



Unmasking Misinformation: A Machine Learning Approach to Detecting Fake News

Maya Hussain

Abstract - In recent times, there has been an increase in misinformation, with misleading information being shared as real news to deceive and manipulate public opinion. The dissemination of misinformation, particularly in areas with global implications such as politics and health, can have severe consequences for society as a whole. For example, recent US elections related widespread misinformation has shown to deepen polarization and erode trust in both democratic institutions and our news media. Misleading reports during crises like the Ebola outbreak or COVID-19 misinformation about vaccines and treatments spread unnecessary fear, created barriers for public health response teams, and resulted in many preventable deaths. Social media further amplifies fake news, making it difficult for fact-checking efforts to keep pace. To distinguish misinformation from credible reporting, this paper aims to apply machine learning techniques to detect fake news with greater accuracy. To research this, we analyze datasets containing both fake and real news articles to uncover linguistic patterns and differences between the two. Natural Language Processing (NLP) techniques such as Term Frequency-Inverse Document Frequency (TF-IDF) are used to convert text data into numerical features for training machine learning models. Several classification algorithms, such as Logistic Regression, Random Forest, and XGBoost, are then trained to differentiate fake from real news. To further explore the differences in the data types, an analysis is done to examine sentiment differences. By leveraging data from everyday news, politics, and health sources, we keep the work grounded in the real-world implications of fake news disguised as fact. The goal is to develop an AI-powered automated fact-checking system to distinguish between real and fake sources, thereby contributing to ongoing efforts to protect the public from the harms of misinformation and uphold their trust in news media.

I. Introduction

Misinformation is the rapid spread of false information, regardless of the intent to deceive¹. Creators of fake news often attempt to spread false information as real news, with the intent to deceive. The impact of disseminating fake news can be extremely harmful, as demonstrated during the COVID-19 pandemic. There was widespread misinformation about vaccines, preventive treatments, and even the origins of the virus. Research indicates that the impact of this misinformation spread is quite severe. A person exposed to just one piece of misinformation saw their likelihood of vaccination fall to 62.9%, and if exposed to six or more such pieces of misinformation, it fell to 52.2%². The data also shows that one-third of the most viewed COVID-19 videos on YouTube, which contained misinformation, had greater than 62 million views worldwide³. These statistics tell the story of the global harmful impact misinformation can have on our society as a whole. Because fact-checking is a manual and time-consuming way to differentiate real from fake information, and most readers do not have the time to go through multiple sources to verify and cross-check information⁴ it is not a viable



way to protect against false data. Furthermore, this traditional method of authenticating sources does not scale in the digital age.

Machine learning (ML) is a branch of artificial intelligence that utilizes data to train models capable of replicating human learning processes. It uses algorithms and statistical models to identify distinct patterns in the data. Machine learning has evolved in numerous ways, driven by improvements in algorithms, advances in computer power, and increased access to data. This evolution has enabled machines to process vast amounts of unstructured information with high accuracy. These advancements are the reason ML is a common facet of our daily lives now. Recommendation systems on platforms like Spotify, Netflix, and Amazon use ML to make suggestions based on previous interests. The image recognition feature on our devices can unlock them using Face ID. Social media platforms utilize recommendation systems to surface relevant posts, and self-driving cars employ computer vision models to navigate safely.

To better equip people to discern the validity of information and its sources, this paper aims to utilize machine learning to provide a reliable and scalable method for slowing the spread of misinformation. We first identify data sources to use for training and then use Term Frequency-Inverse Document Frequency (TF-IDF) to preprocess the data. TF-IDF is a technique that converts text data into numerical feature vectors that a machine learning algorithm can process. Term Frequency determines how frequent a term is within a dataset, and IDF measures how rare that term is across all datasets. We then utilize three datasets for this work: one comprising general fake and real news from various news sources, another containing politicians' statements, and a third with COVID-19 information. We split the data into training, validation, and test sets and then set up our experiment to train and evaluate using Logistic Regression, Random Forest, and XGBoost algorithms. Finally, we compare the models and report on metrics such as their accuracy, precision, recall, and F1-score to determine the best-performing model for this task. This approach ensures that the models are trained in a structured way and are evaluated fairly.

Literature Review

Over the past few years, the rise of misinformation on online platforms has prompted researchers to analyze different approaches to detecting fake news. Numerous studies have addressed this issue through a variety of methods, such as the examination of content using artificial intelligence, behavioral modeling, and hybrid models. These processes provide essential insight into the role of machine learning in detecting fake content.

A content-driven approach to detecting false news is introduced in the work by Choudhary et al.⁵, which combines meaning-focused text analysis with classical machine learning models to identify textual features associated with misinformation. Specific classifiers they apply include Support Vector Machines, Decision Trees, and Naive Bayes. This paper emphasizes the importance of structural language and word-level patterns that identify real from fake news. The authors used multiple datasets to evaluate the performance of these classifiers.

The research by Indu V et al.⁶ highlights the fact that, despite the use of fact-checking tools, users are still highly likely to believe fake news. They use a hybrid framework for misinformation detection that integrates user behavior with emotion analysis. Their work emphasizes the underutilization of emotion as a key indicator of misinformation and shows how user-specific features can serve as strong signals. Although this study demonstrates the value of behavioral and emotional cues in misinformation detection, its dependence on Twitter data limits broad applicability. They also focus only on four primary emotions, which can overlook more subtle dimensions of human reactions. Building on these limitations, our research focuses on more widely available datasets, such as news articles, to train models that can account for a broader range of misinformation areas.

A comprehensive study by Huang and Chen⁷ uses an ensemble learning framework for fake news detection. The similarity between fake and real news makes it challenging to distinguish, so the authors focus on a generalizable model that utilizes textual features and an offline training phase for pre-processing, individual model training, and ensemble optimization, followed by an online phase for real-time detection. This builds on prior works, such as SVM-based satire detection (Rubin et al.⁸), which achieves an 87% F1-score, and RNN models for Twitter data (Ma et al.⁹), which reach 90% accuracy. Multimodal frameworks, such as TI-CNN (Yang et al.,¹⁰) and EANN (Wang et al.,¹¹), also contribute to this foundation. However, Huang and Chen extend them by integrating an ensemble of deep learning models using the Self-Adaptive Harmony Search (SAHS) algorithm for superior performance. While this approach effectively demonstrates how ensemble models can improve detection accuracy, their dependence on heavy computation needs makes it difficult to scale for real-time use. To overcome these limitations, our work focuses on lightweight and context-aware models that adapt more efficiently to dynamic online environments.

II. Methods

Our paper used a research methodology that involved training machine learning models on datasets containing both real and fake news across three domains: politics, health, and general media. Publicly available data was utilized, ensuring the solution was grounded in real-world applications. The data pre-processing involved tokenizing the text, converting it to all lowercase, and removing stop words and punctuation. TF-IDF was applied to transform text into its numerical representation, which downweights common words while highlighting those that are more characteristic of a document. Three supervised learning algorithms were then chosen for training: Logistic Regression, Random Forest, and XGBoost. Logistic Regression was used as the first model to train and serve as a baseline linear model. Then, Random Forest, an ensemble of decision trees, was used to capture nonlinear interactions among features. Finally, XGBoost, a gradient boosting framework, was implemented to handle contextual complexity. Each model was trained separately across the three domains' data sets using a similar three-way train-test split and validated through cross-validation to ensure robustness. We also conducted sentiment analysis on the same datasets using Valence Aware Dictionary and Sentiment Reasoner (VADER¹²), a lexicon-based tool optimized for both short and long texts. To evaluate the performance of the models' accuracy, precision, recall, and F1-score metrics were

used. For sentiment analysis, sentiment scores were computed for each statement, and then these were plotted to output and compare sentiment distributions. This overall experiment setup enabled the assessment of both the effectiveness of each model and the contribution of sentiment features in enhancing misinformation detection.

II.I Datasets

We identified three datasets, comprising both fake and real data, to run the experiments on (see Table 1). We selected these datasets because they cover areas where misinformation is most harmful: health, politics, and general media. This ensures we test our models across domains that young people and the general public frequently encounter online. The datasets include both short-form claims (common in social media and political fact-checking) and long-form articles (as seen in mainstream outlets), allowing for the evaluation of model performance across different formats and content lengths. Each dataset was drawn from sources that use trusted labeling methods. TF-IDF was used to transform each news article into feature vectors, which were further used to train machine learning algorithms. The feature vector was obtained by calculating the frequency of each word within a dataset (TF) and within all datasets (IDF).

- Fake and Real News¹³ dataset – This dataset was collected from real-world sources; the truthful articles were obtained by crawling articles from Reuters.com (a News website). The fake news articles were collected from websites that Politifact and Wikipedia flagged. The dataset comprises articles on a range of topics, with the majority focusing on global news and politics. The Fake vs. Real news is broken down into two datasets, with each containing around 12,600 articles.
- CoAid¹⁴ dataset– This dataset includes COVID-19 and other public health claims, as well as fake news on websites and social platforms, along with users' social engagement regarding such news. Labels are derived from fact-checking groups and datasets compiled by researchers who rely on peer-reviewed medical sources and official health agencies that are backed by scientific consensus and expert review.
- LIAR¹⁵ dataset – This dataset includes short political statements. Labels are based on professional fact-checking organizations such as PolitiFact, which evaluates political statements and campaign claims using a standardized scale. This makes the dataset reliable for distinguishing true from misleading political news.

All three datasets rely on expert-verified and credible labeling rather than crowdsourced or arbitrary annotations, which provides greater confidence in the validity of the results. To maintain consistency, each dataset was split into training and testing sets in a 70:30 proportion, with stratified labels, ensuring balanced representation of real and fake news in both sets.

Dataset	Domain	Type	Total Size	Balance (Fake vs. Real)	Labelling Source	Years & Location
Fake & Real News	General Media & News	Long news articles	44,898	~50/50	Website domains (e.g., reuters.com vs. empireherald.com)	US Only
CoAID	Health	News, Social Media	4,251 news, 296,000 related user engagements, 926 social platform posts	~50/50	Verified by health orgs & fact-checkers	-
LIAR	Politics	Short statements	12,836		PolitiFact (Fact-checkers)	-

Table 1. Summary of datasets used in this study.

II. II Sentiment analysis

Sentiment analysis is a Natural Language Processing (NLP) technique that can be used to determine the emotional tone behind data, and it can be either positive, negative, or neutral. It is another critical factor in detecting the authenticity of news and content. The data typically reflects the emotional tone of the content, and we have observed that fake news often employs extreme sentiment to provoke strong emotions, often in an overwhelmingly positive or overwhelmingly negative manner. Unusual emotional patterns, such as anger or fear in inappropriate contexts, can also be a strong signal. Sentiment analysis can therefore serve as a standalone tool for general analysis or be incorporated directly as a feature within an ML model.

For our purpose, we employed a lexicon-based sentiment analysis methodology known as VADER. It has been widely used for analyzing social media content, online reviews, and short-form text domains that share similarities with language used to convey misinformation online. Unlike other lexicon-based tools, VADER is designed explicitly for analyzing sentiment in short text segments, accounting for factors such as capitalization, punctuation, slang, and even emoticons, which makes it effective in detecting exaggerated emotional tone—a hallmark of fake news. Additionally, it is a computationally efficient algorithm and easy to integrate with other models, which makes it an ideal candidate for our hybrid misinformation detection framework.

II. III Machine Learning algorithms

Having studied the systems and approaches used in the papers described above and analyzed their strengths and weaknesses, we now propose a methodology for identifying the model with the highest accuracy across various datasets. After identifying three data sets to use our methods on, we prepare the data for processing by the following ML models: Logistic Regression, Random Forest, and XGBoost. We chose Logistic Regression due to its simplicity, speed, and interpretability as a linear model for binary classification, making it a suitable fit for the problem at hand. We then run Random Forest, an ensemble model that works by building multiple decision trees and combining their outputs for classification. It handles non-linear data well and is robust, even though it is slower and less interpretable than Logistic Regression. Finally, we run XGBoost, a high-performance implementation of a gradient boost model that builds trees sequentially to correct the errors of previous trees. This model typically achieves high accuracy but is more complex and slower to train.

III. Results

The experiments across the three datasets highlight both the strengths and limitations of traditional machine learning approaches for misinformation detection. Across the three datasets, no single model dominated in every setting. Logistic Regression consistently performed strongly, particularly on long-form articles (Fake & Real News) and short political claims (LIAR), suggesting that linear models work most effectively when the text has a clear and consistent structure. In contrast, XGBoost proved more effective on the CoAID dataset, where misinformation tended to be more context-dependent and noisier, suggesting that tree-based ensemble methods are better suited to handling diverse features. The overall decline in performance on the LIAR dataset highlights the persistent challenge of detecting misinformation in short, nuanced political claims. Together, these results suggest that the choice of model should depend on the nature of the input data, and hybrid approaches may be needed to achieve consistently strong performance across diverse forms of misinformation.

The sentiment analysis results further reinforced these findings. Across the three datasets, fake news showed more extreme sentiment distributions, often leaning toward negative or

fear-inducing tones. Genuine news sentiment, on the other hand, was clustered closer to neutral. This contrast was most noticeable in health-related misinformation, where emotional exaggeration was especially common. Political claims showed significantly less sentiment separation, which helped explain why models struggled more in that domain. Overall, these findings suggest that while textual features drive overall classification performance, sentiment intensity adds a vital signal, particularly in emotionally charged areas such as health misinformation.

- Dataset 1: Fake & Real News: All three models performed very well, with Logistic Regression slightly outperforming Random Forest and XGBoost (Accuracy = 0.987, F1 0.986). The high scores across all models suggest that long-form news articles provide strong linguistic signals that are easier for classifiers to separate.
- Dataset 2: CoAID (COVID-19 health news): Performance was more modest, reflecting the noisier and more varied nature of short-form health misinformation. XGBoost achieved the best overall results (Accuracy = 0.902, F1 = 0.675), driven by stronger recall compared to the other models. Logistic Regression and Random Forest achieved high precision but much lower recall, indicating that they missed a larger proportion of real items.
- Dataset 3: LIAR (short political claims): This was the most challenging dataset, with all models showing significantly lower performance. Logistic Regression again performed best (Accuracy = 0.688, F1 = 0.610), outperforming Random Forest and XGBoost. The drop across all models highlights the difficulty of detecting misinformation in short, context-dependent political statements.
- Sentiment Analysis: Our analysis revealed that fake news consistently carried more extreme sentiment (particularly negative and fear-inducing tones) across all datasets, while real news clustered more toward neutral scores (Figures 4–6). This supports prior research suggesting that emotional intensity is a key signal of misinformation and validates the integration of sentiment features into detection models.

III. I Machine learning results

The results of the experiments are presented in Tables 2, 3, and 4. In Dataset 1 (Fake & Real News), the Logistic Regression model outperformed all other models with an Accuracy of 0.987 and an F1 score of 0.986. In Dataset 2 (CoAid), the XGBoost model outperformed other models, resulting in an Accuracy of 0.986 and an F1 Score of 0.985. In Dataset 3 (LIAR Data), the Logistic Regression model performed the best, achieving an accuracy of 0.688 and an F1 score of 0.610, outperforming both Random Forest and XGBoost on this dataset. Additionally, we provide confusion matrices for each dataset, which lets us see what kind of errors the models make, not just how many. (see Figures 1, 2, and 3). We also provide examples of news misclassified by our models to gain insight into the type of data for which the models do not perform well (see Tables 5 and 6). Finally, the results of the sentiment analysis are presented in Figures 4, 5, and 6.

Dataset 1: Fake & Real News

	Accuracy	Precision	Recall	F1
Logistic Regression	0.987	0.984	0.989	0.986
Random Forest	0.986	0.982	0.988	0.985
XGboost	0.986	0.984	0.986	0.985

Table 2. Results of Logistic Regression, Random Forest, and XGBoost on dataset 1 (Fake & Real News). The best performance is highlighted in bold.

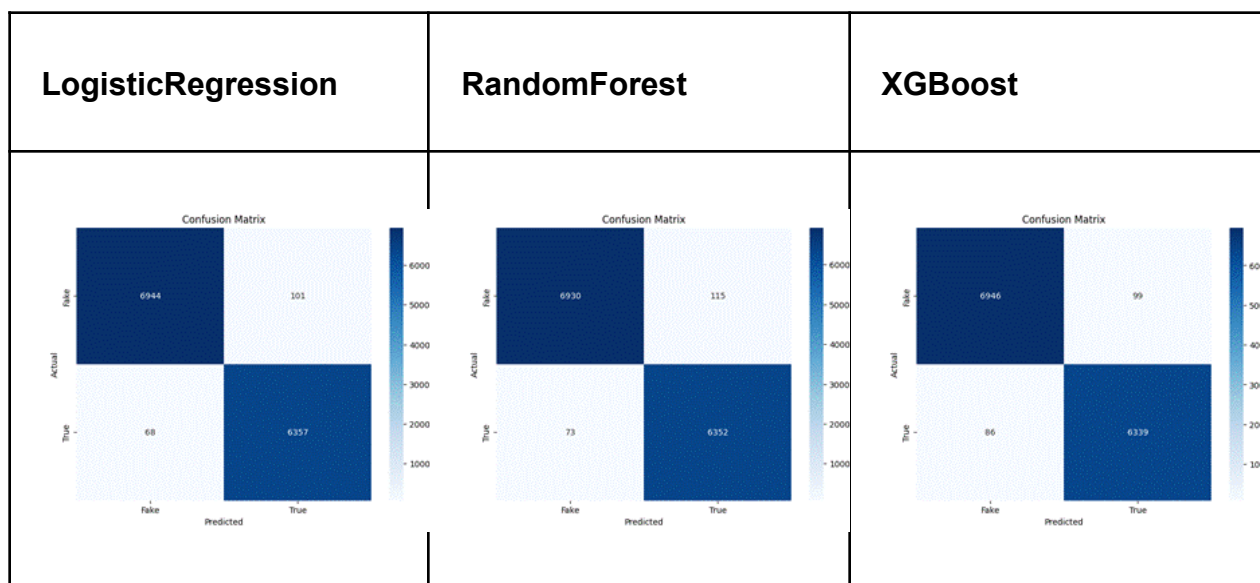


Figure 1. Confusion matrix for Logistic Regression, Random Forest, and XGboost on dataset 1 (Fake & Real News).

Dataset 2: CoAID

	Accuracy	Precision	Recall	F1
Logistic Regression	0.885	0.893	0.427	0.578
Random Forest	0.881	0.945	0.376	0.538
XGboost	0.902	0.873	0.550	0.675

Table 3: Results of Logistic Regression, Random Forest, and XGBoost on dataset 2 (CoAid).

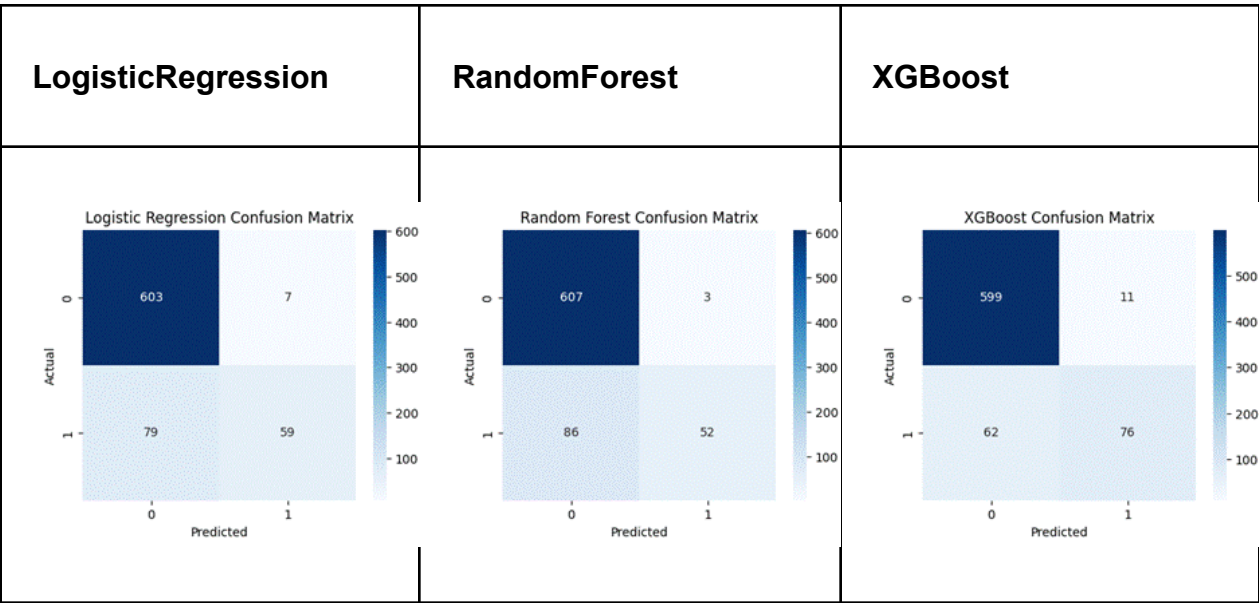


Figure 2. Confusion matrix for Logistic Regression, Random Forest, and XGboost on dataset 2 (CoAid).

Dataset 3: LIAR Data

	Accuracy	Precision	Recall	F1
Logistic regression	0.688	0.649	0.576	0.610
Random Forest	0.678	0.648	0.525	0.580
XGboost	0.669	0.646	0.487	0.555

Table 4: Results of Logistic Regression, Random Forest, and XGBoost on dataset 3 (LIAR).

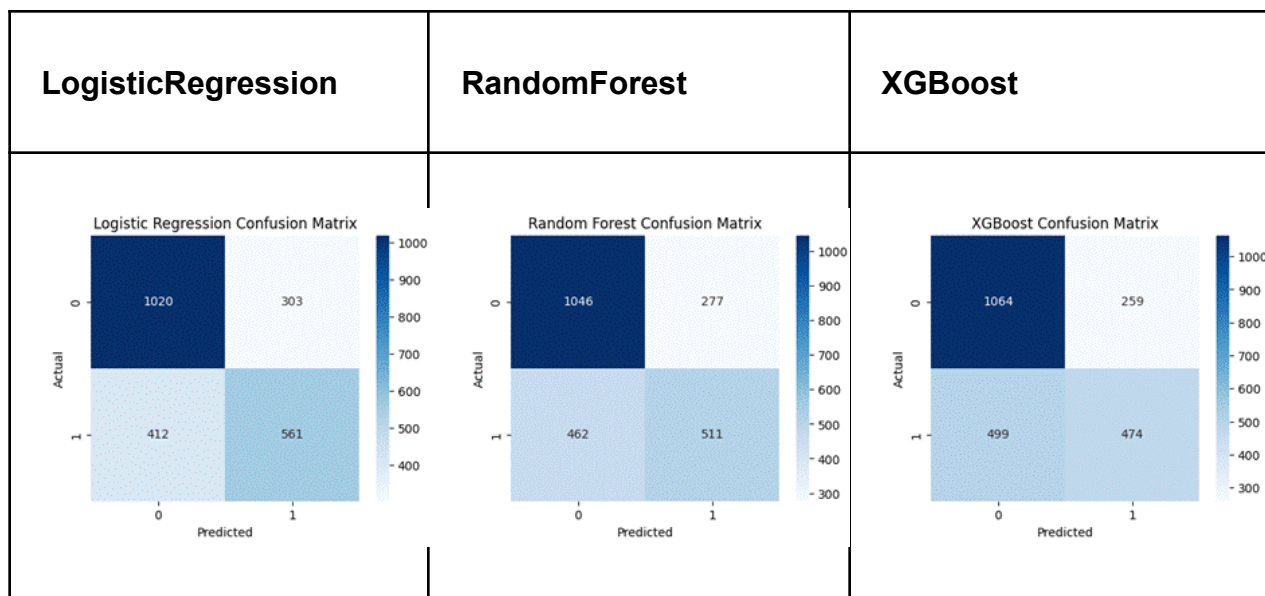


Figure 3. Confusion matrix for Logistic Regression, Random Forest, and XGboost on dataset 3 (LIAR).

Statement	true_label_name	pred_label_name
Says if the AHCA replaces Obamacare, it will "significantly reduce insurance premiums in North Carolina.	Real	Fake

Table 5: LIAR dataset Example: True statement classified as Fake

Statement	true_label_name	pred_label_name
America was the ONLY country that ended slavery	Fake	Real

Table 6: LIAR dataset Example: Fake statement classified as Real

III. II Sentiment Analysis

The sentiment scores for both real and fake news spanned the spectrum from strongly negative to strongly positive. Fake news tended to cluster more at the extreme ends of the sentiment scale, indicating stronger emotional framing. Real news was more balanced, but still shows some spread, though less exaggerated. One takeaway from these results is that fake news in long-form articles often employs extreme sentiment to attract attention, whereas real news tends to remain somewhat closer to neutral.

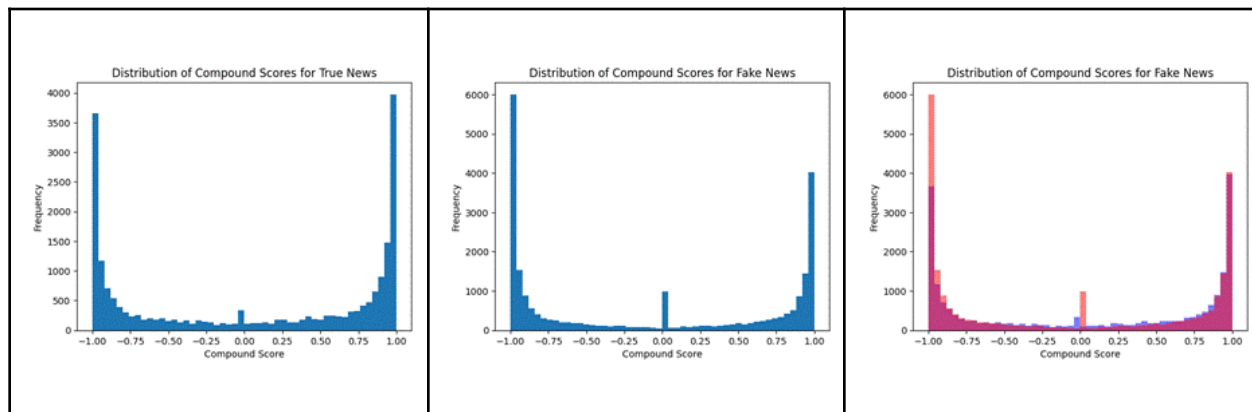


Figure 4: Sentiment distribution for Fake & Real News dataset

Both real and fake content exhibit a substantial spike around neutral sentiment; however, the distribution for fake news is more spread out, with a higher number of instances in the negative sentiment range. Real news is tightly concentrated near neutral, suggesting fact-based reporting with less emotional tone. This can be a sign that misinformation surrounding health is often more fear-inducing and negative in content compared to actual health reporting, which tends to be more neutral in tone.

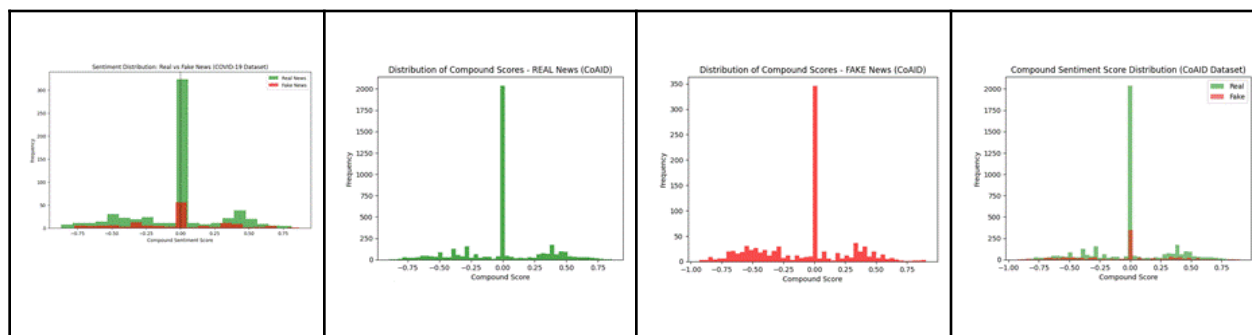


Figure 5: Sentiment distribution for CoAid dataset

Here, we observe that the overlap between fake and real distributions is significantly higher than in the health and general news datasets, and the sentiment for both real and fake claims is heavily clustered around the neutral category. This helps explain why the model's accuracy on the LIAR dataset was lower, as the short and factual style of political claims makes sentiment less discriminative.

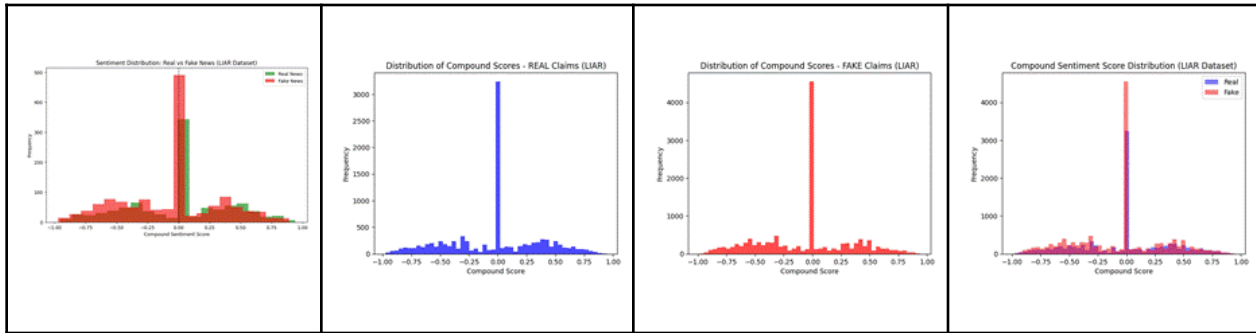


Figure 6: Sentiment distribution for the LIAR dataset

IV. Conclusion

This research shows the true potential of machine learning as a tool to fight the pervasive impact of misinformation. By analyzing three diverse datasets that encompass data from various news sources, politics, and health sources, we identified distinct patterns that distinguish real from fake content. Experimentation showed high detection accuracies, up to 98.7% on general news articles. It also highlighted challenges with data containing shorter or more context-sensitive text, like that in the LIAR dataset, where performance peaked at 68%. Further experimentation with sentiment analysis revealed disparities, with fake news amplifying negative emotions to drive engagement, as evidenced by the skewed distributions.

Prior research provided valuable insights in these areas, and by integrating our findings with them, we developed a scalable AI-powered system to detect fake news. Ultimately, our research shows excellent results from Logistic Regression on long-form data and XGBoost on a dataset with more nuanced context-aware statements. However, our work is not without limitations. First, the datasets we used are limited in their size and scope, which reduces the ability of models to generalize across platforms, languages, or even across rapidly emerging misinformation. Additionally, our approach focuses primarily on textual features, and we could expand to multimodal signals (images, videos, network propagation) to better evaluate which models are suited to each type of problem.. Further research can address these limitations and delve more deeply into issues of bias, as well as examine whether the system performs differently across demographic groups and cultures. It can also facilitate real-time detection that scales more effectively in a dynamic information environment.

V. References

1. Wardle, C., & Derakhshan, H. (2017). *Information disorder: Toward an interdisciplinary framework for research and policymaking*. Council of Europe.
<https://rm.coe.int/information-disorder-report-version-august-2018/16808c9c77>
2. Neely, S. R., Eldredge, C., & Ersing, R. (2022). Vaccine hesitancy and exposure to misinformation: A survey analysis. *Journal of General Internal Medicine*, 37(1), 179–187.
<https://doi.org/10.1007/s11606-021-07171-z>
3. Li, H. O., Bailey, A., Huynh, D., & Chan, J. (2020). YouTube as a source of information on COVID-19: A pandemic of misinformation? *BMJ Global Health*, 5(5), e002604.
<https://doi.org/10.1136/bmjgh-2020-002604>
4. Graves, L. (2018). *Understanding the promise and limits of automated fact-checking*. Oxford University Research Archive.
<https://ora.ox.ac.uk/objects/uuid:f321ff43-05f0-4430-b978-f5f517b73b9b>
5. Choudhary, A., & Arora, A. (2021). Linguistic feature based learning model for fake news detection and classification. *Expert Systems with Applications*, 169, Article 114171.
<https://doi.org/10.1016/j.eswa.2020.114171>
6. Indu, V., and Sabu M. Thampi. "Misinformation detection in social networks using emotion analysis and user behavior analysis." *Pattern Recognition Letters* 182 (2024): 60-66.
7. Huang, Y., & Chen, P. (2020). Fake news detection using an ensemble learning model based on self-adaptive harmony search algorithms. *Expert Systems with Applications*, 159, Article 113584. <https://doi.org/10.1016/j.eswa.2020.113584>
8. Rubin, V. L., Conroy, N., Chen, Y., & Cornwell, S. (2016). Fake news or truth? Using satirical cues to detect potentially misleading news. *Proceedings of the Second Workshop on Computational Approaches to Deception Detection*, 7–17.
<https://doi.org/10.18653/v1/W16-0802>
9. Ma, J., Gao, W., Mitra, P., Zhou, J., & Wong, K.-F. (2016). Detecting rumors using time-aware propagated network embeddings. *Proceedings of the 30th AAAI Conference on Artificial Intelligence (AAAI-16)*, 3042–3048. <https://doi.org/10.1609/aaai.v30i1.10310>
10. Yang, Y., Zheng, L., Zhang, J., Chen, Q., Zhao, Y., & Sun, Y. (2018). TI-CNN: Convolutional neural networks for fake news detection. *arXiv*.
<https://doi.org/10.48550/arXiv.1806.00749>
11. Wang, Y., Ma, F., Jin, Z., Yuan, Y., Xun, G., Jha, K., Su, L., & Gao, J. (2018). EANN: Event adversarial neural networks for multi-modal fake news detection. In *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery & Data Mining* (pp. 849–857). Association for Computing Machinery.
<https://doi.org/10.1145/3219819.3219892>
12. Hutto, C., and E. Gilbert. "VADER: A Parsimonious Rule-Based Model for Sentiment Analysis of Social Media Text". *Proceedings of the International AAAI Conference on Web and Social Media*, vol. 8, no. 1, May 2014, pp. 216-25, doi:10.1609/icwsm.v8i1.14550.
13. Bisailon, C. (2020). *Fake and real news dataset* [Data set]. Kaggle.
<https://www.kaggle.com/datasets/clmentbisailon/fake-and-real-news-dataset>

14. Cui, L., & Lee, D. (2020). CoAID: COVID-19 healthcare misinformation dataset. *arXiv*. <https://doi.org/10.48550/arXiv.2006.00885>
15. Wang, W. Y. (2017). "Liar, liar pants on fire": A new benchmark dataset for fake news detection. *arXiv*. <https://doi.org/10.48550/arXiv.1705.00648>
16. Sanchez, G. R., & Middlemass, K. (2022, July 26). Misinformation is eroding the public's confidence in democracy. *Brookings Institution*. <https://www.brookings.edu/articles/misinformation-is-eroding-the-publics-confidence-in-democracy/>
17. Yasir, M., & Uwishema, O. (2021). Ebola outbreak amid COVID-19 in the Republic of Guinea: Priorities for achieving control. *The American Journal of Tropical Medicine and Hygiene*, 105 (2), 287–289. <https://doi.org/10.4269/ajtmh.21-0228>
18. Vosoughi, S., Roy, D., & Aral, S. (2018). The spread of true and false news online. *Science*, 359 (6380), 1146–1151. <https://doi.org/10.1126/science.aap9559>
19. Shu, K., Sliva, A., Wang, S., Tang, J., & Liu, H. (2017). Fake news detection on social media: A data mining perspective. *ACM SIGKDD Explorations Newsletter*, 19 (1), 22–36. <https://doi.org/10.1145/3137597.3137600>
20. Skafle, I., Nordahl-Hansen, A., Steinsbekk, S., & Engebretsen, E. (2022). Misinformation about COVID-19 vaccines on social media: Rapid review. *Journal of Medical Internet Research*, 24 (8), Article e37367. <https://doi.org/10.2196/37367>
21. Cui, L., & Lee, D. (2020). CoAID: COVID-19 healthcare misinformation dataset. *arXiv*. <https://doi.org/10.48550/arXiv.2006.00885>
22. Gifu, D. (2023). An intelligent system for detecting fake news. *Procedia Computer Science*, 221, 1058–1065. <https://doi.org/10.1016/j.procs.2023.08.088>