

Louvain School of Management

Exploration des Codes de Conduite d'Entreprise par l'IA

Auteurs : Belgnaou Mohammed – En-Naimi Hamza
Promoteur : Vande Kerckhove Corentin
Année académique : 2024 - 2025
Mémoire en vue d'obtenir le titre d'ingénieur de gestion en Business
Analytics
Horaire de jour

Declaration Regarding AI Tool Usage in Master's Thesis

We recognize that AI tools might be valuable aids during the master's thesis work, but they are not infallible. Remember that transparency fosters trust, and acknowledging AI's role enhances the credibility of your work.

Therefore, when deciding to use such a tool, you need to adhere to the following principles of responsible use of AI.

1. Critical Evaluation :

- We critically assessed the AI-generated output, ensuring its alignment with our research objectives.
- Any modifications or corrections were made based on our expertise and domain knowledge.

2. Transparency :

- We acknowledge the use of ChatGPT transparently, emphasizing that it contributed to our work but did not replace human judgment.
- Our commitment to transparency ensures the integrity of this thesis.

3. Ethical Considerations :

- We actively monitored for biases or unintended consequences introduced by the AI tool.
- Our ethical responsibility guided our decisions throughout the research process.

Declaration

During the preparation of this master's thesis, the author(s) utilized ChatGPT 4o for the following purpose:

1. Help with coding: Using this AI to fix errors and improve the performance of some Python code.
2. Help improving the clarity and fluency of the text.

After using ChatGPT, the author(s) diligently reviewed and edited the content produced by the tool. We take full responsibility for the final content presented in this thesis.

By signing this declaration, we affirm that the content of this master's thesis reflects our original work, augmented by the responsible use of AI.

28/05/2025

Belgnaoui  En Naimi Hamza

Résumé

Ce mémoire explore la structuration des codes de conduite des entreprises, en mobilisant des techniques avancées de traitement automatique du langage naturel (NLP). L'objectif est double : (1) identifier les thématiques majeures abordées dans ces documents, et (2) analyser les intentions exprimées, afin de mieux comprendre la manière dont les entreprises formalisent leurs engagements éthiques et communiquent ces messages à leurs collaborateurs.

Le corpus, constitué de plus de 400 codes de conduite collectés via une procédure automatisée. Il a été soigneusement prétraité : nettoyage et filtres, normalisation linguistique, segmentation en paragraphes et enrichissement par des métadonnées contextuelles (secteur, rang, zone géographique). Ces textes ont ensuite été analysés selon deux axes complémentaires.

La première analyse, thématique, s'appuie sur la méthode BERTopic, qui combine embeddings, réduction de dimensionnalité et clustering. Cette analyse thématique a permis d'identifier 16 thèmes dominants dans les codes de conduite, tels que l'éthique et la conformité, la sécurité, ou encore les droits humains. La majorité des entreprises adopte un profil généraliste, caractérisé par une couverture équilibrée de ces thèmes. Cependant, des variations apparaissent : certains secteurs, comme la santé ou l'environnement, tendent vers des profils plus polarisés, marqués par une concentration sur un nombre restreint de thématiques. Ces différences ne sont pas directement liées au classement des Fortune 500 ou au secteur, mais relèvent davantage de logiques contextuelles, tel que l'origine géographique.

La seconde analyse s'intéresse aux intentions des codes, en mobilisant un modèle zero-shot. Ce modèle attribue à chaque paragraphe une intention dominante parmi cinq catégories : protect, guide, discourge, punish et empower. Les résultats révèlent une nette domination de l'intention guide (environ 60% des cas), traduisant une posture explicative et normative, tandis que les intentions punish et empower restent marginales. Des nuances sectorielles et géographiques sont toutefois observées : les entreprises asiatiques présentent une plus forte proportion d'intentions protect, tandis que le secteur des services favorise davantage le registre empower.

L'étude conclut que les codes de conduite se caractérisent par une relative uniformité dans leur structuration thématique et leur intention, indépendamment du secteur ou du rang dans le classement Fortune. Cette homogénéité suggère l'existence de standards globaux en matière de communication éthique, mais masque potentiellement des variations plus fines qui mériteraient d'être explorées par des études qualitatives complémentaires. La démarche adoptée, résolument quantitative, permet de dégager des tendances générales et d'offrir une lecture systématique et reproductible de corpus volumineux.

En définitive, ce mémoire contribue à enrichir la compréhension des pratiques de gouvernance des grandes entreprises, en proposant une méthodologie robuste et adaptable à d'autres types de documents normatifs.

Remerciements

Nous tenons à exprimer notre profonde gratitude à toutes les personnes qui ont contribué, de près ou de loin, à la réalisation de ce mémoire.

Nous tenons également à exprimer notre reconnaissance la plus sincère à nos familles et nos amis, qui nous ont toujours soutenu et encouragé tout au long de la réalisation de ce mémoire. Leur présence et leur bienveillance ont été des piliers essentiels dans ce parcours académique.

Table des matières

1	Introduction	1
2	Revue de littérature	2
2.1	Les codes de conduite en entreprise : enjeux et fonctions	2
2.1.1	Fonctions des codes de conduites	3
2.1.2	Rôle stratégique	4
2.2	Approches classiques d'analyse de textes	4
2.3	Analyse thématique automatisée (topic modeling)	6
2.3.1	Les différentes méthodes d'analyse	6
2.3.2	Approches récentes	7
2.3.3	BERTopic	8
2.4	Dimensions discursives des codes de conduite	9
2.4.1	Distinction entre sentiment, ton et intention	9
2.4.2	Méthode d'analyse des intentions	10
2.5	Conclusion de la revue de littérature	12
3	Méthodologie	14
3.1	Rappel des objectifs de recherche et approche générale	14
3.2	Constitution et préparation du corpus	15
3.2.1	Choix du classement Fortune Global 500	15
3.2.2	Collecte automatisée des codes de conduite	15
3.2.3	Extraction, structuration et format final du corpus	16
3.2.4	Constitution des métadonnées associées	18
3.3	Prétraitement des données textuelles	19
3.3.1	Nettoyage et filtrage des données	19
3.3.2	Normalisation linguistique du corpus	20
3.3.3	Segmentation en unités d'analyse cohérentes	21
3.4	Modélisation des thématiques	21
3.4.1	Embedding des textes	21
3.4.2	Réduction dimensionnelle et clustering	22
3.4.3	Évaluation des clusters : cohérence, diversité et taux d'outliers	24
3.4.4	Post-traitement des clusters et labellisation	25
3.5	Méthode d'analyse thématique	27
3.5.1	Évaluation de la diversité thématique	27
3.5.2	Analyse comparative des profils thématiques	28
3.5.3	Exploration des facteurs explicatifs	28
3.5.4	Discussion de la méthode	29

3.6	Méthode d'analyse des intentions discursives	29
3.6.1	Choix du modèle d'analyse	30
3.6.2	Construction des variables d'analyse	30
3.6.3	Discussion de la méthode	31
4	Résultats	32
4.1	Analyse exploratoire du corpus	32
4.1.1	Répartition par secteur	32
4.1.2	Répartition par continent	33
4.1.3	Répartition par rang	33
4.2	Résultats du modèle Bertopic	34
4.2.1	Configurations testées	34
4.2.2	Comparaison des performances	35
4.2.3	Justification du choix du modèle de référence	35
4.2.4	Regroupement et réduction des clusters	36
4.3	Analyse thématiques des codes de conduites	39
4.3.1	Évaluation de la diversité thématique à l'échelle des entreprises	39
4.3.2	Regroupement des entreprises par profils thématiques	42
4.3.3	Facteurs influençant	45
4.3.4	Discussion des résultats	50
4.4	Résultats d'analyse des intentions	50
4.4.1	Répartition globale des intentions	51
4.4.2	Répartition sectorielle	52
4.4.3	Répartition par continent	55
4.4.4	Répartition par rang	56
4.4.5	Discussion des résultats	57
5	Conclusion	58
5.1	Synthèse des résultats	58
5.2	Pistes d'amélioration	59
	Bibliographie	61

Table des figures

1	Schéma des étapes modulaires du pipeline BERTopic	22
2	Intertopic Distance Map	26
3	Répartition des entreprises du corpus par secteur d'activité.	32
4	Répartition des entreprises par continent.	33
5	Distribution des rangs des entreprises.	33
6	Nuage de mots du topic <i>Workplace Safety and Ethics</i>	37
7	Nuage de mots du topic <i>Data Privacy and Confidentiality</i>	38
8	Distribution de l'entropie thématique des entreprises	40
9	Répartition thématique des codes de conduite pour sept entreprises, en pourcentage de paragraphes par thème.	40
10	Représentation des clusters	42
11	Répartition des thèmes par cluster	43
12	Distribution de l'entropie thématique des codes de conduite selon le quartile de classement.	48
13	Répartition globale des intentions dominantes dans le corpus	51
14	Répartition des intentions par thème	52
15	Répartition des intentions par secteur	54
16	Comparaison des intentions dominantes par continent.	56
17	Répartition des intentions selon le classement des fortune 500	57

Liste des tableaux

1	Comparaison des performances des trois modèles <i>BERTopic</i>	35
2	Les 16 thématiques finales extraites des codes de conduite.	37
3	Valeurs d'entropie mesurant la diversité thématique des codes de conduite par entreprise.	41
4	Répartition sectorielle des entreprises selon les clusters thématiques identifiés. .	46
5	Entropie moyenne des codes de conduite par quartile de classement Fortune Global 500.	47
6	Répartition des clusters thématiques selon le continent d'origine des entreprises	49
7	Entropie moyenne des codes de conduite par continent d'origine	49

1. Introduction

Les codes de conduite occupent une place stratégique croissante dans la gouvernance des entreprises. Ils formalisent les engagements éthiques, encadrent les comportements internes, et véhiculent les valeurs fondamentales que l'entreprise souhaite promouvoir. Loin d'être de simples documents administratifs, ces textes constituent des instruments de régulation interne et de communication externe, ayant un impact sur la culture organisationnelle et la perception du public.

Dans ce mémoire notre objectif est double, d'une part, identifier les principaux thèmes abordés dans ces documents. D'autre part, qualifier le ton ou l'intention qui y est véhiculé.

Le corpus étudié est constitué des codes de conduite issus des entreprises du classement Fortune Global 500. Ce choix se justifie par le caractère influent, international et sectoriellement diversifié de ces entreprises. Il n'est donc pas l'objet d'étude principal, mais un échantillon représentatif de pratiques organisationnelles globales.

Pour mener cette étude, nous mobilisons des techniques avancées de traitement du langage naturel (Natural Language Processing, NLP), en particulier le topic modeling (modélisation thématique) via BERTopic et des outils d'analyse du ton basés sur des modèles pré-entraînés. Cette approche permet une lecture quantitative, systématique et reproductible de documents traditionnellement analysés de façon qualitative.

Ce mémoire s'articule comme suit : une revue de littérature dresse le cadre conceptuel et technologique, une section méthodologique décrit la construction du corpus, le prétraitement et les modèles mobilisés, puis une double analyse est proposée : d'abord thématique, ensuite discursive. Une discussion met en perspective les résultats obtenus et leurs implications pour la gestion des normes éthiques en entreprise.

Il convient de noter que notre démarche s'inscrit résolument dans une approche quantitative. Elle permet de dégager des tendances générales, mais ne prétend pas rendre compte de toutes les subtilités rédactionnelles ou culturelles propres à chaque organisation.

Nous espérons que cette étude offrira des pistes de réflexion utiles pour mieux comprendre comment les grandes entreprises structurent leurs engagements éthiques à travers leurs codes de conduite.

2. Revue de littérature

Les codes de conduite occupent une place stratégique dans le fonctionnement des grandes entreprises, en formalisant leurs engagements éthiques et en orientant les comportements internes. Ces documents, bien que largement répandus, ont fait l'objet d'analyses académiques variées, portant à la fois sur leur contenu, leurs fonctions, et leurs effets sur la gouvernance et la réputation des organisations. Cette littérature, en croissance depuis les années 2000, met en lumière l'importance des codes de conduite comme instruments de gouvernance, mais aussi les limites des approches classiques d'analyse, souvent fondées sur des méthodes qualitatives manuelles.

Face à l'essor de la production textuelle dans les entreprises et à la complexité des enjeux éthiques contemporains, de nouvelles méthodes d'analyse automatisée se développent, mobilisant des techniques avancées de traitement du langage naturel (NLP). Ces approches permettent d'explorer à grande échelle le contenu thématique et discursif des documents normatifs, tout en surmontant certaines limites des approches traditionnelles.

Cette revue de littérature propose donc une vue structurée des travaux existants et des outils mobilisés, afin de situer notre recherche. Elle s'articule en cinq sections. La première présente les enjeux et fonctions des codes de conduite dans les entreprises. La seconde examine les approches classiques d'analyse des textes normatifs, en mettant en lumière leurs apports mais aussi leurs limites. La troisième s'intéresse à l'émergence des méthodes d'analyse thématique automatisée, notamment le topic modeling. La quatrième explore les méthodes d'analyse du ton dans les textes normatifs, un axe complémentaire à l'analyse thématique. Enfin, la cinquième section propose une synthèse critique des limites des approches existantes et des besoins d'innovation méthodologique pour mieux comprendre le rôle des codes de conduite dans les organisations actuelles.

2.1 Les codes de conduite en entreprise : enjeux et fonctions

Un code de conduite en entreprise est généralement défini comme un document écrit et formel énonçant les normes, règles ou principes moraux qui guident le comportement des employés et de l'organisation. Bien que le concept ne soit pas nouveau (des codes existent depuis le début du XX^e siècle), les codes de conduite jouent un rôle clé dans la formalisation des valeurs de l'entreprise et la communication de sa culture organisationnelle (Arrigo, 2006).

2.1.1 Fonctions des codes de conduites

Les fonctions des codes de conduite sont multiples et s'articulent principalement autour de trois grands axes : éthique, régulation interne et communication externe. D'un point de vue normatif, le code vise à formaliser les valeurs fondamentales de l'organisation et à prévenir les comportements contraires à l'éthique. Schwartz (2001) insiste sur la dimension préventive de ces textes, qui ont pour objectif d'« augmenter la résistance morale » de l'entreprise et de « guider le comportement des employés ». Concrètement, les codes énoncent des principes tels que l'intégrité, le respect des lois, la lutte contre la corruption ou la protection de l'environnement, servant de référence dans la prise de décision quotidienne (Stevens, 1996). Ils encadrent ainsi les comportements attendus des collaborateurs tout en traduisant l'engagement éthique de l'organisation.

Coté la régulation interne, le code de conduite incarne un engagement explicite à s'autoréguler et à instaurer des mécanismes de contrôle. Comme le notent Pitt et Grubakmanis (1990), l'adoption d'un code traduit «un engagement à pratiquer l'autorégulation d'entreprise». Cet engagement se manifeste par la mise en place de formations éthiques, de procédures de signalement des comportements inappropriés (whistleblowing) et de dispositifs de sanctions disciplinaires. Le code devient ainsi un outil de gouvernance qui structure les pratiques internes et aligne les comportements sur les standards définis par l'organisation et par la réglementation en vigueur. Il participe à la construction d'un cadre de référence partagé, qui réduit l'incertitude et facilite la prise de décision managériale.

Enfin, le code de conduite remplit une fonction essentielle en matière de communication externe et de gestion de la réputation. Il permet à l'entreprise d'afficher publiquement son engagement éthique et de rendre visible sa volonté de respecter des standards de responsabilité sociale. Comme le soulignent Lugli et al. (2009), « la fonction des codes est de communiquer la position éthique de l'entreprise aux parties externes ». En ce sens, le code contribue à renforcer la confiance des clients, des partenaires et des investisseurs, en démontrant la transparence et l'engagement de l'organisation. Bondy et al. (2004) rappellent que ces textes visent notamment à « protéger et améliorer la réputation » de l'entreprise sur le marché. Le code de conduite devient ainsi un outil stratégique qui participe à la différenciation de l'organisation dans un environnement concurrentiel, tout en renforçant son image de marque et sa crédibilité.

2.1.2 Rôle stratégique

Au-delà de ces fonctions, les codes de conduite occupent un rôle plus large dans la gouvernance des entreprises. Ils sont souvent qualifiés d'outils de gouvernance (Arrigo, 2006), dans la mesure où ils définissent les responsabilités de l'organisation envers ses parties prenantes et encadrent les décisions stratégiques des dirigeants. En ancrant la conduite des affaires dans un cadre de valeurs partagées, ils participent à la structuration de la culture organisationnelle et renforcent la cohérence des pratiques managériales. Sur le plan de la gestion de la réputation, ils fonctionnent également comme des signaux d'engagement éthique : en affichant des principes clairs (intégrité, équité, durabilité), l'entreprise cherche à rassurer ses parties prenantes et à prévenir les crises potentielles (scandales, contentieux). L'existence d'un code de conduite est d'ailleurs devenue un critère d'évaluation de la qualité de la gouvernance et de la responsabilité sociale d'une entreprise, tant pour les investisseurs que pour les agences de notation extra-financière.

Cependant, la littérature souligne également que l'existence d'un code ne garantit pas nécessairement son efficacité dans la pratique. De nombreuses études ont mis en évidence un écart entre le contenu des codes et leur application effective sur le terrain, illustrant le risque d'un « décalage éthique » entre le discours affiché et la réalité des comportements (Schwartz, 2001 ; Stevens, 1996). Cette observation conduit à s'interroger sur les moyens de mieux comprendre et d'analyser ces documents afin d'en évaluer la portée réelle.

C'est dans cette perspective qu'émergent les questions méthodologiques générales sur l'analyse des textes normatifs en entreprise. La section suivante abordera les différentes approches développées pour étudier ces documents, en particulier les méthodes classiques d'analyse qualitative, avant de démontrer l'apport des techniques automatisées modernes basées sur le traitement du langage naturel.

2.2 Approches classiques d'analyse de textes

Comprendre le rôle et la portée des codes de conduite au sein des entreprises nécessite d'en analyser le contenu, les thèmes abordés et la manière dont ils véhiculent des valeurs et des normes. Avant l'avènement des méthodes automatisées, la littérature a principalement mobilisé des approches qualitatives classiques, reposant sur l'interprétation manuelle des documents. Ces méthodes, notamment l'analyse de contenu et l'analyse thématique, ont permis d'explorer

la diversité des textes normatifs produits par les organisations, tout en présentant des limites face aux analyses à grande échelle.

L'analyse de contenu est l'une des techniques les plus répandues pour étudier les textes organisationnels. Définie par L. Bardin comme un « ensemble de techniques d'analyse des communications visant à obtenir des indicateurs quantitatifs ou non, par des procédures systématiques et objectives de description du contenu des messages », elle consiste à classifier le contenu textuel en catégories significatives. Concrètement, le chercheur segmente le texte en unités d'analyse (par exemple, des phrases ou des expressions) et les regroupe dans des catégories correspondant à des thèmes ou des concepts définis à l'avance ou émergeant des données. Cette méthode permet d'identifier des thèmes récurrents, de mesurer la fréquence d'apparition de certaines notions ou encore d'explorer les cooccurrences entre concepts. Appliquée aux textes normatifs tels que les codes de conduite, l'analyse de contenu offre une vue d'ensemble structurée des sujets abordés et des valeurs mises en avant. Par exemple, Kaptein (en 2004) a utilisé cette méthode pour cartographier les thèmes dominants dans les codes d'éthique d'entreprise, tandis que Mercier (en 2002) a développé une typologie des manières dont les organisations formalisent leurs engagements éthiques. Même si cette approche apporte une certaine objectivité dans la structuration des données, elle reste dépendante des choix du chercheur et peut s'avérer lourde à mettre en œuvre sur des corpus volumineux.

L'analyse thématique, proche de l'analyse de contenu, se distingue par son approche plus inductive. Elle vise à identifier les thèmes clés d'un corpus à partir d'une lecture approfondie des documents, sans nécessairement partir d'une grille prédéfinie. Selon Braun(en 2006), cette méthode permet de faire émerger les grandes idées et les motifs récurrents. Elle est particulièrement utile pour explorer de nouveaux terrains de recherche ou pour dégager des tendances à partir d'un nombre limité de textes. Dans le cas des documents normatifs, l'analyse thématique permet d'examiner les valeurs mises en avant dans les chartes éthiques ou les rapports RSE. Cependant, cette méthode repose également sur l'interprétation du chercheur et demande un investissement temporel important, rendant son application à de larges corpus difficilement réalisable sans assistance automatisée.

Dans l'ensemble, ces approches classiques ont contribué à la compréhension des textes normatifs en entreprise. Elles permettent d'explorer à la fois le contenu (les thèmes abordés) et les dimensions plus implicites du discours. Toutefois, elles partagent des limites communes : leur

mise en œuvre prend énormément de temps et est exigeante, leur portée est souvent restreinte à des échantillons réduits, et leur objectivité est partielle, laissant place à une certaine subjectivité analytique. Ces contraintes, combinées à l'augmentation massive des volumes de textes produits par les organisations, ont ouvert la voie au développement de méthodes d'analyse automatisée basées sur le Natural Language Processing (NLP). La section suivante s'intéresse ainsi à ces nouvelles approches et en particulier à l'émergence de techniques d'analyse thématique comme le "topic modeling".

2.3 Analyse thématique automatisée (topic modeling)

Dans un corpus de texte, un topic (ou thème latent) désigne un sujet sous-jacent qui se déclare par un ensemble de termes fréquemment associés. Autrement dit, un topic correspond à une thématique cohérente implicitement présente dans les documents. Identifier ces thèmes est permet de synthétiser le contenu d'un large corpus sans lecture exhaustive. Cependant, réaliser manuellement une telle analyse thématique sur un corpus volumineux (comme une collection de codes de conduite d'entreprise) serait fastidieux et subjectif. C'est pourquoi on recourt à des méthodes automatisées de modélisation thématique (topic modeling) permettant de dégager, de manière non supervisée, les principaux sujets abordés dans un grand ensemble de textes.

2.3.1 Les différentes méthodes d'analyse

Plusieurs méthodes classiques ont été développées pour identifier automatiquement des thèmes latents (topics) dans des corpus textuels.

La méthode Latent Dirichlet Allocation (LDA), introduite par Blei et al. (2003), est la plus connue. C'est une méthode probabiliste qui représente chaque document comme un mélange de thèmes, chaque thème étant défini comme une distribution de probabilité sur les mots. LDA permet d'extraire des listes de mots-clés par thème et d'assigner à chaque document des proportions thématiques. Cependant, la LDA présente des limites. Elle nécessite de fixer le nombre de thèmes à l'avance (un choix souvent arbitraire), repose sur une représentation simplifiée des textes (bag-of-words) et peut être coûteuse en calcul sur de grands corpus (Blei en 2003).

Une autre famille d'approches repose sur le regroupement des documents selon leur similarité, via des techniques de clustering. Ces méthodes regroupent les documents en clusters de contenus similaires, en mesurant la distance entre leurs vecteurs de représentation (Manning en 2008). Une méthode courante est le clustering k-means, qui partitionne le corpus en

un nombre k de groupes définis à l’avance, de sorte que les documents d’un même groupe se ressemblent davantage entre eux qu’avec ceux des autres groupes. Chaque cluster peut être interprété comme un « sujet » dominant. L’avantage de ces méthodes est leur flexibilité : elles peuvent être appliquées à différentes représentations vectorielles (comptes de mots, TF-IDF, ou vecteurs d’embeddings) et sont souvent moins coûteuses en calcul que l’inférence probabiliste de LDA. Cependant, le clustering classique présente des limites : il impose de fixer le nombre de clusters à l’avance, attribue chaque document à un seul groupe (ce qui est réducteur pour des documents multi-thématiques), et nécessite une interprétation manuelle des clusters, ce qui introduit une part de subjectivité.

En résumé, ces approches classiques ont permis de réaliser des analyses exploratoires sur des corpus variés, mais elles présentent des contraintes structurelles : une dépendance à des choix préalables (comme le nombre de thèmes), une incapacité à saisir les nuances sémantiques, et une faible adaptabilité à de grands volumes de textes hétérogènes. Ces limites ont motivé le développement de méthodes récentes exploitant les avancées en *Natural Language Processing* (NLP), notamment les modèles d’embeddings sémantiques, qui seront détaillés dans la suite de cette section.

2.3.2 Approches récentes

Face aux limites des méthodes traditionnelles, des approches plus récentes exploitent les avancées en représentations sémantiques du langage. L’idée est d’utiliser des embeddings (plongements vectoriels) générés par des modèles de langage pré-entraînés sur de vastes corpus. Un embedding est une représentation mathématique d’un texte sous forme de vecteur dense, qui encode le sens des mots ou des phrases dans un espace multidimensionnel. Cette approche permet de dépasser la simple cooccurrence lexicale et de capturer la similarité de sens entre textes. Par exemple, les modèles transformeurs comme BERT (Devlin et al., 2019) produisent des vecteurs qui reflètent le contenu sémantique d’un document. Deux textes abordant un même sujet, même avec des formulations différentes, auront ainsi des embeddings proches dans cet espace. En projetant l’ensemble du corpus dans un tel espace vectoriel, il devient possible de regrouper les documents sur la base de leur proximité sémantique. Des travaux récents (Reimers et Gurevych, 2019) ont montré que cette approche améliore la cohérence des thèmes extraits par rapport aux méthodes basées uniquement sur les mots. Plusieurs méthodes récentes, comme Top2Vec (Angelov, 2020) ou BERTopic (Grootendorst, 2022), adoptent ce paradigme d’« embeddings +

clustering » pour réaliser la modélisation thématique.

2.3.3 BERTopic

Dans ce mémoire, nous avons retenu la méthode BERTopic (BERT-based Topic Modeling) pour explorer les thèmes présents dans les codes de conduite d'entreprise. Cette approche, développée par Grootendorst (2022), combine plusieurs techniques afin d'exploiter la richesse des représentations sémantiques tout en produisant des thèmes interprétables de manière intuitive. Son fonctionnement repose sur quatre étapes principales.

La première consiste à représenter chaque document sous forme de vecteur dense (embedding), en utilisant un modèle de langage pré-entraîné comme Sentence-BERT. Ces vecteurs capturent le sens global des textes, en tenant compte du contexte et des nuances de langage.

La deuxième étape consiste à réduire la dimension des embeddings à l'aide d'un algorithme non linéaire tel qu'UMAP (Uniform Manifold Approximation and Projection). Cette projection permet de simplifier l'espace de représentation tout en conservant la proximité sémantique des documents.

La troisième étape est le regroupement des documents par l'algorithme HDBSCAN (Hierarchical Density-Based Spatial Clustering of Applications with Noise), qui identifie des clusters denses sans nécessiter de fixer leur nombre à l'avance. HDBSCAN permet également de marquer comme "bruit" les documents isolés, ce qui est utile pour traiter des textes très spécifiques ou atypiques.

Enfin, BERTopic extrait pour chaque cluster une liste de mots-clés caractéristiques à l'aide d'une méthode appelée c-TF-IDF (class-based TF-IDF). Cette technique mesure l'importance relative des termes dans un cluster par rapport à l'ensemble du corpus, et permet de produire des étiquettes thématiques compréhensibles. Par exemple, un cluster lié à l'éthique et à la conformité pourra être résumé par des mots comme « éthique », « intégrité » ou « déontologie ».

Cette combinaison d'embeddings, de réduction de dimension, de clustering adaptatif et d'extraction de mots-clés permet à BERTopic d'identifier des thèmes même lorsque les documents utilisent des formulations variées. C'est un atout majeur pour analyser des codes de conduite rédigés dans des styles et des langages différents selon les entreprises. De plus, la capacité d'HDBSCAN à découvrir automatiquement le nombre de clusters et à gérer les cas particuliers améliore la flexibilité et la robustesse de l'analyse. La méthode c-TF-IDF facilite l'interprétation en fournissant des mots-clés clairs pour chaque thème, ce qui est essentiel pour la lecture

humaine des résultats. Enfin, le caractère open-source et bien documenté de l'outil a facilité son implémentation dans ce travail.

Ainsi, la modélisation thématique automatisée par BERTopic permet d'explorer efficacement le contenu de centaines de codes de conduite et d'en faire ressortir les principaux thèmes de manière objective. Après avoir identifié leur contenu, il sera également intéressant d'analyser la manière dont ils le véhiculent. La section suivante portera donc sur l'étude du ton et des intentions discursives à l'œuvre dans ces codes de conduite, complément indispensable à l'analyse thématique.

2.4 Dimensions discursives des codes de conduite

2.4.1 Distinction entre sentiment, ton et intention

L'étude des textes normatifs d'entreprise, tels que les codes de conduite, a traditionnellement mobilisé des approches d'analyse qualitative visant à comprendre le sens des messages véhiculés. Dans ce champ, trois grandes perspectives peuvent être distinguées : l'analyse de sentiment, l'analyse du ton, et l'analyse des intentions discursives. Bien que ces notions soient parfois confondues dans la littérature, elles reposent sur des fondements théoriques et des objectifs distincts.

L'analyse de sentiment vise à déterminer la polarité émotionnelle d'un texte, en le catégorisant comme « positif », « négatif » ou « neutre ». Cette approche est largement utilisée dans des contextes comme l'analyse d'avis clients, de publications sur les réseaux sociaux ou de commentaires en ligne, où l'objectif est d'évaluer l'attitude émotionnelle de l'auteur. Cependant, appliquer une telle grille à des documents formels et normatifs comme les codes de conduite pose question : ces textes ne sont pas conçus pour exprimer des émotions mais pour encadrer des comportements, formuler des directives, ou établir des principes éthiques. Ainsi, une simple catégorisation en sentiments ne permet pas de rendre compte de la complexité des fonctions discursives de ces documents ni d'identifier les messages qu'ils véhiculent à destination des employés.

L'analyse du ton se concentre davantage sur le style rédactionnel et l'attitude de l'auteur, en examinant par exemple si le discours est formel ou informel, encourageant ou autoritaire. Ce niveau d'analyse peut offrir des indices intéressants sur la posture adoptée par l'organisation dans sa communication écrite. Toutefois, cette approche ne permet pas non plus de saisir pleinement

le rôle normatif et prescriptif des codes de conduite : comprendre le style du discours est utile, mais cela ne suffit pas à identifier ce que l'organisation cherche à transmettre comme message clé ou comme objectif comportemental à ses destinataires.

Par rapport à ces limites, l'analyse des intentions apparaît comme une meilleure approche pour explorer le sens des textes normatifs. L'intention se définit comme l'objectif que l'auteur cherche à atteindre en s'adressant à son lecteur : informer, guider, sanctionner, etc. La finalité des codes de conduites est moins d'exprimer une émotion mais plutôt un style qui vise à orienter les comportements des employés, en posant des règles à suivre, en avertissant des risques ou en encourageant certaines pratiques. Comme le souligne la littérature sur l'analyse des discours organisationnels (Abercrombie & Batista-Navarro, 2020 ; Ni et al., 2022). Il en devient intéressant pour distinguer les intentions implicites (par exemple, avertir d'un risque) des simples tonalités émotionnelles, afin de mieux comprendre comment le texte opère comme un outil d'encadrement.

Plusieurs travaux ont tenté de formaliser ces intentions dans des cadres d'analyse. Par exemple, Abercrombie & Batista-Navarro (2020) insistent sur l'importance d'identifier des fonctions discursives comme l'obligation, la prohibition ou la recommandation, en soulignant que ces dimensions structurent la communication réglementaire. Ni et al. (2022) proposent une typologie complémentaire, distinguant des catégories comme « promouvoir », « limiter », ou « sanctionner ». Ces approches ont en commun de considérer que les textes normatifs, tels que les codes de conduite, s'adressent à leurs destinataires non pas pour susciter une émotion, mais pour prescrire un comportement ou une attitude.

Dans cette perspective, l'analyse des intentions permet d'attribuer à chaque segment de texte une fonction précise, en identifiant ce que l'organisation cherche à obtenir de son lecteur. Par exemple, un paragraphe peut avoir pour but de dissuader un comportement jugé inapproprié, d'encourager la prise d'initiative, ou encore d'informer sur un droit ou une obligation. Cette granularité d'analyse permet de mieux refléter la diversité des messages véhiculés dans les codes de conduite et de dépasser une simple évaluation de tonalité ou de sentiment.

2.4.2 Méthode d'analyse des intentions

Plusieurs approches sont envisageables pour l'analyse d'intention, chacune présentant des atouts et des limites spécifiques.

Une première solution consiste à recourir à des modèles supervisés, entraînés sur un corpus annoté où chaque segment de texte est labellisé en fonction des intentions identifiées (par exemple : « protéger », « guider », « punir », etc.). Cette approche a été largement utilisée dans d'autres domaines, notamment l'analyse de textes d'opinion (Pang & Lee, 2008) ou la détection d'intentions dans les forums en ligne (Braun et al., 2017). Toutefois, l'application de cette méthode aux codes de conduite est limitée par l'absence de jeux de données préexistants spécifiquement annotés pour les intentions dans des textes normatifs. La construction d'un tel corpus représenterait un travail considérable, nécessitant l'intervention d'experts et une annotation manuelle systématique, comme le soulignent Abercrombie & Batista-Navarro (2020) dans leur étude sur les textes réglementaires.

Une seconde approche repose sur l'apprentissage par transfert ou fine-tuning. Cette méthode consiste à ajuster un modèle pré-entraîné (par exemple sur la classification d'émotions ou de sentiments) à une nouvelle tâche en utilisant un petit échantillon annoté issu du domaine cible. Bien que cette stratégie ait montré des résultats prometteurs (Howard & Ruder, 2018 ; Devlin et al., 2019), elle reste peu adaptée aux documents formels comme les codes de conduite, dont le style rédactionnel et les objectifs discursifs diffèrent fortement des données courantes utilisées pour l'entraînement des modèles. Enfin, comme le rappelle Ni et al. (2022), les approches par fine-tuning nécessitent des ressources computationnelles importantes.

Face à ces contraintes, les méthodes dites zero-shot sont une bonne alternative. Ces modèles, tels que DeBERTa v3 (He et al., 2021), T5 (Raffel et al., 2020) ou BART-large-MNLI (Lewis et al., 2020), sont pré-entraînés sur des tâches d'inférence logique (Natural Language Inference, NLI), où ils apprennent à évaluer si un texte source implique, contredit ou est neutre par rapport à une hypothèse donnée. Cette capacité d'inférence permet de poser des hypothèses sur les intentions d'un paragraphe, par exemple : « Ce texte cherche-t-il à dissuader un comportement ? ». Cela permet d'obtenir une probabilité d'adéquation sans nécessiter de données annotées dans le domaine d'application. La littérature récente montre que ces modèles zero-shot offrent une robustesse accrue dans des contextes où les annotations manquent, comme le souligne Yin et al. (2019) dans leur évaluation de la transférabilité des modèles NLI, ou Schick & Schütze (2021) dans leurs travaux sur la classification sans supervision directe.

Par ailleurs, cette approche présente un intérêt particulier pour les textes normatifs. Comme le notent Abercrombie & Batista-Navarro (2020) et Ni et al. (2022), les intentions dans ces

textes ne se traduisent pas nécessairement par des émotions explicites, mais par des objectifs discursifs spécifiques : orienter le comportement, signaler des sanctions, ou protéger des droits. La flexibilité des modèles zero-shot permet de formuler des hypothèses directement alignées avec ces fonctions (par exemple : « ce paragraphe cherche-t-il à orienter le lecteur vers un comportement attendu ? »), et d’obtenir un score de correspondance même lorsque le vocabulaire employé varie fortement.

Ainsi, dans le cadre de ce mémoire, le choix d’une méthode zero-shot est une solution pertinente. Elle permet d’analyser un corpus volumineux de codes de conduite sans nécessiter de données annotées spécifiques, tout en offrant une capacité d’adaptation à la diversité des intentions présentes dans les textes. Cette solution permet d’appréhender les fonctions discursives des codes de conduite de manière systématique et reproductible, en complément des analyses thématiques menées.

2.5 Conclusion de la revue de littérature

L’analyse de la littérature montre que les codes de conduite jouent un rôle central dans la gouvernance des entreprises, en formalisant des engagements éthiques et en encadrant les comportements attendus. Les travaux existants ont permis de clarifier les fonctions stratégiques de ces documents – outil de régulation interne, signal d’engagement éthique et vecteur de communication externe – tout en mettant en évidence les limites des approches qualitatives classiques, notamment leur subjectivité et leur faible capacité de traitement à grande échelle.

Face à ces contraintes, les méthodes d’analyse automatisée basées sur le traitement du langage naturel (NLP) se sont progressivement imposées comme des solutions prometteuses pour explorer de grands corpus textuels de manière systématique et reproductible. En particulier, les techniques de *topic modeling* (modélisation thématique) permettent d’identifier les principaux sujets abordés dans les codes de conduite, tandis que l’analyse des intentions discursives offre un angle complémentaire pour comprendre les objectifs sous-jacents véhiculés par ces textes normatifs. La combinaison de ces approches permet de dépasser les limites des méthodes classiques en offrant une analyse plus riche et nuancée des contenus.

Néanmoins, il est important de souligner que ces méthodes automatisées ne sont pas exemptes de limites. Elles peuvent rencontrer des difficultés à capturer les subtilités du langage juridique ou organisationnel, à interpréter correctement les références culturelles ou implicites, et restent

sensibles aux biais des modèles de langage pré-entraînés. Par ailleurs, leur utilisation exige un choix méthodologique réfléchi, notamment concernant la définition des catégories d'analyse (thèmes, intentions) et l'interprétation des résultats produits.

En définitive, la littérature souligne la complémentarité entre approches traditionnelles et méthodes automatisées : si l'automatisation permet un gain d'échelle et une systématique d'analyse, le jugement expert reste indispensable pour interpréter les résultats et contextualiser les enseignements. Ce mémoire adopte ainsi une démarche quantitative centrée sur l'exploration thématique et l'analyse des intentions discursives des codes de conduite, tout en reconnaissant la nécessité d'une lecture critique des résultats obtenus. La section suivante détaillera la méthodologie mise en place pour répondre à ces objectifs.

3. Méthodologie

3.1 Rappel des objectifs de recherche et approche générale

Ce mémoire a pour objectif d’explorer la manière dont les codes de conduite des entreprises sont structurés et rédigés. Comme dit dans l’introduction nous nous focaliserons sur les Fortune Global 500, qui un échantillon représentatif de pratiques organisationnelles globales. Ces documents, qui formalisent les engagements éthiques des entreprises, constituent des outils stratégiques de communication et de gouvernance. L’analyse porte à la fois sur le contenu thématique des codes et sur la forme discursive employée pour transmettre ces messages. Deux questions de recherche principales orientent cette étude :

Q1 : Quels sont les thèmes majeurs abordés dans les codes de conduite des entreprises du Fortune Global 500, et comment leur présence ou leur importance varie-t-elle en fonction de facteurs tels que le secteur d’activité, le rang dans le classement ou la zone géographique ?

Q2 : Comment ces thèmes sont-ils exprimés en termes d’intention et de finalité discursive, c’est-à-dire dans quelle mesure ces codes visent-ils à s’adresser à leurs employés ?

Ces deux questions sont étroitement liées. La première (Q1) s’attache à identifier et comparer les sujets abordés dans les codes de conduite : de quoi parlent-ils ? La seconde (Q2) complète cette analyse en s’intéressant à la manière dont ces sujets sont présentés : comment ces messages sont-ils formulés et avec quelles intentions ?

Pour y répondre, comme introduit précédemment une approche quantitative a été privilégiée. D’une part, l’analyse des thèmes (Q1) repose sur des techniques de topic modeling non supervisé, mises en œuvre à l’aide de l’algorithme BERTopic. Celui-ci permet d’identifier les sujets dominants dans un corpus de textes, en combinant des représentations vectorielles du langage (embeddings) et des méthodes de clustering basées sur la densité. Cette approche permet de regrouper les segments textuels similaires en thèmes cohérents.

D’autre part, l’analyse d’intention et de la finalité discursive (Q2) est également conduite de manière automatisée à l’aide d’outils d’analyse de sentiment et d’intention discursive. Ces outils permettent d’assigner à chaque segment textuel une orientation dominante. L’objectif est de dégager des tendances générales sur la manière dont les entreprises du Fortune Global 500 communiquent leurs attentes et prescriptions éthiques à leurs employés.

3.2 Constitution et préparation du corpus

3.2.1 Choix du classement Fortune Global 500

Le présent travail s'appuie sur un corpus composé des codes de conduite issus des entreprises figurant au classement Fortune Global 500. Ce classement, mis à jour annuellement par le magazine *Fortune*, recense les 500 plus grandes entreprises mondiales en fonction de leur chiffre d'affaires. Il a été retenu pour plusieurs raisons.

Tout d'abord, le Fortune 500 offre un échantillon d'entreprises particulièrement diversifié sur le plan sectoriel, couvrant des domaines tels que la finance, la technologie, l'énergie ou l'industrie, mais également sur le plan géographique, avec des entreprises implantées dans toutes les régions du monde. Cette diversité permet d'obtenir un corpus varié et représentatif des pratiques éthiques des grandes entreprises internationales. Ensuite, la taille du classement garantit un volume de données suffisant pour mener une analyse quantitative robuste tout en restant gérable. Enfin, la disponibilité de ces documents constitue un facteur pratique important : la majorité des entreprises du classement Fortune Global 500 publient leur code de conduite sur leur site institutionnel, ce qui facilite la constitution d'un corpus représentatif.

Dans l'ensemble, le classement Fortune Global 500 constitue donc un cadre d'étude cohérent et pertinent pour analyser la manière dont les grandes entreprises communiquent et structurent leurs engagements éthiques à travers leurs codes de conduite.

3.2.2 Collecte automatisée des codes de conduite

Afin de centraliser, systématiser et fiabiliser la collecte des codes de conduite des entreprises, une procédure automatisée a été développée à l'aide d'un script Python.

Le script repose sur la combinaison de deux outils principaux :

- **SerpAPI**, un service d'interrogation automatisée de Google Search, permet de formuler et d'exécuter des requêtes ciblées telles que : `[nom de l'entreprise] code of conduct filetype:pdf`. Cette API facilite l'accès aux résultats de recherche au format JSON, tout en gérant de manière transparente les limitations typiques du scraping manuel (captchas, proxies, quotas). Elle garantit ainsi l'extraction fiable de liens vers des fichiers PDF publiquement accessibles.
- La bibliothèque **requests** de Python est utilisée pour télécharger les fichiers PDF iden-

tifiés. Elle offre une gestion fine des requêtes HTTP, incluant le contrôle des délais, la gestion des erreurs et des codes de réponse, assurant la robustesse et la continuité du processus de collecte.

Une fois la recherche lancée via SerpAPI, si un lien direct vers un PDF est détecté parmi les premiers résultats, le fichier est automatiquement téléchargé et stocké localement dans un dossier structuré. Chaque fichier est nommé selon une convention standardisée incluant le rang Fortune, le nom de l'entreprise et son pays, facilitant ainsi l'indexation et l'analyse future.

Toutes les opérations sont consignées dans un fichier journal (log.csv), qui enregistre le rang, le nom de l'entreprise, le pays, le lien PDF trouvé (le cas échéant), ainsi que le statut du téléchargement. Ce mécanisme assure une traçabilité complète du processus, permet des reprises manuelles ciblées en cas d'échec.

Cette approche permet d'automatiser efficacement la constitution d'un corpus homogène de documents d'entreprise, tout en réduisant considérablement le besoin d'intervention manuelle. Les entreprises pour lesquelles aucun document n'a pu être identifié ont été temporairement écartées de l'analyse.

3.2.3 Extraction, structuration et format final du corpus

Une fois les fichiers PDF collectés, l'une des principales problématiques méthodologiques rencontrées concerne la nature même des données sources. En effet, les codes de conduites sont généralement pensés pour un affichage visuel, et non pour un traitement automatisé. Chaque code de conduite doit ainsi être transformé en une série d'unités d'analyse cohérentes, permettant de mener des analyses thématiques et discursives sur des segments textuels significatifs. Ce choix méthodologique permet de mieux capturer la diversité des sujets abordés au sein d'un même document, tout en préservant le lien logique avec la structure d'origine.

Cependant, atteindre ce format final ne va pas de soi. L'une des principales problématiques dans ce projet réside dans la nature même des données sources. En effet, les codes de conduite sont conçus pour un affichage visuel, et non pour un traitement automatisé. Le texte qu'ils contiennent est rarement structuré de manière directement exploitable : il peut être réparti sur plusieurs colonnes, entrecoupé de tableaux, d'illustrations ou d'éléments décoratifs, et la hiérarchie des titres et sections n'est pas explicitement encodée.

Dès lors, l'enjeu ne se limite pas à l'extraction du contenu textuel, mais consiste également à restituer une structure logique et hiérarchisée à partir du document d'origine. Il est indispensable d'identifier et de distinguer les titres, sous-titres, paragraphes et sections afin de produire un flux de texte cohérent, prêt à être segmenté en unités d'analyse pour l'étude des thèmes et du ton discursif. Il est donc nécessaire d'identifier et de sélectionner l'outil le plus approprié pour répondre à ce double défi d'extraction et de structuration. Plusieurs approches techniques sont possibles.

Une première solution consiste à utiliser des outils OCR classiques, comme Tesseract, en combinaison avec des bibliothèques telles que PyMuPDF pour extraire le contenu textuel des fichiers PDF. Cette approche permet d'obtenir un flux brut de texte, mais sans structuration : il n'y a pas de distinction claire entre titres, sous-titres et paragraphes, ce qui limite fortement l'exploitabilité des résultats.

Des alternatives plus avancées, comme DocTR (basé sur des modèles de deep learning) ou des services cloud tels que Google Document AI et AWS Textract, offrent de meilleures performances pour la détection des blocs de texte et leur position sur la page. Toutefois, elles présentent des limitations importantes : coût d'utilisation élevé, dépendance à une infrastructure cloud, et difficulté à obtenir une hiérarchisation précise des éléments extraits.

Face à ces contraintes, l'outil *pdf-document-layout-analysis* (Huridocs, 2024), disponible en open source sur GitHub, constitue une solution plus adaptée. Cet outil exploite directement la structure vectorielle des fichiers PDF pour identifier les blocs de texte, analyser leur position relative sur la page, et les regrouper en unités textuelles cohérentes (titres, paragraphes, légendes, etc.). Il permet ainsi d'obtenir un texte structuré qui reflète la logique et la hiérarchie du document d'origine, ce qui est essentiel pour garantir la qualité et la cohérence des analyses ultérieures. L'outil repose sur deux types de modèles complémentaires :

- Le modèle par défaut est un modèle visuel de type Transformer, plus précisément le Vision Grid Transformer (VGT) développé par le groupe de recherche d'Alibaba. Entraîné sur le jeu de données DocLayNet, ce modèle permet d'analyser une page entière, de détecter visuellement les blocs de texte, d'en reconnaître les types (titre, paragraphe, tableau, etc.) et de reconstituer leur agencement spatial. Ce modèle est le plus performant du projet, bien qu'il soit plus exigeant en ressources GPU.

- Une version allégée, reposant sur des modèles LightGBM, est également proposée. Elle permet une extraction plus rapide, compatible avec un traitement sur CPU uniquement, au prix d'une précision légèrement inférieure sur la structuration fine des éléments.

Enfin, pour les documents qui ne contiennent pas de texte vectoriel (scans ou PDF image), le système intègre Tesseract OCR via OCRmyPDF. Cela permet de les rendre compatibles avec les modèles de segmentation même en l'absence de texte « encodé »

À l'issue de cette étape, chaque code de conduite est transformé en un fichier texte brut, dont la structure de base reflète au mieux l'organisation visuelle du document d'origine. Les textes sont extraits sous forme de blocs distincts correspondant à des éléments tels que des titres, des paragraphes ou des sections identifiés lors de l'analyse de la mise en page. Ces blocs constituent une base structurée qui sera affinée et segmentée en unités d'analyse cohérentes lors des étapes de prétraitement décrites ci-après

3.2.4 Constitution des métadonnées associées

En complément des codes de conduite collectés, un ensemble de métadonnées est constitué pour chaque entreprise figurant dans le classement Fortune Global 500. Ces informations visent à contextualiser les entreprises incluses dans l'étude et à permettre des analyses croisées ultérieures. Les variables le constituant sont les suivantes :

- Nom : nom de l'entreprise telle qu'indiquée dans le classement Fortune Global 500.
- Pays : pays d'origine du siège social de l'entreprise ?
- Continent : zone géographique du siège social (Amérique du Nord, Asie, Europe, etc.) ?
- Secteur d'activité : classification économique selon Fortune, ces secteurs ont été uniformisé pour faciliter l'analyse comparative, en simplifiant certaines catégories très spécifiques.
- Rank : position de l'entreprise dans le classement Fortune Global 500.

Ces informations sont collectées manuellement. Les données proviennent des publications officielles du classement Fortune Global 500. Ces métadonnées permettront de filtrer, regrouper ou comparer les résultats du modèle BERTopic selon différents axes.

3.3 Prétraitement des données textuelles

À l'issue de l'extraction, les textes des codes de conduite sont disponibles sous forme brute mais encore imparfaite. Ils contiennent des éléments non pertinents, des anomalies typographiques et des segments hétérogènes susceptibles de perturber l'analyse.

Le prétraitement des données textuelles a pour objectif d'affiner cette base de travail afin de produire un corpus nettoyé, normalisé et segmenté en unités d'analyse cohérentes, à savoir des paragraphes. Ces étapes sont essentielles pour garantir la qualité des représentations textuelles et permettre des analyses thématiques et discursives fiables (ceci sera vu ultérieurement).

Le processus de prétraitement a été structuré en trois phases principales :

3.3.1 Nettoyage et filtrage des données

La première phase consiste à éliminer les éléments parasites issus de la mise en page des documents PDF. Bien que l'outil d'extraction ait permis de récupérer les textes tout en respectant la structure visuelle des fichiers, de nombreux segments non pertinents subsistaient dans le corpus. Ces éléments comprennent notamment :

- Les en-têtes et pieds de page répétitifs.
- Les numéros de page (*Page X of Y*).
- Les mentions légales standardisées (*All rights reserved*, *Confidential*).
- Les tables de matières.
- Les slogans institutionnels.
- Les messages introductifs rédigés par des représentants de la direction (par exemple, des lettres du PDG), dont la formulation est souvent très standardisée et peu informative dans le cadre d'une analyse thématique.

Ces suppressions sont réalisées à l'aide de scripts Python exploitant des expressions régulières et des règles de filtrage fondées sur l'observation de motifs récurrents dans le corpus.

En parallèle, un filtrage linguistique est mis en œuvre afin d'assurer l'homogénéité linguistique du corpus. Bien que la collecte ait ciblé des documents en anglais, certains codes de conduite contiennent des passages rédigés dans d'autres langues, principalement en chinois. Ces segments ont été détectés automatiquement à partir des caractères issus des blocs Unicode spécifiques aux écritures CJK (*Chinese, Japanese, Korean*). Les lignes identifiées ont été supprimées

afin de garantir la compatibilité du corpus avec les modèles NLP utilisés par la suite, qui sont exclusivement entraînés sur des corpus anglophones.

3.3.2 Normalisation linguistique du corpus

Une fois le nettoyage réalisé, une normalisation linguistique est appliquée afin d'uniformiser la forme des textes et d'optimiser leur traitement par les outils d'analyse.

- Tous les textes sont convertis en minuscules afin d'éviter que des mots identiques soient traités comme distincts en raison d'une simple différence de casse.
- Des expressions régulières permettent de supprimer les numéros de sections, les puces, et d'autres symboles typographiques (crochets, tirets décoratifs) n'apportant pas d'information sémantique pertinente.
- Les séquences de lettres artificiellement espacées (par exemple : *C O D E*) sont corrigées pour restaurer leur forme lexicale d'origine.
- Une attention particulière est également portée aux concaténations de mots sans espaces (par exemple : *complianceethicsaudit*), fréquentes dans les textes extraits. Pour traiter ce problème, l'algorithme SymSpell est employé, avec un dictionnaire enrichi de termes spécifiques aux codes de conduite (par exemple : *compliance, stakeholders, misconduct, audit, ethics, integrity*). Cette méthode donne une segmentation fine et fiable de ces séquences tout en préservant les termes correctement formatés.

En complément, une réduction lexicale est effectuée pour regrouper les différentes formes fléchies d'un même mot sous une forme unique.

La suppression des mots vides (stopwords) retire les termes très fréquents en langue naturelle mais peu informatifs sur le plan sémantique (par exemple : *the, and, of*). Cette étape se réalise à l'aide des ressources lexicales proposées par la bibliothèque spaCy, et notamment le modèle préentraîné `en_core_web_sm` adapté à l'anglais.

Enfin, une lemmatisation est là pour réduire chaque mot à sa forme canonique (lemme), en tenant compte de son rôle syntaxique. Par exemple, les formes *running, ran* et *runs* sont ramenées à *run*. Cette opération favorise la détection de similarités sémantiques entre textes et améliore la qualité des regroupements générés par les modèles d'analyse.

3.3.3 Segmentation en unités d’analyse cohérentes

La dernière étape consiste à segmenter le texte en unités d’analyse cohérentes, à savoir des paragraphes. Ce choix repose sur des considérations méthodologiques et pratiques : contrairement à une segmentation par phrase (trop fine) ou par document entier (trop globale), le paragraphe constitue une unité de sens autonome, généralement cohérente sur le plan thématique. Il permet de capturer des idées distinctes tout en offrant une granularité suffisante pour une analyse fine des thèmes et des intentions.

Toutefois, l’extraction automatique des textes génère parfois des fragments très courts (souvent inférieurs à une dizaine de mots), résultant de ruptures de mise en page ou de titres isolés. Pour éviter une fragmentation excessive, une procédure de regroupement est mise en place : les paragraphes trop courts ont été fusionnés avec le segment suivant issu du même document, selon une logique de continuité sémantique.

L’ensemble de ces étapes permet d’obtenir un corpus nettoyé, homogène et structuré, dans lequel chaque paragraphe constitue une unité textuelle riche en informations et prête à être exploitée par les analyses automatisées détaillées dans les sections suivantes.

3.4 Modélisation des thématiques

L’identification des thématiques abordées dans les codes de conduite repose sur la méthode BERTopic (Grootendorst, 2024), dont les principes ont été présentés dans la revue de littérature. Pour rappel, BERTopic combine plusieurs techniques avancées de traitement du langage naturel et de machine learning tel que l’embeddings, réduction de dimensionnalité, clustering adaptatif et extraction de mots-clés. Afin de détecter des regroupements thématiques de manière automatisée et interprétable.

3.4.1 Embedding des textes

Comme expliqué précédemment, la première étape du pipeline BERTopic consiste à représenter chaque texte sous forme de vecteur. Pour réaliser cette opération, le modèle Sentence-BERT (Reimers & Gurevych, 2019) est retenu.

Ce choix s’explique par sa capacité à générer des représentations sémantiques robustes, adaptées à des textes courts ou moyens, comme les paragraphes étudiés dans ce travail.

Sentence-BERT produit des embeddings contextuels qui tiennent compte de l’ensemble du

texte, ce qui permet de capturer avec précision les relations de similarité entre les segments textuels. Ces embeddings constituent la base essentielle du processus d'analyse thématique : ils fournissent une représentation vectorielle du corpus qui sera exploitée dans les étapes suivantes pour identifier des regroupements thématiques. Nous nous situons donc à la première étape de la méthodologie : les embeddings sont générés, et le pipeline BERTopic va désormais dérouler les autres étapes, comme illustré dans la Figure 2 ci-après.

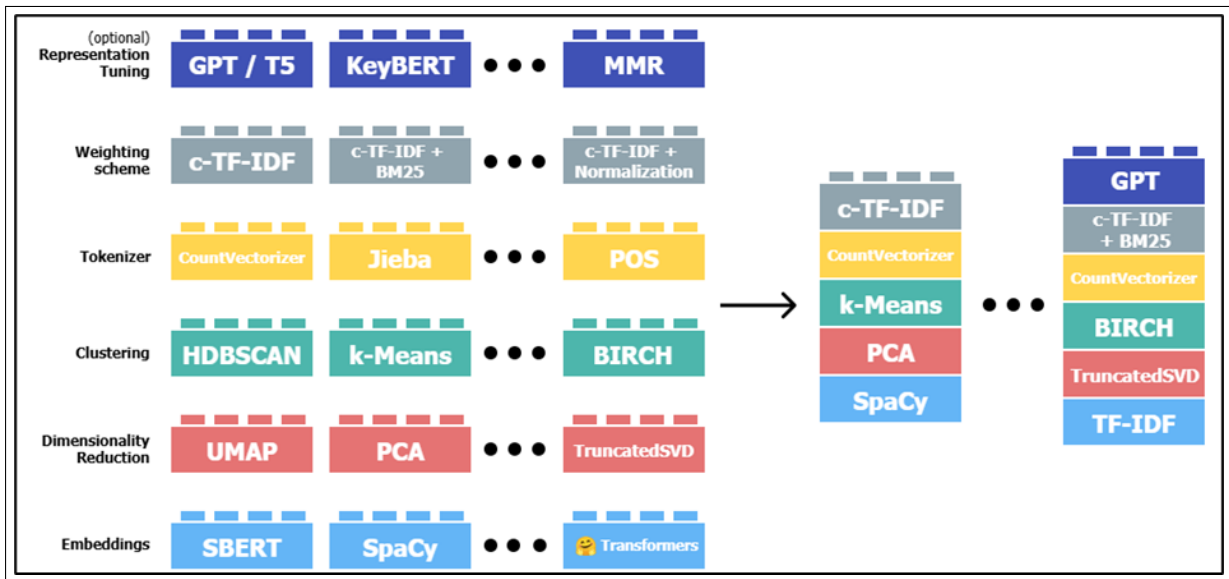


FIGURE 1 – Schéma des étapes modulaires du pipeline BERTopic

3.4.2 Réduction dimensionnelle et clustering

Après la génération des embeddings, la méthodologie BERTopic enchaîne avec deux étapes clés : la réduction de dimensionnalité et le clustering. Ces opérations permettent de simplifier la représentation des textes et de regrouper automatiquement les paragraphes selon leur proximité sémantique.

Les embeddings produits par Sentence-BERT sont des vecteurs de grande dimension (768 dimensions dans la configuration utilisée). Si ces représentations sont très riches, elles sont aussi complexes à manipuler directement. Pour faciliter leur traitement des données, une réduction de dimensionnalité est appliquée à l'aide de l'algorithme UMAP.

UMAP (Uniform Manifold Approximation and Projection) peut être compris intuitivement comme une technique qui "compresse" l'espace des données tout en préservant les relations de proximité entre les points. Il permet de projeter les embeddings dans un espace plus simple (généralement à 2 ou 5 dimensions) où les textes traitant de thématiques proches restent regrou-

pés, tandis que ceux portant sur des sujets différents s'éloignent. Cette compression facilite la détection des motifs et améliore la lisibilité des regroupements à l'étape suivante.

Le comportement de UMAP est guidé par plusieurs paramètres essentiels, dont :

- `n_neighbors` : détermine le nombre de voisins pris en compte pour chaque point. Une valeur faible favorise la détection de petites structures locales et permet d'identifier des thématiques fines. Une valeur plus élevée privilégie la structure globale, en cherchant à préserver l'organisation d'ensemble des clusters, quitte à lisser certaines différences locales.
- `min_dist` : contrôle la distance minimale autorisée entre deux points dans l'espace projeté. Une valeur basse (par exemple 0.0) tend à créer des clusters denses et bien séparés, tandis qu'une valeur plus élevée (par exemple 0.3) favorise une dispersion plus souple des points, avec un risque accru de chevauchement entre groupes.
- `n_components` : fixe le nombre de dimensions finales de projection (par exemple 2 ou 5).

Une fois les données réduites, l'algorithme HDBSCAN (Hierarchical Density-Based Spatial Clustering of Applications with Noise) est utilisé pour réaliser le clustering des textes. HDBSCAN identifie automatiquement les groupes de paragraphes qui partagent des similarités sémantiques, sans nécessiter de spécifier à l'avance le nombre de clusters attendus.

Deux paramètres principaux guident le fonctionnement de HDBSCAN :

- `min_cluster_size` : définit la taille minimale d'un groupe pour être considéré comme un cluster. Une valeur faible permet de détecter des thèmes très spécifiques mais augmente le risque de sur-segmentation. À l'inverse, une valeur plus élevée favorise des regroupements plus larges et stables, au risque de lisser des variations locales.
- `min_samples` : régule la sensibilité de l'algorithme au bruit. Lorsqu'il est égal à `min_cluster_size`, HDBSCAN devient plus strict et rejette davantage de points considérés comme ambigus. À l'inverse, une valeur plus basse rend l'algorithme plus permissif et accepte davantage de points dans les clusters, y compris des textes plus atypiques.

L'un des points forts de HDBSCAN est précisément sa capacité à détecter des clusters de formes et de densités variées, ce qui est particulièrement utile pour des données textuelles hétérogènes comme celles des codes de conduite. De plus, HDBSCAN peut identifier des points considérés

comme des outliers, c'est-à-dire des paragraphes qui ne s'intègrent dans aucun cluster. Cette capacité permet une analyse plus fine du corpus et contribue à la robustesse des résultats.

En combinant UMAP pour simplifier la complexité des données et HDBSCAN pour regrouper automatiquement les textes, la méthode BERTopic permet de détecter des thématiques de manière dynamique et adaptative, en s'ajustant à la structure réelle du corpus.

3.4.3 Évaluation des clusters : cohérence, diversité et taux d'outliers

Une fois les clusters formés par l'algorithme HDBSCAN, il est nécessaire d'évaluer leur qualité et leur pertinence avant d'interpréter les thématiques. Pour cela, plusieurs combinaisons de paramètres (notamment pour UMAP et HDBSCAN) sont testées afin d'optimiser les performances du modèle.

Trois critères principaux permettent de guider cette évaluation et comparer les configurations : le taux d'outliers, la cohérence des clusters, et leur diversité. Ces métriques identifient la configuration la plus adaptée, garantissant des clusters à la fois clairs, distincts, et représentatifs des thématiques du corpus.

- Taux d'outliers : Cette métrique correspond à la proportion de paragraphes pour lesquels l'algorithme HDBSCAN n'a attribué aucun cluster (valeur -1). Ces paragraphes sont considérés comme du bruit ou trop ambigus pour être rattachés à une thématique cohérente. Le score est calculé par la formule suivante :

$$\text{Outlier ratio} = \frac{\text{Nombre de paragraphes non classés}}{\text{Nombre totale de paragraphes}}$$

Un taux d'outliers modéré est généralement recherché : s'il est trop élevé, cela peut indiquer un paramétrage trop strict ou une hétérogénéité excessive des textes ; s'il est trop faible, cela peut révéler un regroupement artificiel de textes peu liés.

- **Cohérence** : Le score de cohérence évalue la proximité sémantique entre les mots-clés d'un même topic. Plus les mots d'un thème apparaissent ensemble dans les mêmes documents, plus le score est élevé. Il a été calculé à l'aide de la classe `CoherenceModel` de la bibliothèque Gensim, avec la méthode C_v , qui combine cooccurrence et similarité contextuelle sur des documents tokenisés. La cohérence permet de vérifier la qualité thématique des clusters : un score élevé indique des thèmes bien délimités et compréhensibles.
- **Diversité** : La diversité mesure le degré de différenciation sémantique entre les topics générés. Elle est basée sur les embeddings de chaque topic, en calculant la moyenne des similarités cosinus entre tous les couples de topics. La formule utilisée est :

$$\text{Diversité} = 1 - \text{moyenne}(\text{similitudes cosinus entre topics})$$

Une diversité élevée indique des groupes bien distincts les uns des autres et une diversité faible suggère une redondance thématique.

Cette approche garantit que les clusters retenus sont à la fois cohérents et distincts.

3.4.4 Post-traitement des clusters et labellisation

Après la formation des clusters par l'algorithme HDBSCAN, une phase de post-traitement est mise en œuvre afin d'affiner, clarifier et enrichir les résultats produits par BERTopic. Cette étape vise à garantir la qualité et l'interprétabilité des thématiques identifiées, en préparant le corpus à une analyse plus approfondie.

Le post-traitement repose sur deux étapes principales.

Dans un premier temps, une analyse visuelle et qualitative des clusters est effectuée à l'aide d'une Intertopic Distance Map (Figure 3). Cette carte positionne chaque cluster sous forme de point, en fonction de sa proximité sémantique avec les autres, tandis que la taille des points reflète le nombre de paragraphes qu'il contient.

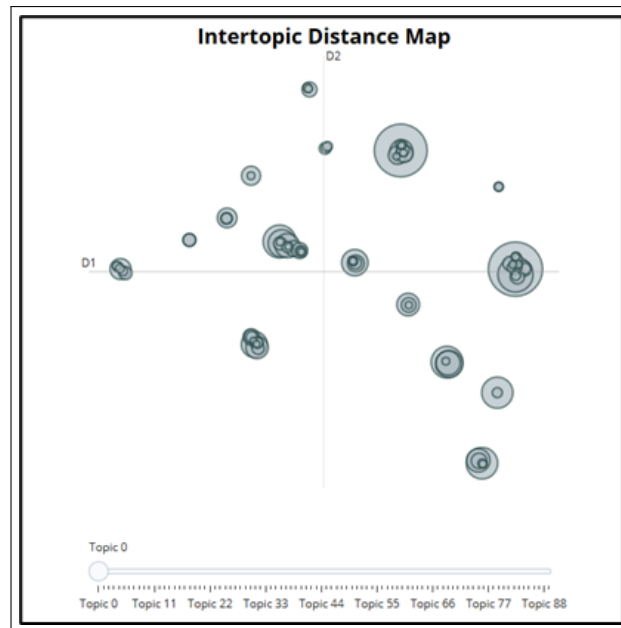


FIGURE 2 – Intertopic Distance Map

L’observation de cette carte permet d’identifier des clusters redondants ou très proches, suggérant des chevauchements thématiques. Lorsque plusieurs groupes abordent un même sujet ou des thématiques similaires sous des angles proches, ils sont fusionnés manuellement afin d’obtenir une segmentation plus claire et pertinente.

Cette étape de consolidation permet de réduire le nombre initial de clusters à un ensemble plus restreint, facilitant ainsi l’interprétation et la lisibilité des résultats. Une fois la structure des clusters consolidée, une phase d’interprétation thématique est réalisée.

Elle combine deux approches complémentaires :

- Une labellisation automatique des clusters est effectuée à l’aide de l’API GPT-3.5 d’OpenAI. Ce modèle génère des intitulés explicites et interprétables pour chaque cluster, en s’appuyant sur ses capacités avancées de traitement du langage naturel.
- En parallèle, la méthode c-TF-IDF (class-based Term Frequency–Inverse Document Frequency) est utilisée pour extraire des mots-clés caractéristiques de chaque cluster. Cette technique identifie les termes qui apparaissent fréquemment dans un cluster donné mais restent rares dans l’ensemble du corpus, permettant ainsi de mettre en évidence les spécificités sémantiques de chaque thématique.

Une vérification croisée est effectuée entre les intitulés générés par GPT-3.5 et les mots-clés extraits par c-TF-IDF afin de s’assurer de leur cohérence et de leur pertinence. Lorsque des

intitulés apparaissent ambigus ou imprécis, une révision manuelle est réalisée sur la base d’une réinspection qualitative des paragraphes concernés.

Ce processus de post-traitement permet de produire des résultats fiables, clairs et directement exploitables pour les analyses ultérieures, notamment l’étude des discours normatifs et des comparaisons sectorielles.

3.5 Méthode d’analyse thématique

L’objectif principal de l’analyse thématique est d’explorer et de quantifier les différences de contenu entre les codes de conduite du Fortune Global 500. Pour cela, il est essentiel d’aller au-delà de la simple extraction de clusters et de fournir des mesures rigoureuses permettant d’interpréter les thématiques identifiées.

La présente sous-section décrit les étapes mises en œuvre pour caractériser, comparer et interpréter ces thématiques. Après l’obtention des clusters thématiques via la méthode BERTopic, une analyse quantitative permet de répondre à la première question de recherche : *Quelles sont les différences thématiques dans les codes de conduite des entreprises du Fortune Global 500 ?*

Cette analyse s’articule autour de trois axes complémentaires : l’évaluation de la diversité thématique, la comparaison des profils thématiques entre entreprises, et l’exploration des facteurs explicatifs susceptibles d’influencer ces différences.

3.5.1 Évaluation de la diversité thématique

La première étape consiste à évaluer la diversité thématique au sein de chaque code de conduite. L’objectif est de déterminer si un document se concentre principalement sur un nombre restreint de thèmes ou s’il aborde un large éventail de sujets. Cette analyse permet de mieux comprendre la couverture thématique des codes et de comparer leur niveau de complexité ou de spécificité. Pour mesurer cette diversité, l’entropie de Shannon est utilisée. Cet indicateur statistique, largement mobilisé en théorie de l’information, permet de quantifier la dispersion d’une distribution. Formellement, l’entropie de Shannon pour une entreprise i est calculée selon la formule suivante :

$$H_i = - \sum_{k=1}^K p_{ik} \cdot \log(p_{ik})$$

où K représente le nombre total de thèmes identifiés dans le corpus et p_{ik} la proportion de paragraphes du code de conduite de l'entreprise i associés au thème k . Une entropie élevée indique une répartition équilibrée des thèmes, tandis qu'une entropie faible suggère une focalisation sur un nombre restreint de sujets.

Cette mesure permet ainsi de comparer la variété des contenus abordés par les différentes entreprises : certaines peuvent couvrir un large spectre de thématiques (éthique, conformité, droits des employés, environnement), tandis que d'autres se concentrent davantage sur des sujets spécifiques. L'analyse de l'entropie offre un premier niveau de compréhension des différences de stratégie éditoriale entre entreprises.

3.5.2 Analyse comparative des profils thématiques

Au-delà de la diversité interne de chaque code de conduite, il est pertinent d'examiner les similarités et différences entre entreprises en termes de structures thématiques. L'objectif est d'identifier des profils types d'organisations partageant des logiques discursives communes.

Pour cela, un clustering non supervisé est appliqué aux profils thématiques des entreprises, représentés par la matrice *entreprise* $\times$ *thème*. L'algorithme *KMeans* est choisi pour sa capacité à regrouper efficacement des observations selon leur proximité dans l'espace thématique. La détermination du nombre optimal de clusters repose sur la méthode du coude, afin d'assurer la robustesse des regroupements.

Une fois les clusters identifiés, une visualisation des résultats est réalisée grâce à *t-SNE* (t-distributed Stochastic Neighbor Embedding), qui permet de projeter les données dans un espace bidimensionnel. Cette représentation facilite l'interprétation des regroupements d'entreprises et permet d'illustrer visuellement la répartition des profils thématiques. L'objectif est de mettre en évidence des logiques discursives dominantes : observe-t-on des entreprises au profil « généraliste », couvrant un large spectre de thèmes, ou des profils plus spécialisés, centrés sur des sujets spécifiques ?

3.5.3 Exploration des facteurs explicatifs

Enfin, pour affiner la compréhension des différences thématiques, des analyses croisées sont menées entre les profils thématiques des entreprises et leurs caractéristiques institutionnelles. L'objectif est de déterminer si certaines variables, comme le secteur d'activité, la position dans le classement Fortune Global 500 ou la zone géographique, influencent la structuration des

codes de conduite.

Cette étape repose sur une combinaison d'analyses descriptives (tableaux croisés, graphiques en barres, heatmaps) et de tests d'indépendance statistiques, notamment le test du χ^2 . Plus précisément, il permet de vérifier si la distribution observée des entreprises dans les clusters est indépendante des catégories étudiées. Plus cette valeur est élevée, plus cela indique une différence marquée entre la répartition observée et celle attendue par hasard, suggérant une association entre la variable testée (par exemple, le secteur d'activité ou le rang Fortune) et la structuration thématique des codes de conduite. Cela permet ainsi de confirmer si certaines caractéristiques institutionnelles des entreprises influencent effectivement leur profil discursif.

3.5.4 Discussion de la méthode

En intégrant ces analyses, la méthodologie vise à fournir un regard global et nuancé sur le contenu des codes de conduite du Fortune Global 500. Ce cadre analytique permet d'identifier à la fois des tendances générales (profils thématiques communs) et des spécificités liées au contexte sectoriel ou géographique des entreprises. Ces résultats, détaillés dans le chapitre suivant, apportent un éclairage sur la structuration des discours normatifs dans le monde des affaires.

3.6 Méthode d'analyse des intentions discursives

La deuxième question de recherche de ce mémoire porte sur l'identification des intentions discursives présentes dans les codes de conduite des entreprises du *Fortune Global 500* :

Comment ces thèmes sont-ils exprimés en termes d'intention et de finalité discursive, c'est-à-dire dans quelle mesure ces codes visent-ils à s'adresser à leurs employés ?

Pour y répondre, il est nécessaire d'aller au-delà d'une simple analyse de sentiment ou de ton, comme discuté dans la revue de littérature.

Cette section détaille la méthodologie développée pour détecter et structurer ces intentions discursives de manière automatisée et reproductible. L'objectif est de produire une variable d'intention pour chaque thème du corpus, afin d'enrichir l'analyse des codes de conduite et de faciliter les comparaisons inter-entreprises et inter-thématiques.

3.6.1 Choix du modèle d'analyse

La détection des intentions repose sur l'utilisation d'un modèle de *zero-shot classification*, qui permet d'attribuer à chaque paragraphe une intention probable sans nécessiter de données annotées au préalable. Comme le souligne la littérature (Abercrombie & Batista-Navarro, 2020 ; Ni et al., 2022), cette approche est particulièrement adaptée aux textes normatifs, où il n'existe pas de corpus d'apprentissage spécifique.

Le modèle sélectionné est *DeBERTa-v3-large-mnli-fever-anli-ling-wanli*, pré-entraîné sur des tâches d'inférence logique (*Natural Language Inference, NLI*). Il évalue la probabilité qu'un texte implique, contredise ou soit neutre par rapport à une hypothèse donnée.

Concrètement, chaque paragraphe est confronté à une série d'hypothèses formulées en anglais, correspondant aux cinq intentions retenues dans ce mémoire suite à la revue de littérature : *protect* (protéger), *guide* (guider), *discourage* (dissuader), *punish* (sanctionner) et *empower* (encourager).

Par exemple : “*This paragraph is intended to protect the employee’s safety, rights, or well-being.*”

Pour chaque hypothèse, le modèle renvoie un score de probabilité indiquant dans quelle mesure le paragraphe exprime cette intention. L'intention dominante est attribuée automatiquement au paragraphe en sélectionnant l'hypothèse avec le score le plus élevé.

Cette approche permet de formaliser les objectifs discursifs des codes de conduite, en identifiant systématiquement les intentions sous-jacentes à chaque segment de texte. Elle répond à la nécessité de dépasser l'analyse émotionnelle classique pour adopter une grille plus adaptée aux textes normatifs, en ligne avec les travaux de la littérature.

3.6.2 Construction des variables d'analyse

Une fois l'attribution des intentions réalisée, chaque paragraphe est enrichi par une variable supplémentaire représentant son intention dominante. Ces paragraphes sont ensuite agrégés par entreprise et par thème, afin de construire une matrice *entreprise × thème × intention*.

Cette structuration permet d'analyser comment les entreprises expriment leurs intentions à travers les différents thèmes abordés dans leurs codes de conduite. Par exemple, on peut observer si un thème comme « santé et sécurité » est majoritairement abordé sous un registre de protection, de dissuasion ou d'empowerment.

Cette approche garantit une cohérence entre l'analyse thématique (présentée dans la section précédente) et l'analyse discursive, en permettant une comparaison croisée des intentions dominantes au sein de chaque thème. Elle évite également les biais liés à la structure des documents (certains codes étant très longs, d'autres plus courts), en ramenant l'analyse au niveau du paragraphe, considéré comme l'unité d'analyse la plus pertinente pour ce type de texte.

3.6.3 Discussion de la méthode

Ce choix méthodologique est motivé par plusieurs constats issus de la littérature. D'une part, les textes normatifs n'expriment pas directement des émotions mais véhiculent plus de intentions telles que guider, protéger ou sanctionner. D'autre part, la méthode *zero-shot* permet de contourner l'absence de corpus annotés pour ces tâches spécifiques et d'adapter l'analyse à la diversité des styles rédactionnels rencontrés dans les codes de conduite.

En combinant l'analyse des thèmes et celle des intentions, cette méthodologie permet d'explorer non seulement le contenu des codes de conduite, mais aussi leur communication. Elle offre ainsi une grille de lecture différente sur la façon dont les entreprises structurent et véhiculent leurs messages normatifs. Les résultats de cette analyse sont détaillés dans le chapitre suivant.

4. Résultats

4.1 Analyse exploratoire du corpus

Avant d’aborder les dimensions thématiques et discursives des codes de conduite, une première exploration statistique du corpus est menée à partir des métadonnées associées à chaque document. Cette analyse permet de décrire la structure globale du dataset, d’identifier les déséquilibres éventuels et de poser un cadre interprétatif solide pour les sections suivantes.

Cette analyse porte sur la composition du corpus à l’échelle des entreprises, en s’appuyant sur les métadonnées disponibles pour chacune d’elles : le secteur d’activité, le continent, et le rang dans le classement *Fortune Global 500*.

Le corpus final comprend 401 codes de conduite d’entreprises collectés. Ce nombre représente environ 80% des entreprises pour lesquelles il a été possible d’obtenir un code de conduite disponible en ligne, reflétant un échantillon significatif et pertinent pour l’analyse.

4.1.1 Répartition par secteur

La répartition sectorielle du corpus (figure 4) montre une bonne couverture de l’ensemble des secteurs d’activité représentés dans le classement *Fortune Global 500*. Aucun secteur n’est totalement absent, bien qu’on observe une certaine concentration autour de branches comme l’énergie, les services financiers, l’industrie et les technologies. Des secteurs moins attendus, tels que les transports, les services postaux ou l’agroalimentaire, sont également représentés, mais dans des proportions moindres.

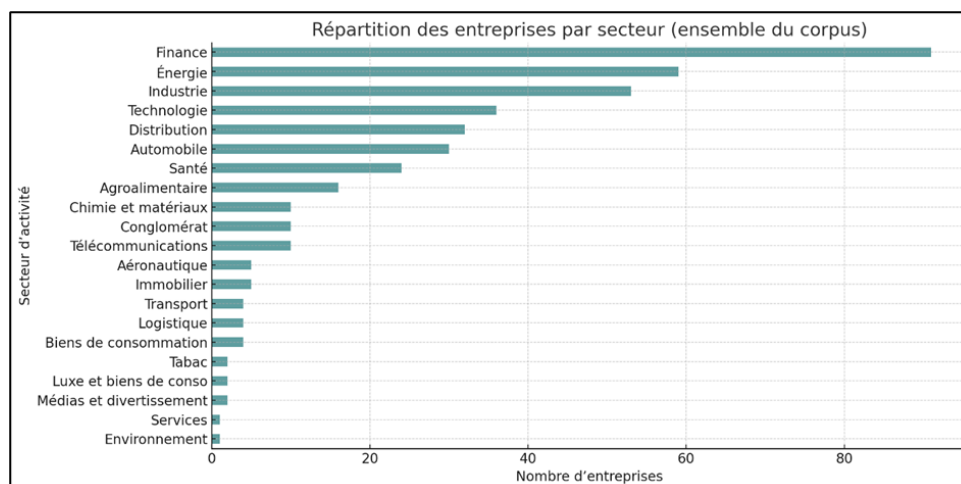


FIGURE 3 – Répartition des entreprises du corpus par secteur d’activité.

4.1.2 Répartition par continent

La répartition des entreprises par continent montre une forte représentation de l'Asie, suivie de près par l'Amérique du Nord et l'Europe. L'Amérique du Sud et l'Océanie apparaissent en revanche très faiblement dans le corpus.

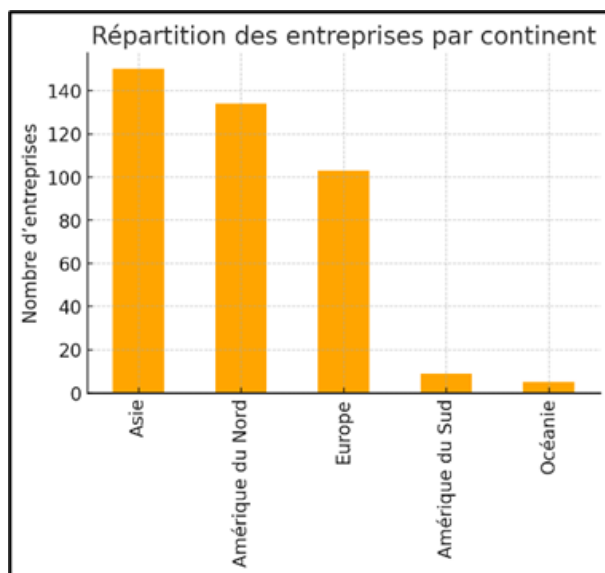


FIGURE 4 – Répartition des entreprises par continent.

4.1.3 Répartition par rang

La distribution des rangs montre que le corpus couvre une part significative du classement *Fortune Global 500*. Toutefois, cette couverture reste légèrement déséquilibrée, en raison de l'inaccessibilité de certains codes de conduite qui n'ont pas pu être collectés malgré leur présence dans le classement.

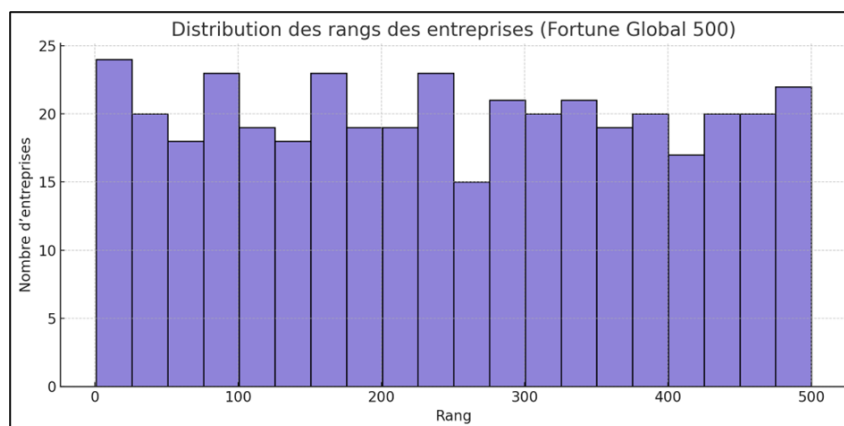


FIGURE 5 – Distribution des rangs des entreprises.

Cette étude présente naturellement certains biais structurels liés à la composition du corpus. La prépondérance de certaines régions du monde (notamment l'Asie) ou de certains secteurs d'activité (comme l'énergie ou la finance) peut influencer la répartition des thèmes et des styles discursifs observés. Toutefois, ces déséquilibres ne relèvent pas d'un biais de sélection méthodologique, mais s'expliquent directement par le contexte de l'étude, qui porte sur les entreprises du classement *Fortune Global 500*, reflet d'un certain équilibre économique mondial.

4.2 Résultats du modèle Bertopic

Dans cette section, nous présentons les résultats obtenus après l'application du modèle BERTopic au corpus des codes de conduite du F500. Notre analyse s'est articulée autour de trois configurations expérimentales : un modèle de référence standard (baseline), un ajustement manuel des hyperparamètres et une optimisation automatique par Optuna.

4.2.1 Configurations testées

Nous avons examiné trois configurations principales du modèle BERTopic sur notre corpus :

- **Modèle de référence (baseline)** : il s'agit de la configuration standard de *BERTopic*, avec les paramètres par défaut pour la réduction de dimension (*UMAP*) et le clustering (*HDBSCAN*).
- **Ajustement manuel des hyperparamètres** : certains hyperparamètres clés ont été ajustés manuellement afin d'améliorer la qualité des thèmes. Par exemple, nous avons modifié *n_neighbors*, *min_samples* et *min_cluster_size* pour influencer la granularité thématique et réduire le nombre d'outliers.
- **Optimisation automatique (Optuna)** : nous avons utilisé la bibliothèque *Optuna* pour effectuer une recherche automatique des meilleurs hyperparamètres. Un espace de recherche a été défini (incluant notamment les paramètres *UMAP* et *HDBSCAN*), et *Optuna* a cherché la combinaison maximisant la cohérence thématique sur un sous-ensemble de validation.

Pour chaque configuration, le modèle a été entraîné sur l'ensemble du corpus avec la même préparation textuelle initiale.

4.2.2 Comparaison des performances

Les performances des trois configurations sont comparées dans le tableau 1 en utilisant les métriques suivantes présentées dans la section méthodologie : cohérence thématique, diversité thématique et taux d’outliers. La cohérence thématique, mesurée par le score C_v , évalue la cohérence sémantique des mots au sein de chaque thème, tandis que la diversité thématique correspond au pourcentage de mots uniques dans l’ensemble des thèmes. Le taux d’outliers indique la proportion de documents du corpus qui n’ont pas été rattachés à un thème.

Configuration	Cohérence (%)	Diversité (%)	Taux d’outliers (%)
Baseline	57.07	37.61	33.30
Ajustement manuel	55.04	37.51	34.50
Optimisation (Optuna)	54.63	15.27	5.00

TABLE 1 – Comparaison des performances des trois modèles *BERTopic*.

Le 1 qui compare les modèles montre que la configuration optimisée par Optuna présente la cohérence thématique la plus faible (54,63%), tandis que le modèle baseline affiche la cohérence la plus élevée (57,07%), suivi du modèle ajusté manuellement (55,04%). En termes de diversité thématique, le modèle baseline se distingue également, avec la valeur la plus élevée (37,61%), légèrement supérieure à celle obtenue avec le modèle ajusté manuellement (37,51%). En revanche, le modèle optimisé par Optuna affiche une diversité nettement plus faible (15,27%), traduisant une perte significative de richesse thématique. Par ailleurs, l’optimisation par Optuna permet de réduire fortement le taux d’outliers, qui passe de 33,3% pour le modèle baseline et 34,5% pour le modèle ajusté manuellement, à seulement 5%.

4.2.3 Justification du choix du modèle de référence

Bien que l’optimisation automatique par Optuna réduise significativement le taux d’outliers (jusqu’à 5,00%), le modèle baseline est retenu comme référence pour les analyses ultérieures. Ce choix repose sur un compromis méthodologique : le modèle baseline affiche la cohérence thématique la plus élevée (0,57), traduisant une structuration claire des thèmes, tout en préservant une diversité plus riche que celle obtenue par les configurations alternatives. À l’inverse, le modèle optimisé, bien qu’il améliore la couverture du corpus en réduisant les outliers, présente une diversité fortement réduite (0,1527), ce qui risque d’appauvrir la richesse des thématiques et de limiter la portée analytique des résultats.

En définitive, le modèle baseline offre un équilibre satisfaisant entre cohérence, diversité et couverture des données. Cette configuration garantit une analyse thématique robuste, préservant la complexité des contenus tout en limitant les pertes d'information, ce qui est essentiel pour l'interprétation qualitative et les analyses discursives approfondies qui suivent.

4.2.4 Regroupement et réduction des clusters

Comme énoncé dans la section méthodologie, après sélection du modèle de référence, une phase de post-traitement est menée sur les clusters extraits afin d'en affiner l'interprétation et de faciliter l'analyse. Cette étape regroupe les thèmes proches.

Nous avons d'abord examiné *l'intertopic distance map* générée à partir du modèle baseline. Cette carte, qui représente la distance sémantique entre les clusters dans un espace à 2 dimensions, a servi à identifier les regroupements potentiels. Les thèmes présentant des positions proches ou des chevauchements ont été jugés sémantiquement similaires et fusionnés manuellement. Cette opération a permis de réduire le nombre total de clusters thématiques de 90 à 16.

Les 16 thèmes finaux identifiés couvrent un large éventail de sujets relatifs aux enjeux des entreprises, allant de la sécurité des employés et la confidentialité des données à la gouvernance d'entreprise et la responsabilité environnementale. Le tableau ci-dessous présente ces thématiques finales.

Thématiques finales identifiées

Workplace Safety and Ethics
Data Privacy and Confidentiality
Fair Competition and Market Practices
Anti-Corruption and Financial Integrity
Gifts, Hospitality, and Political Activities
Conflict of Interest Management
Financial Reporting and Asset Management
Environmental Sustainability and Social Responsibility
Diversity, Equity, and Inclusion
Corporate Governance and Board Oversight
Community Development and Poverty Alleviation
Technical Compliance and Infrastructure Standards
Code of Conduct Enforcement and Reporting
Nuclear Safety and Regulatory Compliance
Environmental Management in Oil and Water Resources
Client Relations and Fiduciary Responsibilities

TABLE 2 – Les 16 thématiques finales extraites des codes de conduite.

Pour illustrer le contenu des thématiques identifiées, des nuages de mots ont été générés à partir des clusters finaux. Ces visualisations mettent en avant les termes les plus fréquemment associés à chaque thème et permettent d’avoir un aperçu rapide des concepts clés abordés dans les codes de conduite.



FIGURE 6 – Nuage de mots du topic *Workplace Safety and Ethics*.

Le nuage de mots ci-dessus montre les termes les plus représentatifs du thème *Workplace Safety and Ethics*, incluant des concepts tels que *employee*, *safety*, *health*, *compliance* et *harassment*. Ces mots-clés reflètent les préoccupations majeures autour de la sécurité et du bien-être des employés.

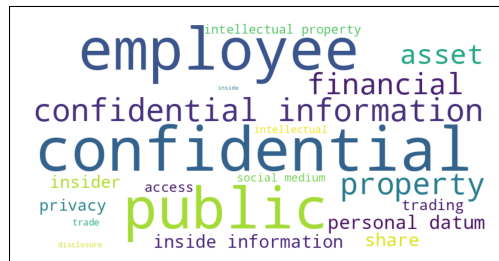


FIGURE 7 – Nuage de mots du topic *Data Privacy and Confidentiality*.

Ce deuxième exemple illustre le thème *Data Privacy and Confidentiality*, dominé par des termes comme *confidential*, *data*, *information*, *access* et *employee*, mettant en évidence les enjeux liés à la gestion des données sensibles et à la protection de la vie privée dans les entreprises.

Pour finir, les thèmes obtenus dans cette section serviront de base aux analyses thématiques et d'intentions présentées dans les sections suivantes du mémoire.

4.3 Analyse thématiques des codes de conduites

Cette section présente les résultats de l'analyse thématique des codes de conduite des entreprises du classement Fortune Global 500, conformément à la méthodologie. L'objectif est de répondre à la première question de recherche :

Quelles sont les différences thématiques dans les codes de conduite des entreprises du Fortune Global 500 ?

L'analyse est structurée selon les trois axes définis dans la méthodologie :

- L'évaluation de la diversité thématique à l'échelle des entreprises, afin de mesurer la richesse et l'équilibre des thèmes abordés.
- La comparaison des profils via des regroupements d'entreprises selon leurs thèmes dominants.
- L'exploration des métadonnées (secteur d'activité, rang Fortune et zone géographique) susceptibles d'influencer ces différences thématiques.

Ces résultats permettent de caractériser les codes de conduite pour mettre en lumière des logiques communes ou spécifiques selon le contexte des entreprises.

4.3.1 Évaluation de la diversité thématique à l'échelle des entreprises

L'analyse débute par une évaluation de la diversité thématique au sein des codes de conduite de chaque entreprise. Cette étape permet de mesurer dans quelle mesure les documents abordent un large éventail de sujets ou, au contraire, se concentrent sur un nombre restreint de thématiques.

Mesure de l'entropie comme indicateur de dispersion

Comme présenté dans la section méthodologique, l'entropie de Shannon est utilisée ici pour évaluer la diversité thématique des codes de conduite, en mesurant la dispersion des paragraphes entre les différents thèmes identifiés. Une entropie élevée traduit une répartition équilibrée entre plusieurs thèmes, tandis qu'une entropie faible indique une concentration du discours sur un nombre limité de sujets.

La figure 8 présente la distribution des valeurs d'entropie au sein du corpus. On observe une variabilité notable : la majorité des entreprises affiche une entropie comprise entre 2.0 et 2.8, reflétant une couverture thématique relativement diversifiée. Toutefois, un nombre non négligeable d'entreprises présente une entropie proche de zéro, suggérant un discours fortement polarisé autour d'un thème unique.

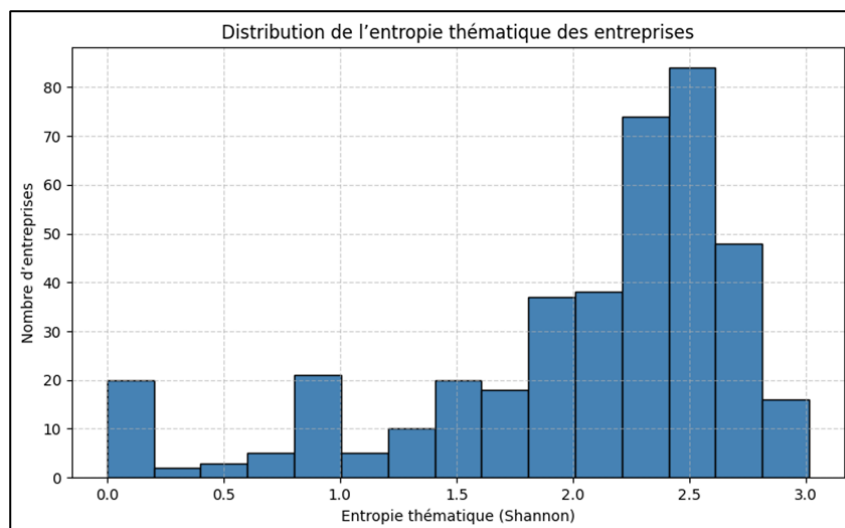


FIGURE 8 – Distribution de l'entropie thématique des entreprises

Étude de cas sur un échantillon d'entreprises

Pour rentrer un peu plus dans le détail de la figure précédente sur la variabilité observée dans la distribution des entropies, un échantillon de sept entreprises a été sélectionné en fonction de leur profil entropique. La figure 9 présente la répartition des paragraphes entre les différents thèmes pour chacune de ces entreprises, exprimée en pourcentage.

Merged_Topic_Name		Application du code de conduite et signalement	Cadeaux, hospitalité et activités politiques	Concurrence loyale et pratiques de marché	Confidentialité des données	Conformité technique et normes d'infrastructure	Diversité, équité et inclusion	Durabilité environnementale et responsabilité sociale	Développement communautaire et réduction de la pauvreté	Gestion des conflits d'intérêts	Gestion environnementale des ressources en eau et pétrole	Gouvernance d'entreprise et surveillance du conseil	Intégrité financière et anticorruption	Rapport financier et gestion des actifs	Relations clients et responsabilités fiduciaires	Sécurité et éthique au travail	Sécurité nucléaire et conformité réglementaire
Sector	Name																
Énergie	ENI	15%	4%	22%	7%	0%	4%	7%	0%	7%	0%	4%	7%	0%	0%	22%	0%
Technologie	Lenovo Group	3%	14%	14%	9%	0%	6%	9%	0%	9%	0%	0%	17%	0%	0%	20%	0%
Finance	Citigroup	0%	9%	9%	21%	0%	3%	6%	0%	12%	0%	0%	18%	0%	6%	18%	0%
	Berkshire Hathaway	0%	9%	9%	36%	0%	0%	0%	0%	9%	0%	0%	0%	0%	0%	36%	0%
Santé	UnitedHealth Group	0%	19%	11%	30%	0%	0%	0%	0%	0%	0%	0%	4%	0%	0%	37%	0%
Automobile	Guangzhou Automobile Industry Group	0%	0%	0%	0%	0%	0%	0%	0%	0%	0%	100%	0%	0%	0%	0%	0%
Biens de consommation	Nike	0%	0%	0%	0%	0%	0%	0%	0%	0%	0%	0%	0%	0%	0%	100%	0%

FIGURE 9 – Répartition thématique des codes de conduite pour sept entreprises, en pourcentage de paragraphes par thème.

En complément, le tableau 3 synthétise les valeurs exactes d'entropie pour ces mêmes entreprises.

Secteur	Entreprise	Entropie
Énergie	ENI	3.01
Technologie	Lenovo Group	3.00
Finance	Citigroup	2.96
Santé	UnitedHealth Group	2.03
Finance	Berkshire Hathaway	2.00
Biens de consommation	Nike	0.00
Automobile	Guangzhou Automobile Industry Group	0.00

TABLE 3 – Valeurs d’entropie mesurant la diversité thématique des codes de conduite par entreprise.

Ces résultats permettent de distinguer trois profils thématiques types :

- **Codes à entropie élevée** (ENI, Lenovo Group, Citigroup) : ils sont caractérisés par une couverture homogène et transversale d’un large spectre de thématiques.
- **Codes à entropie intermédiaire** (UnitedHealth Group, Berkshire Hathaway) : ils sont marqués par une focalisation sur quelques thématiques dominantes, tout en maintenant une certaine diversité.
- **Codes à entropie nulle** (Nike, Guangzhou) : ils sont concentrés exclusivement sur un thème unique, traduisant une spécialisation très marquée.

Pour résumer, l’analyse de l’entropie a permis de quantifier la diversité thématique des codes de conduite, en mesurant la répartition des paragraphes entre les différents thèmes. La distribution des entropies met en évidence une variabilité importante entre entreprises : si la majorité d’entre elles présentent une couverture relativement équilibrée (entropie entre 2,0 et 2,8), un nombre non négligeable se concentre sur un très petit nombre de thèmes, traduisant une structuration plus spécialisée du discours. Cette diversité se retrouve également dans l’étude de cas sur un échantillon d’entreprises, qui illustre concrètement les écarts de profils : certains codes affichent une approche généraliste et homogène, tandis que d’autres adoptent un discours plus ciblé, centré sur quelques thèmes dominants.

4.3.2 Regroupement des entreprises par profils thématiques

Après avoir évalué la diversité des codes de conduite au niveau individuel, on élargit un peu l'analyse en concentrant la recherche sur les structures partagées entre entreprises. L'objectif est d'identifier des regroupements d'entreprises qui, indépendamment de leur secteur ou de leur taille, adoptent des logiques similaires dans la manière dont elles abordent les thèmes dans leur code de conduite. Pour ce faire, une approche de clustering non supervisé est appliquée sur la matrice entreprise $\times$ thème, permettant de détecter des profils thématiques récurrents dans le corpus.

Construction du profil thématique

Pour mettre en évidence des structures communes entre les entreprises, une matrice entreprise $\times$ thème est construite.

Cette représentation vectorielle permet d'appliquer l'algorithme de clustering non supervisé KMeans sur les données. Le choix du nombre optimal de clusters a été déterminé à l'aide de la méthode du coude, qui met en évidence un palier autour de trois groupes. Cette configuration a été confirmée visuellement grâce à une projection t-SNE : la figure 10 montre que trois regroupements cohérents se dessinent nettement dans l'espace réduit.

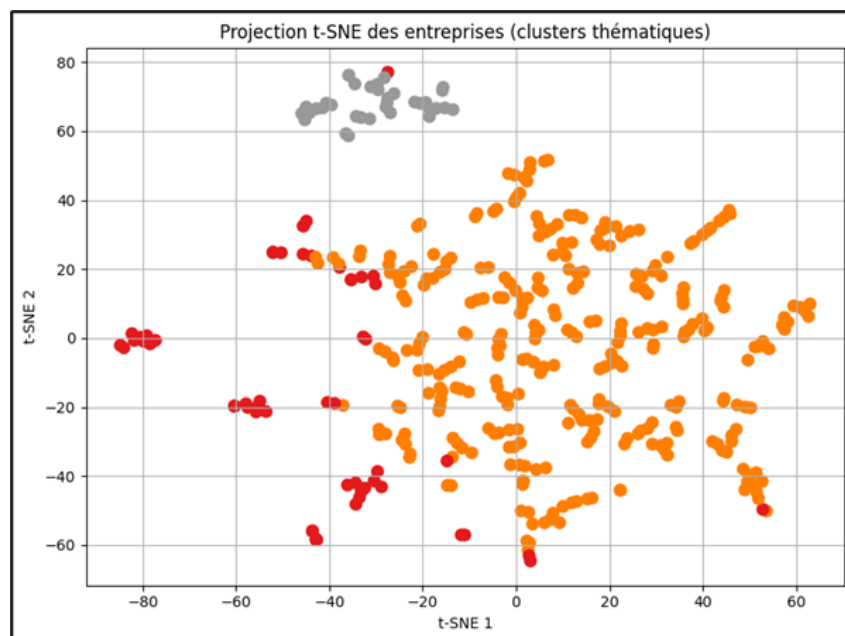


FIGURE 10 – Représentation des clusters

Après avoir visualisé la structuration générale des clusters dans l'espace réduit (Figure 11), il est essentiel d'examiner de manière plus détaillée la composition thématique de chacun d'entre eux afin de mieux comprendre les logiques sous-jacentes à ces regroupements.

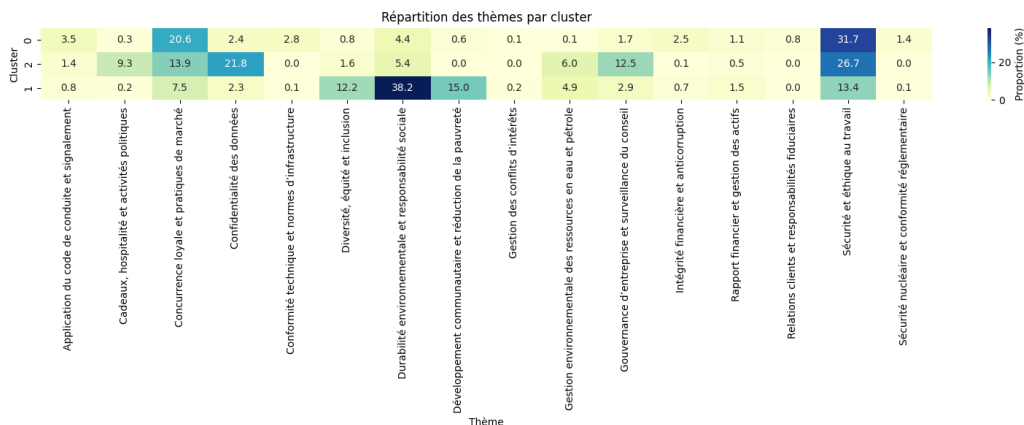


FIGURE 11 – Répartition des thèmes par cluster

La Figure 11 présente la répartition des thèmes identifiés dans les codes de conduite au sein de chaque cluster. Chaque ligne correspond à un cluster, et chaque colonne indique la proportion moyenne de paragraphes d'un cluster donnée abordant un thème particulier. Cette heatmap permet de dégager des tendances générales dans la structuration thématique des entreprises du corpus.

L'analyse révèle trois profils distincts :

- **Le cluster 2 (orange)** se distingue par une couverture thématique large et équilibrée, abordant de manière significative des sujets variés tels que la sécurité au travail, la confidentialité des données, ou encore la diversité et l'inclusion. Ce profil généraliste est majoritaire dans le corpus, regroupant 337 entreprises. Sa richesse thématique suggère une approche de communication éthique plus globale et transverse.
- **Le cluster 0 (gris)** regroupe des codes de conduite au profil polarisé. Ces entreprises concentrent leur discours sur un nombre limité de thématiques, notamment la sécurité au travail et la gouvernance d'entreprise. Ce profil représente 9 entreprises et traduit une orientation plus spécialisée, souvent focalisée sur des enjeux clés spécifiques à leur secteur ou à leur contexte réglementaire.
- **Le cluster 1 (rouge)** se caractérise par une prédominance marquée des thématiques liées à la durabilité environnementale et à la responsabilité sociale (environ 38% des paragraphes), mais conserve néanmoins une certaine diversité, en abordant également des

sujets comme la diversité, l'équité et l'inclusion ou la sécurité au travail. Ce profil mixte, bien que minoritaire (55 entreprises), reflète une attention particulière portée aux enjeux environnementaux, tout en maintenant un socle thématique plus large.

Cette analyse préliminaire met en évidence la diversité des structurations thématiques au sein des codes de conduite. Mais ces observations nécessitent d'être approfondies et mises en relation avec des variables institutionnelles ou contextuelles, telles que le secteur d'activité, le rang dans le classement Fortune Global 500 ou la localisation géographique des entreprises. Ces facteurs potentiels d'explication seront examinés de manière systématique dans la section suivante.

4.3.3 Facteurs influençant

L'analyse descriptive et comparative réalisée jusqu'ici a permis de caractériser la structuration thématique des codes de conduite à travers deux dimensions complémentaires : d'une part, la diversité interne des thèmes abordés au sein de chaque code, mesurée par l'entropie de Shannon ; d'autre part, l'appartenance des entreprises à des profils discursifs distincts, identifiés par clustering non supervisé sur la base des distributions thématiques. Ces deux angles d'analyse fournissent une première cartographie des différences thématiques entre entreprises. Toutefois, afin de mieux comprendre les facteurs pouvant expliquer ces variations, il convient désormais d'examiner systématiquement l'influence de variables institutionnelles et contextuelles, telles que le secteur d'activité, le rang dans le classement Fortune Global 500 ou la localisation géographique.

Afin d'affiner la compréhension de ces variations, cette section explore l'hypothèse selon laquelle des facteurs exogènes pourraient influencer le positionnement thématique des entreprises. Plus précisément, il s'agit d'examiner si des variables institutionnelles ou contextuelles – telles que le secteur d'activité, la position dans le classement *Fortune Global 500*, ou la zone géographique d'origine – sont associées aux profils discursifs identifiés.

Trois axes d'analyse sont considérés :

1. **Le secteur d'activité**, qui pourrait structurer le contenu des codes en fonction des risques métiers ou des normes.
2. **Le rang dans le classement Fortune Global 500**, qui pourrait refléter la taille ou la visibilité attendu dans les codes.
3. **La zone géographique d'origine**, en lien avec les cultures de gouvernance ou les environnements juridiques locales.

Chacun de ces facteurs sera analysé à travers des croisements de données, des représentations graphiques et des tests statistiques, dans le but de valider ou non leur influence sur la structuration thématique des codes de conduite.

Influence du secteur d'activité

L'analyse commence par une étude de la distribution des profils de thèmes selon le secteur d'activité. Le tableau 4 présente la répartition des entreprises dans les trois clusters identifiés

(généraliste, polarisé, marginal) au sein des 21 secteurs étudiés.

Secteur	Cluster 0	Cluster 1	Cluster 2
Agroalimentaire	0 (0%)	1 (6%)	15 (94%)
Automobile	0 (0%)	5 (17%)	25 (83%)
Aéronautique	0 (0%)	0 (0%)	5 (100%)
Biens de consommation	0 (0%)	0 (0%)	4 (100%)
Chimie et matériaux	1 (10%)	2 (20%)	7 (70%)
Conglomérat	0 (0%)	1 (10%)	9 (90%)
Distribution	0 (0%)	4 (12%)	28 (88%)
Environnement	0 (0%)	0 (0%)	1 (100%)
Finance	2 (2%)	10 (11%)	79 (87%)
Immobilier	0 (0%)	0 (0%)	5 (100%)
Industrie	4 (8%)	15 (28%)	34 (64%)
Logistique	0 (0%)	1 (25%)	3 (75%)
Luxe et biens de conso	0 (0%)	0 (0%)	2 (100%)
Médias et divertissement	0 (0%)	0 (0%)	2 (100%)
Santé	0 (0%)	0 (0%)	24 (100%)
Services	0 (0%)	0 (0%)	1 (100%)
Tabac	0 (0%)	0 (0%)	2 (100%)
Technologie	0 (0%)	4 (11%)	32 (89%)
Transport	0 (0%)	0 (0%)	4 (100%)
Télécommunications	0 (0%)	0 (0%)	10 (100%)
Énergie	2 (3%)	12 (20%)	45 (76%)

TABLE 4 – Répartition sectorielle des entreprises selon les clusters thématiques identifiés.

L'observation générale met en évidence une domination nette du **cluster 2**, qui regroupe la majorité des entreprises dans presque tous les secteurs. Ce profil « généraliste » traduit une couverture thématique étendue et équilibrée, indépendamment du secteur d'activité. Les clusters 1 et 0 apparaissent de manière plus marginale, sans logique sectorielle stricte : même lorsque plusieurs entreprises d'un même secteur se retrouvent dans le même cluster, leurs socles de thèmes varient. Ces constats suggèrent que l'appartenance sectorielle n'exerce pas de réelle

influence sur la structuration des codes de conduite.

Pour valider cette observation, un test d'indépendance du χ^2 a été effectué sur le tableau de contingence croisant les 21 secteurs d'activité et les trois clusters. Les résultats sont les suivants : $\chi^2 = 42.66$, **Degrés de liberté = 40**, **p-value = 0.3575**.

La non-significativité statistique confirme l'absence de lien robuste entre le secteur d'activité et la structuration des thèmes dans les codes, renforçant l'idée d'une logique transversale.

Il devient, alors, intéressant d'observer d'autres facteurs potentiellement explicatifs, tels que le rang Fortune ou l'origine géographique des entreprises.

Influence du rang

Après l'examen du secteur d'activité, l'analyse se poursuit en explorant l'influence du *rang* des entreprises dans le classement *Fortune Global 500* sur la diversité thématique des codes de conduite. L'hypothèse sous-jacente est que les entreprises les mieux classées, en raison de leur taille et de leur visibilité, pourraient adopter des discours plus diversifiés, couvrant un spectre plus large de thèmes. Pour tester cette hypothèse, les entreprises ont été regroupées en quatre quartiles, de Q1 (les 125 premières entreprises) à Q4 (les 125 dernières), et la diversité des thèmes a été mesurée par l'entropie moyenne de Shannon.

Pour rappel cette entropie moyenne de Shannon permet de quantifier la dispersion des thèmes au sein des codes de conduite. Plus cette valeur est élevée, plus le discours est diversifié et équilibré entre différents sujets ; à l'inverse, une entropie faible indique une concentration sur un nombre restreint de thèmes.

Le **tableau 5** et la **figure 12** ci-dessous présentent la distribution de l'entropie pour chaque quartile. Ces résultats permettent d'examiner si des variations significatives apparaissent en fonction du rang des entreprises.

Quartile	Entropie moyenne	Écart-type	Nombre d'entreprises
Q1 (1–125)	2.06	0.72	101
Q2 (126–250)	2.02	0.71	100
Q3 (251–375)	2.05	0.71	100
Q4 (376–500)	2.00	0.72	100

TABLE 5 – Entropie moyenne des codes de conduite par quartile de classement Fortune Global 500.

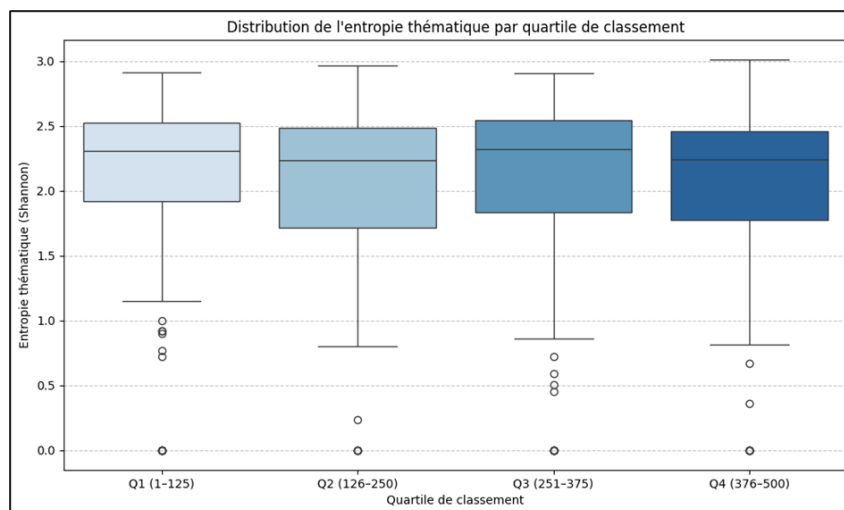


FIGURE 12 – Distribution de l’entropie thématique des codes de conduite selon le quartile de classement.

Les résultats mettent en évidence une remarquable stabilité de l’entropie moyenne entre les différents quartiles de classement, oscillant entre 2,00 et 2,06. Cette constance suggère que le rang des entreprises n’a pas d’effet significatif sur la diversité thématique des codes de conduite. Autrement dit, les entreprises les mieux classées ne formulent pas nécessairement des discours plus riches ou plus variés que celles de rang inférieur.

Afin de confirmer l’absence d’influence du rang des entreprises sur la structuration thématique des codes de conduite, un test d’indépendance du χ^2 a été réalisé. Les résultats obtenus sont les suivants : $\chi^2 = 8,93$, degrés de liberté = 6, p -value = 0,1773.

La non-significativité statistique de ce test indique qu’aucune association robuste n’est observée entre le rang des entreprises et leur appartenance à un profil thématique donné. Cette absence de lien vient corroborer les observations descriptives et suggère que la structuration des codes de conduite n’est pas directement déterminée par le classement Fortune. Toutefois, un dernier facteur explicatif, l’origine géographique des entreprises, reste à examiner.

Influence géographique

Enfin, la dernière analyse explore l’influence de la **zone géographique** d’origine des entreprises. Le tableau 6 montre la répartition des clusters par continent, et le tableau 7 présente l’entropie moyenne.

Continent	Cluster 0	Cluster 1	Cluster 2
Amérique du Nord	0 (0%)	3 (2%)	131 (98%)
Amérique du Sud	0 (0%)	0 (0%)	9 (100%)
Asie	9 (6%)	48 (32%)	93 (62%)
Europe	0 (0%)	4 (4%)	99 (96%)
Océanie	0 (0%)	0 (0%)	5 (100%)

TABLE 6 – Répartition des clusters thématiques selon le continent d’origine des entreprises

Continent	Entropie moyenne	Écart-type	Nombre d’entreprises
Amérique du Sud	2.30	0.34	9
Amérique du Nord	2.22	0.61	134
Europe	2.07	0.70	103
Asie	1.85	0.77	150
Océanie	1.42	0.98	5

TABLE 7 – Entropie moyenne des codes de conduite par continent d’origine

Les résultats confirment que la localisation géographique constitue un facteur différenciant : si les entreprises d’Amérique du Nord, d’Europe, d’Amérique du Sud et d’Océanie présentent une structuration homogène et généraliste (cluster 2, entropie élevée), l’Asie se distingue par une plus grande hétérogénéité, avec une proportion plus élevée d’entreprises polarisées ou atypiques et une entropie moyenne plus faible (1.85).

Afin de vérifier si l’origine géographique des entreprises influence la structuration thématique de leurs codes de conduite, un test d’indépendance du χ^2 a été réalisé. Les résultats obtenus sont les suivants : $\chi^2 = 145,28$, degrés de liberté = 12, p-value < 0,0001.

La significativité très forte de ce test indique qu’il existe une association statistiquement robuste entre la zone géographique et la répartition des entreprises au sein des clusters thématiques. Autrement dit, la structuration des codes de conduite semble partiellement déterminée par l’origine géographique des entreprises, ce qui suggère l’existence de logiques régionales ou culturelles influençant la manière dont les engagements éthiques sont formulés. Ces résultats contrastent avec l’absence d’effet significatif du secteur d’activité et du rang dans le classement Fortune Global 500, et soulignent l’importance du contexte géographique dans l’élaboration des discours normatifs des grandes entreprises.

4.3.4 Discussion des résultats

L'analyse des résultats met en lumière plusieurs constats majeurs. Tout d'abord, les tests d'indépendance et les analyses croisées confirment l'absence d'influence significative du secteur d'activité et du rang dans le classement Fortune Global 500 sur la structuration des codes de conduite. Les entreprises, quels que soient leur domaine d'activité ou leur position dans le classement, ne présentent pas de différences notables en matière de diversité thématique ou d'appartenance à un profil discursif donné.

En revanche, la dimension géographique émerge comme un facteur explicatif différenciant.

L'analyse statistique, renforcée par un test du χ^2 très significatif ($\chi^2 = 145,28$, $p < 0,0001$), met en évidence une association robuste entre la zone géographique et la structuration des codes de conduite. Les entreprises asiatiques, notamment, affichent des profils thématiques plus polarisés et une entropie moyenne plus faible, traduisant une certaine concentration des discours sur des enjeux spécifiques. À l'inverse, les entreprises d'Amérique du Nord, d'Europe et d'Amérique du Sud adoptent davantage des profils généralistes, marqués par une plus grande diversité thématique et une homogénéité dans la structuration des engagements.

4.4 Résultats d'analyse des intentions

L'analyse des intentions discursives a permis de caractériser la manière dont les entreprises du *Fortune Global 500* expriment leurs normes et recommandations dans leurs codes de conduite. En s'appuyant sur le modèle *DeBERTa-v3-large-mnli-fever-anli-ling-wanli*, chaque thème du corpus a été classifié selon l'une des cinq intentions définies :

- *protect* : protéger les droits, la sécurité ou le bien-être des employés,
- *guide* : orienter les employés vers un comportement attendu, en les aidant à faire les bons choix dans des situations professionnelles,
- *discourage* : dissuader des comportements jugés inappropriés ou risqués,
- *punish* : évoquer ou annoncer des sanctions en cas de non-respect des règles,
- *empower* : encourager l'initiative individuelle, la prise de parole ou la responsabilité personnelle.

Cette approche permet de dépasser une simple analyse de thèmes, en explorant les intentions qui se cachent derrière les codes.

4.4.1 Répartition globale des intentions

Commençons par une analyse globale sur l'ensemble du corpus.

Le premier graphique (Figure 13) le montre. On observe une prédominance de l'intention « guide », qui représente plus de la moitié des cas. Cela reflète un positionnement largement normatif mais non contraignant, dans lequel les entreprises cherchent à orienter les comportements attendus à travers des formulations non menaçantes. Cette stratégie suppose une adhésion douce du salarié aux règles internes, sans recourir à un discours punitif.

En deuxième position, l'intention « discouragement » occupe environ un quart des cas. Elle correspond à des formulations qui cherchent à prévenir les comportements inappropriés, sans forcément formuler de sanctions explicites.

Les intentions « protect », « empower » et surtout « punish » sont nettement plus rares. Cela suggère que les entreprises recourent relativement peu à un discours strictement répressif (punish), ou à l'inverse, à un registre responsabilisant (empower) qui mettrait l'accent sur la capacité d'action du salarié, tout comme l'intention qui vise à protéger l'employé (protect).

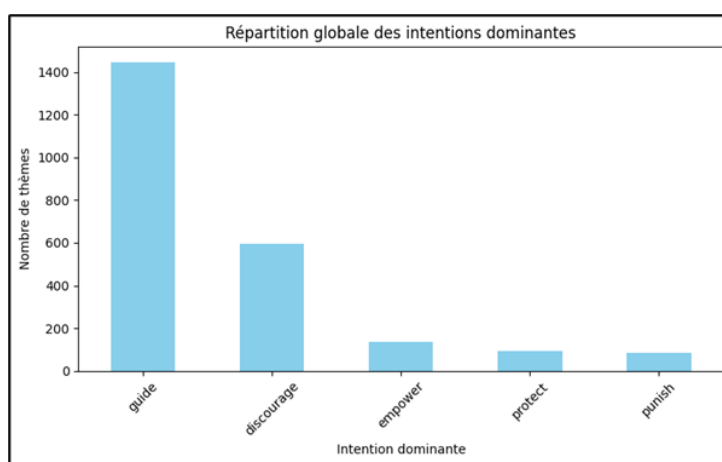


FIGURE 13 – Répartition globale des intentions dominantes dans le corpus

Le deuxième graphique (Figure 14) affine cette analyse en représentant la distribution des intentions dominantes par thème (ou topic). Il révèle que certains thèmes sont très homogènes, tandis que d'autres font l'objet de traitements variés, certainement selon les entreprises.

C'est le cas, par exemple, de « l'application du code de conduite », des « cadeaux, hospitalité et activités politiques », ou encore du « développement communautaire et de la pauvreté », qui sont tous très majoritairement associés à l'intention « *guide* ».

À l'inverse, d'autres thèmes témoignent d'une diversité plus marquée dans le ton employé. Par exemple, les codes abordant la « gestion des conflits d'intérêts », la « protection des données », ou la « sécurité nucléaire et conformité réglementaire » présentent une plus grande répartition entre plusieurs intentions. Ces sujets sont parfois traités sur le mode de la prévention (discourage), parfois dans une logique de responsabilisation (empower), ou même de sanction (punish), ce qui révèle des choix stratégiques différents selon les entreprises ou les secteurs.

Enfin, on note que des domaines liés à la conformité réglementaire, la surveillance, ou l'intégrité financière suscitent des réponses discursives plus équilibrées, mêlant *guide*, *discourage* et, dans une plus faible proportion, *punish*. Cela pourrait suggérer que lorsque les enjeux sont fortement liés au risque juridique, les entreprises adoptent des styles plus nuancés, combinant information, prévention et rappel des conséquences.

Ce graphique met donc en évidence le fait que les codes de conduite ne sont pas uniformes dans leur intentions, et que le thème traité joue un rôle dans cette intention exprimée.

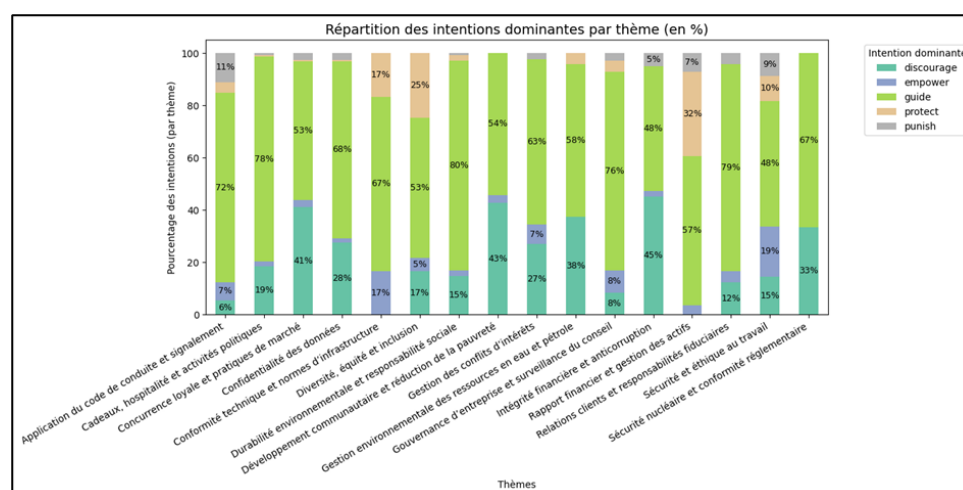


FIGURE 14 – Répartition des intentions par thème

4.4.2 Répartition sectorielle

Le graphique ci-dessous (Figure 15) présente la distribution des intentions dominantes selon les secteurs d'activité des entreprises du corpus. Chaque barre représente la répartition, en pourcentage, des intentions dominantes exprimées par les entreprises d'un même secteur, tous thèmes confondus.

Comme dans l'analyse thématique, on observe une nette domination de l'intention *guide*, quelle que soit la branche d'activité. Ce constat confirme que, de manière transversale, les entreprises

préfèrent adopter un registre orienté vers l'explication des comportements attendus plutôt que vers la répression. Toutefois, des différences notables apparaissent lorsque l'on entre dans le détail des secteurs.

Certains secteurs affichent un profil particulièrement homogène. C'est notamment le cas des secteurs du transport (71 %), de la technologie (66 %), de la logistique (62 %) ou encore de l'énergie (65 %), où l'intention *guide* domine largement et où les autres intentions sont peu mobilisées. Ce type de profil suggère une volonté d'uniformiser les discours internes autour de normes explicites, sans insister fortement sur la dissuasion ou la sanction.

À l'inverse, certains secteurs présentent une plus grande diversité. Par exemple :

- Le secteur des services affiche une part élevée d'intentions *empower* (33 %), traduisant une responsabilisation accrue des collaborateurs.
- Les secteurs de la santé, de l'environnement ou de la chimie laissent plus de place à des intentions *discourage*, signe d'une posture plus préventive, sûrement liée à des enjeux réglementaires ou à la gestion des risques.
- Le secteur des biens de consommation ou des conglomérats se distingue par une présence non négligeable de l'intention *punish* (8 à 11 %), ce qui peut traduire un recours plus à un registre coercitif, pour renforcer l'autorité de certaines règles.

Ces différences soulignent que la stratégie adoptée dans les codes de conduite varie en fonction des contraintes propres à chaque secteur : degré de régulation, exposition au risque, pression juridique, attentes sociétales, etc.

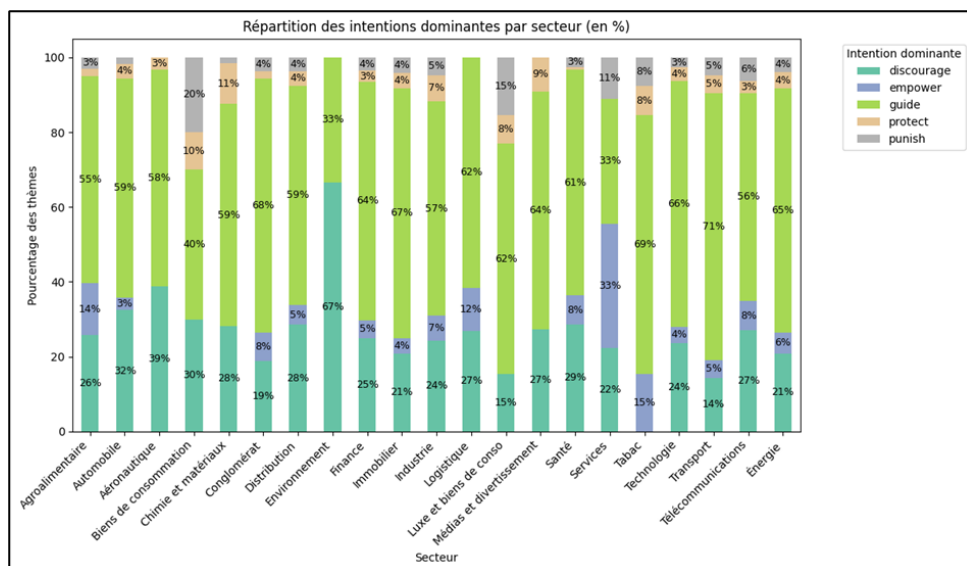


FIGURE 15 – Répartition des intentions par secteur

4.4.3 Répartition par continent

Le graphique ci-dessous (Figure 16) compare les profils des entreprises selon leur continent (siège social). Chaque ligne représente la distribution relative des intentions dominantes exprimées par les entreprises européennes, asiatiques et nord-américaines, toutes thématiques confondues. L'Amérique du Sud et l'Océanie, étant donné qu'elles apparaissent très faiblement dans le corpus, n'ont pas été prises en compte pour cette analyse. Le radar plot permet ainsi de mettre en évidence les différences, ou les similitudes, dans les intentions des codes de conduite selon le contexte géographique.

Dans l'ensemble, quel que soit le continent, le registre *guide* reste dominant. Toutefois, on observe des nuances significatives dans la manière dont ce registre s'articule avec les autres intentions.

Les entreprises européennes et nord-américaines présentent des profils relativement similaires. Elles proposent un usage massif de l'intention *guide*, combiné à une proportion non négligeable de registres *discourage*, et à une très faible utilisation des registres *punish* ou *protect*.

Enfin, on note que dans les trois continents, les intentions *empower* et *punish* restent extrêmement marginales. Cela confirme une tendance globale à éviter les postures extrêmes, qu'il s'agisse de la sanction ou de la responsabilisation, au profit d'un ton intermédiaire centré sur la guidance.

En revanche, des différences intéressantes émergent toutefois. Le profil asiatique se distingue par une proportion plus élevée d'intentions *protect*, relativement à l'Europe et à l'Amérique du Nord. Cette spécificité peut s'interpréter comme la volonté d'adopter un ton plus bienveillant, mettant l'accent sur la sécurité, la stabilité ou la préservation des salariés, plutôt que sur la prévention ou la sanction.

Cette comparaison met en évidence l'influence du contexte culturel et institutionnel sur le choix des intentions. Tandis que les entreprises nord-américaines et européennes semblent privilégier des postures à la fois prescriptives et préventives, les entreprises asiatiques tendent un peu plus leur communication autour de normes protectrices et structurantes. Même si toutefois, il n'y a pas grande différence entre ces trois continents.

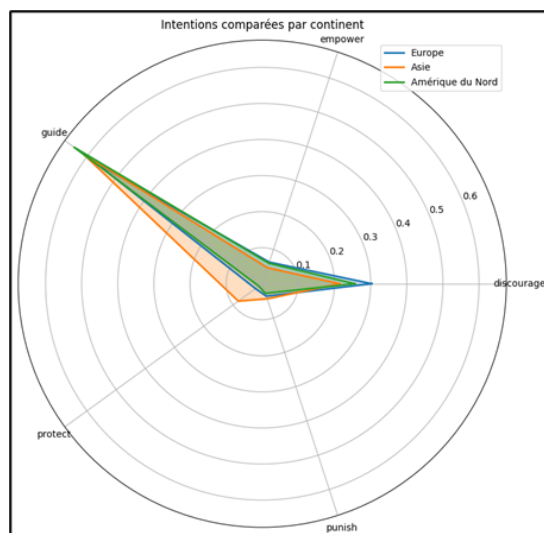


FIGURE 16 – Comparaison des intentions dominantes par continent.

4.4.4 Répartition par rang

Le graphique ci-dessous (Figure 17) définit dans quelle mesure le classement des entreprises dans le *Fortune Global 500* influence les intentions dominantes exprimées dans leurs codes de conduite. Pour cette analyse, les entreprises ont été réparties en cinq groupes équivalents selon leur rang, du groupe Top 100 au Top 500. Chaque barre représente la répartition des intentions dominantes pour les entreprises situées dans chaque tranche de classement (en %).

Dans l'ensemble, l'intention *guide* domine systématiquement, avec des proportions comprises entre 58 % et 63 %, ce qui confirme sa position centrale dans la stratégie des grandes entreprises, indépendamment de leur taille ou de leur statut.

Les autres intentions présentent des variations, mais celles-ci restent minimales. L'intention *discourage* se maintient autour de 24 % à 28 %, sans tendance nette selon le classement. L'intention *protect*, bien que faible, atteint un léger maximum dans le Top 200 (7 %), tandis que *empower* et *punish* ne dépassent jamais les 7 %, et ne présentent aucune progression significative d'un groupe à l'autre.

En résumé, le classement dans le *Fortune Global 500* n'apparaît pas comme un facteur définissant une intention dominante. Les entreprises, qu'elles soient parmi les plus puissantes ou non, semblent adopter des postures similaires, centrées sur la *guidance*, avec des variations marginales sur les autres registres. Ce constat suggère que la communication éthique à travers les codes de conduite obéit à des standards, indépendants du poids économique ou du prestige d'une entreprise.

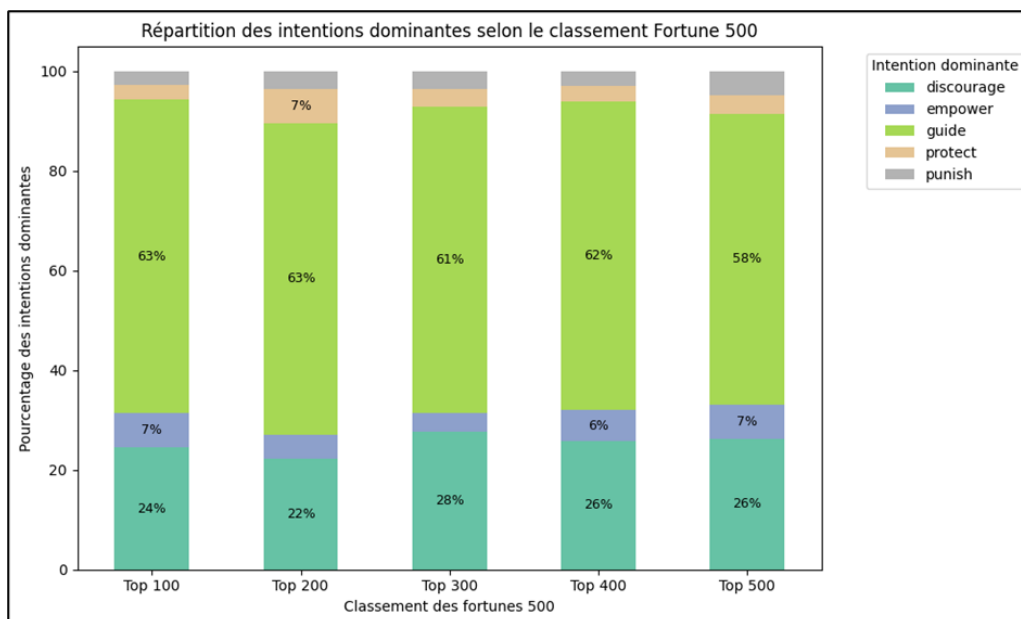


FIGURE 17 – Répartition des intentions selon le classement des fortune 500

4.4.5 Discussion des résultats

L'analyse quantitative des intentions a permis de mettre en évidence plusieurs tendances transversales au sein des codes de conduite des entreprises du Fortune Global 500. Tout d'abord, le registre guide domine largement, quelle que soit la variable considérée (thème, secteur, zone géographique ou rang). Cette domination du ton guidant révèle une volonté partagée d'encadrer les comportements sans recourir à la contrainte.

En complément, des nuances ont pu être observées. Certains secteurs comme la santé ou les services se distinguent par une plus grande présence de registres empower ou protect, traduisant une logique de soutien ou de responsabilisation. D'autres, comme la finance ou la chimie, mobilisent davantage le registre discourage, traduisant un ton préventif, orienté vers la maîtrise du risque. D'un point de vue géographique, les différences entre continents sont réelles mais relativement modérées, et le classement Fortune 500 ne semble pas constituer une variable explicative déterminante du ton adopté permettant de les distinguer.

Il convient toutefois de rappeler que cette étude adopte une approche quantitative. Elle vise à repérer des tendances globales dans la manière dont les entreprises formulent leurs normes et intentions à travers leurs codes de conduite. D'autres niveaux de lecture, plus fins ou plus contextuels, nécessiteraient une analyse qualitative, par exemple à travers une étude de cas, ou une lecture critique des choix lexicaux.

5. Conclusion

Ce mémoire a pour objectif d’analyser les codes de conduite des entreprises du classement Fortune Global 500 à travers les techniques avancées de traitement automatique du langage. En mobilisant des outils de NLP de dernière génération, notamment les modèles basés sur les Transformers tels que BERT et BERTopic, l’étude a permis d’extraire et de comparer, à grande échelle, les thématiques et les tonalités dominantes portées par ces documents.

La démarche suivie a combiné une revue de littérature, une phase de préparation et de structuration du corpus, ainsi qu’une série d’analyses quantitatives centrées sur la détection automatique des thèmes et l’identification des intentions. Les résultats ont mis en lumière une diversité thématique significative entre les entreprises de continents différents. Par ailleurs, l’introduction des intentions (protéger, guider, dissuader, sanctionner, responsabiliser) a offert un autre cadre d’interprétation, permettant de dépasser la seule présence de thèmes pour explorer la posture adoptée par les entreprises.

5.1 Synthèse des résultats

Plus concrètement, en réponse à la première problématique, l’analyse des thèmes a révélé que la structuration thématique des codes de conduite ne s’explique que partiellement par des variables institutionnelles classiques telles que le secteur d’activité ou le rang dans le classement Fortune. Ni l’analyse croisée, ni les tests statistiques ne permettent d’identifier une corrélation forte entre ces facteurs et l’appartenance à un profil discursif donné. En revanche, la dimension géographique se distingue comme un facteur plus différenciant. Les entreprises asiatiques, en particulier, affichent une plus grande hétérogénéité dans la forme de leurs codes (clusters) et une entropie moyenne plus faible, suggérant des dynamiques locales plus contrastées. À l’inverse, les entreprises d’Amérique du Nord, d’Europe et d’Amérique du Sud convergent davantage vers des profils discursifs standardisés et diversifiés.

En résumé, la construction des codes de conduite semble moins déterminée par des critères sectoriels ou de classement que par des logiques transversales et culturelles, potentiellement liées à la gouvernance interne, à l’histoire institutionnelle, ou au degré d’alignement sur des standards internationaux.

En réponse à la deuxième problématique, l'analyse quantitative des intentions a permis de mettre en évidence plusieurs tendances. Le registre guide domine largement dans les codes de conduite, quelle que soit la variable considérée (thème, secteur, région, ou rang), traduisant une volonté commune d'encadrer les comportements sans recourir à la sanction. Toutefois, des nuances apparaissent : certains secteurs, comme la santé ou les services, se caractérisent par une plus grande présence des registres empower et protect, traduisant une logique de soutien ou de valorisation. D'autres, comme la finance ou la chimie, mobilisent plus fréquemment le registre discourge, dans une perspective plus préventive et orientée vers le contrôle des risques. À l'échelle géographique, les différences entre continents restent modérées, tandis que le rang dans le classement Fortune 500 n'a pas montré d'impact significatif sur les postures adoptées.

5.2 Pistes d'amélioration

Plusieurs pistes permettraient de renforcer l'interprétation des résultats et d'enrichir la portée de l'analyse.

Tout d'abord, la dimension géographique gagnerait en précision en passant du niveau continental à celui des pays ou des familles juridiques (common law, droit civil, etc.). Cela permettrait d'identifier des logiques normatives plus fines, potentiellement liées à des traditions institutionnelles différenciées.

Ensuite, au-delà du secteur et du rang, d'autres variables structurelles pourraient être intégrées, telles que le statut public/privé, la présence internationale, ou encore des indicateurs de contexte comme l'indice de perception de la corruption (CPI). Ces dimensions sont susceptibles d'influencer la formulation du code de conduite de manière plus significative que les critères purement économiques.

Sur le plan méthodologique, les regroupements thématiques ont été construits à l'aide de KMeans, une méthode efficace mais sensible au choix du nombre de clusters. L'utilisation de techniques non paramétriques comme HDBSCAN ou le clustering hiérarchique permettrait de valider ou d'enrichir la typologie actuelle sans imposer de structure a priori.

Enfin, plusieurs choses peuvent être envisagées pour approfondir l'analyse des intentions discursives présentes dans les codes de conduite. Une approche qualitative sur un sous-corpus permettrait d'interpréter plus finement les formulations, les nuances de ton, au-delà de la classi-

fication automatique par intentions. Une annotation manuelle experte pourrait également venir valider ou ajuster les catégories d'intention retenues (protect, guide, discouragement, punish, empower), et évaluer la précision effective du modèle zero-shot utilisé.

Par ailleurs, l'attribution d'une seule intention dominante par paragraphe, telle qu'opérée ici, limite l'analyse dans les cas où plusieurs intentions coexistent. Un même passage peut chercher à responsabiliser tout en dissuadant ou en protégeant. Un modèle multi-label, ou une analyse des co-occurrences d'intentions, pourrait enrichir l'étude pour mieux refléter la complexité des textes produits par les entreprises.

En définitive, cette recherche confirme la pertinence d'une approche quantitative fondée sur le NLP pour analyser à grande échelle les logiques portées par les codes de conduite d'entreprise. Ce travail ouvre des perspectives pour de futurs travaux mêlant data science et sciences sociales.

Références

- [1] Abercrombie, G., & Batista-Navarro, R. (2020). Sentiment and position-taking analysis of parliamentary debates : A systematic literature review. *Journal of Computational Social Science*, 3(1), 245–270. <https://link.springer.com/article/10.1007/s42001-019-00060-w>
- [2] Angelov, D. (2020). Top2Vec : Distributed representations of topics. *arXiv preprint*. <https://arxiv.org/abs/2008.09470>
- [3] Bardin, L. (2013). L’analyse de contenu. Presses universitaires de France.
- [4] Blei, D. M., Ng, A. Y., & Jordan, M. I. (2003). Latent Dirichlet allocation. *Journal of Machine Learning Research*, 3, 993–1022. <https://www.jmlr.org/papers/volume3/blei03a/blei03a.pdf>
- [5] Braun, V., & Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), 77–101. <https://www.tandfonline.com/doi/abs/10.1191/1478088706qp063oa>
- [6] Braun, V., Clarke, V., & Hayfield, N. (2017). Using thematic analysis in qualitative research. In P. Liamputtong (Ed.), *Handbook of research methods in health social sciences* (pp. 1–18). Springer. https://doi.org/10.1007/978-981-10-2779-6_103-1
- [7] Devlin, J., Chang, M.-W., Lee, K., & Toutanova, K. (2019). BERT : Pre-training of deep bidirectional transformers for language understanding. *Proceedings of NAACL-HLT 2019*. <https://arxiv.org/abs/1810.04805>
- [8] Grootendorst, M. (2022). BERTopic : Neural topic modeling with a class-based TF-IDF procedure. *arXiv preprint*. <https://arxiv.org/abs/2203.05794>
- [9] He, P., Liu, X., Gao, J., & Chen, W. (2020). DeBERTa : Decoding-enhanced BERT with disentangled attention. *arXiv preprint*. <https://arxiv.org/abs/2006.03654>
- [10] Howard, J., & Ruder, S. (2018). Universal language model fine-tuning for text classification. *Proceedings of ACL 2018*. <https://arxiv.org/abs/1801.06146>
- [11] Kaptein, M. (2004). Business codes of multinational firms : What do they say? *Journal of Business Ethics*, 50(1), 13–31. <https://link.springer.com/article/10.1023/B:BUSI.0000021051.53452.dd>
- [12] Lewis, M., Liu, Y., Goyal, N., Ghazvininejad, M., Mohamed, A., Levy, O., Stoyanov, V., & Zettlemoyer, L. (2020). BART : Denoising sequence-to-sequence pre-training for

- natural language generation, translation, and comprehension. *Proceedings of ACL 2020*. <https://aclanthology.org/2020.acl-main.703/>
- [13] Manning, C. D., Raghavan, P., & Schütze, H. (2008). Introduction to information retrieval. Cambridge University Press. <https://nlp.stanford.edu/IR-book/>
- [14] Mercier, S. (2002). Du discours à la réalité : Théorie et pratique du management éthique. *Cadres-CFDT*, 401–402(novembre), 39–44. [PDF non disponible en ligne]
- [15] Ni, J., Hernandez-Abrego, G., Constant, N., Ma, J., Hall, K., Cer, D., & Yang, Y. (2022). Sentence-T5 : Scalable sentence encoders from pre-trained text-to-text models. *Findings of ACL 2022*, 1864–1874. <https://aclanthology.org/2022.findings-acl.147/>
- [16] Pang, B., & Lee, L. (2008). Opinion mining and sentiment analysis. *Foundations and Trends in Information Retrieval*, 2(1–2), 1–135. <https://www.nowpublishers.com/article/Details/INR-011>
- [17] Racine, J.-B. (1996). Les codes de conduite à l’épreuve du droit international économique. In M. Doucin (Dir.), *La responsabilité sociale des entreprises : de la déclaration à la mise en œuvre* (pp. 67–78). Éditions La Découverte.
- [18] Raffel, C., Shazeer, N., Roberts, A., Lee, K., Narang, S., Matena, M., Zhou, Y., Li, W., & Liu, P. (2020). Exploring the limits of transfer learning with a unified text-to-text transformer (T5). *Journal of Machine Learning Research*, 21(140), 1–67. <https://www.jmlr.org/papers/v21/20-074.html>
- [19] Reimers, N., & Gurevych, I. (2019). Sentence-BERT : Sentence embeddings using Siamese BERT-networks. *Proceedings of EMNLP-IJCNLP 2019 (Hong Kong)*. <https://aclanthology.org/D19-1410/>
- [20] Resources for the Future. (2005). Corporate social responsibility and codes of conduct : The importance of international standards. <https://media.rff.org/documents/RFF-DP-05-09.pdf>
- [21] Saïd, Y. (2013). Responsabilité sociale de l’entreprise, code de conduite et charte d’éthique : aspects juridiques. [Mémoire de Master, Université Panthéon-Assas Paris II].
- [22] Schick, T., & Schütze, H. (2021). It’s not just size that matters : Small language models are also few-shot learners. *Proceedings of NAACL-HLT 2021*. <https://aclanthology.org/2021.naacl-main.164/>

- [23] Vallée, G., Murray, G., Coutu, M., Rocher, G., & Giles, A. (2001). Les codes de conduite des entreprises multinationales canadiennes : aux confins de la régulation privée et des politiques publiques du travail. *Commission du droit du Canada*. https://publications.gc.ca/collections/collection_2008/lcc-cdc/JL2-54-2001F.pdf
- [24] Veilleux, S., & Bachand, R. (2011). L'autorégulation des entreprises transnationales en matière de droits humains : Le cas des codes de conduite. *Études internationales*, 42(1), 9–37. <https://doi.org/10.7202/1006366ar>

