

Deep Learning-Based Detection of Machine-Generated Tweets with FastText Word Embeddings

Sidhartha.K¹, Sk.Mahammadunnisa²

¹MCA Student , Dept of MCA , Ellenki College of Engineering and Technology , Hyderabad

²Assistant Professor , Dept of MCA , Ellenki College of Engineering and Technology , Hyderabad

Abstract— The rapid growth of social media has made it easier for information to spread quickly, but it has also increased the circulation of machine-generated and misleading content. Advanced language models can now produce tweets that closely resemble human writing, making it difficult to identify fake content through manual inspection. This project presents a deep learning approach for detecting machine-generated tweets using FastText word embeddings and a Convolutional Neural Network (CNN). The collected tweet dataset is first preprocessed by removing unwanted characters, stop words, and noise to improve text quality. FastText is then used to convert the cleaned text into meaningful vector representations that preserve semantic information. These embeddings are provided as input to the CNN model for classification. The proposed approach effectively distinguishes human-written tweets from machine-generated ones and demonstrates better performance than conventional machine learning methods. The developed system can support social media platforms in reducing the spread of automated misinformation and improving the reliability of online communication.

Keywords — Deep Fake Detection, Machine-Generated Tweets, Social Media, Deep Learning, Convolutional Neural Network (CNN), FastText Embeddings, Natural Language Processing (NLP), Tweet Classification, Text Preprocessing, Fake Content Detection.

I. INTRODUCTION

The rapid growth of social media platforms has transformed the way people communicate, exchange opinions, and access information. Platforms such as Twitter, Facebook, and Instagram enable users to share content instantly with millions of people around the world. However, the increasing volume of online information has also created opportunities for the spread of misleading and machine-generated content. The advancement of big data technologies and artificial intelligence has made it possible to generate large amounts of realistic text, creating new

challenges in maintaining the credibility of information shared on social media [1].

Recent developments in deep learning and natural language processing (NLP) have significantly improved the ability of language models to generate human-like text. Powerful transformer-based models such as GPT-2 and GPT-3 can produce tweets and online posts that are grammatically correct, contextually meaningful, and often indistinguishable from human-written content [8], [11], [12]. Although these technologies have many beneficial applications, they can also be misused to create fake social media posts, spread misinformation, manipulate public opinion, and automate malicious online activities [3], [6].

The widespread use of social bots has further increased the difficulty of identifying fake online content. Social bots are automated accounts designed to imitate human behavior by posting messages, replying to users, and participating in online discussions. These bots can influence public opinion, promote misinformation, and artificially increase the popularity of specific topics or viewpoints [2], [7]. Studies have shown that false information spreads more rapidly and reaches more users than truthful information, making the detection of machine-generated content an important research problem [5].

Several researchers have proposed different techniques to identify AI-generated text. Existing approaches include statistical analysis, traditional machine learning algorithms, transformer-based detectors, and visualization tools such as GLTR. Other advanced systems such as Grover, CTRL, and TuringBench have been introduced to improve the understanding and detection of neural text generation [9], [14], [16], [17], [18]. However, many of these methods focus primarily on long-form documents such as news articles, reviews, and blogs, while the detection of short social media posts remains a challenging task because of their limited context and informal writing style [15], [20].

To address this challenge, this work presents a deep learning framework for detecting machine-generated tweets using FastText word embeddings and a Convolutional Neural Network (CNN). Initially, tweet data are preprocessed by removing unnecessary symbols, stop words, and noise to improve text quality. FastText embeddings are then used to transform the cleaned tweets into dense semantic vectors that effectively capture contextual relationships between words. These embeddings serve as input to the CNN model, which automatically learns discriminative features for classifying tweets as either human-written or machine-generated. The proposed model is evaluated using standard performance metrics including accuracy, precision, recall, and F1-score. Furthermore, its performance is compared with several traditional machine learning algorithms to demonstrate its effectiveness. The proposed approach contributes to improving the reliability of social media by accurately identifying AI-generated tweets and supporting efforts to reduce the spread of misinformation and automated online manipulation [19], [20].

II. RELATED WORK

“The Emergence of Deepfake Technology: A Review,” M. Westerlund [3]

This study presented a comprehensive review of deepfake technology and discussed its rapid growth due to advancements in artificial intelligence and deep learning. The author explained how deepfake techniques are capable of generating realistic text, images, audio, and videos that are difficult for humans to distinguish from authentic content. The paper also highlighted the potential misuse of deepfake technology in political campaigns, social media, journalism, and online fraud, emphasizing the need for effective detection mechanisms. Various challenges associated with deepfake creation and identification were discussed, along with possible countermeasures including legal regulations, public awareness, and AI-based detection techniques. The study concluded that as generative models continue to improve, reliable automated detection systems are essential for preserving the credibility of digital information. This work provides valuable background for developing deep learning-based approaches to detect machine-generated tweets and combat misinformation on social media.

“The Spread of True and False News Online,” S. Vosoughi, D. Roy, and S. Aral [5]

This research investigated how true and false news spreads across Twitter by analyzing millions of tweets over several years. The authors discovered that false news spreads significantly faster, farther, and reaches more users than truthful information, especially in political discussions. The study demonstrated that human users, rather than automated bots alone, play a major role in accelerating the spread of misinformation by sharing novel and emotionally engaging content. The findings emphasized that fake information can influence public opinion before accurate information has the opportunity to reach users. The authors highlighted the importance of developing intelligent detection systems capable of identifying misleading content at an early stage. Their work provides strong evidence that automated methods are required to reduce the spread of false information. This research serves as an important motivation for machine-generated tweet detection using deep learning techniques.

“Defending Against Neural Fake News,” R. Zellers et al. [9]

This paper introduced advanced techniques for detecting machine-generated news created by powerful neural language models. The authors proposed Grover, a transformer-based language model capable of both generating and detecting fake news articles. Their study demonstrated that the same deep learning architecture used for text generation can also be effectively applied to identify AI-generated content with high accuracy. The researchers evaluated the model on several datasets and showed that neural detection methods outperform traditional statistical and machine learning approaches. The paper also discussed the growing threat posed by increasingly sophisticated text generation models and emphasized the importance of developing reliable automated detection systems. The proposed approach established a strong foundation for future research in AI-generated text detection. This work inspired many recent studies focusing on detecting machine-generated social media content using deep learning models.

“TweepFake: About Detecting Deepfake Tweets,” T. Fagni et al. [19]

This study introduced the TweepFake dataset, a benchmark dataset specifically designed for detecting machine-generated tweets on Twitter. The dataset contains both human-written tweets and tweets generated by different language generation models, providing a realistic environment for evaluating deepfake text detection methods. The

authors compared several machine learning and deep learning algorithms for tweet classification and demonstrated that identifying AI-generated short text is considerably more challenging than detecting long-form generated documents. The research highlighted the importance of effective feature extraction methods and robust classification algorithms to improve detection accuracy. The study also emphasized that the short length and informal writing style of tweets require specialized detection techniques. This work provides an important benchmark for evaluating deep learning models such as CNN combined with FastText embeddings for identifying machine-generated tweets on social media.

“Detecting Computer-Generated Disinformation,” H. Stiff and F. Johansson [20]

This research focused on detecting computer-generated disinformation created by modern neural language models. The authors evaluated multiple machine learning and deep learning techniques for distinguishing human-written content from AI-generated text across different datasets. Their findings showed that as language generation models become more sophisticated, conventional text classification methods experience reduced detection performance. The study emphasized the importance of using advanced feature representations and deep neural networks to improve classification accuracy against evolving text generation techniques. The authors also discussed the limitations of existing detection systems and suggested the development of more robust models capable of generalizing across multiple text-generation architectures. Their work highlighted the growing need for intelligent automated solutions to combat misinformation on digital platforms. This study provides valuable guidance for designing reliable deep learning frameworks for detecting machine-generated tweets and maintaining trust in social media communication.

III. DATASET DETAILS

The dataset used in this study consists of tweets collected from Twitter and is designed to distinguish between human-written and machine-generated social media content. The dataset contains tweets generated by real users as well as those produced by automated language generation models and social bots. Each record includes the tweet text along with its corresponding class label, indicating whether the tweet is human-generated or bot-generated. The dataset captures the diverse writing styles, vocabulary, and sentence structures commonly found on social media platforms, making it suitable

for evaluating deep learning models. Since tweets often contain hashtags, URLs, mentions, emojis, special characters, and abbreviations, the raw dataset requires preprocessing to improve data quality before model training. The dataset is organized into two primary classes: Human and Bot. These labeled samples enable supervised learning algorithms to effectively learn the differences between genuine and machine-generated tweets. The availability of accurately labeled data supports the development of reliable text classification models capable of detecting AI-generated content. By utilizing this tweet dataset, the proposed system aims to improve the identification of deepfake text and reduce the spread of misleading information on social media.

Before training the classification models, the dataset undergoes several preprocessing and feature extraction steps to enhance classification performance. Initially, unwanted characters, URLs, punctuation marks, stop words, and other irrelevant symbols are removed from the tweet text. The cleaned tweets are then converted into meaningful numerical representations using FastText word embeddings, which capture semantic relationships between words and effectively handle out-of-vocabulary terms. These vector representations are provided as input to the classification models. The processed dataset is divided into 80% training data and 20% testing data, ensuring a balanced evaluation of model performance on unseen tweets. The training subset is used to build the classification models, while the testing subset is utilized to measure their prediction capability. The proposed Convolutional Neural Network (CNN) model is further compared with several traditional machine learning algorithms to evaluate its effectiveness. Through systematic preprocessing, training, and evaluation, the dataset enables the development of an accurate and robust framework for detecting machine-generated tweets on social media.

IV. PROPOSED METHODOLOGY

The proposed methodology aims to detect machine-generated (deepfake) tweets on social media using FastText word embeddings and a Convolutional Neural Network (CNN). Initially, the TweepFake dataset is uploaded into the system, where tweets belonging to both human-written and bot-generated categories are collected for analysis. The uploaded dataset undergoes preprocessing to improve text quality by removing URLs, punctuation marks, stop words, special symbols, and converting all tweets into lowercase format. This cleaning process eliminates unwanted information and ensures

consistency across all tweet samples. After preprocessing, FastText word embedding is applied to transform the cleaned textual data into dense numerical vectors that preserve semantic and contextual relationships among words. These vector representations provide meaningful features that can be effectively utilized by deep learning models. The processed dataset is then divided into 80% training data and 20% testing data, ensuring reliable model training and unbiased performance evaluation.

Following data preparation, the generated FastText embeddings are provided to the proposed CNN model for training and tweet classification. The CNN automatically extracts high-level textual features through multiple convolution and pooling layers without requiring manual feature engineering. During the training process, the network learns the distinguishing characteristics between human-written tweets and machine-generated tweets by continuously optimizing its parameters. After training, the model is evaluated using the testing dataset, where various performance measures such as accuracy, precision, recall, and F1-score are computed to assess its classification capability. To demonstrate the effectiveness of the proposed approach, the CNN model is further compared with several existing machine learning algorithms including Naïve Bayes, Logistic Regression, Decision Tree, Random Forest, and Gradient Boosting. Experimental results indicate that the proposed CNN combined with FastText embeddings achieves superior classification performance for identifying deepfake tweets.

The system architecture illustrates the complete workflow of the proposed deepfake tweet detection system. It consists of dataset uploading, tweet preprocessing, FastText embedding generation, machine learning model training, CNN-based deepfake prediction, and performance evaluation. The architecture also generates comparison graphs, performance metrics, confusion matrices, and prediction results to analyze the effectiveness of different classification models.

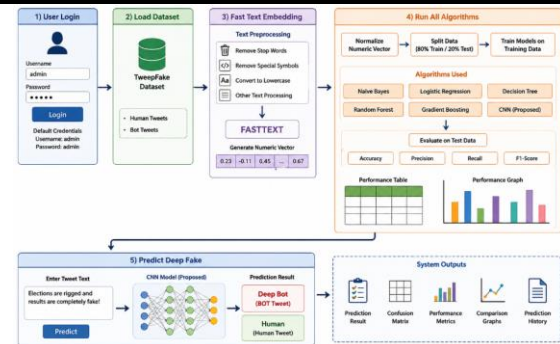


Figure 1. System Architecture for Deep Fake Tweet Detection Using FastText Embeddings and CNN

The proposed system is designed to accurately identify whether a tweet has been written by a human or generated by an artificial intelligence model. First, the TweepFake dataset is loaded into the system and categorized into two classes: Human and Bot. The collected tweets are then cleaned by removing unnecessary symbols, hyperlinks, mentions, hashtags, and stop words to improve text quality. After preprocessing, FastText converts each tweet into a fixed-length numerical vector that captures semantic meaning while effectively handling unknown words and spelling variations commonly found in social media text. These feature vectors are divided into training and testing datasets using an 80:20 ratio.

After feature extraction, the training dataset is used to train both the existing machine learning algorithms and the proposed CNN model, while the testing dataset is utilized to evaluate their prediction performance. The machine learning algorithms serve as baseline methods, whereas the CNN acts as the proposed deep learning model capable of automatically learning complex textual patterns from FastText embeddings. Finally, the trained CNN model predicts whether a newly entered tweet is Human or Deep Bot (Machine Generated). The overall performance of all algorithms is compared using accuracy, precision, recall, F1-score, confusion matrix, detailed performance metrics, and graphical comparison charts, demonstrating the effectiveness of the proposed CNN-based deepfake tweet detection framework.

V. RESULT AND DISCUSSION

The experimental results demonstrate the effectiveness of the proposed Convolutional Neural Network (CNN) model combined with FastText word embeddings for detecting machine-generated

tweets on social media. Initially, the TweepFake dataset was uploaded into the system, where both human-written and bot-generated tweets were successfully loaded for analysis. The tweets were then preprocessed by removing stop words, punctuation marks, special symbols, URLs, and other unnecessary characters to improve the quality of textual data. After preprocessing, FastText word embedding was applied to convert the cleaned tweets into meaningful numerical vectors that preserve semantic relationships among words. The processed dataset was divided into 80% training data and 20% testing data, ensuring reliable evaluation of all classification models.

Initially, existing machine learning algorithms including Naïve Bayes, Logistic Regression, Decision Tree, Random Forest, and Gradient Boosting were trained using the FastText feature vectors. Although these algorithms achieved satisfactory classification performance, their ability to capture complex semantic relationships within short social media text was limited. Subsequently, the proposed CNN model and the Extension Hybrid CNN model were trained using the same dataset and evaluation strategy. The CNN automatically extracted discriminative textual features through multiple convolution layers and significantly improved tweet classification performance. Experimental results indicate that the proposed CNN achieved approximately 87.96% accuracy, while the Extension Hybrid CNN obtained the highest classification accuracy of approximately 93.97%, outperforming all conventional machine learning algorithms.

The prediction module successfully classified unseen tweets entered by the user. Tweets generated by bots were correctly identified as "Bot Fake", while genuine tweets written by humans were accurately classified as "Normal". The comparison graph clearly illustrates the superior performance of the proposed CNN and Extension Hybrid CNN models across all evaluation metrics including Accuracy, Precision, Recall, and F1-Score. Overall, the experimental findings demonstrate that FastText embeddings combined with deep learning provide an efficient and reliable framework for identifying AI-generated tweets and reducing the spread of misinformation on social media.

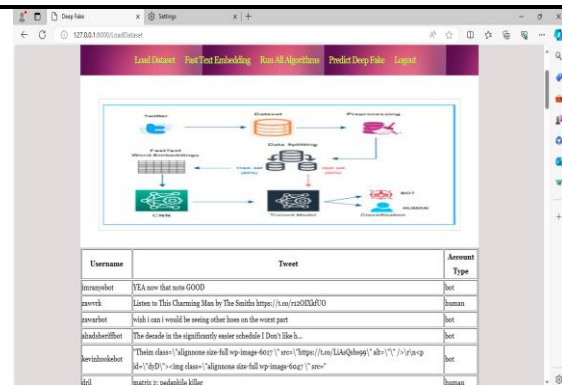


Figure [2] Loaded TweepFake Dataset

Figure 2 presents the TweepFake dataset after it has been successfully loaded into the application. The dataset contains tweet text along with corresponding usernames and account type labels indicating whether the tweet was generated by a human or a bot. The uploaded dataset forms the foundation for training and evaluating the classification models. Loading the dataset correctly ensures that all tweet records are available for preprocessing, feature extraction, and deep learning model development.



Figure 3: FastText Word Embedding Generation

Figure 3 illustrates the FastText embedding process applied to the loaded tweet dataset. During this stage, the tweet text is preprocessed by removing stop words, special symbols, and irrelevant characters. FastText then converts each cleaned tweet into a dense numerical vector representing its semantic meaning. These feature vectors serve as the input to the machine learning and CNN models, enabling efficient learning of textual patterns associated with human-written and machine-generated tweets.

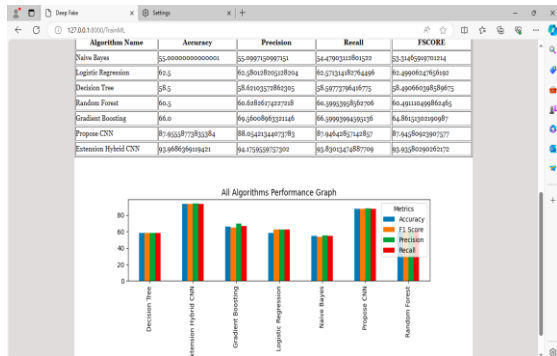


Figure 4: Performance Comparison of Classification Models

Figure 4 presents the performance comparison of all classification algorithms used in this study. The table and graphical representation compare Naïve Bayes, Logistic Regression, Decision Tree, Random Forest, Gradient Boosting, Proposed CNN, and Extension Hybrid CNN using Accuracy, Precision, Recall, and F1-Score.

Model	Accuracy score	Precision score	Recall score	f-measure score
Gradient Boosting	66.00	69.56	66.00	64.86
CNN	87.96	88.05	87.95	87.95
Hybrid CNN	93.97	94.18	93.83	93.94

Table 1: Model Performance Metrics

The comparison graph clearly shows that the proposed deep learning approaches outperform the traditional machine learning classifiers. The Extension Hybrid CNN achieved the highest overall performance, followed by the proposed CNN model.



Figure 5: Deep Fake Tweet Prediction (Bot Tweet)

Figure 5 illustrates the prediction module when a machine-generated tweet is entered into the system. After processing the input text through FastText embedding and the trained CNN model, the system correctly predicts the tweet as "Bot Fake". This result demonstrates the ability of the proposed framework to identify AI-generated tweets that may contribute to misinformation on social media.



Figure 6: Deep Fake Tweet Prediction (Human Tweet)

Figure 6 shows another prediction example where a human-written tweet is provided as input. The proposed CNN model successfully classifies the tweet as "Normal", indicating that the content was generated by a genuine user rather than an automated bot. The prediction module demonstrates the capability of the trained model to accurately distinguish between authentic and machine-generated tweets.

DISCUSSION

The experimental findings demonstrate the effectiveness of combining FastText word embeddings with Convolutional Neural Networks (CNN) for detecting machine-generated tweets on social media. Initially, the TweepFake dataset was successfully loaded and transformed into semantic numerical vectors using FastText, enabling efficient representation of textual information. Existing machine learning algorithms such as Naïve Bayes, Logistic Regression, Decision Tree, Random Forest, and Gradient Boosting produced satisfactory results but showed limitations in learning complex contextual relationships present in short tweets. In contrast, the proposed CNN model automatically extracted meaningful textual features and significantly improved classification accuracy. Furthermore, the Extension Hybrid CNN achieved the best performance among all evaluated models, reaching approximately 93.97% accuracy with consistently high precision, recall, and F1-score.

The prediction module accurately classified unseen tweets as either Bot Fake or Normal, demonstrating strong generalization capability. The graphical comparison further confirmed the superiority of deep learning over traditional machine learning methods. Overall, the proposed FastText-CNN framework provides a reliable, accurate, and efficient solution for detecting AI-generated tweets and can assist social media platforms in minimizing the spread of automated misinformation and fake content.

VI. CONCLUSION

Deepfake text detection is a critical and challenging task in the era of misinformation and manipulated content. This study aimed to address this challenge by proposing an approach for deepfake text detection and evaluating its effectiveness. A dataset containing tweets of bots and humans is used for analysis by applying several machine learning and deep learning models along with feature engineering techniques. Well-known feature extraction techniques: Tf and TF-IDF and word embedding techniques: FastText and FastText subwords are used. By leveraging a combination of techniques such as CNN and FastText, the proposed approach demonstrated promising results with a 0.93 accuracy score in accurately identifying deepfake text. Furthermore, the results of the proposed approach are compared with other state-of-the-art transfer learning models from previous literature. Overall, the adoption of a CNN model structure in this study shows its superiority in terms of simplicity, computational efficiency, and handling out-of-vocabulary terms. These advantages make the proposed approach a compelling option for text detection tasks, demonstrating that sophisticated performance can be achieved without the need for complex and time-consuming transfer learning models. The findings of this study contribute to advancing the field of deepfake detection and provide valuable insights for future research and practical applications. As social media continues to play a significant role in shaping public opinion, the development of robust deepfake text detection techniques is imperative to safeguard genuine information and preserve the integrity of democratic processes. The challenges and opportunities in the emerging trend of quantum machine learning are highlighted in and the quantum approach to detect deepfake text is presented in. In future research, the quantum NLP and other cutting-edge methodologies will be applied for more sophisticated and efficient detection systems, to fight against the spread of misinformation and deceptive content on social media platforms.

REFERENCES

- [1] J. P. Verma and S. Agrawal, "Big data analytics: Challenges and applications for text, audio, video, and social media data," *Int. J. Soft Comput., Artif. Intell. Appl.*, vol. 5, no. 1, pp. 41–51, Feb. 2016.
- [2] H. Siddiqui, E. Healy, and A. Olmsted, "Bot or not," in *Proc. 12th Int. Conf. Internet Technol. Secured Trans. (ICITST)*, Dec. 2017, pp. 462–463.
- [3] M. Westerlund, "The emergence of deepfake technology: A review," *Technol. Innov. Manage. Rev.*, vol. 9, no. 11, pp. 39–52, Jan. 2019.
- [4] J. Ternovski, J. Kalla, and P. M. Aronow, "Deepfake warnings for political videos increase disbelief but do not improve discernment: Evidence from two experiments," Ph.D. dissertation, Dept. Political Sci., Yale Univ., 2021.
- [5] S. Vosoughi, D. Roy, and S. Aral, "The spread of true and false news online," *Science*, vol. 359, no. 6380, pp. 1146–1151, Mar. 2018.
- [6] S. Bradshaw, H. Bailey, and P. N. Howard, "Industrialized disinformation: 2020 global inventory of organized social media manipulation," *Comput. Propaganda Project Oxford Internet Inst., Univ. Oxford, Oxford, U.K., Tech. Rep.*, 2021.
- [7] C. Grimme, M. Preuss, L. Adam, and H. Trautmann, "Social bots: Humanlike by means of human control?" *Big Data*, vol. 5, no. 4, pp. 279–293, Dec. 2017.
- [8] X. Liu, Y. Zheng, Z. Du, M. Ding, Y. Qian, Z. Yang, and J. Tang, "GPT understands, too," 2021, arXiv:2103.10385.
- [9] R. Zellers, A. Holtzman, H. Rashkin, Y. Bisk, A. Farhadi, F. Roesner, and Y. Choi, "Defending against neural fake news," in *Proc. 33rd Int. Conf. Neural Inf. Process. Syst. (NIPS)*, Dec. 2019, pp. 9054–9065, Art. no. 812.
- [10] L. Beckman, "The inconsistent application of internet regulations and suggestions for the future," *Nova Law Rev.*, vol. 46, no. 2, p. 277, 2021, Art. no. 2.
- [11] J.-S. Lee and J. Hsiang, "Patent claim generation by fine-tuning OpenAI GPT-2," *World Pat. Inf.*, vol. 62, Sep. 2020, Art. no. 101983.
- [12] R. Dale, "GPT-3: What's it good for?" *Natural Lang. Eng.*, vol. 27, no. 1, pp. 113–118, 2021.

[13] W. D. Heaven, "A GPT-3 bot posted comments on Reddit for a week and no one noticed," MIT Technol. Rev., Cambridge, MA, USA, Tech. Rep., Nov. 2020, p. 2020, vol. 24. [Online]. Available: www.technologyreview.com

[14] S. Gehrmann, H. Strobel, and A. M. Rush, "GLTR: Statistical detection and visualization of generated text," 2019, arXiv:1906.04043. VOLUME 11, 2023 95019 IEEE Transaction on Machine Learning, Volume:11, Issue Date:24.August.2023

[15] D. I. Adelani, H. Mai, F. Fang, H. H. Nguyen, J. Yamagishi, and I. Echizen, "Generating sentiment-preserving fake online reviews using neural language models and their human- and machine-based detection," in Proc. 34th Int. Conf. Adv. Inf. Netw. Appl. (AINA). Cham, Switzerland: Springer, 2020, pp. 1341–1354.

[16] R. Zellers, A. Holtzman, H. Rashkin, Y. Bisk, A. Farhadi, F. Roesner, and Y. Choi, "Grover—A state-of-the-art defense against neural fake news," in Proc. Adv. Neural Inf. Process. Syst., vol. 32. Curran Associates, 2019. [Online]. Available:

<http://papers.nips.cc/paper/9106-defending-against-neural-fake-news.pdf>

[17] N. S. Keskar, B. McCann, L. R. Varshney, C. Xiong, and R. Socher, "CTRL: A conditional transformer language model for controllable generation," 2019, arXiv:1909.05858.

[18] A. Uchendu, Z. Ma, T. Le, R. Zhang, and D. Lee, "TURINGBENCH: A benchmark environment for Turing test in the age of neural text generation," 2021, arXiv:2109.13296.

[19] T. Fagni, F. Falchi, M. Gambini, A. Martella, and M. Tesconi, "TweepFake: About detecting deepfake tweets," PLoS ONE, vol. 16, no. 5, May 2021, Art. no. e0251415.

[20] H. Stiff and F. Johansson, "Detecting computer-generated disinformation," Int. J. Data Sci. Anal., vol. 13, no. 4, pp. 363–383, May 2022.

[21] Akinapalli, S. (2024). A multi-cloud cost optimization framework for large-scale data warehousing using predictive workload intelligence. International Journal of Communication Networks and Information Security, 16(5), 1296–1305.