

Text Characteristics Reveal AI-Written News

Felix Thielen^{1,2}
Selina Gabric¹

Ahmed Mahmoud¹
Anastasia Yablokova¹

¹University of Trier, Department of Computational Linguistics

²Trier Center for Digital Humanities

{thielenf, s2ahmahm, s2segabr, s2anyabl}@uni-trier.de

Abstract

Distinguishing AI-generated from human-written text has become an important challenge. This study investigates whether linguistic features extracted with the Text Characterization Toolkit (TCT) can capture systematic differences between human-written and AI-generated news articles. We compiled a custom dataset of human-written news articles and generated AI counterparts. We trained an XGBoost classifier using these metrics and compared it to a RoBERTa baseline. Our analysis reveals distinct linguistic patterns between human-written and AI-generated news articles. We show that fundamental linguistic properties captured by TCT provide a more robust signal for distinguishing AI-generated text across domains than learned word embeddings.

1 Introduction

Text generation has become one of the most important and fast-developing areas of Natural Language Processing (NLP). The transformer architecture, introduced by Vaswani et al. (2017), enabled substantial improvements in large language models (LLMs), which now produce coherent, high-quality text that is difficult to distinguish from human-written content.

A recent survey (Fraser et al., 2025) also shows that an increasing number of people who participated in research experiments to identify AI-generated and human-written text are unable to reliably distinguish between these two types of materials. While these developments are outstanding from a technical viewpoint, they also raise several concerns. In particular, the domain of news and media can be strongly affected by the misuse of such technologies.

AI-generated texts can be used to spread misinformation, influence public opinion, and harm a reputation by creating misleading content. Hanley and Durumeric (2023) show that the amount

of machine-generated news articles published online has increased noticeably between early 2022 and mid-2023. Because of this, it becomes increasingly vital to develop methods that can distinguish between AI-generated and human-written text.

In this study, we investigate whether linguistic characteristics can systematically distinguish AI-generated from human-written news. We construct paired human-written and AI-generated articles, compare a TCT-based XGBoost classifier with a RoBERTa baseline in both in-domain and cross-domain evaluations, and analyse the linguistic differences between the two types of text. The accompanying code, datasets, and experimental results will be made available in a repository¹.

2 Related Work

Our study builds on a growing body of research in AI-generated text detection. The survey by Wu et al. (2025) provides a comprehensive overview of contemporary approaches, identifying key challenges such as out-of-distribution performance, data ambiguity, and evaluation framework limitations.

Hanley and Durumeric (2023) demonstrated the increasing prevalence of synthetic news articles, with significant growth observed in both mainstream and misinformation websites. Their findings underscore the importance of our research, as the ability to distinguish AI-generated content becomes increasingly critical for media integrity.

Previous work has also explored feature-based methods and benchmark datasets for this task. Fröhling and Zubiaga (2021) proposed a classifier using manually designed features to distinguish human-written from machine-generated text. They evaluated the method on texts generated by GPT-2, GPT-3, and Grover and found that it performed competitively with more computationally expensive detectors. Uchendu et al. (2021) introduced

¹<https://github.com/thielenf/tct-ai-news>

TURINGBENCH, a benchmark for distinguishing human-written from machine-generated text and identifying which model generated a given text. In our study, we focus on paired news articles and combine classification with a feature-level analysis of linguistic differences.

The Text Characterization Toolkit serves as the foundation for our feature extraction approach. Unlike traditional benchmark-based evaluations that report single-value performance scores, TCT enables a deeper analysis of textual features and their correlations with authorship.

3 Methodology

3.1 Dataset Collection and Preparation

The dataset was combined from Cornell Newsroom (Grusky et al., 2018) and BBC News Archive (Greene and Cunningham, 2006) (see Table 1). We decided on a cutoff point in 2016 to ensure that no AI content was accidentally included in the samples. The main body had to be at least 120 words to ensure that all metrics can be meaningfully computed. All headlines had to be at least 4 words so that the LLMs had enough information to follow instructions.

For each human-written article, we generated AI counterparts using mistral-small3.2 (24B) (AI, 2025) and gemma2 (27B) (Gemma et al., 2024). The models were run locally through Ollama (Ollama Team, 2025) using its self-hosted API. Each model received the following prompt:

```

Headline: {headline}
Date: {date}
Paragraphs: {paragraph count}
Words: {word count}
Language: English

```

Rather than using a single fixed target length, the prompt included a separate word target derived from the word count of the corresponding human article. It also included a paragraph target; for the BBC subset, this value was uniformly set to one. By providing the models exclusively with high-level metadata, this generation setup reduces the risk of the models recalling the original articles from their pre-training.

3.2 Feature Extraction

We applied the Text Characterization Toolkit to extract 68 linguistic features covering lexical, syntactic, and semantic dimensions. These features include metrics for lexical diversity, syntactic complexity, readability, and part-of-speech distribution.

Beyond this we extended TCT to compute Shannon entropy on character level and word level.

3.3 Statistical Analysis and Visualization

Due to the paired nature of our data, we performed statistical significance testing using Wilcoxon signed-rank tests, comparing each original human text against each AI-generated counterpart. We decided against a paired t-test because during normality testing, the differences in paired samples were not normally distributed. For each metric, comparative box plots were created to visually assess the magnitude of the differences and similarities. Effect size was computed via rank-biserial correlation. A selection can be seen in Figures 1, 2, and 3.

3.4 Model Training and Evaluation

We trained an XGBoost classifier (Chen and Guestrin, 2016) on the 68 numerical features extracted via TCT and fine-tuned a roberta-base model (Liu et al., 2019) for sequence classification as a baseline. For in-domain experiments on our custom dataset, data was split using stratified sampling (XGBoost: 80% train, 20% validation; RoBERTa: 70% train, 15% validation, 15% test). To account for class imbalance during XGBoost training, the loss function was weighted using the sample ratio of human to AI texts (scale_pos_weight). RoBERTa was fine-tuned for up to 5 epochs with a learning rate of 2×10^{-5} and early stopping based on validation macro F1 (patience of 4 checks). Beyond in-domain testing, both models were evaluated zero-shot on an external news dataset (Roy et al., 2025, “gsingh-1”) and a non-news benchmark (Yatsenko, 2025, “ai-detection-pile”) containing student essays and fiction to assess cross-domain generalization (see Table 1).

4 Results

Our analysis revealed distinct linguistic patterns between human-written and AI-generated news articles. Specifically, 61 of the 68 extracted text characteristics showed statistically significant differences across models ($p < 0.001$, Wilcoxon signed-rank test). Effect size analysis via rank-biserial correlation revealed that structural and stylistic differences between human and LLM-generated texts are pronounced across most evaluated metrics. **gemma2-27b** yielded medium-to-large effect sizes on 85.2% of these metrics (42 large, 10 medium, 9 small out

	Our dataset			gsingh-1		ai-detection-pile	
	Human ¹	Human ²	AI	Human	AI	Human	AI
Articles	2,171	8,000	20,342	7,295	21,925	25,000	25,000
Article Length*	374	668	501	424	431	225	344
Sentence Length*	23	23	23	23	25	16	23
Type-Token Ratio*	0.54	0.51	0.51	0.51	0.48	0.57	0.48

Table 1: Descriptive statistics of the sources for our custom dataset and the evaluation datasets. Lengths in tokens. *Median. ¹BBC News. ²Cornell Newsroom.

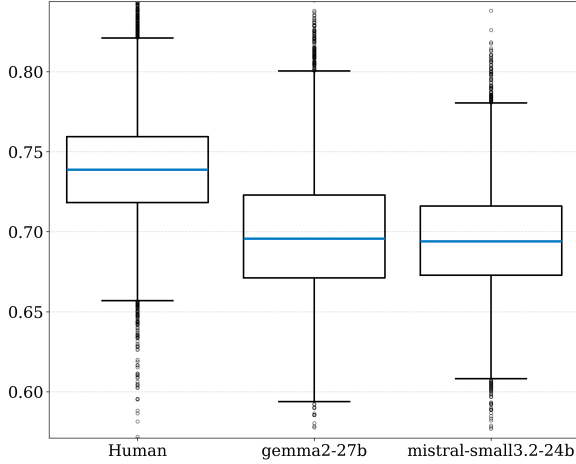


Figure 1: SYNMEDpos: This metric quantifies syntactic variation by measuring Levenshtein distance of POS tags between consecutive sentences. Human texts show significantly higher values compared to gemma2-27b ($p < 0.001$, $r_{tb} = 0.74$) and mistral-small3.2-24b ($p < 0.001$, $r_{tb} = 0.83$).

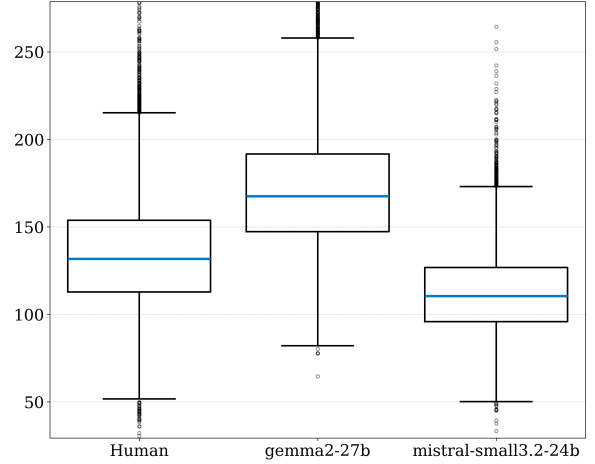


Figure 2: LDMTLD: Measure of Textual Lexical Diversity (MTLD). This metric assesses the variety of vocabulary in a text. Statistically significant differences were observed compared to gemma2-27b ($p < 0.001$, $r_{tb} = -0.76$) and mistral-small3.2-24b ($p < 0.001$, $r_{tb} = 0.57$).

of 61). **mistral-small3.2-24b** exhibited a similar pattern, with 82.0% of metrics showing medium-to-large effect sizes (38 large, 12 medium, 11 small) (see Table 2). This confirms systematic variation in writing style between sources.

The RoBERTa baseline model demonstrated near-perfect performance on news datasets ($F1 = 0.99$) but showed dramatic performance degradation on the non-news domain ($F1 = 0.35$). This indicates specialization in the trained model, which may be a symptom of deep neural classifiers’ domain overfitting and shortcut learning when fine-tuned on specific corpora. In contrast, our XGBoost model showed a less drastic decline, achieving $F1 = 0.46$ on non-news data, a 31% relative improvement over the baseline (see Table 3). This suggests that fundamental linguistic properties captured by text characteristic metrics provide a more robust signal for recognizing AI-generated text across domains than learned word embeddings.

5 Discussion

The distinct linguistic patterns we observed provide valuable insight into the differences between human and AI writing styles. They may indeed be useful in identifying AI-written text. We will discuss a selection of metrics from different categories for illustration. The complete set of plots, raw data and source code will be made available in a repository.

5.1 Syntactic Complexity

Syntactic complexity analysis reveals noteworthy findings for the SYNLE and SYNNP metrics. The SYNLE metric (see Figure 3) quantifies the number of left-embedded non-verb tokens preceding the main verb of a sentence (Zou et al., 2022). Consequently, a higher SYNLE value indicates an increased structural complexity and longer dependency distances. Conversely, a lower SYNLE value depicts a shorter pre-verbal distance, a common

Model	Large	Medium	Small	Total
gemma2-27b	42	10	9	61
mistral-small3.2-24b	38	12	11	61

Table 2: Distribution of effect size magnitudes via rank-biserial correlation. Thresholds used: Large ($|r_{rb}| \geq 0.5$), Medium ($0.3 \leq |r_{rb}| < 0.5$), and Small ($|r_{rb}| < 0.3$).

	Our dataset			gsingh-1			ai-detection-pile		
	P	R	F1	P	R	F1	P	R	F1
RoBERTa	0.99	0.99	0.99	0.99	0.99	0.99	0.35	0.35	0.35
XGBoost	0.98	0.98	0.98	0.90	0.88	0.88	0.54	0.52	0.46

Table 3: Classification scores on all datasets as Precision, Recall and F1.

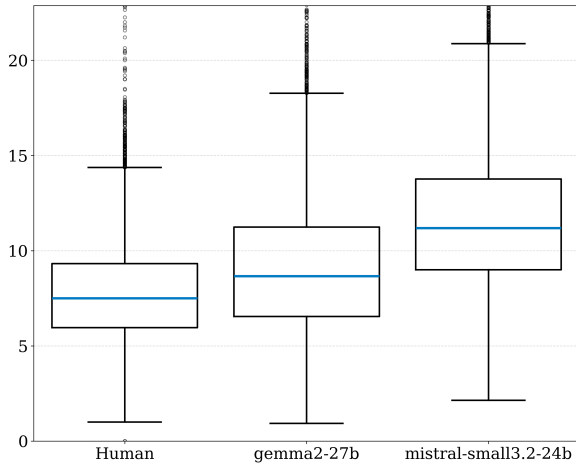


Figure 3: SYNLE: Left embeddedness: average words before main verb. Human texts differ significantly compared to gemma2-27b ($p < 0.001$, $r_{rb} = -0.33$) and mistral-small3.2-24b ($p < 0.001$, $r_{rb} = -0.78$).

feature of Subject-Verb-Object (SVO) languages such as English. Our finding demonstrates that AI-generated texts exhibit a higher SYNLE score compared to their human-authored counterparts, suggesting a systematically higher structural complexity.

A similar pattern is observable with SYNNP, which measures the average number of pre- and post-nominal elements that elaborate a noun phrase (NP) beyond the head noun. The SYNNP of AI-generated texts exhibits an increased score compared to their human-authored counterparts. An increased SYNNP value indicates that the analyzed NP encompass an increased amount of descriptive structures, e.g. stacked adjectives, thus further increasing the structural complexity.

These findings could be indicators of a general trend in LLM text, and its effects may be impacting

other scores. For instance, we measure strongly reduced readability scores for AI-generated texts, which aligns with the idea that readers need to process more embedded information before reaching the core meaning of a sentence as measured in SYNLE and SYNNP. Additionally, the measured lower Entropy of AI-generated news articles suggests that models tend to use non-informative fillers rather than adding meaningful content. This interpretation is in alignment with previous findings on the differences in Entropy between human-authored and AI-generated texts (Widjaja et al., 2023).

When comparing the SYNMEDpos (see Figure 1) and SYNSTRUTa metrics, the results seem contradictory. SYNMEDpos quantifies syntactic variation by measuring the similarity of POS tags between consecutive sentences while SYNSTRUTa measures the syntactical similarity between neighbouring sentences. Our analysis shows that Human-written texts show an increased SYNMEDpos value while the SYNSTRUTa score is lower than for the AI-generated counterpart. This finding may reflect the stylistic choices human writers utilize while maintaining good readability through consistent syntactical structures.

5.2 Lexical Diversity

While it could be intuitively assumed that lexical diversity measurements show general trends like many of the other metrics (humans scoring categorically higher or lower), the data seems to show individual model characteristics. In LDMTLD (see Figure 2), one model scored much higher than the human baseline and the other scored much lower. The differences suggest that lexical diversity be-

haves like a model-specific fingerprint instead of a universal property of LLMs. This may ultimately be an indication of the models’ training data and how they learned to comply with the instruction to write a news article.

The near-perfect performance on our custom news dataset may partly reflect how its AI-generated articles were produced. The models received the headline, date, and length targets, but not the source article body. This setup may have produced systematic differences between the generated and human-written articles in how journalistic details, quotations, and source attribution were expressed, thereby making the classification task easier. Benchmarks constructed using similar generation procedures may therefore produce optimistic estimates of detection performance on naturally occurring machine-generated news.

6 Conclusion

This study investigated whether or not linguistic features extracted through the Text Characterization Toolkit can capture systematic differences between human-authored and AI-generated news articles. The analysis revealed statistically significant differences across 61 text characteristics. While the present study focuses on the detection of AI-generated news articles using text characteristics, the findings contribute to the detection approaches at large.

The trained XGBoost model demonstrated greater generalizability compared to the RoBERTa baseline model, achieving a F1-score of 0.46 on non-news data, a 31% relative improvement over the baseline. This result suggests that text characteristic metrics provide a more robust signal for the identification of AI-generated text across domains than learned word embeddings.

Nonetheless, the present findings on the identification of AI-generated texts using text characteristic are needed to test the applicability of the proposed approach to other text genres, and languages.

7 Ethical Considerations

Our research raises important ethical considerations regarding AI-generated content detection. Although our methods demonstrate effectiveness in distinguishing AI from human writing, we acknowledge potential misuse of such technologies for censorship or surveillance purposes. The development

of detection methods must be accompanied by responsible deployment guidelines and transparency about limitations.

As AI models continue to improve, the detection arms race presents challenges. Our findings suggest that current detection methods may become less effective as AI-generated text becomes more human-like. This underscores the need for ongoing research and adaptive detection approaches.

8 Limitations

While our study provides valuable insights into AI-generated text detection, several limitations should be considered. Our training data consists exclusively of English news articles from two major outlets, which may limit generalization to other languages, genres, and cultural contexts. The performance decline on the non-news dataset indicates that models trained on news should not necessarily be expected to transfer reliably to other domains. The 2016 cutoff for human-written articles may not fully represent contemporary journalistic styles. Although the generation prompts did not include the source article bodies, some source articles may have been present in the models’ pretraining data, and we did not conduct a dedicated contamination analysis. The effect of any such prior exposure on the generated texts and the observed feature differences therefore remains unknown. Rapid AI model evolution also means that our findings may not apply to newer model generations.

Because generation was based on the original headline and date, different prompt formulations or generation settings could produce different linguistic patterns and therefore affect classification performance. We evaluated a subset of the metrics available in TCT, and future work could examine additional linguistic measures. Both our XGBoost classifier and the RoBERTa baseline have inherent limitations in capturing subtle linguistic patterns and generalizing across domains. We did not compare XGBoost with other feature-based classifiers, and alternative models may perform better, particularly on non-news texts.

Future research could explore multilingual studies, temporal analysis of AI text evolution, and adversarial testing. As AI text generation advances, detection methods may need to incorporate multi-modal analysis and contextual information beyond purely linguistic features. Despite these limitations, our findings demonstrate the potential of feature-

based approaches to distinguish AI-generated content.

References

- Mistral AI. 2025. [Mistral small 3.2 \[Version 2506, 24B Parameters\]](#).
- Tianqi Chen and Carlos Guestrin. 2016. [XGBoost: A Scalable Tree Boosting System](#). In *Proceedings of the 22nd ACM SIGKDD International Conference on Knowledge Discovery and Data Mining, KDD '16*, pages 785–794. ACM.
- Kathleen C. Fraser, Hillary Dawkins, and Svetlana Kiritchenko. 2025. [Detecting ai-generated text: Factors influencing detectability with current methods](#). *Journal of Artificial Intelligence Research*, 82:2233–2278.
- Leon Fröhling and Arkaitz Zubiaga. 2021. [Feature-based detection of automated language models: Tackling GPT-2, GPT-3 and Grover](#). *PeerJ Computer Science*.
- Team Gemma, Morgane Riviere, Shreya Pathak, Pier Giuseppe Sessa, Cassidy Hardin, Surya Bhupatiraju, Léonard Hussenot, Thomas Mesnard, Bobak Shahriari, Alexandre Ramé, Johan Ferret, Peter Liu, Pouya Tafti, Abe Friesen, Michelle Casbon, Sabela Ramos, Ravin Kumar, Charline Le Lan, Sammy Jerome, and 179 others. 2024. [Gemma 2: Improving Open Language Models at a Practical Size](#). *arXiv preprint*. ArXiv:2408.00118 [cs].
- Derek Greene and Pádraig Cunningham. 2006. [Practical solutions to the problem of diagonal dominance in kernel document clustering](#). In *Proceedings of the 23rd International Conference on Machine Learning, ICML '06*, pages 377–384. ACM.
- Max Grusky, Mor Naaman, and Yoav Artzi. 2018. [Newsroom: A Dataset of 1.3 Million Summaries with Diverse Extractive Strategies](#). In *Proceedings of the 2018 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies, Volume 1 (Long Papers)*, pages 708–719, New Orleans, Louisiana. Association for Computational Linguistics.
- Hans W. A. Hanley and Zakir Durumeric. 2023. [Machine-Made Media: Monitoring the Mobilization of Machine-Generated Articles on Misinformation and Mainstream News Websites](#). *arXiv preprint*.
- Yinhan Liu, Myle Ott, Naman Goyal, Jingfei Du, Mandar Joshi, Danqi Chen, Omer Levy, Mike Lewis, Luke Zettlemoyer, and Veselin Stoyanov. 2019. [Roberta: A robustly optimized BERT pretraining approach](#). *CoRR*, abs/1907.11692.
- Ollama Team. 2025. [Ollama](#). GitHub Repository. Version 0.11.3.
- Rajarshi Roy, Nasrin Imanpour, Ashhar Aziz, Shashwat Bajpai, Gurpreet Singh, Shwetangshu Biswas, Kapil Wanaskar, Parth Patwa, Subhankar Ghosh, Shreyas Dixit, Nilesch Ranjan Pal, Vipula Rawte, Ritvik Garimella, Gaytri Jena, Amit Sheth, Vasu Sharma, Aishwarya Naresh Reganti, Vinija Jain, Aman Chadha, and Amitava Das. 2025. [A Comprehensive Dataset for Human vs. AI Generated Text Detection](#). *arXiv preprint*. ArXiv:2510.22874 [cs].
- Adaku Uchendu, Zeyu Ma, Thai Le, Rui Zhang, and Dongwon Lee. 2021. [TURINGBENCH: A benchmark environment for Turing test in the age of neural text generation](#). In *Findings of the Association for Computational Linguistics: EMNLP 2021*, pages 2001–2016, Punta Cana, Dominican Republic. Association for Computational Linguistics.
- Ashish Vaswani, Noam Shazeer, Niki Parmar, Jakob Uszkoreit, Llion Jones, Aidan N. Gomez, Łukasz Kaiser, and Illia Polosukhin. 2017. [Attention is all you need](#). In *Advances in Neural Information Processing Systems 30*, pages 5998–6008. Curran Associates, Inc.
- Henry Widjaja, Shreya Das, George Daniel Dixon, and Farah Basher. 2023. [Significance of entropy in distinguishing between ai-driven content](#). STEM2SHEM 2023, Stanford Compression Forum, Stanford University.
- Junchao Wu, Shu Yang, Runzhe Zhan, Yulin Yuan, Lidia Sam Chao, and Derek Fai Wong. 2025. [A Survey on LLM-Generated Text Detection: Necessity, Methods, and Future Directions](#). *Computational Linguistics*, 51(1):275–338.
- Artem Yatsenko. 2025. [artem9k/ai-text-detection-pile](#).
- Longhui Zou, Michael Carl, Mehdi Mirzapour, H  l  ne Jacquet, and Lucas Nunes Vieira. 2022. [Ai-based syntactic complexity metrics and sight interpreting performance](#). In *Intelligent Human Computer Interaction*, pages 534–547, Cham. Springer International Publishing.

9 Appendix

9.1 XGBoost Configuration

The XGBoost classifier was trained on the 68 numerical features extracted via TCT. An 80/20 stratified train/validation split was applied to the main dataset. To address class imbalance, `scale_pos_weight` was dynamically set to the ratio of human to AI samples in the training split ($\sum y_{\text{train}} = 0 / \sum y_{\text{train}} = 1$). Feature columns across validation and held-out sets were aligned to match the training feature schema.

Parameter	Value
Objective	binary:logistic
$n_{\text{estimators}}$	100
Learning rate	0.1
Max depth	5
scale_pos_weight	Human / AI sample ratio
Eval metric	logloss
Train / Val split	80 / 20 (Stratified)

Table 4: XGBoost hyperparameters and configuration.

9.2 RoBERTa Fine-Tuning Configuration

roberta-base was fine-tuned using a 70/15/15 stratified train/validation/test split. Evaluation was performed every 200 steps, and training was stopped early if validation macro F1 did not improve after 4 consecutive evaluation checks (patience = 4).

Parameter	Value
Base Model	roberta-base
Max Sequence Length	512 tokens
Epochs	5 (Max)
Per-device Batch Size	8
Learning Rate	2×10^{-5}
Warmup Steps	50
Weight Decay	0.01
Optimizer	AdamW
FP16 Precision	Enabled
Early Stopping Patience	4 checks (200 steps)
Best Model Metric	Macro F1
Train / Val / Test split	70 / 15 / 15 (Stratified)

Table 5: RoBERTa fine-tuning hyperparameters.

9.3 Overview of Extracted Text Characteristics

All 68 extracted metrics were evaluated using Wilcoxon signed-rank tests. A total of 61 features showed statistically significant differences ($p < 0.001$), categorized below:

Descriptive (10) DESPL, DESPLw, DESSC, DESSL, DESSLd, DESWC, DESWL1t, DESWL1td, DESWLsy, DESWLsyd.

Lexical Diversity (3) LDMTLD, LDTTRa, LDTTRc.

Repetition (10) REPETITION_FRACTION_LINE_CNT, REPETITION_FRACTION_LINE_LEN, REPETITION_FRACTION_NGRAM_2_LEN, REPETITION_FRACTION_NGRAM_4-10_LEN.

Syntactic (6) SYNLE, SYNMEDlem, SYNMEDpos, SYNMEDwrd, SYNNP, SYNSTRUTa.

Token Attributes (5)

TOKEN_ATTRIBUTE_RATIO_ALPHA, TOKEN_ATTRIBUTE_RATIO_DIGIT, TOKEN_ATTRIBUTE_RATIO_EMAIL, TOKEN_ATTRIBUTE_RATIO_PUNCT, TOKEN_ATTRIBUTE_RATIO_URL.

Word Properties (11) WORD_PROPERTY_WRDADJ, WORD_PROPERTY_WRDADV, WORD_PROPERTY_WRDFRQa, WORD_PROPERTY_WRDFRQc, WORD_PROPERTY_WRDFRQmc, WORD_PROPERTY_WRDHYPn, WORD_PROPERTY_WRDHYPnv, WORD_PROPERTY_WRDHYPv, WORD_PROPERTY_WRDNOUN, WORD_PROPERTY_WRDPOLc, WORD_PROPERTY_WRDVERB.

Word Set Incidence (12)

WORD_SET_INCIDENCE_C4_COMMON_WORDS, WORD_SET_INCIDENCE_CNCAAdd, WORD_SET_INCIDENCE_CNCCaus, WORD_SET_INCIDENCE_CNCLogic, WORD_SET_INCIDENCE_CNCNeg, WORD_SET_INCIDENCE_CNCPos, WORD_SET_INCIDENCE_CNCTemp, WORD_SET_INCIDENCE_WRDPRP1p, WORD_SET_INCIDENCE_WRDPRP1s, WORD_SET_INCIDENCE_WRDPRP2, WORD_SET_INCIDENCE_WRDPRP3p, WORD_SET_INCIDENCE_WRDPRP3s.

Readability (2) flesch, flesch-kincaid.

Entropy (2) entropy_char, entropy_word.

Non-significant (7)

DESPLd, LDHDD, SYNSTRUTt, REPETITION_FRACTION_NGRAM_3_LEN, REPETITION_FRACTION_PARA_CNT, REPETITION_FRACTION_PARA_LEN.