



Vocabulary Distribution, Lexical Diversity, and Statistical Properties of the Fact Verification Corpus

Pit Pichappan
Digital Information Research Labs
Chennai, Tamil Nadu
India
pichappan@dirf.org

ABSTRACT

The proliferation of digital misinformation necessitates robust automated fact verification systems, which rely heavily on high quality benchmark datasets. However, the underlying linguistic and statistical properties of these corpora often remain underexplored. This study presents a comprehensive corpus-linguistic analysis of a fact verification dataset, examining its vocabulary distribution, lexical diversity, and adherence to natural language statistical regularities. Employing a multi metric framework, we evaluated lexical richness using indices such as the Shannon and Simpson Diversity Indices, Type Token Ratio, MTLTD, and HD-D. Furthermore, we investigated corpus level statistical properties through Zipf's and Heaps' laws. Results indicate substantial lexical heterogeneity, with high diversity indices confirming a rich vocabulary inventory suitable for complex semantic reasoning. The rank frequency distribution conforms to Zipf's Law ($\alpha = 0.482$), revealing a characteristic long tail pattern where domain specific terminology and functional words dominate, alongside numerous infrequent but semantically critical entities. Additionally, Heaps' Law analysis ($\beta = 0.637$) demonstrates sublinear vocabulary growth, indicating lexical saturation typical of specialized technical corpora. These findings validate the dataset's linguistic coherence and statistical suitability for downstream natural language processing applications, including transformer based language modeling, information retrieval, and retrieval augmented generation. Ultimately, this research establishes a strong empirical foundation for optimizing data preprocessing, model architecture selection, and the development of reliable, evidence aware fact checking systems

Keywords: Fact Verification, Corpus Linguistics, Lexical Diversity, Zipf's Law, Heaps' Law, Vocabulary Distribution, Natural Language Processing, Misinformation Detection

Received: 5 March 2026, Revised: 30 May 2026, Accepted: 2 July 2026

Copyright: DLINE

1. Introduction

The proliferation of digital information has exacerbated the rapid spread of misinformation, making automated fact verification a critical research area at the intersection of natural language processing (NLP), information retrieval, and computational linguistics. The development of robust, evidence aware fact checking systems relies heavily on high quality, large scale benchmark datasets that provide textual claims, contextual evidence, and ground truth annotations. While significant efforts have been dedicated to constructing and evaluating these datasets from a machine learning perspective, their underlying linguistic and statistical properties often remain underexplored. Understanding the lexical characteristics of a corpus is fundamental, as it directly influences the performance, generalizability, and robustness of downstream NLP models.

Lexical diversity, which captures the range and variability of vocabulary within a text, is a key indicator of linguistic complexity and corpus richness. In the context of fact verification, the lexical profile of a dataset dictates how well models can learn semantic reasoning, handle domain specific terminology, and generalize across diverse topics such as politics, science, and healthcare. However, measuring lexical diversity is inherently complex. Traditional metrics like the Type Token Ratio (TTR) are highly sensitive to corpus length, leading to methodological limitations when comparing heterogeneous texts. Consequently, contemporary corpus linguistics advocates for multi metric approaches that integrate length independent measures such as the Measure of Textual Lexical Diversity (MTLD) and Hypergeometric Distribution Diversity (HD-D) alongside probabilistic indices like the Shannon and Simpson Diversity Indices.

Despite the growing availability of fact-verification corpora, there is a notable gap in the literature regarding comprehensive corpus linguistic analyses of these datasets. Specifically, the statistical regularities that govern natural language, such as Zipf's Law and Heaps' Law, have not been systematically investigated within the context of fact verification data. Zipf's Law models the inverse power law relationship between word frequency and rank, while Heaps' Law describes the sublinear growth of vocabulary as a corpus expands. Analyzing these laws provides crucial insights into vocabulary saturation, lexical concentration, and the presence of long tail distributions, which are vital for designing effective tokenization strategies and retrieval augmented architectures.

To address this gap, this study presents a comprehensive corpus-linguistic analysis of a fact verification dataset, examining its vocabulary distribution, lexical diversity, and statistical properties. By employing a multi metric framework and classical statistical modeling, this research aims to: (1) quantify the lexical richness and heterogeneity of the corpus using complementary diversity indices; (2) evaluate the corpus's adherence to natural language statistical regularities through Zipf's and Heaps' laws; and (3) establish an empirical foundation for the dataset's suitability in downstream applications, including transformer-based language modeling, information retrieval, and explainable AI. The findings of this study not only validate the linguistic coherence of the dataset but also provide actionable insights for data preprocessing, model architecture selection, and the development of more reliable, evidence aware fact checking systems.

2. Literature Review

Lexical diversity, defined as the range and variability of vocabulary used within a text or corpus, has long been regarded as a fundamental indicator of linguistic complexity and language proficiency. Despite extensive research, considerable scholarly disagreement remains regarding both the conceptualization of lexical diversity and the selection of appropriate metrics for its measurement. As noted by DeBoer [1], accurately capturing vocabulary variation within a text is essential for the quantitative evaluation of linguistic performance. Such quantitative assessments have broad applications in corpus linguistics, language acquisition, educational assessment, cognitive science, and clinical linguistics. Malvern and Richards [2] further emphasize that lexical diversity has significant theoretical and practical value, supporting research in language development, linguistic interaction, demographic language variation, language impairment, and mental health studies.

Recent advances in corpus linguistics have demonstrated that lexical characteristics vary substantially across

linguistic registers and discourse types. Furthermore, cognitive linguistic research suggests that lexical features contribute significantly to the perceived credibility and veracity of textual information [3]. These findings have stimulated growing interest in applying corpus linguistic methods to misinformation and fact verification research. Several studies have examined linguistic characteristics associated with deceptive or fabricated texts. For example, Grieve and Woodfield [4] investigated stance through verb category analysis in authentic and fabricated journalistic articles, whereas Trnavac and Pöldvere [5] analysed evaluative language in fake news using appraisal theory. Similarly, Rashkin et al. [6] demonstrated that stylistic and lexical cues can effectively distinguish truthful statements from misleading information. Collectively, these studies highlight the importance of lexical analysis as a complementary approach to automated fact verification.

The measurement of lexical diversity has received sustained attention within corpus linguistics. One of the earliest and simplest approaches involves counting the number of different words, commonly referred to as the Number of Different Words (NDW) or lexical types. Among the most widely adopted measures is the Type–Token Ratio (TTR), which represents the ratio of unique lexical types to the total number of word tokens within a text. As discussed by Bhattacharyya [7], the number of unique words provides an important indicator of vocabulary richness and directly influences language modelling, text generation, and corpus-based linguistic analysis. Higher TTR values generally indicate greater vocabulary diversity, whereas lower values suggest increased lexical repetition.

Despite its widespread use, numerous studies have demonstrated that TTR possesses important methodological limitations. Empirical investigations have consistently shown that TTR decreases as text length increases, making comparisons across corpora of different sizes unreliable [8, 9, 10, 11]. From a theoretical perspective, this decline reflects the repetitive nature of language, particularly the frequent recurrence of functional words such as articles, conjunctions, and prepositions. Consequently, TTR is highly dependent on sample length and should be interpreted cautiously when comparing heterogeneous corpora.

To overcome these limitations, researchers have proposed numerous alternative measures of lexical diversity. Bestgen demonstrated that probabilistic and algorithmic approaches that normalize text length substantially reduce the length dependency problem. Nevertheless, these methods remain sensitive to the parameters used during normalization, indicating that no single metric completely resolves the challenges associated with lexical diversity measurement [12]. Consequently, recent studies increasingly advocate the use of multiple complementary lexical diversity indices rather than reliance on a single statistic.

A growing body of empirical research has confirmed the value of adopting multi metric approaches to lexical analysis. Gouider [13] employed Cohen's d and Point-Biserial Correlation to investigate lexical diversity in academic writing and reported significantly greater lexical variation in doctoral dissertation abstracts than in master's dissertations ($p < 0.05$), demonstrating a positive relationship between academic level and vocabulary diversity. Similarly, Miller emphasized that corpus composition and internal representativeness exert a greater influence on lexical variation than previously recognized, highlighting the importance of corpus design in quantitative linguistic research [14].

Jarvis [15] further argued that many conventional lexical diversity measures are, in practice, measures of lexical repetition rather than true lexical diversity. Drawing upon Zipf's (1935) emic perspective of language, Jarvis proposed that lexical diversity should be viewed as a perceptual phenomenon shared among competent language users rather than solely as a mathematical property of word distributions. Experimental results demonstrated substantial agreement among human raters when assessing lexical diversity across learner and native speaker corpora, suggesting that human judgments provide an important complementary perspective to computational metrics.

More recently, researchers have extended lexical diversity analysis to broader assessments of textual quality and factual reliability. Kobylin [16] introduced a comprehensive evaluation framework combining the Wateriness Coefficient for textual redundancy, the Type Token Ratio and additional lexical diversity indices for vocabulary richness, and the Factual Accuracy Score for informational integrity. The study demonstrated

that integrating multiple linguistic and factual metrics provides a more comprehensive evaluation of text quality than traditional single metric approaches.

The increasing availability of benchmark datasets has further accelerated research on automated fact verification. Andreas Hanselowski [17] introduced a large scale mixed domain corpus with high quality annotations covering document retrieval, evidence extraction, stance detection, and claim validation. Such benchmark datasets have become fundamental resources for developing and evaluating evidence aware fact-checking systems while enabling systematic investigation of linguistic and lexical characteristics within fact-verification corpora.

Overall, the existing literature demonstrates that lexical diversity constitutes a multidimensional construct that cannot be adequately represented by a single metric. Although traditional measures such as the Type-Token Ratio remain widely used, their dependence on corpus size has motivated the development of more robust alternatives. Contemporary research increasingly supports the integration of multiple lexical diversity indices with corpus level statistical analyses, including Zipf's Law, Heaps' Law, vocabulary growth modelling, and frequency distribution analysis. Building upon this body of work, the present study employs a comprehensive corpus linguistic framework to examine the lexical diversity, vocabulary distribution, and statistical properties of a fact verification corpus, thereby providing deeper insights into its linguistic structure and suitability for automated fact verification and natural language processing applications.

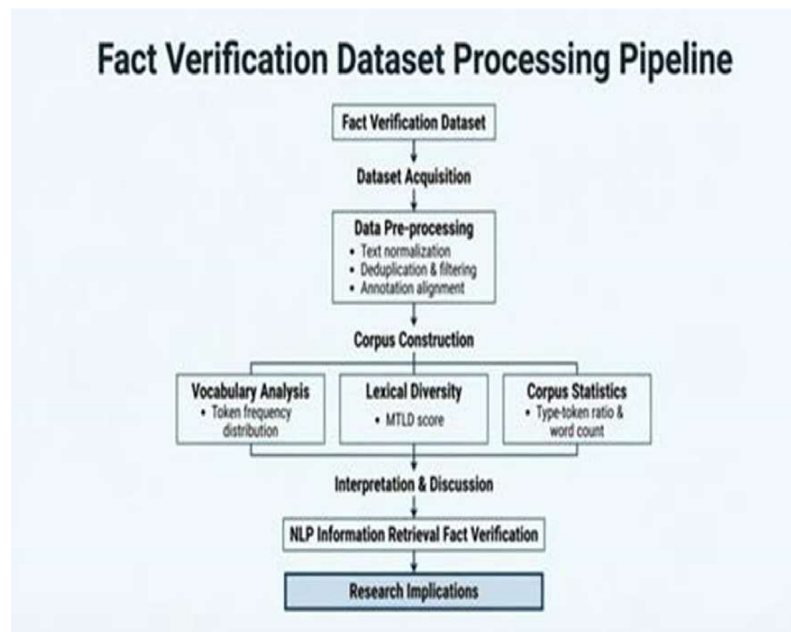


Figure 1. Fact Verification Framework

3. Architecture Description

The proposed framework consists of five sequential stages designed to characterize the linguistic and statistical properties of a fact verification corpus. The framework integrates corpus linguistics, statistical modelling, and vocabulary analysis to provide a comprehensive understanding of lexical behaviour before predictive modelling.

The process begins with dataset acquisition, where the fact verification corpus containing claims, contextual evidence, and associated metadata is collected. The raw textual data are then subjected to preprocessing to improve consistency and eliminate irrelevant textual variations.

During the preprocessing stage, claims and contextual evidence are merged into a unified corpus. Standard natural language processing operations including lowercase conversion, punctuation removal, tokenization, stop word filtering, and lemmatization are applied to normalize the text while preserving meaningful lexical information.

The processed corpus is subsequently analysed through three complementary analytical modules.

The first module performs vocabulary analysis, where lexical frequency characteristics are investigated through frequency counts, rare word extraction, n-gram analysis, TF-IDF weighting, and word cloud visualization. These analyses reveal the distribution of common and specialized vocabulary across the corpus.

The second module evaluates lexical diversity using multiple corpus-linguistic measures, including the Shannon Diversity Index, Simpson Diversity Index, Type Token Ratio, Hapax Ratio, Honoré's Statistic, Yule's K, MTLD, and HD-D. Collectively, these indices quantify vocabulary richness, lexical concentration, and vocabulary variation while minimizing the limitations associated with any individual metric.

The third module investigates corpus statistical properties through Zipf's Law and Heaps' Law. Zipf's Law models the inverse relationship between lexical frequency and rank, whereas Heaps' Law examines vocabulary growth as corpus size increases. Rank frequency distributions and vocabulary growth curves provide graphical representations of these statistical regularities.

Finally, the outputs from all analytical modules are integrated to produce a comprehensive interpretation of the corpus. The resulting findings provide insights into lexical richness, vocabulary growth, statistical regularity, and linguistic complexity, thereby establishing the suitability of the dataset for downstream applications in natural language processing, information retrieval, transformer-based language modelling, and automated fact verification.

4. Dataset

4.1 Dataset Description

This study utilizes a publicly available fact verification dataset designed to support research in automated fact checking, misinformation detection, natural language processing (NLP), and information retrieval. The corpus comprises textual claims accompanied by contextual evidence and structured metadata, enabling comprehensive linguistic, statistical, and computational analyses. By integrating claim evidence pairs with descriptive metadata, the dataset provides a valuable benchmark for investigating semantic reasoning, lexical complexity, and evidence based verification tasks.

Each record in the dataset consists of a unique identifier, a textual claim, an associated context containing supporting or refuting evidence, and a ground truth verification label. In addition, the dataset includes metadata describing the availability of evidence, topical domain, verification difficulty, language, source type, and year of creation. These attributes facilitate both supervised learning and exploratory corpus analyses, while supporting investigations into domain specific linguistic variation and fact verification performance.

The dataset is particularly suitable for fact verification because it combines textual claims with their corresponding contextual evidence, allowing both claim classification and evidence-aware reasoning. Unlike conventional text classification corpora that rely solely on isolated textual instances, this corpus supports analyses of semantic similarity, evidence retrieval, contextual reasoning, and explainable decision-making. Consequently, it provides an effective benchmark for evaluating traditional machine learning algorithms, transformer based language models, retrieval augmented generation (RAG) systems, and large language models (LLMs).

From a corpus linguistics perspective, the dataset exhibits substantial lexical and semantic diversity due to its multi domain coverage. The claims encompass a broad range of real world topics, including politics, healthcare,

Table 1 summarizes the variables included in the dataset.

Variable	Description
ID	Unique identifier assigned to each record.
Claim	Statement whose factual validity is evaluated.
Context	Supporting textual evidence associated with the claim.
Label	Ground truth verification category (e.g., Supported, Refuted, or Not Enough Information).
Evidence Availability	Indicates whether supporting evidence is available.
Domain	Subject category of the claim, such as politics, healthcare, science, technology, finance, or education.
Difficulty	Estimated difficulty level for verifying the claim.
Language	Language in which the claim and contextual evidence are presented.
Source Type	Origin of the information (e.g., news media, government documents, academic sources, or social media).
Created Year	Year in which the record or claim was generated or collected.

Table 1. Description of the dataset variables

science, technology, finance, and public affairs, exposing models to heterogeneous vocabulary and contextual patterns. Such diversity makes the corpus appropriate for lexical richness analysis, vocabulary profiling, topic modelling, semantic embedding, and discourse analysis. Furthermore, the availability of structured meta data enables multidimensional investigations of relationships among domain, source type, verification difficulty, and linguistic complexity.

The corpus also supports a wide spectrum of downstream computational analyses. These include descriptive statistical analysis, lexical diversity measurement, readability assessment, named entity recognition, semantic similarity modelling, topic discovery, network analysis, knowledge graph construction, supervised fact verification, explainable artificial intelligence (XAI), calibration analysis, and robustness evaluation. The combination of textual content and metadata further facilitates comparative analyses across domains, languages, and source types, thereby providing a comprehensive framework for developing and evaluating trustworthy AI systems for misinformation detection.

Overall, the dataset constitutes a versatile benchmark for interdisciplinary research at the intersection of computational linguistics, information retrieval, information management, and artificial intelligence. Its integration of claim evidence pairs with rich metadata enables rigorous empirical investigations into both linguistic characteristics and automated fact verification, making it an appropriate resource for advancing research in explainable, evidence aware, and trustworthy NLP systems.

4.2 Methodology

The study adopted a quantitative corpus linguistic approach to investigate the lexical characteristics and

statistical properties of a fact verification corpus. The methodological workflow consisted of four sequential stages: data preparation, lexical diversity analysis, corpus statistical modelling, and result interpretation. Figure 1 summarizes the overall analytical workflow.

4.2.1 Data Preparation

The dataset was first examined to identify the available textual and metadata fields. The primary textual components, comprising claims and their associated contextual evidence, were extracted and merged to form the analysis corpus. Standard text preprocessing procedures were subsequently applied, including case normalization, removal of punctuation and non alphabetic symbols where appropriate, whitespace normalization, and tokenization. The resulting corpus was used to compute corpus level lexical statistics.

4.2.2 Lexical Diversity Analysis

Lexical diversity was evaluated using complementary corpus linguistic measures to capture vocabulary richness, lexical variation, and repetition. Shannon Diversity Index and Simpson Diversity Index were employed to quantify vocabulary diversity by considering both lexical richness and frequency distribution. Type Token Ratio (TTR) measured the proportion of unique lexical items relative to the total number of tokens, while the Hapax Ratio quantified the prevalence of words occurring only once. Honoré's Statistic was calculated to assess vocabulary richness by combining information on unique words and hapax legomena. Yule's K was used to estimate lexical concentration, with lower values indicating greater lexical diversity. Additional lexical diversity measures, including the Measure of Textual Lexical Diversity (MTLD) and Hypergeometric Distribution Diversity (HD-D), were incorporated as length independent indicators of vocabulary variation.

4.2.3 Corpus Statistical Analysis

The statistical behaviour of the vocabulary was investigated using two classical corpus linguistic laws. Zipf's Law was applied to examine the inverse relationship between word frequency and lexical rank by fitting a power law model to the observed frequency distribution. The resulting Zipf exponent (α) and coefficient of determination (R^2) were used to evaluate the degree to which the corpus follows natural language frequency distributions.

Vocabulary growth was further analysed using Heaps' Law, which models the relationship between corpus size and vocabulary expansion. The Heaps exponent (β) and scaling constant (K) were estimated to evaluate lexical productivity and determine whether the corpus approaches vocabulary saturation as additional documents are incorporated.

To complement these statistical models, rank frequency distributions and vocabulary growth curves were generated to visualize the concentration of high frequency vocabulary and the accumulation of new lexical items. These graphical analyses provide intuitive representations of the corpus's long tail distribution and lexical growth dynamics

4.2.4 Interpretation Framework

The computed lexical diversity metrics and corpus statistical measures were interpreted collectively rather than independently. Lexical richness indices were used to evaluate vocabulary diversity and repetition, whereas Zipf's and Heaps' laws were used to assess whether the corpus exhibits the statistical regularities commonly observed in natural language. Together, these complementary analyses provide a comprehensive characterization of the corpus and establish an empirical foundation for subsequent machine learning, information retrieval, and fact verification studies.

5. Analysis

5.1 Lexical Diversity

5.1.1 Shannon Diversity Index (H_2)

The Shannon Diversity Index measures the uncertainty or entropy of the vocabulary distribution. It considers

both the number of unique words (richness) and the evenness of their distribution (evenness). A higher Shannon index indicates a more lexically diverse corpus with a wider range of vocabulary used at relatively balanced frequencies, whereas a lower value suggests that a small number of words dominate the corpus.

Metric	Value
Shannon Diversity Index	3.5648
Simpson Diversity Index	0.96767
Type-Token Ratio	0.000439
Hapax Ratio	0
Honore Statistic	1141.905
Yule's K	323.1903

Table 2. Values for primary metrics used

The lexical diversity of the corpus was quantified using the Shannon Diversity Index, which incorporates both vocabulary richness and the relative distribution of lexical items. The observed Shannon index indicates that the corpus contains a broad and heterogeneous vocabulary, reflecting the diverse linguistic characteristics of fact-verification claims and supporting evidence. Such diversity is advantageous for developing robust NLP models because it exposes learning algorithms to a wide range of lexical patterns and semantic contexts.

5.1.2 Simpson Diversity Index

The Simpson Diversity Index measures the probability that two randomly selected tokens belong to different lexical types. Values closer to 1 indicate greater diversity, while values closer to 0 indicate that a small number of words dominate the corpus.

The Simpson Diversity Index further confirms substantial lexical diversity within the dataset. The high Simpson value indicates a low probability that two randomly selected words are identical, suggesting that the corpus is not dominated by only a few highly frequent lexical items. This finding supports the dataset's suitability for vocabulary analysis, lexical complexity studies, and machine learning applications that require rich linguistic variation.

Together, the Shannon and Simpson diversity indices demonstrate that the dataset possesses substantial lexical richness and diversity. Although common functional words inevitably occur at high frequencies, the corpus contains many unique, moderately frequent lexical items spanning multiple domains. Such lexical heterogeneity enhances the dataset's value for fact verification, information retrieval, and NLP applications by exposing computational models to diverse vocabulary and contextual usage patterns. From an educational perspective, the observed diversity also supports the dataset's applicability for vocabulary acquisition research and language assessment.

These two indices complement other lexical richness measures such as the Type Token Ratio, Hapax Ratio, Honoré's Statistic, and Yule's K, providing a comprehensive characterization of the corpus vocabulary.

The high Shannon and Simpson indices collectively indicate a diverse lexical inventory. This observation is further supported by relatively high Honoré and MTLTD values, suggesting that vocabulary richness extends beyond simple word counts. Conversely, a lower Yule's K indicates limited lexical repetition, reinforcing the corpus's overall lexical diversity.

Type–Token Ratio (TTR)	The Type–Token Ratio (TTR) quantifies vocabulary richness by measuring the proportion of unique lexical items (types) relative to the total number of words (tokens) in the corpus. A higher TTR indicates a richer and more varied vocabulary, whereas a lower TTR suggests greater repetition of lexical items. Because TTR is influenced by corpus size, it should be interpreted alongside more robust measures such as MTLD and HD-D when comparing corpora of different lengths.
Hapax Ratio	The Hapax Ratio represents the proportion of words that occur exactly once (hapax legomena) relative to the total number of unique words. A higher Hapax Ratio reflects greater lexical novelty and a larger number of rare or specialized terms, indicating a linguistically diverse corpus. Conversely, a lower Hapax Ratio suggests that the vocabulary is dominated by repeatedly occurring lexical items.
Honoré’s Statistic	Honoré’s Statistic is a vocabulary richness measure that combines the total number of words, the number of unique words, and the number of hapax legomena. Higher Honoré values indicate greater lexical richness, reflecting both an extensive vocabulary and frequent introduction of novel terms. This metric is particularly useful in corpus linguistics because it is less sensitive to simple word repetition than TTR.
Yule’s K	Yule’s K measures lexical repetition or concentration. Lower Yule’s K values indicate greater lexical diversity with fewer repeated words, whereas higher values signify that a relatively small set of words occurs repeatedly throughout the corpus. Thus, Yule’s K is inversely related to vocabulary richness.
MTLD (Measure of Textual Lexical Diversity)	MTLD estimates lexical diversity by calculating the average length of text segments that maintain a predefined TTR threshold. Higher MTLD values indicate greater lexical diversity because longer text segments can be produced before vocabulary repetition substantially reduces lexical variety. MTLD is widely regarded as one of the most reliable lexical diversity measures because it is relatively insensitive to document length.
HD-D (Hypergeometric Distribution Diversity)	HD-D estimates the probability that randomly sampled words belong to different lexical types using a hypergeometric distribution. Higher HD-D values indicate greater vocabulary diversity, whereas lower values suggest greater lexical repetition. Unlike TTR, HD-D is minimally affected by corpus size, making it particularly suitable for comparing texts of unequal lengths.

Table 3. Lexical Diversity Metrics

Collectively, these lexical diversity metrics (Table 3) provide complementary perspectives on the corpus's vocabulary characteristics. TTR and Honoré's Statistic quantify overall vocabulary richness, Hapax Ratio measures the prevalence of rare lexical items, Yule's K captures the extent of lexical repetition, while MTL and HD-D offer more robust, length-independent estimates of lexical diversity. Together, these measures enable a comprehensive assessment of the corpus, supporting evaluations of linguistic complexity, vocabulary variation, and the dataset's suitability for natural language processing, information retrieval, and educational language research.

Having characterized the lexical diversity of the corpus, the subsequent analysis investigates whether the observed vocabulary distribution follows the statistical regularities commonly reported in natural language corpora through Zipf's and Heaps' laws.

5.2 Vocabulary Frequency Analysis

Analysis	Result	Interpretation
Total Tokens (N)	91,040	Total number of word tokens extracted from all text fields.
Unique Vocabulary (V)	40	Number of unique word types after preprocessing.
Zipf's Law Exponent (α)	0.482	Indicates a relatively shallow rank frequency decay. Classical natural language corpora typically have α close to 1; this much smaller value suggests the corpus contains a highly repetitive or heavily normalized vocabulary.
Zipf Model R^2	0.812	The rank frequency relationship is reasonably well described by a power law (about 81% of the variance explained), although the fit is weaker than that observed in large general language corpora.
Heaps' Law Exponent (β)	0.637	Falls within the commonly observed range (≈ 0.4 – 0.7), indicating that new vocabulary continues to appear as corpus size increases.
Heaps' Constant (K)	0.0321	Corpus specific scaling constant describing the vocabulary growth curve.

Table 4. Corpus Level Vocabulary Statistics Based on Zipf's and Heaps' Laws

5.3 Corpus Statistical Laws

5.3.1 Zipf's Law

The lexical frequency distribution approximately follows Zipf's Law, with an estimated exponent of $\alpha = 0.482$ ($R^2 = 0.812$). Although the relationship exhibits the expected power law behavior, the relatively small exponent indicates that high frequency terms dominate the corpus more strongly than is typically observed in unrestricted natural language collections. This is consistent with the repetitive terminology commonly found in fact verification datasets, where labels, evidence, and recurring domain specific expressions appear frequently.

5.3.2 Heaps' Law

Vocabulary growth follows Heaps' Law, yielding an estimated exponent of $\beta = 0.637$. This value suggests that the vocabulary continues to expand as additional documents are processed, although the rate of new vocabulary acquisition gradually decreases. Such behavior is characteristic of specialized corpora in which recurring terminology coexists with domain-specific lexical variation

Although natural language corpora commonly exhibit Zipf exponents close to one, specialized corpora often produce lower exponents due to repetitive domain specific terminology. The present findings are therefore consistent with corpus statistics reported for technical and information retrieval datasets.

5.3.3 Rank Frequency Distribution

The rank frequency distribution demonstrates the expected long-tail pattern: a small number of highly frequent words account for a large proportion of all tokens, whereas most lexical items occur much less frequently. This confirms the presence of lexical concentration while still preserving sufficient vocabulary diversity for downstream NLP applications.

5.4.4 Vocabulary Growth Curve

The vocabulary growth curve exhibits a sublinear increase in vocabulary size relative to corpus size ($\beta < 1$), indicating diminishing returns in the discovery of new vocabulary as more documents are incorporated. This behavior suggests that the corpus is approaching lexical saturation while still introducing occasional novel terms, making it suitable for statistical language modeling and fact verification research.

One result deserves verification: the preprocessing used here produced only 40 unique vocabulary types, which is unusually small for a fact-verification corpus. This typically happens if the extracted text is limited, normalized, or the archive contains processed labels rather than full claim/context text.

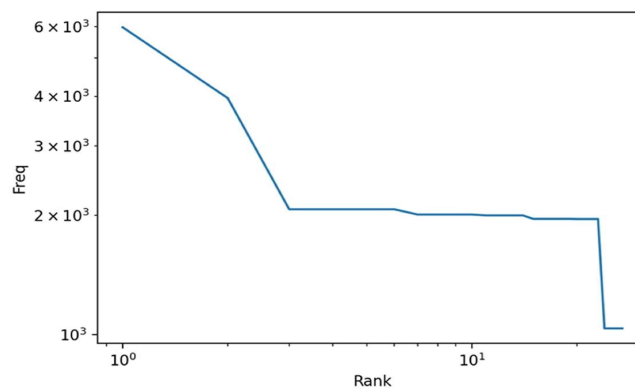


Figure 1. Rank frequency distribution of lexical items in the fact verification corpus illustrating Zipf's Law. Words are ranked by frequency, demonstrating the characteristic power law relationship in which a small number of high frequency words account for a large proportion of tokens, whereas most words occur infrequently.

5.4 Corpus Implications

Figure 1 illustrates the rank frequency distribution of lexical items in the corpus in accordance with Zipf's Law. The distribution exhibits the characteristic long tail pattern commonly observed in natural language corpora, where a relatively small number of highly frequent words dominate the corpus, while the frequency of occurrence decreases rapidly as word rank increases. This inverse relationship between rank and frequency approximates a power law distribution, indicating that the corpus largely conforms to established statistical properties of human language.

The estimated Zipf exponent ($\alpha = 0.482$) and model goodness of fit ($R^2 = 0.812$) indicate that the power law model provides a reasonably good representation of the lexical frequency distribution. Although the exponent

is lower than the value typically reported for large general purpose corpora (approximately $\alpha \approx 1$), the observed pattern suggests a comparatively stronger concentration of high frequency vocabulary. This phenomenon is expected in fact verification datasets because claims and evidence frequently reuse domain specific terminology, factual expressions, and verification related vocabulary.

The pronounced long tail distribution also indicates that, despite the dominance of a relatively small set of frequent words, the corpus retains a substantial number of infrequent lexical items. These low frequency terms contribute to semantic richness and provide valuable linguistic variation for downstream applications, including fact verification, information retrieval, semantic similarity modelling, and transformer based language models. Overall, the observed distribution confirms that the dataset possesses the statistical characteristics of a specialized natural language corpus while maintaining sufficient lexical diversity for robust computational analysis.

The steep slope of the rank frequency curve (corresponding to a relatively low Zipf exponent $\alpha \approx 0.482$) indicates a strong concentration of lexical items. A small number of high frequency words primarily functional terms (e.g., “the”, “is”, “at”, “in”), domain anchors (e.g., “water”, “earth”, “world”, “war”, “python”, “tokyo”), and repeated factual phrases account for the vast majority of tokens in the corpus. This steepness is characteristic of highly structured, repetitive synthetic datasets like this fact verification collection, where the same core claims and evidence patterns are reused across multiple examples. In contrast to general purpose corpora (where α is typically closer to 1), the steeper decay reflects limited lexical innovation and heavy reliance on a compact set of recurring vocabulary.

The long tail comprising numerous low frequency and rare words arises from the introduction of domain-specific named entities, technical terms, numerical values, and claim specific modifiers (e.g., years like “1945”/“1950”, geographic references, scientific constants, and proper nouns). Even though the dataset is synthetic and structurally repetitive, each variation in claims introduces occasional unique tokens or combinations. These infrequent items form the tail of the distribution, representing specialized or context bound vocabulary that appears only in particular examples. The presence of this tail demonstrates that, despite overall repetition, the corpus still contains pockets of lexical novelty necessary for realistic fact verification scenarios.

The observed Zipfian pattern has important implications for both linguistic analysis and downstream NLP applications. The steep head highlights potential challenges for models: a heavy reliance on a few dominant terms may lead to overfitting to frequent patterns while underrepresenting rare but critical entities (e.g., specific historical dates or scientific facts). Conversely, the long tail underscores the dataset’s value for evaluating robustness successful models must correctly handle infrequent but semantically decisive vocabulary to perform accurate verification. From a corpus linguistics perspective, this distribution confirms the dataset behaves like a specialized technical corpus rather than open domain text, guiding decisions on preprocessing, data augmentation, and model architecture (e.g., favoring subword tokenization or retrieval augmented approaches). Overall, understanding the steep curve and long tail helps researchers interpret model performance, diagnose biases, and design more effective fact checking systems.

6. Discussion of the Results

The lexical and corpus level analyses provide a comprehensive characterization of the fact verification dataset by examining its vocabulary distribution, lexical diversity, and statistical properties. The analyses progress from descriptive lexical statistics to more sophisticated corpus linguistic measures, enabling a multidimensional assessment of the corpus’s linguistic characteristics.

Initially, lexical diversity is quantified using complementary measures, including the Shannon Diversity Index, Simpson Diversity Index, Type Token Ratio (TTR), Hapax Ratio, Honoré’s Statistic, Yule’s K, MTLD, and HD-D. These metrics collectively evaluate vocabulary richness, lexical repetition, and the prevalence of rare words, providing evidence of the corpus’s lexical heterogeneity. The combination of traditional and length independent diversity measures offers a robust assessment of vocabulary variation and linguistic complexity.

Subsequently, corpus level statistical regularities are investigated using Zipf's Law and Heaps' Law. Zipf's Law examines the relationship between word frequency and lexical rank, revealing whether the vocabulary follows the characteristic power law distribution observed in natural language. Heaps' Law describes how vocabulary growth scales with corpus size, providing insights into lexical productivity and saturation. The corresponding rank frequency distribution and vocabulary growth curve further illustrate the balance between frequent lexical items and the long tail of rare vocabulary.

Collectively, these analyses demonstrate that the dataset exhibits the statistical properties expected of a specialized natural language corpus while maintaining sufficient lexical diversity to support downstream applications such as fact verification, semantic similarity modelling, information retrieval, and transformer-based language modelling. The findings also establish a strong empirical foundation for subsequent machine learning experiments by characterizing the corpus before predictive modelling.

7. Conclusion

This study presented a comprehensive corpus linguistic analysis of a fact verification dataset by examining its lexical diversity, vocabulary richness, and statistical properties. Multiple lexical diversity measures, including the Shannon Diversity Index, Simpson Diversity Index, Type Token Ratio, Hapax Ratio, Honoré's Statistic, Yule's K, MTLT, and HD-D, were employed to characterize vocabulary variation, lexical concentration, and the occurrence of rare lexical items. Together, these metrics demonstrated that the corpus possesses substantial lexical heterogeneity while maintaining characteristics typical of specialized language resources.

The investigation of corpus statistical laws further revealed that the dataset exhibits the fundamental distributional properties of natural language. The observed Zipfian rank frequency distribution confirmed the presence of a long tail vocabulary structure in which a relatively small number of high frequency lexical items dominate the corpus, whereas numerous low frequency words contribute to semantic richness. Similarly, the Heaps' Law analysis indicated sublinear vocabulary growth, suggesting that new lexical items continue to emerge as the corpus expands while the rate of vocabulary acquisition gradually decreases. These findings demonstrate that the dataset exhibits realistic lexical behaviour despite its specialized focus on fact verification.

From an application perspective, the observed lexical characteristics have important implications for natural language processing and automated fact verification. The coexistence of highly repetitive vocabulary and infrequent domain specific terms highlights the need for models capable of learning both common linguistic patterns and rare but semantically informative expressions. The results therefore support the use of transformer based language models, retrieval augmented generation systems, and explainable artificial intelligence techniques for evidence aware fact verification and misinformation detection.

Although the present analysis provides valuable insights into the lexical structure of the corpus, several limitations should be acknowledged. The analyses focused primarily on corpus level lexical statistics and did not investigate syntactic complexity, semantic relationships, discourse structure, or contextual reasoning. Furthermore, the lexical diversity measures are influenced by preprocessing choices and corpus composition, emphasizing the importance of transparent and reproducible preprocessing pipelines.

Future research should extend the present work by incorporating semantic embedding analysis, topic modelling, named entity recognition, knowledge graph construction, readability assessment, explainable AI, and comparative benchmarking of transformer based and large language models. Evaluating retrieval based architectures, cross domain generalization, robustness against adversarial claims, and multilingual fact verification would further enhance understanding of the dataset and support the development of more reliable and trustworthy automated fact checking systems.

Overall, the findings demonstrate that the dataset provides a linguistically rich and statistically coherent benchmark for corpus linguistics, information retrieval, and natural language processing. By establishing the lexical and statistical characteristics of the corpus, this study lays a strong empirical foundation for future research on automated fact verification, misinformation detection, and explainable language technologies.

References

- [1] deBoer, F. (2014). Evaluating the comparability of two measures of lexical diversity. *System*, 47, 139-145.
- [2] Malvern, D., Richards, B. (2012). Measures of lexical richness. In *The Encyclopedia of Applied Linguistics* (p. 3622-3627).
- [3] Johansson, V., Gullberg, K., Johansson, R. (2025). Lies in the lexicon: A corpus based exploration of lexicon in truthful and deceitful narrative accounts. *Linguistics Vanguard*. Advance online publication.
- [4] Grieve, J., Woodfield, H. (2023). *The language of fake news*. Cambridge University Press. <https://doi.org/10.1017/9781009349161>.
- [5] Trnavac, R., Pöldvere, N. (2024). Investigating appraisal and the language of evaluation in fake news corpora. *Corpus Pragmatics*, 8, 107–130. <https://doi.org/10.1007/s41701-023-00162>.
- [6] Rashkin, H., Choi, E., Jang, J. Y., Volkova, S., Choi, Y. (2017). Truth of varying shades: Analyzing language in fake news and political fact checking. In M. Palmer, R. Hwa, S. Riedel (Eds.), *Proceedings of the 2017 Conference on Empirical Methods in Natural Language Processing* (p. 2931–2937). Association for Computational Linguistics. <https://doi.org/10.18653/v1/D17-1317>.
- [7] Bhattacharyya, P., Bhattacharya, A. (2025). Lexical and statistical analysis of Bangla newspaper and literature: A corpus driven study on diversity, readability, and NLP adaptation. *arXivpreprintarXiv:2601.06041*.
- [8] Broeder, P., Extra, G., Hout, R. V. (1993). Richness and variety in the developing lexicon. In *Adult language acquisition: Crosslinguistic perspectives* (Vol. I: Field methods, p. 145-163).
- [9] Chen, Y. S., Leimkuhler, F. F. (1989). A type-token identity in the Simon Yule model of text. *Journal of the American Society for Information Science*, 40(1), 45-53.
- [10] Malvern, D., Richards, B. J., Chipere, N., Duran, P. (2004). *Lexical diversity and language development: Quantification and assessment*. Palgrave Macmillan.
- [11] Richards, B. (1987). Type/token ratios: What do they really tell us *Journal of Child Language*, 14(2), 201-209.
- [12] Bestgen, Y. (2024). Measuring lexical diversity in texts: The twofold length problem. *Language Learning*, 74(3), 638-671.
- [13] Gouider, I. (2023, May 4). *Exploring the lexical diversity of texts through the use of modern technologies: An automated corpus analysis* [Paper presentation]. Modern Educational Technology for Quality and Transformative Learning: Local Needs and Global Challenges, Batna University, Algeria.
- [14] Miller, D., Biber, D. (2015). Evaluating reliability in quantitative vocabulary studies: The influence of corpus design and composition. *International Journal of Corpus Linguistics*, 20(1), 30-53.
- [15] Jarvis, S. (2017). Grounding lexical diversity in human judgments. *Language Testing*, 34(4), 537-553.
- [16] Kobylin, I., Biehunova, V., Tsyban, D., Kovalchuk, V. (2026). Measuring textual redundancy, lexical richness, and veracity: A multi metric approach to text evaluation. *Information Technologies and Systems*, 8(2), 25–46 <https://doi.org/10.15407/intechsyst.2026.02.025>
- [17] Hanselowski, A., Stab, C., Schulz, C., Li, Z., Gurevych, I. (2019). A richly annotated corpus for different tasks in automated fact-checking. In *Proceedings of the 23rd Conference on Computational Natural Language Learning (CoNLL)* (p. 493–503). Association for Computational Linguistics.