs GloVe. Word2Vec peut être entraîné par des méthodes telles que le modèle Continuous SkipGram ou CBOW.

Les algorithmes d'apprentissage profond courants comprennent : Le réseau neuronal convolutif (CNN) est un réseau neuronal à propagation directe (feed-forward neural Network) avec calcul convolutif [13]. CNN est né dans le domaine de la vision par ordinateur et s'est ensuite étendu à de nombreux domaines, tels que le PNL. Le réseau neuronal récurrent (RNN) [14] a été largement utilisé en SA car il garantit que les informations relatives à une longue

séquence étaient capturées et mémorisées [15]. L'avantage le plus remarquable du RNN est qu'il utilise les connaissances antérieures et peut ainsi mémoriser les informations précédentes. Les réseaux à mémoire à long terme (LSTM) sont un type particulier de RNN, qui peut résoudre les problèmes d'explosion de gradient et de disparition [16]. De plus, le LSTM intègre un mécanisme de filtrage pour résoudre la dépendance à longue distance que les RNN ordinaires ne peuvent pas résoudre. Un LSTM mis en [17] peut stocker des informations loin de leur emplacement actuel dans une séquence. Les structures des réseaux d'unités récurrentes fermées (GRU) et du LSTM sont similaires, et leurs effets sont presque satisfaisants. L'avantage de GRU est que le modèle est plus simple et la vitesse de convergence est plus rapide. Mais lorsque la quantité de données est importante, le LSTM fonctionne mieux car il y a plus de portes et plus de paramètres [18].

Le Transformer a été proposé par Google en 2017 et est l'un des modèles d'apprentissage profond les plus puissants [19]. Transformer résout le problème de séquence à séquence en remplaçant LSTM par une structure d'attention complète, ce qui permet d'obtenir de meilleurs résultats et de réduire la complexité de calcul. Transformer utilise le mécanisme d'auto-attention pour la modélisation. Sans limitation de mémoire et de puissance de calcul, le Transformer peut théoriquement coder du texte infiniment long. Cependant, comme la quantité de calcul d'attention est énorme et que la complexité de calcul et la longueur de la séquence sont $O(n^2)$, la longueur de la séquence augmente et la consommation de mémoire et de calcul augmente rapidement.

1.4.4 Approche hybride

L'approche hybride combine des approches basées sur le lexique, l'apprentissage automatique traditionnel et l'apprentissage profond [20]. De cette manière populaire, l'analyse des sentiments combine l'analyse linguistique et la sémantique contextuelle des mots. [51] a proposé deux nouvelles méthodes de sélection de caractéristiques, MCPD (Modified Categorical Proportional Difference) et BCF (Balanced Category Feature), pour améliorer la classification des sentiments dans des textes à haute dimension.

Le tableau 2.1 montre quelques comparaisons de méthodes d'analyse des sentiments [21], [22], [23], [24] :

Approche	Principe général	Avantages	Inconvénients
Basée sur un dictionnaire ou lexique	Utilise un lexique prédéfini de mots associés à des sentiments positifs ou négatifs.	- Simple à mettre en œuvre - Ne nécessite pas de données annotées	- Dépend fortement de la qualité du lexique - Faible précision sémantique - Ne gère pas bien le contexte
Apprentissage automatique (Machine Learning)	Utilise des algorithmes (SVM, Naïve Bayes...) entraînés sur des textes étiquetés pour prédire le sentiment.	- Généralement plus précis qu'une approche lexicale - Capable de capter des motifs complexes	- Requiert des données annotées - Moins performant face à des textes très variés ou ambigus
Apprentissage profond (Deep Learning)	Utilise des réseaux de neurones (RNN, CNN, Transformers) pour apprendre automatiquement les représentations sémantiques.	- Haute performance - Capacité à comprendre le contexte et la complexité du langage - Bonne généralisation	- Requiert beaucoup de données annotées - Long à entraîner - Peu interprétable
Approche hybride	Combine les méthodes basées sur le dictionnaire et l'apprentissage automatique et profond.	- Améliore la robustesse et la précision - Réduit les faiblesses des approches individuelles	- Plus complexe à mettre en œuvre - Peut nécessiter un ajustement fin selon les cas d'usage

Tableau 2.1 comparaisons de méthodes d'analyse des sentiments

1.5 Niveaux d'analyse des sentiments

La portée de l'analyse sentiments peut être généralement divisée en niveaux de document, de phrase et d'aspect, selon la portée du texte [25] (figure 3):

a) SA au niveau du document

Ce niveau d'analyse consiste à étudier l'émotion de l'ensemble du document. L'analyse des sentiments au niveau du document traite chaque document comme un objet indépendant, et un document n'a qu'une seule polarité émotionnelle. Cette tâche est donc à granularité grossière.

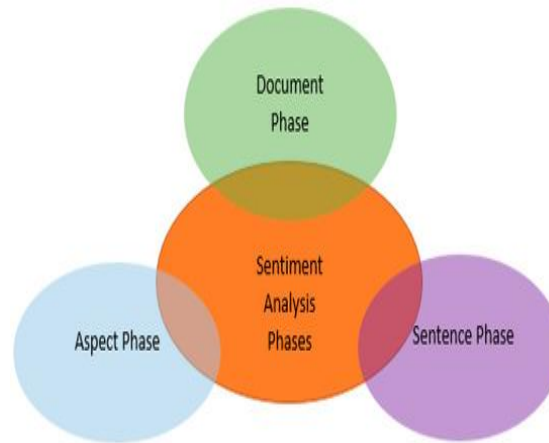


Figure 2.1 : Différentes phases/niveaux de l'analyse des sentiments [26]

La contribution dans [25] a mené une SA au niveau du document en utilisant un réseau LSTM bidirectionnel basé sur l'attention et CNN. [27] on 2020 proposition d'un modèle d'analyse sentiments spéculatif et ont émis l'hypothèse que des avis similaires étaient plus susceptibles d'être rédigés par des utilisateurs partageant les mêmes sentiments. Ils ont donc utilisé des documents similaires pour améliorer la précision de l'analyse sentiments.

b) SA au niveau des phrases

L'AS au niveau de la phrase (Sentence-level SA) vise à classer la polarité du sentiment d'une phrase en positif, négatif ou neutre [25]. Cette tâche classait chaque phrase en phrases objectives (subjective sentences) et subjectives. Une phrase objective était une phrase qui ne transmettait aucune opinion. Les phrases subjectives (**Subjective sentences**) décrivant les pensées et les idées de la critique peuvent être divisées en sentiments positifs ou négatifs.

c) SA au niveau des aspects

Elle consiste à évaluer la polarité des sentiments de caractéristiques ou aspects spécifiques dans un contexte donné.

L'analyse sentiments au niveau de l'aspect est une analyse plus fine, et l'algorithme vise principalement à modéliser la relation entre le terme d'aspect, la catégorie d'aspect, le terme d'opinion et la polarité du sentiment. [28] D'une manière générale, l'analyse des sentiments au niveau des aspects (ou ABSA) nécessite deux étapes : premièrement, extraire les termes d'aspect et les termes d'opinion, et par la suite, identifier la polarité sentimentale des aspects. Considérant la phrase : « This restaurant's steak is delicious. » Le « steak » est un terme

d'aspect de la catégorie d'aspect « nourriture », « délicieux » est le terme d'opinion et la polarité du sentiment de « steak » est positive. [29] développé une méthode de génération de texte pour effectuer une ABSA et les performances sont meilleures que la méthode de classification ML.[30] L'étude utilisé des informations contextuelles de localisation spécifiques à l'aspect pour attribuer différents poids et réduire l'erreur dans le jugement de la polarité sentimentale.[31] construit un CNN déformable pour incorporer une attention de corrélation croisée et des phrases contextuelles pour SA au niveau des aspects.

1.6 Les applications de l'analyse des sentiments

Le traitement de texte axé opinion possède de nombreuses applications importantes comme la détermination des opinions concernant un certain produit ou service via la classification des critiques en ligne ou l'enregistrement de l'évolution des attitudes du public à l'égard d'un parti politique via l'extraction de sites d'actualités en ligne ou de contenu de blogs [32]. Bien que les applications basées sur l'opinion ou sur le feedback sont plus populaires, le domaine du traitement du langage naturel s'intéresse actuellement aux SA ainsi qu'aux systèmes d'exploration d'opinion. Les principales applications de SA et de l'extraction d'opinions sont données comme suit:

- *Achat de produits ou de services* : La prise de décision précise pour l'achat des produits ou des services n'est plus un travail difficile. Les demandeurs ou clients peuvent évaluer les opinions et les expériences du grand public concernant tous les produits et services et comparer les marques concurrentes.
- *Amélioration de la qualité du produit ou du service* : grâce à l'exploration d'opinions et à l'analyse des sentiments, les fabricants peuvent recueillir les opinions des critiques positives concernant leurs produits ou services et ainsi améliorer la qualité de leurs services.
- *Recherche marketing* : les résultats des analyses de sentiment peuvent être utilisés à des fins d'étude de marché. Grâce à des méthodes d'analyse des sentiments, les tendances récentes des clients concernant des produits ou services particuliers peuvent être examinées. De même, les attitudes actuelles d'un public à l'égard des nouvelles politiques prises par des institutions gouvernementales peuvent également être facilement examinées.

- *Systèmes de recommandation* : Grâce à la classification des opinions des utilisateurs comme positives ou négatives, le système peut déterminer laquelle quel produit ou service est recommandé et lequel ne l'est pas.

- *Surveillance des sites sur les médias sociaux*: La surveillance des sites d'information, des blogs et des réseaux sociaux devient plus simple grâce à l'analyse des sentiments. Cette technique d'AS permet d'identifier automatiquement les propos agressifs, haineux, provocateurs ou les insultes présents dans les courriels, les articles de blog ou sur les plateformes médiatiques.

- *Détection d'opinion Spam* : Comme Internet est accessible à tout le monde, n'importe qui peut télécharger du n'importe quoi. Cela signifie que la probabilité que le contenu soit du spam augmente de jour en jour. Les particuliers pourraient télécharger du contenu de spam dans le but d'induire les gens en erreur. L'exploration d'opinions ainsi que les analyses de sentiment sont capables de classer le contenu Internet en spam et en non spam.

- *Elaboration de politiques* : en utilisant des analyses de sentiment, les décideurs politiques sont en mesure de prendre en considération les perspectives des citoyens concernant certaines politiques et ces connaissances peuvent être utilisées pour créer de nouvelles politiques en faveur des citoyens.

- *Prise de décision* : les opinions et les expériences du public sont des facteurs très utiles lors de la prise de décisions. Les sites de réseaux sociaux tels que Twitter ou Facebook sont l'un des moyens les plus exploitables pour la communication et le partage des opinions et de sentiments envers une variété de sujets et domaines. L'AS des avis partagés et leur classification automatique en classes positives, négatives ou neutres peut générer des données cruciales aidant les entreprises à prendre de décisions, sous la forme d'études de marché.

1.7 Processus d'analyse des sentiments

- *Etape 1 - Extraction de données (Data collection)*: Elle peut provenir de diverses sources en ligne, y compris les plateformes de réseaux sociaux, le scraping web, les sites d'actualités, les forums, les blogs et les sites de commerce électronique. Cette phase initiale de l'analyse des sentiments correspond à la collecte de données textuelles, bien que d'autres types de données puissent être inclus selon la tâche à accomplir, comme des données audio, vidéo ou de localisation. Les principales sources de collecte de données incluent :

Les étapes du processus sentimental sont illustrées ci-dessus :

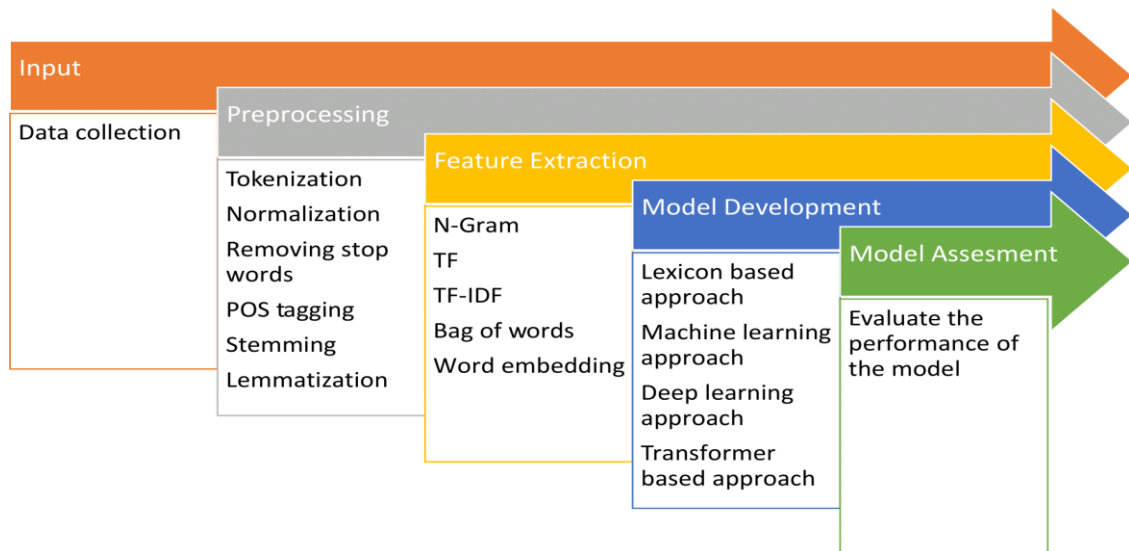


Figure 2.2 Etapes fondamentales de l'analyse des sentiments

- . Réseaux sociaux : Cela inclut les données collectées à des réseaux sociaux, mettant particulièrement en évidence l'engagement des consommateurs envers les produits lors de l'accès, de la publication ou du partage de contenu. De plus, les données des réseaux sociaux restent un objet de recherche vivant pour étudier le comportement des individus et de leurs groupes.

- . Forums : Les participants interagissent avec d'autres utilisateurs lors de discussions, posent des questions, ou font des suggestions et des demandes par le biais de messages textuels dans les forums. Ces plateformes fournissent une importante quantité d'informations, utiles pour l'analyse des sentiments, en particulier lorsque l'analyse est réalisée dans des catégories ou domaines spécifiques.

- . Blogs : Ils se composent de courtes entrées, incluant des opinions, des faits, des journaux personnels ou des liens organisés de manière chronologique. Les blogs constituent des ressources précieuses pour l'analyse des sentiments sur divers sujets.

- . Sites de commerce électronique : Ces plateformes permettent aux utilisateurs d'évaluer et d'exprimer leurs opinions sur des entreprises ou des organisations. Les sites de commerce électronique avec des avis produits, ainsi que les sites de critiques professionnelles, offrent une grande quantité d'avis exploitables pour l'analyse.

- o *Etape 2 - Prétraitement* : Le prétraitement est le processus de nettoyage ou de préparation des sources de données non structurées ou textes pour la classification. Les textes en ligne contiennent généralement beaucoup de bruits et des éléments inutiles tels que

des balises, des scripts. Le prétraitement des données réduit le bruit dans lequel contribue à améliorer les performances du classificateur. Le prétraitement accélère également le processus de classification, aidant ainsi en temps réel SA. Il permet de supprimer les informations parasites, de réduire le contenu des données et d'améliorer la précision de la classification. La stratégie de prétraitement est la suivante :

- . Tokenization (Segmentation) : Le texte du jeu de données est prétraité pour être segmenté en unités cohérentes appelées tokens, telles qu'une tabulation, une virgule, un espace ou tout autre délimiteur présent dans les données.

- . Stemming : Ce processus normalise chaque mot en le convertissant à sa forme de base en supprimant les suffixes, selon des règles morphologiques spécifiques.

- . Normalisation : De nombreux utilisateurs préfèrent utiliser des abréviations ou des ponctuations incorrectes lorsqu'ils s'expriment. La normalisation transforme ces termes en leurs formes écrites canoniques afin de permettre au système de distinguer les différentes utilisations de mots qui, bien que réécrits, signifient la même chose.

- . Suppression du bruit non pertinent : Les réseaux sociaux contiennent souvent beaucoup de bruit dans les données, ce qui peut nuire à l'efficacité du modèle de classification. Parmi ces bruits, on trouve des ponctuations, des chiffres, des symboles, des noms d'utilisateurs et des adresses de sites web. Pour atteindre l'objectif souhaité, toutes les données non pertinentes doivent être supprimées à l'avance afin de nettoyer les données:

- a) Suppression des mots : Certaines approches considèrent que certains mots très communs n'apportent aucune information utile pour l'analyse du texte. Dans ce cas, il est courant, d'utiliser une liste de ces mots pour filtrer le texte. La suppression de ces mots présente l'avantage de réduire la phrase à des mots pleins.

- b) Structuration des phrases : Structurer une phrase permet de mieux appréhender la sémantique de chaque mot. Certaines techniques d'analyse de sentiments utilisent la structure des phrases afin d'identifier l'opinion.

- c) Suppression des localisateurs de ressources uniformes (URL), hashtags, références, caractères spéciaux : Le nettoyage des données des hashtags, références, caractères spéciaux, aidera à réduire la plupart des bruits.

- d) Traduction de mots d'argot : Pour cela, nous prenons l'aide du dictionnaire d'argot Internet et remplaçons les mots d'argot dans leur format significatif.

- e) Suppression des lettres supplémentaires des mots : Les mots qui ont la même lettre plus de deux fois et qui ne sont pas présents dans le lexique sont réduits au mot

avec la lettre répétitive n'apparaissant qu'une seule fois. Par exemple, le mot exagéré «Happyyyyyy» est réduit à « Happy».

f) Enracinement : L'enracinement donne le mot racine, est fait à l'aide de Natural Language Tool Kit (NLTK). Par exemple, des mots tels que «waiting», «waits», «waited» sont remplacés par le mot « wait».

. Suppression des mots vides (Stop words removal) : Il n'est pas conseillé de supprimer des mots utiles, mais il est bénéfique de se débarrasser des "stop words", c'est-à-dire des mots comme les conjonctions et les prépositions qui ne transmettent pas d'informations utiles car ils servent principalement de connecteurs. Dans le cadre de SA ou de la reconnaissance des émotions, les stop words ont généralement un impact très limité sur le déplacement du sentiment d'une catégorie (par exemple, positif ou heureux) vers une autre (négatif ou triste).

○ *Etape 3 - Extraction des caractéristiques (Feature Extraction)*: L'extraction des caractéristiques est une étape clé en AS, transformant le texte brut en représentations structurées pour améliorer les performances des modèles. Elle identifie les éléments linguistiques porteurs de sentiment tout en réduisant la dimensionnalité des données. Le processus d'extraction des caractéristiques (feature extraction) prend un texte en entrée et produit des caractéristiques dans différents formats tels que lexico-syntaxique ou stylistique, syntaxique, et basé sur le discours [33].

. Fréquence du Terme (Term Frequency - TF): Il s'agit d'une méthode consistant à comptabiliser les occurrences de chaque terme (mot) dans un document ou un ensemble de données. Cette méthode constitue une approche fondamentale pour représenter la connotation d'un terme au sein d'un document. Cette méthode analyse des mots individuels (unigrammes), des combinaisons de deux mots (bigrammes), ou trois mots (trigrammes), en comptant leurs occurrences, qui servent ainsi de caractéristiques [32]. La présence d'un terme se voit attribuer une valeur binaire de 1 (présent) ou 0 (absent) . En revanche, le TF est représentée par une valeur entière correspondant à son nombre d'occurrences dans le document analysé.

Le TF-IDF évalue l'importance d'un mot dans un document par rapport à sa fréquence dans l'ensemble d'un corpus. Il permet d'identifier les termes fréquents dans un document spécifique mais rares dans la collection complète, fournissant ainsi des caractéristiques plus significatives pour l'analyse.

- **Bag of words:** La méthode du Bag of words (BoW) reste une procédure simple d'extraction de caractéristiques textuelles à partir d'un document de référence. Pour un document particulier, elle décrit les fréquences de mots en distinguant les phrases sous forme de vecteurs basés sur le dictionnaire lexical [34]. Cependant, elle présente un inconvénient majeur : les représentations obtenues ne sont absolument pas sensibles à la structure syntaxique du texte.

- **Word embedding:** Word embeddings consiste à traduire des mots en vecteurs, les comparaisons de sens étant effectuées à l'aide de structures entraînées de manière similaire à celle des réseaux de neurones. Les algorithmes SG et CBOW utilisent des méthodes de type fenêtre (window type methods) pour prédire soit le contexte à partir du mot, soit les mots à partir du contexte. Ils sont censés extraire des significations du contexte local, utiles à des opérations telles que l'analyse des sentiments.

Il existe principalement quatre types d'embeddings de texte, classés en catégories [35]: word-based embeddings, phrase-based embeddings, sentence-based embeddings et document-based embeddings.

D'autres sous-catégories importantes sont également utilisées, telles que: Word2vec, GloVe, FastText, ELMo, CBOW, Skip-gram, Sentiment-Specific Word Embedding (SSWE), GLoMo, Sent2Vec, OpenAI Transformer, BERT, Context2Vec, Universal Language Model Fine-tuning (ULMFiT) [36].

1.8 Les défis de l'analyse des sentiments

Différents facteurs affectent le processus d'AS et doivent être traités correctement pour obtenir le rapport final de classification ou de regroupement, parmi ces facteurs, on peut citer [16]:

- **Résolution de coréférence :** Ce problème se réfère principalement à savoir ce qu'indique un pronom ou un proverbe ? » Par exemple, dans la phrase "Après avoir regardé le film, nous sommes partis manger; c'était bien." A quoi se réfère le mot « c'était » ; que ce soit le film ou la nourriture ? Ainsi, lorsque l'analyse du film est en cours, si la phrase concerne le film ou la nourriture ? C'est une préoccupation pour l'analyste. Ce type de problème se produit principalement dans le cas d'une SA orientée aspect.

- **Association avec une période :** le moment de la collecte d'avis est une question importante dans l'AS. Le même utilisateur ou groupe d'utilisateurs peut donner une réponse

positive pour un produit à un moment donné, et il peut y avoir un cas où il peut donner une réponse négative. C'est donc un défi pour l'analyseur de sentiment à un autre moment. Ce type de problème survient principalement dans la SA comparative.

- Gestion du sarcasme : l'utilisation de mots qui signifient le contraire de ce qu'ils informent sont surtout connus comme des mots de sarcasme. Par exemple, la phrase « Quel bon batteur il est, il marque zéro dans toutes les autres manches. » Dans ce cas, le mot positif « bon » a un sens négatif. Ces phrases sont difficiles à trouver et par conséquent, elles affectent l'analyse du sentiment.

- ⊖ Indépendance de domaine : Dans la SA, les mots sont principalement utilisés comme fonction d'analyse. Mais, le sens des mots n'est pas fixé de bout en bout. Il y a peu de mots dont la signification change d'un domaine à l'autre. En dehors de cela, certains mots portent des significations opposées selon le contexte; on les appelle contronyme. Ainsi, il est difficile de connaître le contexte pour lequel le mot est utilisé.

- Négations : Les mots négatifs présents dans un texte peuvent totalement changer le sens de la phrase dans laquelle il est présent. Ainsi, lors de l'analyse des critiques, ces mots doivent être pris en compte. Par exemple, les phrases « Ceci est un bon livre. » Et « Ce n'est pas un bon livre. » ont une signification opposée, mais lorsque l'analyse est effectuée en utilisant un seul mot à la fois, le résultat peut être différent. Pour gérer ce type de situations, une analyse en n-gramme est préférable.

- Détection de spam : le Web contient à la fois des contenus authentiques et du spam. Pour une classification efficace des sentiments, ce contenu de spam doit être éliminé avant le traitement. Cela peut être fait en identifiant les doublons, en détectant les valeurs aberrantes et compte tenu de la réputation du critique.

- Mots orthographiques : les gens utilisent des mots orthographiques pour exprimer leur excitation, leur bonheur, par exemple : le mot Sooo, Sweeettt, Haappy ou s'ils sont pressés, ils insistent sur les mots par exemple : comeeeee, fasssssst , waittttnggg...ect

- Manière d'exprimer son sentiment : les gens n'expriment pas toujours leurs sentiments de la même manière. Le sentiment de chaque individu est différent de la façon de penser, où la manière de s'exprimer varie d'une personne à l'autre.

- Asymétrie dans la disponibilité des outils d'extraction d'opinion : le logiciel d'extraction d'opinion est très coûteux et actuellement abordable uniquement pour les grandes organisations et le gouvernement. Cela dépasse les attentes des citoyens ordinaires. Il est donc essentiel que ce soit accessible à tous, afin que chacun puisse en bénéficier [37].

1.9 Conclusion

Ce chapitre présente une revue de la littérature sur l'analyse des sentiments appliquée aux textes, en examinant les principales approches : les méthodes basées sur le lexique, sur l'apprentissage automatique, sur l'apprentissage profond, sur les transformateurs, ainsi que les approches hybrides. Il aborde également des aspects essentiels tels que le prétraitement des données, les niveaux de sentiment, l'extraction et la sélection des caractéristiques et l'intégration du texte.

Chapitre 2

Analyse des Sentiments à base d'aspects

2.1 Introduction	21
2.2 Analyse des sentiments à base d'aspects	21
2.3 Tâches ABSA	22
2.3.1 Tâches simples de l'ABSA (Single ABSA Tasks)	22
2.3.1.1 Extraction des termes d'aspect (ATE)	23
2.3.1.2 Catégorisation des termes d'aspect	23
2.3.1.3 Extraction de termes d'opinion	23
2.3.2 Tâches composées de l'ABSA	24
2.3.2.1 Extraction de paires aspect-opinion (AOPE)	24
2.3.2.2 Extraction de triplets aspect-opinion-sentiment (ASTE)	24
2.3.2.3 Prédiction de quadruplets aspect-opinion-sentiment (ASQP)	24
2.4 Évolution de l'ABSA	25
2.4.1 Modèles basés sur des transformateurs	28
2.5 Jeux de données de référence pour ABSA	30
2.6 Métriques de performance	32
2.7 Défis de l'ABSA	35
2.8 Large Modèles de Langue (LLM)	35
2.8.1 Principaux constituants des LLM	36
2.8.2 Fonctionnement des LLM	36
2.8.3 Apprentissage par transfert à partir des LLM	37
2.8.4 Méthodes de fine-tuning des LLM	38

2.8.5 Avantages et inconvénients de l'apprentissage par transfert	39
2.8.6 Domaines d'application des LLM	40
2.9 Relation entre LLM et apprentissage automatique	40
2.9.1 Facteurs clés de l'IA générative et des LLM	41
2.10 LLM dans l'analyse des sentiments	42
2.11 LLM dans l'apprentissage en contexte	42
2.11.1 Facteurs clés pour le choix des LLM	42
2.12 Conclusion	43

2-1 Introduction

Ces dernières années, les travaux de recherche sur l'AS ont connu des avancées significatives, accompagnées d'une complexification progressive des tâches liées à la détection et à l'identification des opinions. Cette évolution a progressivement transformé les objectifs et les méthodes du domaine. A ses débuts, l'AS se concentrait principalement sur des tâches de classification, telles que la détection de la polarité binaire ou l'analyse globale au niveau des phrases et des textes. Les approches reposaient alors sur des algorithmes d'apprentissage automatique classiques, associés à l'extraction de caractéristiques pertinentes.

Cependant, un défi persiste lorsqu'il s'agit d'étudier les opinions associées à une cible spécifique ou à ses caractéristiques (ou aspects) au sein d'un texte. En effet, un même document peut contenir des passages subjectifs évoquant des entités différentes, mais aussi des aspects distincts d'une même cible. La tâche visant à identifier et classer les opinions relatives à ces aspects est appelée Aspect-Based Sentiment Analysis, ou ABSA. Il s'agit d'une analyse fine dont l'objectif est d'inférer la polarité des sentiments exprimés sur des entités ou leurs propriétés.

2-2 Analyse des Sentiments à base d'aspects

Dans l'ABSA, le sentiment fait référence aux ressentis subjectifs d'une personne à propos d'un aspect ou caractéristique spécifique. L'opinion publique peut connaître des hauts et des bas à tout moment concernant certaines entités spécifiques, entraînant un changement d'attitude vis-à-vis d'un aspect donné. Cette transformation illustre la capacité d'adaptation du comportement humain, l'autonomie de choix et la pensée innovante. Les auteurs dans [1] le définissent ainsi : « Un sentiment est une opinion qu'une personne exprime à propos d'un aspect, d'une entité, d'une personne, d'un événement, d'une caractéristique, d'un objet ou d'une cible donnée ». Le terme *caractéristique* est d'ailleurs souvent utilisé de manière interchangeable avec *aspect* par la majorité des chercheurs dans le monde.

L'ABSA est la tâche qui consiste à identifier les éléments de sentiment pertinents dans un texte donné, soit sous la forme d'un élément unique, soit comme un ensemble d'éléments liés entre eux par des relations de dépendance. Il s'agit donc de détecter les aspects et leurs sentiments associés dans un contenu textuel.

L'ABSA repose principalement sur l'identification de quatre éléments fondamentaux du sentiment (figure 4) [38] :

- La *catégorie d'aspect* (*c*) définit un aspect unique d'une entité et est censée appartenir à un ensemble de catégories, prédéfini pour chaque domaine spécifique.
- Le *terme d'aspect* (*a*) est la cible de l'opinion qui apparaît explicitement dans le texte,
- Le *terme d'opinion* (*o*) est l'expression utilisée par l'émetteur de l'opinion pour exprimer son sentiment à l'égard de la cible.
- La *polarité du sentiment* (*p*) décrit l'orientation du sentiment à l'égard d'une catégorie ou d'un terme d'aspect. Elle est généralement positive, négative ou neutre.

En combinant les éléments *a* (aspect), *c* (catégorie), *o* (opinion) et *s* (sentiment), l'ABSA se décline en plusieurs sous-tâches, chacune visant à capturer les différentes dimensions des opinions exprimées dans les textes.

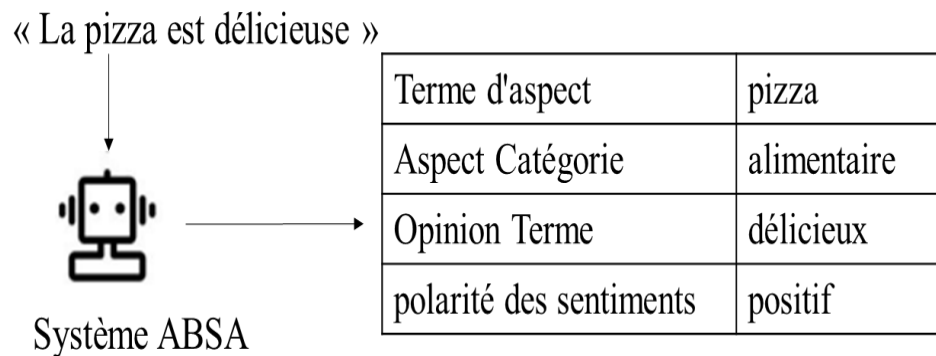


Figure 2.1 : Quatre éléments de sentiment clés en ABSA

2.3. Tâches ABSA

Les tâches d'ABSA sont généralement divisées en deux grandes classes : les tâches simples et les tâches composées. Le tableau 2.1 présente en détail les diverses sous-tâches associées à l'ABSA [39] :

2.3.1. Tâches simples de l'ABSA (Single ABSA Tasks)

Les tâches simples visent principalement à prédire un facteur émotionnel unique. Les principales tâches simples de l'ABSA se divisent en:

Type	Abrév.	Nom de la tâche	Entrée	Sortie	Méthode
Tâches simples de l'ABSA					
Simple	ATE	Extraction de termes d'aspect	P : Phrase	a	Extraction
	OTE	Extraction de termes d'opinion	P	o	Extraction
	ACD	Détection de catégories d'aspect	P	c	Classification
	AOCE	Co-extraction aspect-opinion	P	a, o	Extraction
	AOOE	Extraction d'opinion orientée aspect	P + a	o	Extraction
	ABSC	Classification de sentiment basée sur l'aspect	P + a	s	Classification
	COSC	Classification de sentiment orientée catégorie	P + c	s	Classification
Tâches composées de l'ABSA					
Paire	AOPE	Extraction de paires opinion-aspect	P	(a, o)	Extraction & Classification
	ASPE	Extraction de paires aspect-sentiment	P	(a, s)	Extraction & Classification
	CSPE	Extraction de paires catégorie-sentiment	P	(c, s)	Extraction & Classification
Triplet	ACSTE (TASD)	Extraction de triplets catégorie-aspect-sentiment ou Détection de sentiment ciblé sur un aspect	P	(a, c, s)	Extraction & Classification
	AOSTE (ASTE)	Extraction de triplets opinion-aspect-sentiment ou aspect-sentiment-opinion	P	(a, o, s)	Extraction & Classification
Quad	ACOSQE	Extraction de quadruplets catégorie-opinion-aspect-sentiment	P	(a, c, o, s)	Extraction & Classification

Tableau 2.1 : Tâches simples et composées en ABSA

2.3.1.1. Extraction des termes d'aspect (Aspect Term Extraction - ATE)

L'ATE est une tâche clé de l'ABSA, qui consiste à extraire les aspects explicites

mentionnés par les utilisateurs dans un texte donné. Le processus d'ATE peut être classé selon trois catégories, en fonction de la disponibilité des données : supervisée, semi-supervisée et non supervisée.

Comme les termes d'aspect visés sont généralement des tokens ou expressions dans une phrase donnée, la version supervisée de l'ATE est souvent formulée comme une tâche de classification au niveau des tokens, à partir de données annotées.

De nombreux travaux de recherche se concentrent sur l'amélioration de la représentation des mots, car l'ATE nécessite des informations spécifiques au domaine pour détecter les caractéristiques pertinentes. Pour l'étiquetage de séquences et dans les formulations séquence-à-séquence, l'objectif est de capter le sens global de l'énoncé et de mieux prédire les aspects grâce à une contextualisation enrichie.

L'extraction des termes d'aspect consiste à identifier et extraire les aspects ou caractéristiques spécifiques évoqués dans un texte donné. L'objectif est de reconnaître automatiquement les aspects pertinents et d'en fournir une représentation structurée, mettant en évidence les attributs ou caractéristiques abordés dans les propos analysés. L'ATE peut être divisée en ATE explicite et implicite:

. ATE explicite : Un aspect explicite est directement mentionné dans le texte. Par exemple, dans la phrase « *Les sushis du restaurant sont merveilleux* », « sushis » représente un aspect explicite, ainsi qu'une entité explicite liée à la nourriture.

. ATE implicite : Un aspect implicite est un aspect non exprimé directement dans le texte. Par exemple, dans la phrase « *Le téléphone est beau mais il est lent au chargement* », « chargement » renvoie à un aspect implicite lié à l'entité *téléphone*, en particulier à sa batterie, bien que ce composant ne soit pas mentionné explicitement.

2.3.1.2 Catégorisation des termes d'aspect

La deuxième sous-tâche de l'ABSA consiste à regrouper les termes d'aspect similaires en catégories (ATC), chacune représentant un aspect spécifique. Ces catégories peuvent être considérées comme des groupes d'aspects partageant une signification ou une fonction commune.

2.3.1 Extraction de termes d'opinion

L'extraction de termes d'opinion (OTE) désigne le processus d'identification des expressions subjectives associées à un aspect donné. Ainsi, les termes d'aspect apparaissent fréquemment dans la majorité des travaux existants sur l'OTE. Cette tâche peut être divisée en deux sous-tâches, selon que le terme d'aspect est fourni en entrée ou prédit en sortie : AOCE (Aspect-Oriented Contextual Extraction) et TOTE (Target-Oriented Term Extraction). La figure 6 illustre les principales caractéristiques de cette tâche dans le cadre de l'ABSA.

L'extraction ciblée de termes d'opinion, quant à elle, se concentre sur l'identification des opinions associées à un aspect spécifique. Le cœur des recherches en TOTE porte sur la modélisation de représentations contextuelles centrées sur l'aspect dans le texte, afin de permettre une extraction précise des opinions correspondantes.

2.3.2 Tâches composées de l'ABSA

Les tâches composées de l'ABSA sont généralement réparties en quatre sous-catégories : l'extraction de paires aspect-opinion, l'ABSA de bout en bout, l'extraction de triplets aspect-opinion-sentiment, et la prédiction de quadruplets enrichis de sentiment :

2.3.2.1 Extraction de paires aspect-opinion (AOPE)

Dans le cadre de l'ABSA, l'extraction conjointe des termes d'aspect et des termes d'opinion constitue une étape essentielle. Pour la tâche AOPE, l'identification de l'un de ces éléments peut faciliter celle de l'autre. Le résultat attendu se présente généralement sous forme de paires comprenant un aspect et une opinion. Cette observation constitue un fondement clé de l'approche AOPE, qui cherche à extraire simultanément les termes d'aspect et les expressions d'opinion qui leur sont associées, en tant que paires cohérentes aspect-opinion.

2.3.2.2 Extraction de triplets aspect-opinion-sentiment (ASTE)

La tâche ASTE consiste à extraire, à partir d'une phrase donnée, des triplets composés de l'aspect concerné, du terme d'opinion associé, et de la polarité du sentiment exprimé. Elle permet ainsi d'identifier non seulement l'objet de l'opinion, mais aussi la direction du sentiment à son égard et les raisons de cette orientation, exprimées par les termes d'opinion.

2.3.2.3 Prédiction du quadruplet aspect-sentiment (ASQP)

L'objectif principal des tâches composées de l'ABSA est d'obtenir une analyse des sentiments plus fine et précise au niveau des aspects, que ce soit par extraction de paires ou de triplets. La structure de sentiment la plus complète à ce niveau est fournie par un modèle capable de prédire simultanément les quatre éléments clés du sentiment, correspondant aux différentes composantes des tâches ABSA.

Cela conduit à la problématique récente de la prédiction du quadruplet aspect-sentiment (ASQP), qui vise à extraire, à partir d'un texte ou d'un avis, un quadruplet complet comprenant les quatre informations essentielles du sentiment. Par exemple, pour la phrase « *l'écran est remarquable* », un modèle ASQP retournerait le quadruplet suivant : (*affichage, écran, positif, remarquable*).

2.4. Evolution De L'ABSA

Avec l'essor des algorithmes d'apprentissage et ses performances remarquables dans divers traitements du langage naturel (NLP), les modèles basés sur CNN, les modèles RNN utilisant LSTM et GRU, les modèles LSTM basés sur l'attention, et les modèles à base de transformateurs sont largement utilisés pour l'ABSA dans des ensembles de données de pointe [40]. Les deux grandes catégories de modèles d'apprentissage ABSA sont les modèles d'apprentissage profond (DL) et les modèles d'apprentissage automatique (ML) [41] , [42] :

L'apprentissage automatique -ML- relève du domaine de l'intelligence artificielle (IA) et se concentre sur le développement de modèles et d'algorithmes statistiques. L'apprentissage automatique permet aux ordinateurs d'apprendre et de faire des prédictions ou de prendre des décisions sans nécessiter de programmation explicite. L'application des techniques de ML a marqué une avancée significative dans l'ABSA. Des modèles d'apprentissage supervisé, tels que Naïve Bayesian, Support Vector Machine (SVM) et Artificial Neural Networks (ANN), ont été utilisés pour la classification des sous-tâches de l'ABSA. L'ingénierie des caractéristiques a joué un rôle crucial dans ces modèles, les chercheurs extrayant des caractéristiques pertinentes du texte, telles que les n-grammes, les sacs de mots, parties de discours, les structures syntaxiques et les lexiques de sentiments.

Les performances de ces méthodes dépendent fortement des caractéristiques élaborées manuellement, qui malheureusement demandent beaucoup de travail et sont comparativement moins efficaces. C'est pourquoi les chercheurs ont exploré des techniques révolutionnaires de DL pour les sous-tâches ABSA au cours des dernières années.

Les Techniques d'apprentissage -DL - Propulsés par les progrès rapides des techniques de réseaux neuronaux ont atteint un succès remarquable dans diverses applications. Ce succès a conduit les recherches en ABSA à passer des techniques basées sur les caractéristiques aux méthodes neuronales profondes. L'essor des architectures de réseaux de neurones profonds (DNN), telles que les réseaux neuronaux récurrents (RNN), les réseaux neuronaux convolutifs (CNN) et les transformateurs, a entraîné un changement de paradigme dans l'ABSA. Ces modèles se sont révélés plus performants pour capturer les informations contextuelles et pour apprendre des modèles complexes dans les données textuelles. Les modèles LSTM et CNN basés sur l'attention ont été largement proposés pour l'ABSA.

Ce passage des méthodes traditionnelles basées sur les caractéristiques à diverses architectures de réseaux neuronaux reflète les progrès substantiels de la recherche sur les réseaux neuronaux et leur applicabilité aux sous-tâches de l'ABSA. Chacune de ces approches basées apporte des atouts uniques, contribuant ainsi à l'évolution du domaine de l'ABSA. Toutefois, les modèles LSTM et CNN présentaient deux limites principales. Tout d'abord, le modèle formé avec un ensemble de données n'était pas performant pour d'autres ensembles de données et d'autres domaines. Deuxièmement, ces modèles peinaient à atteindre des performances remarquables en termes de précision. C'est pourquoi les modèles à base de transformateurs sont apparus ces dernières années comme des solutions proposées pour les tâches ABSA.

2.4.2.1 Modèles Bases Sur Des Transformateurs

Le développement d'architectures de transformateurs, comme le démontrent des modèles tels que BERT (Bidirectional Encoder Representations from Transformers), a constitué une avancée significative dans le domaine de l'ABSA. Les modèles de transformateurs ont donné d'excellents résultats dans l'analyse des sentiments et d'autres tâches de traitement du langage naturel, surpassant les méthodes antérieures dans la capture des dépendances à longue portée et des informations contextuelles. L'affinement des modèles de transformateurs pré-entraînés pour

les tâches d'ABSA est devenu une pratique courante, permettant aux modèles de tirer parti d'un pré-entraînement à grande échelle sur diverses données textuelles. Les transformateurs possèdent un mécanisme d'auto-attention qui leur permet de reconnaître les dépendances entre les mots dans un texte. Contrairement aux RNN, qui traitent les séquences de manière séquentielle, les transformateurs peuvent traiter les mots en parallèle. Le mécanisme d'auto-attention peut être représenté par l'équation suivante [46] :

$$Attention(Q, K, V) = softmax\left(\frac{QK^T}{\sqrt{d_k}}\right)V$$

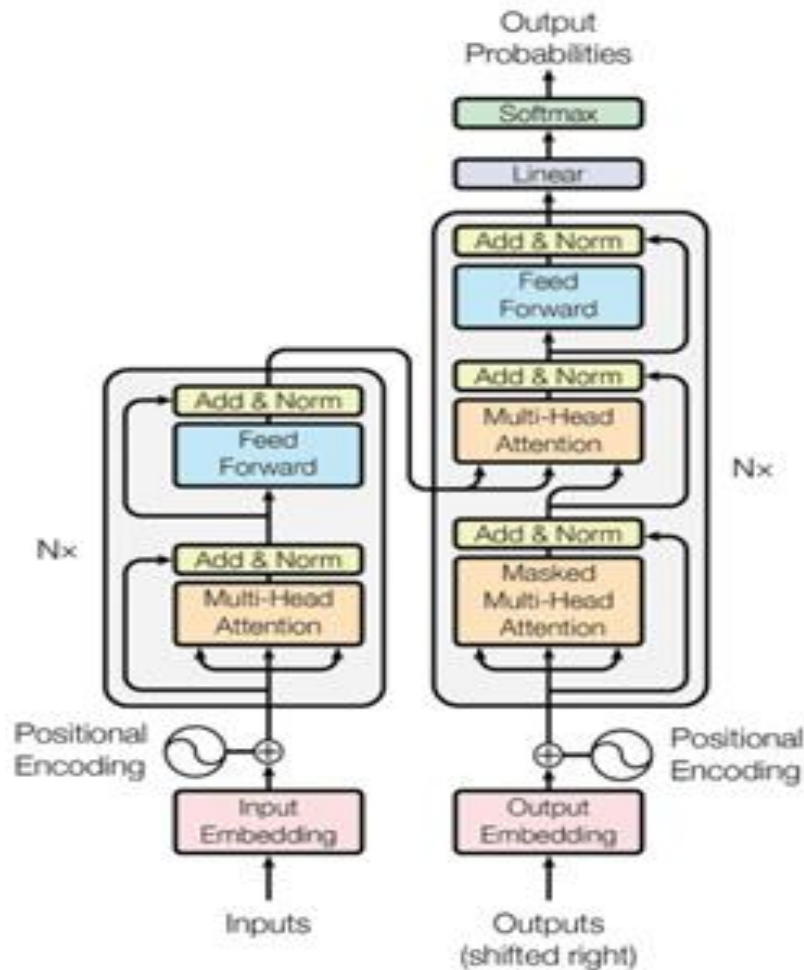


Figure 2.2 Architecture du transformateur [44]

La figure 2.2 illustre comment un codeur et un décodeur constituent l'architecture d'un transformateur. Dans le contexte de l'ABSA, le codeur prend le texte (avis) et l'aspect en entrée et

les code dans une représentation à haute dimension. D'autre part, le décodeur produit la sortie, telle que la polarité du sentiment. Le mécanisme d'auto-attention des transformateurs permet au codeur et au décodeur d'enregistrer les dépendances globales de l'examen et de faciliter un flux d'informations efficace.

BERT est un modèle puissant basé sur les transformateurs qui a révolutionné les tâches de traitement du langage naturel. Il comprend la compréhension bidirectionnelle du contexte grâce à l'architecture des transformateurs. BERT a atteint des performances exceptionnelles dans divers domaines du NLP tels que l'analyse des sentiments, la classification des textes, la réponse aux questions et la traduction automatique. La compréhension contextuelle du langage par BERT a permis d'améliorer les moteurs de recherche, les chatbots et les assistants virtuels, permettant des réponses plus précises et de meilleures expériences pour l'utilisateur. Bien que les modèles basés sur les BERT aient connu un succès remarquable dans l'ABSA, des défis persistent, notamment la nécessité d'un pré-entraînement spécifique au domaine, le traitement des contextes ambigus et la prise en compte de la rareté des données étiquetées pour des domaines spécifiques. Par conséquent, pour surmonter ces limitations, les grands modèles de langage (LLM) tels que PaLM API, GPT-3 et flan-t5 peuvent être utilisés pour les sous-tâches ABSA [43], [44].

GPT est un LLM de pointe et probablement un moment clé dans le domaine du traitement du langage naturel (NLP). Développé par l'OpenAI, le GPT utilise une architecture de transformation pour générer un texte cohérent et contextuellement précis. Le GPT est capable de produire des textes fluides tels que la traduction, la génération de texte et la réponse à des questions d'une manière semblable à celle des humains, car il a été affiné en utilisant un vaste corpus de données textuelles pour diverses tâches de TAL. GPT1, GPT-2 et GPT-3 ont été mis à disposition par OpenAI. En outre, la dernière innovation, ChatGPT, a stupéfié tout le monde par ses caractéristiques sophistiquées et s'est rapidement hissée au sommet des médias sociaux et des systèmes d'information. Le ChatGPT est capable d'accomplir diverses tâches à partir d'ensembles de données d'avis de clients, telles que l'analyse des sentiments, le résumé de texte, la classification de texte, la détection de mots offensants, etc. Bien que le ChatGPT ait démontré des capacités remarquables de génération de langage et qu'il ait été appliqué à diverses tâches de NLP, sa mise en œuvre pour ABSA reste limitée et nécessite une étude scientifique plus

approfondie. Les derniers modèles LLM n'ont pas encore été explorés en profondeur dans le domaine ABSA dans les études susmentionnées [43] .

2.5 Jeux de données de référence pour ABSA

Dans le domaine de l'ABSA, plusieurs jeux de données de référence sont largement utilisés pour l'entraînement et l'évaluation des modèles. Parmi les plus connus, on trouve SemEval, Twitter Airlines, MAMS, Yelp ou Amazon Reviews, jeux de données en chinois et pour la finance (ABSAbank-Imm) [52]:

Jeu de données	Domaine	Langue	Utilisation
SemEval (2014–2016)	Restaurants, ordinateurs, hôtels	Anglais (+ autres)	Référence majeure pour ABSA ; annotations d'aspects, polarité et opinion expressions
Twitter Airlines	Compagnies aériennes	Anglais	Annoté pour aspects liés au service (ponctualité, personnel, etc.)
REST (Rest. Reviews)	Restauration	Anglais	Basé sur SemEval ; annotations d'aspects spécifiques aux avis de restaurants
Laptop Reviews (SemEval)	Informatique (laptops)	Anglais	Annoté pour aspects techniques : écran, batterie, performance, etc.
MAMS Dataset	Restauration	Anglais	Multi-aspect et multi-sentiment dans un seul avis
Hotel Review Datasets	Hôtellerie	Multiling.	Données issues de TripAdvisor, Booking.com, etc. ; aspects comme propreté, service, etc.
Yelp Challenge Dataset	Multi-domaines (services)	Anglais	Grand volume ; annotations possibles pour aspects extraits automatiquement
Amazon Reviews Dataset	E-commerce	Anglais (+ autres)	Avis clients ; nécessite annotation ou extraction automatique d'aspects
ChnSentiCorp / NLPCC	Restauration, e-commerce	Chinois	Jeux de données chinois annotés pour ABSA
ABSAbank-Imm Dataset	Finance, banques	Anglais	ABSA orienté finance ; aspects : services en ligne, frais, relations client, etc.

Tableau 2.2 Jeux de données de référence

2.6. Métriques de la performance

Les indicateurs de performance les plus couramment utilisés sur l'ABSA sont l'exactitude, la précision, le rappel et le score F1. Ces mesures quantitatives restent les références standards pour évaluer l'efficacité des modèles. Le choix des mesures peut influencer considérablement la manière dont la performance et l'efficacité d'un modèle sont évaluées et comparées. Une matrice de confusion fondamentale prend généralement la forme d'une matrice 2x2, comme illustré à la figure 4 résumant les prédictions correctes et incorrectes du classificateur. Toutefois, la plupart des recherches doivent encore s'appuyer sur un protocole d'évaluation cohérent et reconnu à l'échelle de la communauté. Dans ce contexte, nous présenterons brièvement les principales métriques utilisées en ABSA, y compris certaines moins répandues telles que la perte de classement ou l'erreur absolue moyenne. L'exactitude, notamment, représente la proportion de prédictions correctes sur l'ensemble des prédictions effectuées. Elle constitue l'un des indicateurs les plus utilisés pour juger de la justesse d'un modèle de classification. Les méthodes modernes d'analyse des sentiments s'appuient souvent sur accuracy, F1-score et la précision comme principaux indicateurs d'évaluation. Cependant, une étude récente sur l'analyse des sentiments utilisant des architectures d'apprentissage profond met en évidence l'utilisation du rappel et de la précision (recall and accuracy) pour évaluer les performances. Ces mesures sont décrites ci-dessous :

ACTUAL	Positive	True Positive	False Negative
	Negative	False Positive	True Negative
		Positive	Negative
		PREDICTED	

Figure 2.2 Matrice de confusion

Le Vrais positifs (VP) : le nombre d'avis positifs qui ont été correctement classés.

Le Vrai négatif (VN) : nombre d'avis négatifs correctement classés comme négatifs.

Le Faux positifs (FP) : Nombre d'avis positifs mal classés.

Le Faux négatifs (FN) : nombre d'avis négatifs mal classés.

Les principales métriques quantitatives sont :

Précision

La précision [45] est également appelée valeur prédite positive, mesure la justesse du modèle. Une précision plus élevée indique moins de FP. Mathématiquement, il est défini comme:

$$\text{Précision} = TP / (TP + FP)$$

Rappel

Le rappel [45] est également connu sous le nom de sensibilité, mesure les cas positifs correctement classés par le modèle, une valeur de rappel élevée signifie que peu de cas positifs sont mal classés comme négatifs. Le rappel peut être calculé à l'aide de la formule suivante.

$$\text{Rappel} = TP / (TP + FN)$$

Exactitude

L'exactitude est utilisée comme mesure pour les techniques de catégorisation. Les valeurs d'exactitudes sont beaucoup moins réticentes aux variations du nombre de décisions correctes que la précision et le rappel. Cette technique est représentée sous la forme suivante [52] :

$$\text{Exactitude} = (TP + TN) / (TP + FP + TN + FN)$$

Score F1

Le score F1 ou la mesure F1 [38] [45] est la moyenne harmonique de la précision et du rappel. Le score F peut être calculé comme suit :

$$\text{Score F1} = 2 * (\text{Précision} \times \text{Rappel}) / (\text{Précision} + \text{Rappel})$$

Autres métriques utiles :

- **Ranking Loss** : La Ranking Loss (ou perte de classement) est une fonction de perte utilisée pour apprendre un modèle à préserver l'ordre relatif entre paires d'éléments. Elle

est souvent utilisée dans les systèmes de recommandation, le ranking d'informations, ou dans le cas de la classification multilabel: $L_{rank} = \max(0, 1 - (s_i - s_j))$

- MAE (Erreur Absolue Moyenne) : également connue sous le nom de perte L1, est une fonction de perte utilisée dans les tâches de régression qui calcule les différences absolues moyennes entre les valeurs prédites par un modèle d'apprentissage automatique et les valeurs cibles réelles: $MAE = (1/n) * \sum |y_i - \bar{y}|$
- NDCG (Gain Cumulé Actualisé Normalisé) : mesure la pertinence d'une liste d'aspects retournée. Les équipes de machine learning utilisent souvent le NDCG pour évaluer les performances d'un moteur de recherche, d'une recommandation ou d'un autre système de recherche d'informations: $NDCG_k = IDCG_k / DCG_k$

Il est conseillé de donner la priorité à la (F1-score) pour les jeux de données déséquilibrés, car elle offre un bon compromis entre précision et rappel, ce qui est essentiel lorsque les classes ne sont pas réparties équitablement.

Pour les jeux de données équilibrés, l'exactitude peut constituer un indicateur pertinent, mais il est préférable de la compléter avec la précision et le rappel pour une évaluation plus complète.

L'erreur absolue moyenne (MAE) est quant à elle recommandée pour les tâches impliquant des scores de sentiment continus, indépendamment de l'équilibre du jeu de données, car elle permet de mesurer directement les écarts de prédiction.

Ces recommandations visent à aider les chercheurs à choisir les métriques d'évaluation les plus adaptées à la nature de leurs données et à leurs objectifs analytiques.

2.7 Défis de l'ABSA

Les principaux problèmes et difficultés auxquels fait face l'ABSA sont essentiellement liés aux éléments qui composent ses définitions fondamentales:

- L'étude de la classification des sentiments dépend fortement du domaine. Cela signifie qu'un modèle de classification des sentiments entraîné dans un domaine, avec des paramètres adaptés, risque de mal performer dans un autre domaine. En effet, un même mot émotionnel peut

avoir des significations différentes selon le contexte.

- Chaque domaine a son propre langage spécifique, ce qui affecte la détection des sentiments et des aspects. L'ajustement des modèles sur des données spécifiques à un domaine ou l'utilisation de lexiques spécialisés permet d'améliorer la précision.

Par ailleurs, le développement de corpus annotés plus riches et variés, propres à différents domaines, représente une étape clé pour surmonter ces obstacles dans l'analyse du sentiment au niveau des aspects.

2.8 Large Modèle de Langue

Les Grandes Modèles de Langues (LLM) ont transformé le domaine de l'intelligence artificielle (IA). Ces modèles, caractérisés par leur grand nombre de paramètres et leurs architectures sophistiquées, ont démontré une capacité sans précédent à comprendre, générer et manipuler du texte de type humain.

Les LLM sont des systèmes d'IA actuelle conçus pour comprendre et générer du texte de type humain en fonction des informations qu'ils reçoivent. Ces modèles utilisent des techniques d'apprentissage profond, en particulier des réseaux de neurones avec de nombreux paramètres.

Les LLM fonctionnent sur la base d'une architecture de transformateur. L'architecture du transformateur a été introduite par Vaswani et al [46].

2.8.1 Principaux constituants des LLMs

- Encodage positionnel : Etant donné que les transformateurs traitent les données en parallèle, ils ont besoin d'un moyen de comprendre l'ordre des jetons dans une séquence. L'encodage positionnel injecte des informations sur la position du jeton dans l'entrée, ce qui permet au modèle de conserver une compréhension de la structure de la séquence, même si elle est traitée en parallèle.

- Cadre encodeur-décodeur : Le modèle original du transformateur est basé sur une structure encodeur-décodeur. L'encodeur prend une séquence d'entrée, la traite à travers plusieurs couches et crée une représentation interne. Le décodeur, à son tour, utilise cette

représentation pour générer une séquence de sortie, qui peut être une traduction, une classification ou un autre type de prédiction.

- Mécanisme d'attention multi-têtes : L'auto-attention permet au modèle de se concentrer sur les parties pertinentes de la séquence. L'attention multi-têtes exécute plusieurs opérations d'attention en parallèle, ce qui permet au modèle d'apprendre différents aspects de la séquence à la fois. Chaque tête peut se concentrer sur différentes parties de l'entrée, ce qui donne au transformateur plus de flexibilité et de précision.

- Couches d'anticipation : Après les couches d'attention, chaque jeton passe par un réseau neuronal à anticipation entièrement connecté. Ces couches aident le modèle à affiner sa compréhension de la relation de chaque jeton au sein de la séquence.

2.8.2 Fonctionnement des LLMs :

Les LLMs utilisent une architecture de transformateur, qui repose sur des mécanismes d'auto-attention. Les transformateurs peuvent traiter efficacement les données d'entrée, ce qui les rend adaptés à l'entraînement de grands modèles sur de vastes ensembles de données :

Pré-entraînement : Le modèle est pré-entraîné sur de grandes quantités de données textuelles. Au cours du pré-entraînement, le modèle apprend à prédire le mot suivant d'une phrase ou à combler les lacunes dans le texte. Ce processus permet au modèle de capturer la grammaire, la syntaxe, le contexte et les relations sémantiques dans le langage.

Paramètres : Le terme « Large/Grand » dans les LLM fait référence au grand nombre de paramètres. Les paramètres sont les variables internes que le modèle ajuste pendant l'entraînement pour comprendre et générer du texte.

Contexte et mécanisme d'attention : Le mécanisme d'attention permet au modèle de se concentrer sur différentes parties du texte d'entrée lorsqu'il génère une sortie. Cela aide le modèle à comprendre le contexte et les dépendances à longue portée dans les données.

Réglage fin : Après la préformation, le modèle peut être affiné pour des tâches spécifiques. Cela implique d'entraîner le modèle sur un ensemble de données plus petit lié à la tâche cible et d'ajuster ses paramètres.

Inférence : Pendant l'inférence, le modèle formé prend du texte d'entrée et génère du texte de sortie en fonction des modèles qu'il a appris lors du pré-entraînement et du réglage fin. Le modèle peut être utilisé pour des tâches telles que la complétion de texte, la traduction, la synthèse, la réponse à des questions, etc.

2.8.3 Apprentissage par transfert à partir de LLM

L'apprentissage par transfert consiste à exploiter les connaissances acquises en pré-entraînant un modèle sur un grand ensemble de données, puis en l'affinant sur une tâche spécifique. Les aspects essentiels de cet apprentissage des LLM :

- Pré-entraînement : Un LLM est formé sur un vaste ensemble de données de texte et de code, ce qui lui permet d'apprendre la compréhension et la représentation générales du langage.
- Réglage fin : Ce LLM pré-entraîné est ensuite adapté à une tâche spécifique en l'entraînant sur des ensembles de données pertinents plus petits. Cela repose sur des connaissances générales du pré-entraînement tout en se spécialisant dans le nouveau domaine.

L'apprentissage par transfert des LLM peut être utilisé à diverses fins, notamment les chatbots personnalisés, la synthèse automatique de texte, la traduction automatique et la création de contenu.

2.8.4 Méthodes de fine-tuning de LLM

L'ajustement fin des LLM est un processus d'apprentissage supervisé qui utilise un ensemble de données d'exemples étiquetés pour mettre à jour les pondérations des LLM et améliorer la capacité du modèle à effectuer des tâches spécifiques. Parmi les méthodes les plus efficaces pour le finetuning des LLM [47], [48] :

Méthode 1 : Instruction fine-tuning Affinement des instructions est une méthode d'adaptation d'un modèle, consistant à apprendre à partir d'exemples structurés en paires instruction-réponse. Ces exemples sont sélectionnés pour atteindre un objectif spécifique (ex. : générer un résumé ou traduire un texte). Le modèle apprend ainsi à répondre aux instructions en s'inspirant des exemples du corpus utilisé pour son affinage [42].

Méthode 2 : Full fine-tuning Le fine-tuning complet consiste à ajuster tous les paramètres

du modèle (ses poids) pour l'adapter à une tâche spécifique. Ce processus aboutit à une version entièrement mise à jour du modèle, avec des paramètres optimisés pour la nouvelle cible. Cependant, comme lors du pré-entraînement, cette méthode nécessite des ressources conséquentes (mémoire GPU/TPU, puissance de calcul) pour gérer les gradients, les optimiseurs et les autres éléments mis à jour pendant l'entraînement [42].

Méthode 3 : Parameter-efficient fine-tuning (PEFT) Les méthodes d'adaptation paramétrique efficace (PEFT), telles que LoRA , QLoRA , Adapters Layers , et Prefix/Prompt Tuning, visent à ajuster les grands modèles linguistiques (LLMs) avec un coût computationnel et une empreinte mémoire réduite. Ces techniques modifient uniquement une fraction des paramètres du modèle initial :

- LoRA insère des matrices de faible rang pour approximer les mises à jour des poids ;
- QLoRA combine quantification et LoRA pour optimiser davantage les ressources ;
- Adapters ajoutent de petites couches légères entre les modules existants ;
- Prefix/Prompt Tuning introduit des tokens apprenables en début de séquence pour guider la génération. Ces approches permettent de spécialiser les LLMs pour des tâches spécifiques tout en préservant la majorité des paramètres préentraînés, rendant le fine-tuning accessible même avec des modèles extrêmement volumineux (ex : Llama 65B). Elles sont particulièrement utiles dans des contextes à ressources limitées ou pour éviter le surapprentissage sur de petits jeux de données [42].

Méthode 4 : Reinforcement learning from human feedback (RLHF) L'apprentissage par renforcement à partir du feedback humain (RLHF) est une approche innovante qui vise à améliorer les modèles de langage en intégrant des retours humains au processus d'entraînement. Grâce à cette méthode, les modèles apprennent à produire des réponses plus précises et adaptées au contexte en s'ajustant aux préférences et corrections fournies par des évaluateurs humains. En exploitant l'expertise humaine, le RLHF permet non seulement une optimisation continue des performances du modèle, mais aussi une adaptation dynamique aux exigences réelles des utilisateurs, garantissant ainsi une efficacité accrue et une meilleure fidélité aux attentes humaines [43]. Cependant, le fine-tuning présente des limites. Si la tâche ou le domaine cible diffèrent

significativement des données de pré-entraînement, le modèle peut avoir du mal à s'adapter efficacement. De plus, l'ajustement fin peut entraîner des oublis catastrophiques : le modèle oublie une partie de ses connaissances générales au profit de la tâche spécialisée.

2.8.5 Avantages et inconvénients de l'apprentissage par transfert à partir de LLM

Les grands modèles de langage sont utiles pour une variété de tâches, et l'apprentissage par transfert aide ces modèles à fonctionner plus efficacement. Dans le même temps, ce processus comporte certaines limites à prendre en compte : plus vous comprendrez ces avantages et ces inconvénients, mieux vous serez préparé à travailler avec des LLM.

Avantages

- Réduction du temps de formation
- Utilisation plus efficace des ressources
- Sortie plus précise

Inconvénients

- Nécessite des données de haute qualité exemptes de biais de l'utilisateur
- Transfert accidentel de biais à partir de données d'entraînement existantes
- Préoccupations concernant la sécurité des données

2.8.6 Domaines d'application des LLM

Les entreprises de divers secteurs utilisent l'apprentissage par transfert des LLM pour créer des prévisions, prédire des tendances et automatiser des tâches [49]:

- **Marketing** : Certaines entreprises utilisent les LLM pour créer des supports marketing, tels que des textes publicitaires, des blogs et d'autres contenus. Les développeurs continuent d'améliorer les LLM existants pour mieux comprendre comment aligner le contenu sur la marque de l'entreprise et engager le public cible.

- Centres de contact : Les LLM offrent aux clients une communication humaine et des réponses personnalisées, qui sont la clé d'un bon service client. Les LLM permettent également aux clients de recevoir une réponse immédiate du support client en ligne et des chatbots, de sorte qu'ils n'ont pas à attendre l'aide.

- Santé : À l'aide d'un logiciel LLM, les professionnels de la santé peuvent organiser plus efficacement les données des patients. Il peut également être utile pour l'imagerie médicale pour détecter des anomalies et dans la recherche en soins de santé, où il est utilisé pour étudier des maladies rares ou les effets de maladies sur des sous-groupes.

- Education : Un LLM peut agir en tant que tuteur ou enseignant, en fournissant un contenu personnalisé à l'apprenant. Les éducateurs peuvent également utiliser les LLM pour différencier les leçons afin de répondre aux besoins des élèves ayant des capacités d'apprentissage différentes et de répondre aux besoins spécifiques des élèves à leur niveau d'enseignement.

2.9. Relation entre les LLM et l'apprentissage automatique

Les LLM sont un type de modèle d'apprentissage automatique. Plus précisément, ils entrent dans la catégorie de l'apprentissage profond, qui est un sous-ensemble de l'apprentissage automatique :

- Apprentissage automatique : Un vaste domaine de l'IA qui implique le développement d'algorithmes et de modèles pour apprendre des modèles à partir de données. L'objectif est de permettre aux systèmes de faire des prédictions et des décisions ou d'effectuer des tâches même s'ils ne sont pas explicitement programmés pour le faire.

- Deep Learning : sous-domaine de l'apprentissage automatique qui se concentre sur les réseaux neuronaux comportant plusieurs couches (réseaux neuronaux profonds). Ces réseaux sont particulièrement efficaces pour apprendre des représentations hiérarchiques et capturer des modèles complexes dans les données.

- LLM : ils sont basés sur les principes du Deep Learning. Ils utilisent un type spécifique d'architecture de réseau neuronal appelé transformateurs. L'architecture, combinée à un grand nombre de paramètres, permet à ces modèles de comprendre et de générer un texte de type humain en apprenant à partir de grandes quantités de données de pré-entraînement.

2.9.1 Facteurs clés de l'IA générative et des LLM

Les LLM sont également des exemples d'IA générative. L'IA générative fait référence à des modèles ou des systèmes qui ont la capacité de générer de nouveaux contenus, tels que du texte, des images ou d'autres types de données, en fonction des modèles et des informations qu'ils ont appris pendant la formation.

- Génération de texte : Les LLM sont conçus pour générer du texte de type humain.
- Applications diverses : Ces modèles sont polyvalents et peuvent être appliqués à diverses tâches de traitement du langage naturel, notamment la complétion de texte, la traduction et la synthèse.
- Formation sur des données diverses : Au cours de leur phase d'entraînement, ces modèles sont exposés à des ensembles de données divers et étendus contenant des exemples de langage humain.
- Génération conditionnelle et inconditionnelle : les LLM peuvent effectuer la génération de texte conditionnelle et inconditionnelle. La génération conditionnelle génère du texte en fonction d'une invite ou d'un contexte donné, tandis que la génération inconditionnelle génère du texte sans invite spécifique.
- Réglage fin pour des tâches spécifiques : Ces modèles peuvent être affinés pour des tâches spécifiques, ce qui leur permet de générer du contenu adapté à des applications ou à des domaines particuliers.

2.10 LLM dans l'analyse de sentiments

L'exploitation des grands modèles de langage (LLM) tels que BERT, RoBERTa ou GPT-4 représente une avancée majeure pour l'analyse des sentiments, notamment dans le cadre de l'analyse des sentiments à base d'aspects (ABSA). Grâce à leur entraînement sur d'immenses volumes de textes, ces modèles sont capables de comprendre le langage naturel de manière fine et contextuelle, permettant ainsi d'identifier non seulement l'opinion globale exprimée dans un avis, mais aussi les sentiments associés à des aspects spécifiques, comme la qualité des plats, le service ou l'ambiance dans le domaine de la restauration. Plusieurs stratégies peuvent être mobilisées

selon les ressources disponibles, allant du fine-tuning supervisé sur des données annotées, au zero-shot learning via des prompts soigneusement conçus, permettant d'exploiter les capacités des LLM sans phase d'apprentissage supplémentaire. En outre, leur adaptabilité à des domaines spécifiques peut être renforcée par des techniques comme le pré-entraînement adaptatif (DAPT) ou l'adaptation de domaine non supervisée (UDA). Bien que leur utilisation soulève certaines limites, telles que les coûts computationnels ou la sensibilité à la formulation des requêtes, les LLM offrent un potentiel considérable pour analyser automatiquement et efficacement des opinions complexes exprimées dans des textes issus des réseaux sociaux ou des plateformes d'avis [50].

2.11 LLM dans l'apprentissage en contexte

Les LLM présentent des capacités d'apprentissage en contexte améliorées par rapport aux modèles plus petits. L'apprentissage contextuel fait référence à la capacité du modèle à comprendre et à générer du texte en réponse à une invite ou à une entrée donnée.

- Compréhension contextuelle : les modèles plus grands, avec des millions ou des milliards de paramètres, ont une plus grande capacité à capturer et à comprendre le contexte.
- Compréhension sémantique : Les LLM ont une meilleure compréhension de la sémantique, ce qui signifie qu'ils peuvent comprendre les nuances et les subtilités du langage.
- Adaptabilité contextuelle : les LLM peuvent s'adapter plus efficacement aux changements de contexte au sein d'une conversation ou d'une invite.
- Conversations à plusieurs cycles : dans le contexte de l'IA conversationnelle, les LLM peuvent gérer plus efficacement les conversations à plusieurs cycles.

2.11.1 Facteurs clés pour le choix de LLM

- Adéquation des tâches : Evaluation des performances du logiciel sur des tâches spécifiques pertinentes pour votre application. Certains modèles peuvent exceller dans les tâches NLP, telles que la complétion de texte, la synthèse, la traduction ou la réponse à des questions.

- Capacités de fine tuning : Vérification si le modèle permet un réglage fin. Le réglage fin vous permet d'adapter le modèle pré-entraîné à votre domaine ou à votre tâche spécifique, en améliorant ses performances de manière ciblée.
- Communauté et soutien : une communauté solide et une bonne documentation peuvent être des ressources précieuses pour le dépannage, l'apprentissage et la mise à jour.
- Licence et coût : certains modèles peuvent avoir des limites d'utilisation ou exiger un paiement pour un accès au-delà d'un certain seuil.

2.12 Conclusion

En résumé, l'ABSA représente une approche fine et puissante de l'analyse des sentiments, permettant de capturer non seulement l'opinion globale d'un texte, mais aussi les sentiments associés à des aspects précis. Grâce aux approches avancées par DL, et à la richesse des jeux de données disponibles, l'ABSA continue de progresser et de s'adapter à des contextes variés. Ces avancées ouvrent la voie à des applications exploitant l'IA générative, les LLMs et d'autres techniques de plus en plus performantes, dans des domaines tels que le commerce, l'éducation, la santé, entre autres.

Chapitre 3

Implémentation et Résultats

3.1 Introduction	45
3.2 Formulation du problème ABSA	45
3.3 Choix techniques utilisés	46
3.4 Architecture globale du système ABSA	46
3.5 Structure globale du pipeline ABSA	47
3.6 Structure technique du pipeline ABSA	48
3.7 Jeu de données utilisé	50
3.8 Processus séquentiel détaillé de la solution ABSA	51
3.9 Résultats expérimentaux	53
3.9.1 Dataset après formatage	56
3.9.2 Présentation et évaluation des résultats	56
3.9.3 Perte d'entraînement	57
3.10 Conclusion	59

3.1 Introduction

La mise en œuvre de la tâche ABSA repose sur l'utilisation des capacités avancées d'un LLM pré-entraîné en lui appliquant une technique de fine-tuning via l'apprentissage par transfert. Cette technique nous permet de bénéficier des larges connaissances linguistiques pré-acquises et de les adapter efficacement au domaine spécifique de la restauration.

L'objectif principal est, donc, d'effectuer une analyse fine et contextualisée des sentiments exprimés dans des fils de discussion concernant les avis de clients via un LLM finement ajusté.

3.2 Formulation du problème ABSA

Etant donné un ensemble de phrases ou avis textuels $D = \{S_1, \dots, S_N\}$. Chaque phrase " S_i " contient potentiellement plusieurs mots ou tokens $S_i = \{w_1, w_2, \dots, w_n\}$, où w_i est le i -ème mot de la séquence, annotées avec des quadruplets $Q_i = \{(a_j, c_j, s_j, o_j)\}$. Le but est de construire une fonction linguistique finetunée.

$$LLM_{\theta} : S_i \rightarrow Q_i = \{(a_j, c_j, s_j, o_j)\} \text{ (pour } j=1, k \text{ et } k \leq N),$$

Où l'entrée du modèle LLM_{θ} est une phrase " S_i " souvent accompagnée d'un prompt système pour guider le format de la réponse : Input = (SystemPrompt, UserQuery) et la sortie est une séquence de tuples de sentiment en quadruplets au format structuré $\{(a_j, c_j, s_j, o_j)\}$ (pour $j=1, k$), séparés par « ; » :

- Aspect a_j : Entité ou caractéristique explicite mentionnée dans S (ex : "service", "nourriture").
- Catégorie c_j : Classe sémantique de A_j (ex : "service général", "qualité des plats").
- Opinion o_j : Expression subjective associée à A_j (ex : "lent", "délicieux").
- Polarité s_j : Sentiment $\in \{\text{positif, négatif, neutre}\}$.

Le modèle finetuné est fondé sur LLM (Llama-3.1), fine-tuné via LoRA et entraîné avec SFT, tel que $LLM_{\theta}(S_i) \approx Q_i$, où θ représente les paramètres ajustés du modèle.

3.3 Choix techniques utilisées

Technique	Description	Avantage
LLM	Meta-Llama-3.1-8B-Instruct	Modèle puissant et adapté aux conversations
Quantification 4-bit	load_in_4bit=True	Réduction de la consommation mémoire
LoRA	r=32, modules cibles définis	Fine-tuning efficace avec peu de paramètres
Template de chat	llama-3.1via Unsloth	Format compatible avec le modèle instruct
SFT	TRL/SFTTrainer	Entraînement supervisé sur les réponses
Masking	train_on_responses_only()	Focus uniquement sur la partie générée
Evaluation	Precision / Recall / F1	Mesure objective de la performance
Visualisation	Matplotlib / Seaborn	Compréhension approfondie des résultats

3.3 Architecture globale du système ABSA

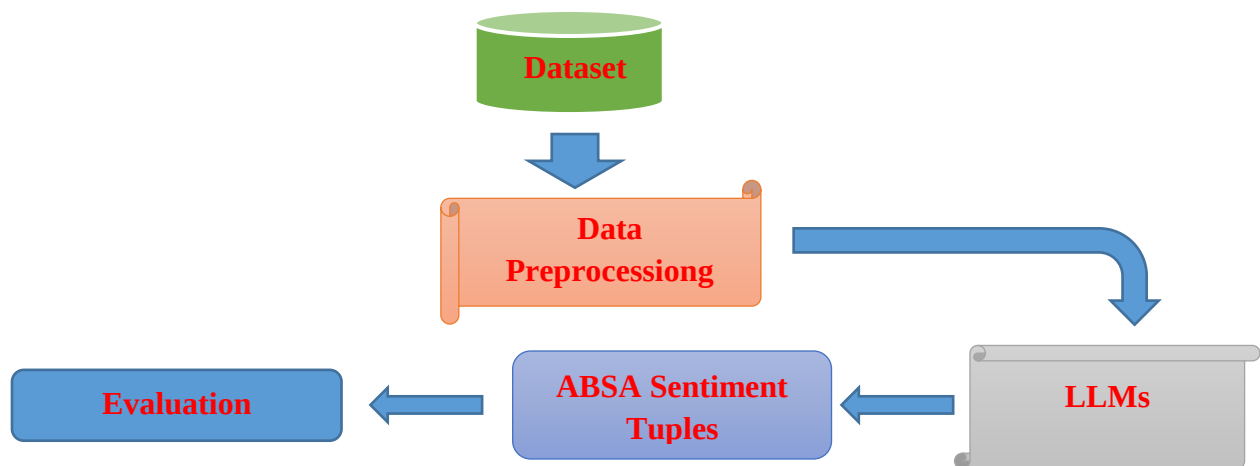


Figure 3.1 Architecture globale du système ABSA

Cette figure décrit un pipeline complet pour l'ABSA dans le domaine de la restauration. Chaque étape joue un rôle essentiel dans le traitement des données, depuis la collecte jusqu'à l'évaluation finale des résultats. L'utilisation de LLM avancés permet d'extraire des informations exploitables, précises et structurées à partir des avis clients.

3.5 Structure globale du pipeline ABSA

1. Prétraitement

Les données brutes sont converties en prompts structurés utilisant le template de chat Llama-3.1 Instruct .

2. Fine-tuning du modèle

Un modèle pré-entraîné M pretrained est fine-tuné via LoRA (Low-Rank Adaptation) :

$$M_{\theta} = \text{LoRA} (M_{\text{pretrained}}, r=32)$$

Où r est le rang de la matrice d'adaptation.

3. Entraînement supervisé

Utilisation de SFTTrainer avec supervision uniquement sur la partie réponse de l'assistant (masquage des tokens utilisateur) :

$$L = \sum_{t=1}^T \log P_{\theta} (y_t | x_{\text{response}})$$

Où x_{response} est la partie générée par l'assistant dans le prompt.

4. Evaluation du modèle

Pour chaque phrase S_i , on extrait les quadruplets prédits Q_p et les quadruplets réels Q_i . Les metriques classiques de Precision, Recall et F1-score sont alors calculées pour mesurer la performance du modèle.

3.6 Structure technique du pipeline ABSA

1) Modèle de langage utilisé (LLM)

Le modèle choisi est Meta-Llama-3.1-8B-Instruct, un modèle de langage pré-entraîné de grande taille (8 milliards de paramètres), optimisé via la bibliothèque Unsloth pour accélérer le fine-tuning et réduire la consommation mémoire.

Caractéristiques clés pour le choix du modèle :

- Taille modérée adaptée au matériel grand public.
- Version "Instruct" alignée pour la génération conversationnelle.
- Support de la quantification en 4 bits pour réduire l'empreinte mémoire.

2) Fine-tuning avec LoRA (Low-Rank Adaptation)

Pour réduire la complexité computationnelle, on utilise LoRA avec un rang $r=32$, ce qui permet de ne mettre à jour qu'une fraction des paramètres du modèle initial.

Paramètres :

- Rang LoRA : $r=32$
- Modules cibles : ["q_proj", "k_proj", "v_proj", "o_proj", "gate_proj", "up_proj", "down_proj"]
- Dropout LoRA : 0

3) Prétraitement des données et mise en forme des prompts

Les données sont extraites du dataset HuggingFace NEUDM/absa-quad ¹, contenant des annotations sous forme de listes imbriquées.

La fonction `formatting_prompts_func()` convertit ces annotations en prompt structuré utilisant le template Llama-3.1 Instruct :

¹ [NEUDM/absa-quad · Datasets at Hugging Face](#)

4) Entraînement du modèle (SFT - Supervised Fine-Tuning)

Utilisation de SFTTrainer de TRL (Transformers Reinforcement Learning) avec les paramètres suivants :

Configuration d'entraînement :

- Batch size par GPU : 2
- Accumulation de gradients : 4
- Longueur maximale de séquence : 2048 tokens
- Optimiseur : AdamW 8-bit
- Taux d'apprentissage : 2×10^{-4}
- Époques : 5
- Pas maximaux : 60 (pour tests rapides)

Masquage des entrées utilisateur : Seule la réponse de l'assistant est prise en compte pour le calcul de la perte grâce à la fonction `train_on_responses_only()`.

5) Evaluation du modèle

Afin d'évaluer les performances du modèle après l'étape de fine-tuning, des métriques standards de classification sont à exploiter : la précision (Precision), le rappel (Recall) et le F1-score. Ces mesures sont particulièrement adaptées au cadre de l'ABSA, où l'objectif est de générer des tuples structurés de sentiments d'ABSA à partir d'avis non structurés, en l'occurrence des avis de clients dans le domaine de la restauration.

La précision indique dans quelle mesure les tuples générés par le modèle sont corrects. Plus précisément, elle représente la proportion de tuples exacts (c'est-à-dire correctement identifiés dans les trois composantes : aspect, opinion, polarité) parmi l'ensemble des tuples produits par le modèle. Par exemple, si le modèle génère ("service", "rapide", "positif"), la précision évalue à quelle fréquence ce genre de sortie correspond bien à une annotation de référence.

Le rappel mesure la capacité du modèle à détecter l'ensemble des tuples pertinents présents dans un texte donné. Il s'agit de la proportion de tuples corrects identifiés par le modèle parmi tous ceux que l'on aurait dû retrouver selon les annotations de référence. Un rappel élevé signifie que le modèle n'omet que peu d'informations importantes, ce qui est crucial dans des domaines

riches en nuances comme la restauration, où un même avis peut évoquer plusieurs aspects (par exemple : "qualité des plats", "temps d'attente", "ambiance").

Le F1-score constitue une mesure synthétique qui équilibre précision et rappel. Il est calculé comme la moyenne harmonique des deux et permet d'avoir une vue d'ensemble sur les performances du modèle, en particulier lorsqu'un compromis entre exactitude et couverture est nécessaire. Dans notre contexte, un F1-score élevé suggère que le modèle affiné est à la fois précis et complet dans l'extraction des tuples de sentiment.

3.7 Jeu de données utilisé (Dataset):

Le jeu de données d'entraînement a été formaté spécifiquement pour l'ABSA. Chaque ligne du dataset représente une instance d'entraînement et contient plusieurs informations clés :

- **input** : La phrase originale de l'avis client, servant de texte source pour l'analyse.
- **output** : La sortie attendue du modèle, qui semble être une structure complexe combinant l'aspect (**aspect**), le sentiment (**sentiment**) et potentiellement un terme d'opinion (**opinion**) ou un détail (**holder**) lié à l'aspect. Cette colonne guide le modèle sur la forme de la réponse à générer.
- **instruction** : Une consigne explicite décrivant la tâche à accomplir (par exemple, "Extracting aspect terms and their corresponding sentiment polarity"), utile pour les modèles de type prompt-based ou instruction-tuned.
- **aspect** : Le terme spécifique ou l'entité (l'aspect) extrait(e) de la phrase.
- **opinion** : Le terme ou la phrase qui exprime l'opinion sur l'aspect.
- **sentiment** : La polarité de sentiment associée à l'aspect (par exemple, "negative").
- **holder** : Une information supplémentaire, potentiellement un descripteur ou une propriété liée à l'aspect ou à l'opinion.
- Les colonnes **dataset**, **situation**, **label**, et **extra** fournissent des métadonnées additionnelles sur l'origine ou la catégorisation de l'entrée."

Ce jeu de données est organisé en trois sous-ensembles principaux :

- Train : 2098 exemples.

- Validation : 525 exemples.
- Test : 1081 exemples.

	task_type	dataset	input	output	situation	label	extra	instruction	aspect	opinion	sentiment	holder
0	generation	absa-quad	['The wait here is long for dim sum, but if y...	['wait', 'service general', 'negative', 'long...]	none			Task: Extracting aspect terms and their corres...	wait	service general	negative	long
1	generation	absa-quad	['The wait here is long for dim sum, but if y...	['wait', 'service general', 'negative', 'long...]	none			Task: Extracting aspect terms and their corres...	atmosphere	ambiance general	negative	raucous
		absa-	['The wait here is	['wait', 'service				Task: Extracting aspect				restaurant

Figure 3.3 dataset public NEUDM/absa-quad

3.8 Processus séquentiel détaillé de la solution ABSA

Pour inférer des tuples de sentiments ABSA exploitant la technique du fine-tuning, appliquée aux avis de restauration, la démarche à suivre est structurée en plusieurs étapes clés formant un pipeline cohérent. Le processus séquentiel commence par la collecte de données, suivie de la préparation du jeu de données. Ensuite, un LLM approprié est choisi et configuré avant son affinage. Après l'affinage, le modèle est exploité pour la génération de tuples ABSA. Les résultats sont ensuite évalués, et à la fin, la solution est déployée :

a. Collecte de données : Cette phase de préparation des données est fondamentale pour la performance de tout modèle d'apprentissage automatique. Le volume et la diversité des données sont cruciaux pour garantir la robustesse de notre modèle.

b. formatage de dataset : La figure montre le flux de données depuis le dataset ABSA original (avec les colonnes 'input', 'output', 'instruction') à travers les étapes de la fonction de formatage (formatting_prompts_func) et de l'application du chat template du tokenizer (tokenizer.apply_chat_template), jusqu'au dataset formaté final (avec les colonnes 'text' et 'conversations') prêt pour l'entraînement du LLM.

Les étapes de processus de formatage de dataset

- Format Data Function : Conversion des chaînes en objets Python
- Parse Input/Output : Conversion des chaînes en objets Python
- Build Conversation : Création du dialogue
- Apply Chat Template : Formatage selon le modèle
- Create Dataset : Application du formatage

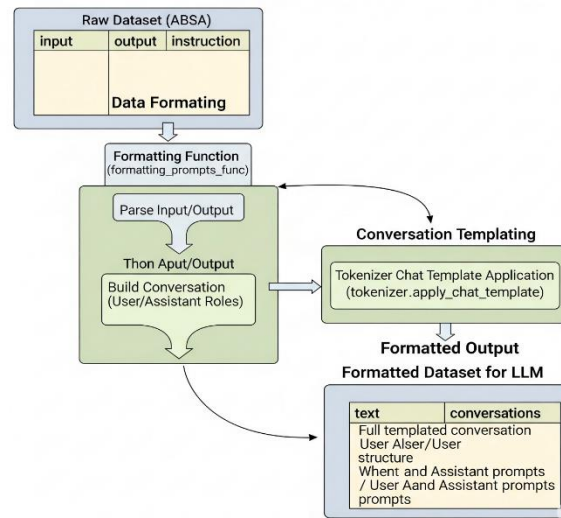


Figure 3.2 processus de formatage des données

c. Choix du LLM : Le choix du LLM (Large Language Model) est une étape déterminante. Après une revue des modèles disponibles, nous avons exploité un modèle basé sur l'architecture Transformer, pré-entraîné sur de vastes corpus textuels. Des considérations telles que la taille du modèle, ses capacités de génération et de compréhension, ainsi que sa disponibilité pour l'affinage, ont guidé notre choix. Nous avons examiné des modèles tels que **Meta-Llama-3.1-8B-Instruct** pour leur équilibre entre performance et ressources nécessaires.

d. Choix de la méthode d'amélioration du LLM : Pour adapter le LLM choisi à la tâche spécifique de l'ABSA, nous avons privilégié la méthode de l'affinage (Fine-tuning). Plus précisément, nous avons utilisé des techniques d'affinage léger (**Low-Rank Adaptation - LoRA**)

pour optimiser les performances du modèle tout en minimisant les ressources de calcul nécessaires et le risque de sur-apprentissage. Cette approche permet de mettre à jour un petit nombre de paramètres supplémentaires ou de couches du réseau, plutôt que l'ensemble du modèle, rendant l'entraînement plus efficace et plus rapide.

e. Evaluation des résultats : L'évaluation des résultats est essentielle pour mesurer la performance de notre solution ABSA. Nous utilisons un ensemble de métriques standard pour évaluer à la fois l'extraction des aspects et la classification des sentiments. Ces métriques incluent la précision (**Precision**), le rappel (**Recall**), et le score F1 (**F1-score**) pour l'identification des aspects et la classification de leur polarité. Nous avons également recours à des matrices de confusion pour une analyse plus détaillée des erreurs.

3.9 Résultats expérimentaux :

La figure ci-dessous illustre la distribution des sentiments dans le jeu de données d'entraînement utilisé pour notre tâche d'ABSA. On observe un déséquilibre marqué entre les différentes classes. La majorité des échantillons sont étiquetés comme positifs, avec environ 1 450 occurrences, ce qui représente une part prédominante du corpus. En comparaison, la classe négative ne compte qu'environ 550 échantillons, soit près de trois fois moins. La classe neutre, quant à elle, est très peu représentée, avec moins de 100 échantillons, ce qui la rend statistiquement marginale. Ce déséquilibre pourrait introduire un biais d'apprentissage en faveur des sentiments positifs, rendant plus difficile la détection correcte des sentiments négatifs ou neutres. Il est donc crucial de prendre en compte cette répartition lors de l'entraînement du modèle, en envisageant, par exemple, des techniques de rééchantillonnage ou d'ajustement de la fonction de perte afin d'atténuer l'impact de ce déséquilibre sur les performances globales du système.

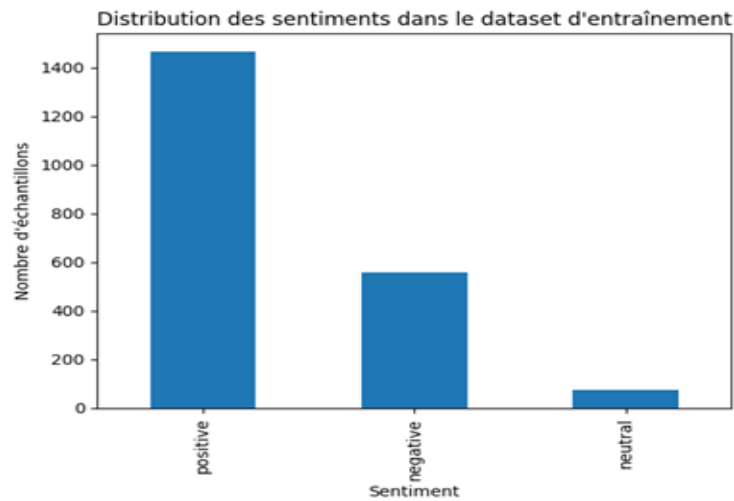


Figure 3.4 : Distribution des sentiments dans le jeu de données d'entraînement

. Les occurrences de chaque catégorie dans la colonne category du DataFrame df_train se présentent comme suit :

Category	
food quality	744
restaurant general	440
service general	348
ambiance general	201
restaurant miscellaneous	96
food style_options	82
restaurant prices	54
drinks quality	42
location general	32
drinks style_options	28
food prices	27
drinks prices	4

Figure 3.4 Tableau de distribution des catégories prédéfinies

La Figure 3.5 montre une répartition détaillée des sentiments (positif, négatif, neutre) pour chaque catégorie d'aspect identifiée dans le corpus. Il apparaît clairement que la catégorie "food quality" domine en termes de volume, avec 542 mentions positives, 175 négatives et 27 neutres, ce qui reflète l'importance accordée par les clients à la qualité de la nourriture. La catégorie "restaurant general" arrive en deuxième position, avec 354 commentaires positifs et 79 négatifs, indiquant qu'une évaluation globale du restaurant est fréquemment formulée. La catégorie "service general" suit, avec une forte présence de sentiments, notamment 177 positifs contre 164 négatifs, soulignant la polarisation fréquente des opinions sur le service. À l'inverse, certaines catégories comme "drinks prices", "drinks style_options" ou "location general" présentent très peu d'occurrences, ce qui suggère qu'elles sont moins fréquemment mentionnées ou jugées moins pertinentes par les utilisateurs. Enfin, la distribution globale est fortement biaisée vers les sentiments positifs, bien que certaines catégories révèlent une proportion non négligeable de commentaires négatifs. Cette analyse est précieuse pour orienter les priorités d'amélioration dans un contexte de satisfaction client.

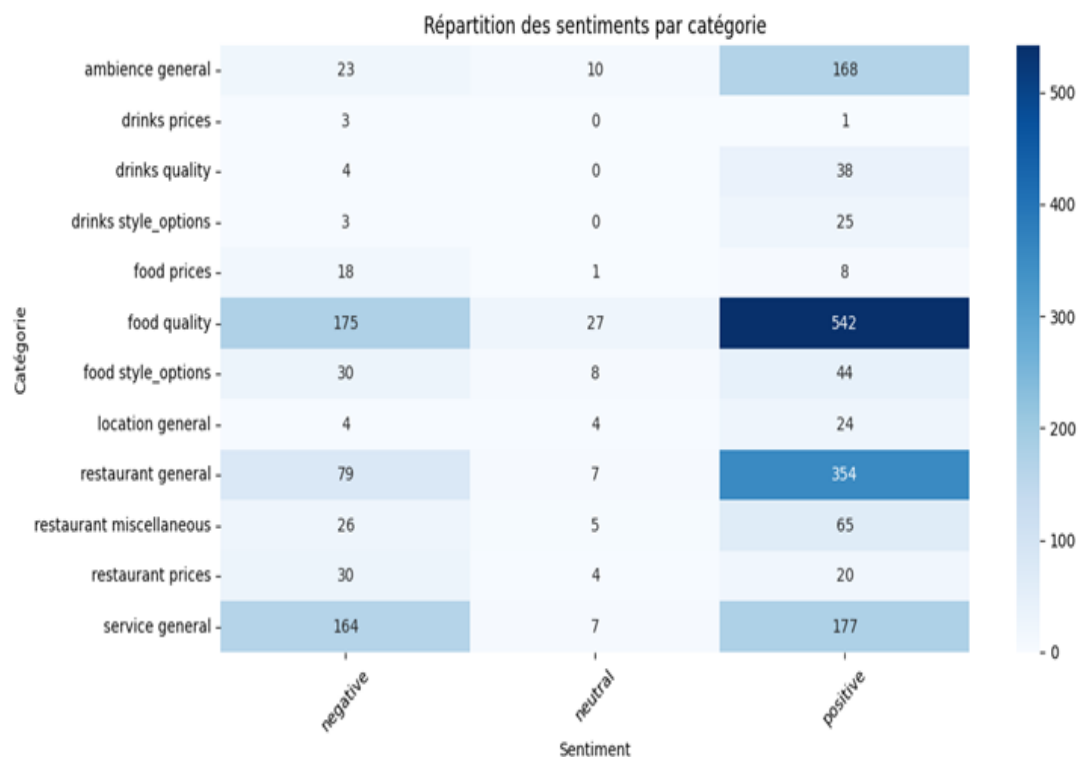


Figure 3.5 Répartition des sentiments par catégorie

3.9.1 Dataset après formatage :

Initialement, le dataset se présente sous la forme Dataset ({features: ['task_type', 'dataset', 'input', 'output', 'situation', 'label', 'extra', 'instruction'], num_rows: 2098}). Chaque entrée contient des informations brutes, notamment la phrase d'entrée (input), la sortie ABSA attendue sous forme de chaîne (output), et une instruction (instruction).

Après l'étape de formatage, effectuée par la fonction **formatting_prompts_func** et l'application du chat_template du tokenizer, le dataset est transformé en Dataset({features: ['text', 'conversations'], num_rows: 2098}). Toutes les colonnes d'origine sont remplacées par deux nouvelles colonnes essentielles :

- **conversations** : Cette colonne contient une représentation structurée du dialogue entre un utilisateur et un assistant. Comme illustré par dataset[0], il s'agit d'une liste de dictionnaires, où chaque dictionnaire spécifie un rôle ('user' ou 'assistant') et le contenu correspondant ('content'). Le message de l'utilisateur inclut généralement l'instruction et la phrase à analyser, tandis que le message de l'assistant contient la réponse attendue, c'est-à-dire les quads ABSA formatés.
- **text** : C'est la colonne finale et la plus importante pour l'entraînement du LLM. Elle contient une chaîne de caractères unique qui est la concaténation des messages de la conversation, mais formatée avec les balises spécifiques du modèle (comme <|begin_of_text|>, <|start_header_id|>, <|eot_id|>). Ce format spécifique permet au LLM de comprendre la structure du dialogue et de différencier les rôles et les contenus, ce qui est essentiel pour qu'il apprenne à générer la réponse correcte (output) lorsqu'il est confronté à un input donné.

3.9.2 Présentation et Évaluation des résultats

Les résultats présentés détaillent la performance du LLM fine-tuné avec LoRA pour la tâche ABSA :

	Precision	Recall	F1-Score
Train Metrics	0.6929	0.6263	0.6580
Validation Metrics	0.6810	0.6131	0.6453
Test Metrics	0.6310	0.5728	0.6005

Tableau 5. Métriques de performances

Les performances du modèle ont été évaluées sur les ensembles de validation, de test et d'entraînement en utilisant la Précision, le Rappel et le Score F1. Ces métriques sont cruciales pour évaluer la capacité du modèle à identifier correctement les aspects et leurs sentiments associés :

Performance Générale : Le modèle démontre une capacité raisonnable à effectuer la tâche d'ABSA, avec des scores F1 autour de 0.60 à 0.65 sur les différents ensembles de données.

Généralisation : Les métriques sur l'ensemble de validation (F1: 0.6453) et l'ensemble de test (F1: 0.5997) sont légèrement inférieures à celles de l'ensemble d'entraînement (F1: 0.6580). Cela est attendu et suggère que le modèle généralise relativement bien aux données non vues, bien qu'il y ait une légère baisse de performance sur les données de test, ce qui est courant.

Équilibre Précision/Rappel : Les scores de Précision et de Rappel sont relativement proches, indiquant un bon équilibre entre la capacité du modèle à éviter les faux positifs (Précision) et à trouver tous les éléments pertinents (Rappel).

3.9.3 Perte d'entraînement:

La perte d'entraînement montrent une diminution significative et constante de la perte au fil des étapes. Partant d'une perte initiale élevée (3.706400 à l'étape 1), la valeur diminue rapidement pour atteindre des niveaux beaucoup plus bas (par exemple, 0.058900 à l'étape 60). Cette tendance indique que le modèle apprend efficacement à partir des données d'entraînement et que le processus de fine-tuning converge bien.

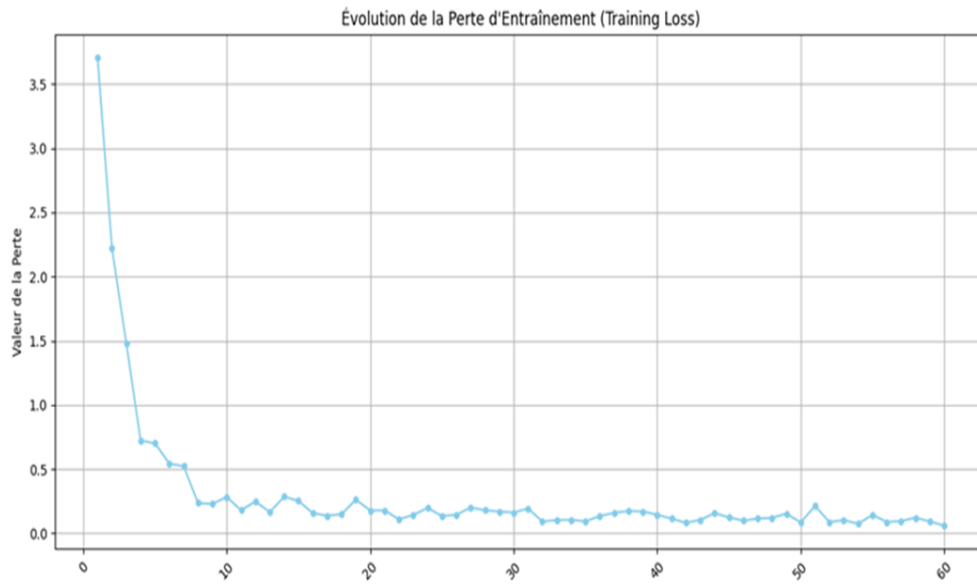


Figure 3.6 Évolution de la Perte d'Entraînement (Training Loss) :

En conclusion, les résultats indiquent que le fine-tuning du LLM avec LoRA a été efficace pour adapter le modèle à la tâche d'ABSA, comme en témoigne la diminution de la perte d'entraînement et les scores de performance obtenus sur les ensembles de validation et de test. Ces performances sont prometteuses pour l'application de ce modèle dans l'analyse des sentiments basée sur les aspects

3.10 Conclusion

Ce chapitre a permis de concrétiser la solution proposée en mettant en œuvre les différentes étapes nécessaires à sa réalisation. La construction du dataset a impliqué la collecte de données d'un dataset de référence (NEUDM/absa-quad), structurées sous forme de textes annotés. Pour la modélisation, un modèle de langage de grande taille (LLM) approprié a été sélectionné, en utilisant la méthode du Fine-Tuning (plus spécifiquement l'affinage léger LoRA). Ensuite, l'implémentation s'est déroulée selon un pipeline séquentiel bien défini, qui a d'abord transformé les données brutes en un format conversationnel structuré (étapes de formatage et de construction du nouveau dataset avec les colonnes 'text' et 'conversations'), puis a configuré l'entraînement du LLM affiné avec LoRA. Des tests ont été réalisés pour évaluer les performances des modèles

sous des métriques standards telles que la Précision, le Rappel et le score F1, suivis d'une analyse critique des résultats. Grâce à cette solution, il sera possible d'aller au-delà de l'analyse statique des sentiments pour offrir une compréhension dynamique et relationnelle des interactions sociales en ligne. Ce projet contribue ainsi à enrichir les méthodes actuelles d'ABSA avec une dimension sociale innovante, ouvrant la voie à des analyses plus pertinentes et actionnables.

Conclusion générale

L'arrivée du web 2.0 a profondément transformé la manière dont les entreprises perçoivent et interagissent avec leurs clients. Les avis en ligne, notamment dans le secteur de la restauration, représentent une source de connaissances précieuses permettant de mieux comprendre les attentes, les préférences et les insatisfactions des consommateurs. Dans ce contexte, l'ABSA s'est imposée comme un outil essentiel pour capter des opinions fines et contextualisées, en associant chaque sentiment exprimé à un aspect spécifique du service ou du produit.

Ce mémoire a abordé une approche moderne et efficace pour résoudre le problème la tâche d'ABSA en exploitant les capacités avancées du LM. Plutôt que de recourir à des architectures complexes et spécialisées, notre méthode repose sur la solution d'adaptation par fine-tuning de modèles pré-entraînés. Cette stratégie permet de tirer parti des connaissances linguistiques pré-acquises dans ces modèles tout en les adaptant aux spécificités de la restauration.

Les expérimentations menées et les résultats obtenus démontrent que les LLM, une fois adaptés, sont capables d'effectuer une analyse fine des sentiments avec un bon niveau de précision.

Cependant, ce travail ouvre également la voie à plusieurs perspectives d'amélioration. Il serait pertinent, par exemple, d'explorer l'intégration de techniques de détection automatique d'aspects dans des contextes non annotés, de renforcer la robustesse des modèles face au langage informel typique des réseaux sociaux, ou encore d'évaluer l'approche sur d'autres domaines et d'autres langues que la restauration pour tester son pouvoir de généralisation. Par ailleurs, l'usage de modèles multilingues ou la combinaison avec des signaux multimodaux (images, vidéos) pourrait enrichir davantage l'analyse.

En conclusion, cette étude montre que les LLM fine-tuné représentent une avancée significative pour l'ABSA appliquée aux avis du dataset public NEUDM. Elle offre une nouvelle voie vers des analyses plus intelligentes, flexibles et contextuellement riches, au service d'une meilleure compréhension de la voix de l'utilisateur.

Bibliographies

- [1] Pang, B., Lee, L., 2008. Opinion mining and sentiment analysis. *Found. Trends Inf. Retr.* 2(1–2), 1–135.
- [2] Merriam-Webster, n.d. College – definition. Available at: <https://www.merriam-webster.com/dictionary/college>
- [3] Madhoushi, Z., Hamdan, A.R., Zainudin, S., 2015. Sentiment analysis techniques in recent works. In: *2015 Science and Information Conference (SAI)*, IEEE, 288–291.
- [4] Taboada, M., Brooke, J., Tofiloski, M., Voll, K., Stede, M., 2011. Lexicon-based methods for sentiment analysis. *Comput. Linguist.* 37, 267–307.
- [5] Mitra, A., 2020. Sentiment analysis using machine learning approaches (Lexicon based on movie review dataset). *J. Ubiquitous Comput. Commun. Technol.* 2, 145–152.
- [6] Agarwal, B., Mittal, N., 2016. Machine learning approach for sentiment analysis. In: *Prominent Feature Extraction for Sentiment Analysis*, 21–45.
- [7] Kang, H., Yoo, S.J., Han, D., 2012. Senti-lexicon and improved Naïve Bayes algorithms for sentiment analysis of restaurant reviews. *Expert Syst. Appl.* 39, 6000–6010.
- [8] Ahmad, M., Aftab, S., Ali, I., 2017. Sentiment analysis of tweets using SVM. *Int. J. Comput. Appl.* 177, 25–29.
- [9] Nigam, K., Lafferty, J., McCallum, A., 1999. Using maximum entropy for text classification. In: *IJCAI-99 Workshop on Machine Learning for Information Filtering*, 61–67.
- [10] Myles, A.J., Feudale, R.N., Liu, Y., Woody, N.A., Brown, S.D., 2004. An introduction to decision tree modeling. *J. Chemom.* 18, 275–285.
- [11] Daeli, N.O.F., Adiwijaya, A., 2020. Sentiment analysis on movie reviews using information gain and K-nearest neighbor. *J. Data Sci. Its Appl.* 3, 1–7.
- [13] Chen, Y., 2015. Convolutional neural network for sentence classification. University of Waterloo.
- [14] Liu, P., Qiu, X., Huang, X., 2016. Recurrent neural network for text classification with multi-task learning. arXiv preprint arXiv:1605.05101.
- [16] Hochreiter, S., Schmidhuber, J., 1997. Long short-term memory. *Neural Comput.* 9(8), 1735–1780. <https://doi.org/10.1162/neco.1997.9.8.1735>
- [15] Zaremba, W., Sutskever, I., Vinyals, O., 2014. Recurrent neural network regularization. arXiv preprint arXiv:1409.2329.
- [17] Xu, J., Chen, D., Qiu, X., Huang, X., 2016. Cached long short-term memory neural networks for document-level sentiment classification. arXiv preprint arXiv:1610.04989.

- [18] Abdelgwad, M.M., Soliman, T.H.A., Taloba, A.I., Farghaly, M.F., 2022. Arabic aspect based sentiment analysis using bidirectional GRU based models. *J. King Saud Univ. Comput. Inf. Sci.* 34, 6652–6662.
- [19] Parmar, N., Vaswani, A., Uszkoreit, J., Kaiser, L., Shazeer, N., Ku, A., Tran, D., 2018. Image transformer. In: *Proc. Int. Conf. Mach. Learn.*, PMLR, 4055–4064.
- [20] Appel, O., Chiclana, F., Carter, J., Fujita, H., 2016. A hybrid approach to the sentiment analysis problem at the sentence level. *Knowl. Based Syst.* 108, 110–124.
- [21] Liu, B., 2012. *Sentiment Analysis and Opinion Mining*. Morgan & Claypool Publishers.
- [22] Medhat, W., Hassan, A., Korashy, H., 2014. Sentiment analysis algorithms and applications: a survey. *Ain Shams Eng. J.* 5(4), 1093–1113.
- [23] Cambria, E., Schuller, B., Xia, Y., Havasi, C., 2013. New avenues in opinion mining and sentiment analysis. *IEEE Intell. Syst.* 28(2), 15–21.
- [24] Zhang, L., Wang, S., Liu, B., 2018. Deep learning for sentiment analysis: a survey. *WIREs Data Min. Knowl. Discov.* 8(4), e1253.
- [25] Behdenna, S., Barigou, F., Belalem, G., 2018. Document level sentiment analysis: a survey. *EAI Endorsed Trans. Context-Aware Syst. Appl.* 4, e2.
- [26] Çano, E., Morisio, M., 2019. Word embeddings for sentiment analysis: a comprehensive empirical survey. *arXiv preprint arXiv:1902.00753* <https://arxiv.org/abs/1902.00753>
- [27] Wen, J., Wu, Y., Li, Y., Zhang, J., 2020. Speculative text mining for document-level sentiment classification. *Neurocomputing* 412, 52–61. <https://doi.org/10.1016/j.neucom.2020.06.024>.
- [28] Wu, L., Morstatter, F., Liu, H., 2018. SlangSD: building, expanding and using a sentiment dictionary of slang words for short-text sentiment classification. *Lang. Resour. Eval.* 52, 839–852.
- [29] Zhang, W., Li, X., Deng, Y., Bing, L., Lam, W., 2021. Towards generative aspect-based sentiment analysis. In: *Proc. 59th Annu. Meeting Assoc. Comput. Linguist. (Vol. 2: Short Papers)*, 504–510.
- [30] Huang, H., Yu, P., Tan, W., Niu, W., Shi, B., 2022. Aspect-location attention networks for aspect-category sentiment analysis in social media. *J. Intell. Inf. Syst.* 61, 395–419. <https://doi.org/10.1007/s10844-022-00760-2>
- [31] Zunic, A., Corcoran, P., Spasic, I., 2020. Sentiment analysis in health and well-being: systematic review. *JMIR Med. Inform.* 8, e16023.
- [32] Sharma, A., Lyons, J., Dehzangi, A., Paliwal, K.K., 2013. A feature extraction technique using bi-gram probabilities of position-specific scoring matrix for protein fold recognition.
- [33] Das, A., Bandyopadhyay, S., 2010. Topic-based Bengali opinion summarization. In: *COLING 2010: Posters*, 232–240.

- [34] Zhang, Y., Jin, R., Zhou, Z.-H., 2010. Understanding bag-of-words model : a statistical framework. *Int. J. Mach. Learn. Cybern.* 1, 43–52. <https://doi.org/10.1007/s13042-010-0001->
- [36] Sharma, H.D., Goyal, P., 2023. An analysis of sentiment: methods, applications, and challenges. *Eng. Proc.* 59(1), 68. <https://doi.org/10.3390/engproc2023059068>.
- [35] GeeksforGeeks, n.d. Word embedding techniques in NLP. Available at: <https://www.geeksforgeeks.org/word-embedding-techniques-in-nlp/>
- [37] Jiang, D., Luo, X., Xuan, J., Xu, Z., 2016. Sentiment computing for the news event based on the social media big data. *IEEE Access* 5, 2373–2382.
- [38] Zhang, H., Cheah, Y.N., Alyasiri, O.M., An, J., 2023. A survey on aspect-based sentiment quadruple extraction with implicit aspects and opinions.
- [39] Wankhade, M., Kulkarni, C., Rao, A.C.S., 2024. A survey on aspect base sentiment analysis methods and challenges. *Appl. Soft Comput.* 167(A), 112249. <https://doi.org/10.1016/j.asoc.2024.112249>.
- [40] Mughal, N., et al., 2024. Comparative analysis of deep neural networks and large language models for aspect-based sentiment analysis. *IEEE Access* 12, 60943–60959. <https://doi.org/10.1109/ACCESS.2024.3386969>.
- [41] Zhang, W., Li, X., Deng, Y., Bing, L., Lam, W., 2022. A survey on aspect-based sentiment analysis: tasks, methods, and challenges. *IEEE Trans. Knowl. Data Eng.* 35, 11019–11038.
- [42] Mao, Y., Liu, Q., Zhang, Y., 2024. Sentiment analysis methods, applications, and challenges: a systematic literature review. *J. King Saud Univ. Comput. Inf. Sci.* 36(4), 102048. <https://doi.org/10.1016/j.jksuci.2024.102048>.
- [43] Mughal, N., Mujtaba, G., Shaikh, S., Kumar, A., Daudpota, S.M., 2024. Comparative analysis of deep neural networks and large language models for aspect-based sentiment analysis. *IEEE Access* 12, 60943–60959. <https://doi.org/10.1109/ACCESS.2024.3386969>.
- [44] Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A.N., Kaiser, L., Polosukhin, I., 2017. Attention is all you need. *arXiv preprint arXiv:1706.03762*.
- [45] Kundi, F.M., Khan, A., Ahmad, S., Asghar, M.Z., 2014. Lexicon-based sentiment analysis in the social web. *J. Basic Appl. Sci. Res.* 4(6), 238–248.
- [46] Vaswani, A., et al., 2017. Attention is all you need. *Adv. Neural Inf. Process. Syst.* 30.
- [47] SuperAnnotate, n.d. LLM Fine-Tuning: Methods, Best Practices, and Use Cases. SuperAnnotate Blog. Available at: <https://www.superannotate.com/blog/llm-fine-tuning>
- [48] Zhou, C., Shu, R., Zhou, X., 2024. LLaMA Pro: Progressive LLaMA with Block Expansion. *arXiv preprint arXiv:2403.14608*. <https://arxiv.org/pdf/2403.14608>
- [49] Zhao, L., et al., 2024. When to do what: A taxonomy of LLM prompting methods. *arXiv preprint arXiv:2410.12837*. <https://arxiv.org/pdf/2410.12837>

[50] OpenAI, 2023. GPT-4 Technical Report. arXiv preprint arXiv:2303.08774.
<https://arxiv.org/abs/2303.08774>

[51] Chang, J.-R., Liang, H.-Y., Chen, L.-S., Chang, C.-W., 2020. Novel feature selection approaches for improving the performance of sentiment classification. *J. Ambient Intell. Hum. Comput.* 1–14.

[52] Chebolu, S. U. S., Dernoncourt, F., Lipka, N., & Solorio, T. (2022). Survey of aspect-based sentiment analysis datasets. arXiv preprint arXiv:2204.05232.