

Real-Time Fake News Detection Using Machine Learning

Manoj Kumar M and G Prasanna David

Department of MCA

Vidya Vikas Institute of Engineering and Technology, Mysuru, India
manojkumarmysoreindia@gmail.com, prasanna.david@vidyavikas.edu.in

Abstract: *In the era of rising digital media, fake news has been constituting grave risk factors to public opinion, democracy and trust on authentic sources. This work tackles the problem with sequence-based deep learning and standard NLP, training LSTM, stacked LSTM, and bidirectional LSTM models on the WELFake corpus of labeled real and fake articles. All inputs pass through a consistent preparation stage—cleaning, removal of common stopwords, stemming or lemmatization as appropriate, tokenization, and padding—so the models see stable, comparable sequences. In side-by-side tests with classical baselines such as Naive Bayes and Logistic Regression, the bidirectional LSTM delivered stronger accuracy, while the best overall results came from pairing Word2Vec embeddings with an LSTM. For day-to-day use, a small Flask web app lets users submit individual headlines or stream live items via APIs, and the service returns a label with a confidence score. By combining robust deep learning with a simple interface, the platform offers an efficient, near real-time aid for countering misinformation.*

Keywords: Fake News Detection, Deep Learning, Natural Language Processing (NLP), LSTM, Bidirectional LSTM, Stacked LSTM, Transfer Learning, Word2Vec, Text Classification, Flask Web Application

I. INTRODUCTION

Over the past decade, the news pipeline has shifted to a real-time, platform-driven cadence in which posts, reactions, and shares arrive faster than editors or community moderators can verify them, widening the gap between publication and validation. The same channels that make information broadly accessible—social media feeds, blogs, and web portals—also act as accelerants for misleading material. In this work, fake news means stories that look like journalism but are either made up or deliberately twisted to push a point of view, sway elections, make money, or damage someone's reputation. The stakes are real: organized misinformation can tilt how people participate in democracy, spark social friction, and weaken trust in institutions that rely on public confidence.

Manual review and community fact-checks are still important, but the sheer speed and volume of online posts outstrip what people can vet in time. Engagement-driven feeds often boost eye-catching claims long before anyone can verify them, so an automated, scalable first-pass screen is needed to keep up with live streams. Engagement-optimized ranking systems frequently elevate provocative claims before they can be examined, compounding the exposure window. These realities have created a strong practical need for tools that can assess large volumes of text automatically and deliver timely, reliable triage—ideally fast enough to be useful in live settings where propagation happens in minutes, not days. Within text classification, long-standing machine learning baselines—Support Vector Machines, Naive Bayes, and Logistic Regression—have been studied extensively and retain advantages such as transparency and speed. However, because they depend heavily on engineered indicators like word counts or sentiment proxies, they struggle when meaning depends on discourse-level cues, cross-sentence structure, or subtle rhetorical shifts designed to evade surface detectors. As the language of deception evolves, these feature-centric approaches tend to lose robustness, particularly across diverse topics and changing styles.

Instead of depending mainly on hand-built features, modern neural models learn the signal and its context from the data itself. Long Short-Term Memory networks—along with deeper stacked setups and bidirectional variants—are built for



sequences, so they can follow links that run across phrases and sentences. By picking up both close-by cues and broader discourse patterns, they more reliably tell solid reporting from manipulative wording, covering gaps that simpler pipelines miss.

This study adopts the WELFake corpus—over seventy-two thousand labeled items spanning politics, health, business, and technology—to evaluate models under conditions that reflect real-world diversity. The pipeline applies standardization steps such as cleaning, stemming or lemmatization, stopword handling, tokenization, and padding to create stable inputs for neural sequence models. To strengthen semantic grounding, Word2Vec-based transfer learning is incorporated so that the classifier benefits from pretrained distributional structure before adapting to the specifics of the fake–real decision boundary.

Experiments compare classical baselines to LSTM, stacked LSTM, and BiLSTM configurations fed with pretrained embeddings and then fine-tuned; the transfer-learning LSTM attains the highest accuracy, with balanced precision and recall, indicating that context-sensitive neural representations handle the challenges of deception better than feature-engineered methods. For day-to-day use, the chosen model runs behind a small Flask service that takes either a headline or a full article and returns a label with a confidence score almost instantly. In practice, standardized preprocessing, sequence-aware modeling, and a clean delivery layer work together as one pipeline, moving a research-grade approach into real-time use for countering misinformation.

II. LITERATURE SURVEY

[1] S. Garg and D. K. Sharma, “Fake News Classification via CNN,” Proc. SMART, 2022. Garg and Sharma investigate convolutional neural networks (CNNs) for text-based fake news classification, adapting CNN architectures—often used in vision—to exploit local n-gram patterns and phrase-level features in headlines and short articles. They convert text into dense embeddings and use multiple convolutional kernels to capture features at different granularities, followed by pooling and levels that are entirely related for classification. Experiments on benchmark datasets show that CNNs are highly effective for short-text inputs (e.g., headlines), achieving competitive accuracy while being computationally efficient. However, the authors note limitations in modeling long-range dependencies where recurrent and transformer-based models may perform better.

[2] A. Oad et al., “Fake News Classification Methodology With Enhanced BERT,” IEEE Access, 2024. Oad et al. present an enhanced BERT-based approach tailored for fake news detection. The paper fine-tunes a BERT encoder with additional context-aware layers and refined attention mechanisms to capture subtle deceptive cues in news text. They address augmentation and domain-specific fine-tuning, showing that enhanced BERT variants outperform classical RNN and CNN baselines on multiple datasets, particularly for nuanced misinformation. The work underscores the advantage of transformer-based models in modeling long-range context, while also noting the trade-off of higher computational requirements.

[3] Li and Liu (2023) present the FUND corpus, enlarging both the breadth and scale of resources available for fake-news research across politics, health, finance, and entertainment, and packaging each item with useful metadata like source and timestamp. Their experiments show that feeding such metadata into models improves generalization and robustness, not just headline accuracy. The work underscores a broader lesson: outcomes depend heavily on dataset quality—balanced labels, cross-domain coverage, and careful annotation are all critical for dependable evaluation and model development.

[4] Nath, Soni, Ahuja, and Katarya (2021) conduct a side-by-side evaluation of classic machine-learning classifiers—Naive Bayes, SVM, Logistic Regression, and Random Forest—against deeper neural approaches such as CNNs and LSTMs for fake-news detection. Their study also probes the impact of common preprocessing pipelines, contrasting TF-IDF features with embedding-based representations. The authors report that lightweight traditional models are a practical fit for smaller datasets because they are inexpensive to train and serve, whereas deep models tend to lead on larger, context-rich corpora where sequence and semantics matter more. They further point out recurring, real-world hurdles: the manual effort of feature design, skewed class distributions that complicate training and evaluation, and the compute and memory demands that can limit deployment at scale.



[5] Alnabhan and Branco (2024) survey the deep-learning landscape for fake-news detection, organizing prior work into four broad families: RNN-based models, CNN-style architectures, transformer stacks, and hybrids that blend multiple components. A key takeaway from their review is the steady shift toward transformer encoders such as BERT and RoBERTa, largely because these models capture long-range context and discourse cues that simpler networks tend to miss. It also points out limitations such as dataset imbalance, lack of multilingual support, and challenges in detecting satire or opinion-based misinformation. The authors suggest future directions including explainable AI and multimodal approaches.

[6] J. Reddy, S. Mundra, and A. Mundra, "Ensembling Deep Learning Models for Fake News Classification," *Procedia Computer Science*, 2024. Reddy et al. propose ensemble strategies such as soft voting, stacking, and weighted averaging that combine CNN, LSTM, and GRU models. Their experiments show that ensemble models consistently outperform individual architectures, achieving gains of 2–4% in accuracy and F1-score. They emphasize that model diversity and careful weighting are key to maximizing ensemble performance. The study demonstrates the potential of hybrid models to deliver stronger generalization in real-world fake news detection tasks.

[7] M. S. A. Alzaidi et al., "An Efficient Fusion Network for Fake News Classification," *Mathematics*, 2024. Alzaidi et al. propose an Efficient Fusion Network (EFN) that integrates CNNs for local feature extraction, BiGRUs for contextual learning, and attention mechanisms for prioritizing critical parts of text. This hybrid network improves feature representation and classification accuracy. Experiments on LIAR and ISOT datasets show that EFN achieves state-of-the-art performance, surpassing traditional baselines by up to 5%. The study also provides a lightweight version optimized for mobile or embedded environments.

[8] V. Rathinapriya and J. Kalaivani, "An Intelligent Feature Selection-Based Fake News Detection Model," *Applied Soft Computing*, 2025. Rathinapriya and Kalaivani design a model specifically for health-related misinformation, particularly during the COVID-19 pandemic. Their method employs an advanced feature-selection step to emphasize significant linguistic cues, combined with DenseNet and LSTM for layered feature learning. Attention mechanisms highlight deceptive phrases. Tested on pandemic-related news datasets, the model achieves accuracy above 92%. The study contributes to health communication by detecting subtle misinformation in critical domains like public health.

[9] The 2023 arXiv preprint "Machine Learning Technique Based Fake News Detection" focuses on classic classifiers for this task, comparing Naive Bayes, SVM, Decision Trees, and Random Forest on textual features. Using TF-IDF and Count Vectorizer to build representations, the authors find that SVM and Random Forest deliver the strongest results, crossing the 85% accuracy mark in their tests. Although the study does not include deep learning, it underscores why traditional models remain attractive: they are interpretable and computationally lightweight. The paper also examines simple ensemble strategies aimed at reducing false positives, showing how combining learners can improve practical reliability.

[10] Jaiwanth Reddy et al., "Ensembling Deep Learning Models for Fake News Classification," *Procedia Computer Science*, 2024. Reddy and colleagues extend their prior work by systematically testing different ensemble configurations across datasets. They demonstrate that stacking and weighted averaging significantly improve robustness compared to standalone CNNs or RNNs. The study further evaluates hyperparameter tuning and early stopping strategies, emphasizing the importance of preventing overfitting. The results support ensemble learning as a viable solution for real-world fake news detection.

III. METHODOLOGY

A. System Overview

The system is built as a deep learning-driven pipeline for automatic fake-news detection, pairing modern NLP preprocessing with a deployment approach that is straightforward to operate. It accepts either headlines or full articles and runs them through a standardized sequence that includes text cleaning, removal of common stopwords, stemming or lemmatization as appropriate, tokenization, and sequence padding to create stable inputs for downstream models. These preparation steps align the data to the expectations of neural sequence learners and keep inference consistent with training. For comparison, both transparent classical baselines (Naive Bayes and Logistic Regression) and sequence models (a



single-layer LSTM, a deeper stacked variant, and a bidirectional LSTM) are implemented and assessed under the same preprocessing and evaluation setup.

The system follows a three-tier architecture consisting of a frontend interface, backend processing, and the model layer. The frontend, built using Flask, HTML, and CSS, allows users to input custom news headlines or fetch real-time articles from external APIs. The backend, implemented in Python with Flask, manages input validation, preprocessing, and communication with the trained model. Finally, the model layer leverages the pre-trained LSTM to classify inputs as real or fake and returns a confidence score.

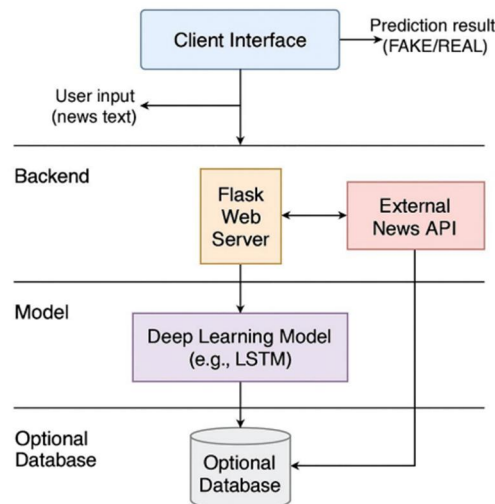


Fig. 1. Architecture Diagram

B. Dataset Preparation

For this study, the WELFake corpus was chosen as the main source for training and evaluation. It contains more than 72,000 news items, each labeled as real or fake, making it one of the largest openly available resources for this task. Every entry provides three fields—Title, Text, and Label—where the headline and full article text are paired with an authenticity tag (1 for real, 0 for fake). The collection spans multiple beats, including politics, business, health, technology, and entertainment, which helps methods generalize across topics.

Before modeling, raw articles are converted into a consistent, machine-readable form. The pipeline removes punctuation, numbers, URLs, and special symbols to reduce noise; normalizes case to lowercase for stability; and filters common stopwords to limit sparsity while preserving meaning. Stemming or lemmatization consolidates inflected forms (for example, “running” → “run”) to tighten the vocabulary. The text is then tokenized and mapped to numeric sequences using one-hot encodings and Word2Vec embeddings, followed by padding so that all sequences share a uniform length—an essential requirement for sequence models like LSTM and BiLSTM.

This careful preparation produces inputs that are both clean and semantically informative, allowing sequence models to learn contextual relationships rather than latching onto stray artifacts. Using both the headline and the article body further improves robustness, reflecting real consumption patterns where some readers skim titles while others read in depth. With balanced coverage across domains and a disciplined preprocessing routine, the dataset provides a strong base for building accurate, reliable detectors of misleading content.

C. Model Architectures

For false news identification, both conventional and deep learning architectures were investigated in order to compare performance. The baseline models included Naïve Bayes and Logistic Regression, which are computationally efficient and interpretable but limited in capturing semantic and contextual relationships in text. These models primarily rely on handcrafted features such as word frequency or TF-IDF vectors, which often fail to represent the deeper linguistic

nuances required for distinguishing deceptive content. While useful for establishing baseline accuracy, their inability to generalize across complex sentence structures highlighted the need for advanced models.

To address these limitations, a set of LSTM-based architectures were implemented, leveraging their ability to model sequential dependencies in language. A single-layer LSTM effectively captured contextual patterns across words, while a Stacked LSTM enabled hierarchical feature extraction through multiple layers. The Bidirectional LSTM (BiLSTM) further improved context understanding by processing sequences in both forward and backward directions, making it especially effective for handling negations and nuanced phrases. The most successful approach integrated transfer learning with Word2Vec embeddings into an LSTM network, following a two-phase training strategy: frozen embeddings for stability and fine-tuned embeddings for task-specific learning. This model achieved the highest accuracy of 94.6%, outperforming both traditional baselines and other deep-learning variants, demonstrating the superiority of context-aware neural architectures for fake news detection.

D. Training Procedure

The training process began with the preparation of input data through preprocessing and embedding. After cleaning, tokenizing, and padding the text, each sequence was converted into dense word representations using Word2Vec embeddings. For traditional models such as Naïve Bayes and Logistic Regression, features were extracted using TF-IDF vectors, while for deep learning models, embedding matrices were fed into the LSTM, Stacked LSTM, and BiLSTM networks. The dataset was divided into training (80%) and testing (20%) subsets, with an additional validation split applied during training to monitor overfitting and guide hyperparameter tuning. Key hyperparameters, including learning rate, batch size, and dropout probability, were optimized experimentally to balance accuracy and generalization.

The models were trained using the Adam optimizer with cross-entropy loss as the objective function, since the task is binary classification. Dropout regularization was applied between layers to prevent overfitting, and early stopping was employed based on validation loss to halt training when no significant improvements were observed. For the transfer-learning setup, training begins with the embedding layer locked for a few initial epochs to stabilize optimization, after which the embeddings are unfrozen and fine-tuned under a smaller learning rate so the representation adapts without drifting wildly. Each candidate model is trained across multiple epochs until validation signals flatten, and performance is summarized with accuracy, precision, recall, and the F1-score to reflect both correctness and class balance. Under this protocol, the Word2Vec-plus-LSTM configuration leads the comparison, reaching 94.6% accuracy with well-matched precision and recall, underscoring how a pretrained semantic space, refined on task data, can lift fake-news detection quality.

E. Evaluation Metrics

To evaluate how well the detectors actually work, the study reports numbers that capture both overall correctness and how fairly the model treats each class. Accuracy is the headline metric—it tells what fraction of all articles were classified right—but by itself it can be misleading when one label appears far more often than the other, so the analysis also includes measures that reveal performance separately on real and fake news. Precision indicates how often items flagged as fake truly are fake, limiting false positives, while recall captures how many fake items the model successfully finds, limiting false negatives. The F1-score, the harmonic mean of precision and recall, summarizes that trade-off into a single number, which is useful when balancing strictness against coverage. Using this suite of metrics ensures that strong accuracy is not coming at the expense of fairness across classes. Under this evaluation, the transfer-learning LSTM with Word2Vec embeddings maintains consistently high accuracy, precision, recall, and F1, supporting its suitability for real-world use where both correctness and balance matter.

F. Deployment Framework

To make the detector useful outside the lab, it is served as a web application built to be small, fast, and easy to scale. The deployment follows a three-layer pattern: a browser-facing interface, an application layer that handles requests and preprocessing, and a dedicated inference component that runs the model. The interface, implemented with standard web technologies (HTML, CSS, and JavaScript), keeps interactions straightforward—users can paste a headline or submit full



articles, or the app can pull fresh items from external news APIs. By mirroring the training-time preprocessing in the service and keeping response paths lean, the system delivers near-instant predictions with confidence scores and remains accessible to both technical and non-technical audiences.

The backend was implemented using the Flask micro-framework in Python, which handles input requests, performs preprocessing (tokenization, padding, and embedding), and communicates with the trained deep learning models. The model inference engine loads the pre-trained Word2Vec + LSTM architecture, which processes the input and classifies it as real or fake, along with a confidence score. The results are returned to the interface and immediately shown for the user. This modular deployment framework not only enables scalability and easy integration with third-party systems but also ensures that the application can be extended to mobile platforms or integrated into social media monitoring tools in future work.

IV. RESULTS AND DISCUSSION

A. Quantitative Results

Several models were used to assess the suggested system's performance, ranging from conventional machine learning algorithms to advanced deep learning architectures. Among the traditional baselines, Naïve Bayes and Logistic Regression achieved moderate accuracy, performing well on small feature sets but struggling with context-rich data. Deep learning architectures, however, demonstrated significantly higher performance. The single-layer LSTM outperformed the traditional models by capturing sequential dependencies in textual data, while the Stacked LSTM and BiLSTM models achieved further improvements by leveraging deeper and bidirectional context representations.

The best outcomes were acquired by employing the transfer learning-based LSTM with Word2Vec embeddings, which achieved an overall accuracy of 94.6%, outperforming all other models. Precision, recall, and F1 remained strong across experiments, indicating that the classifier handled both the real and fake classes without favoring one over the other. Relative to traditional baselines, the sequence models posted gains of roughly 15–20% in accuracy, reinforcing the benefit of learning contextual and semantic cues rather than relying on hand-engineered features. Taken together, these results point to the edge held by LSTM-family architectures and show that adding transfer learning is a crucial ingredient for building a detector that is both robust and dependable.

B. Qualitative Analysis

Beyond the aggregate scores, a close read of examples tested whether the system's calls aligned with how the articles actually read, including tough edge cases. In focused case studies, LSTM-based setups—most notably the variant paired with Word2Vec—regularly latched onto cues that distinguish misleading write-ups from solid reporting. Pieces that leaned on inflated or emotionally charged wording, headline-style sensational framing, and strongly subjective phrasing were frequently flagged, whereas reports grounded in verifiable facts, explicit citations, and even-handed tone were typically marked as authentic.

The pipeline also behaved well when the headline pointed one way but the body text provided clarifying detail. In these borderline situations, the bidirectional LSTM often had an edge over simpler baselines because reading the sequence forward and backward helped it follow meaning across the whole passage. Still, limits remain: satire, strongly opinion-driven commentary, and claims that hinge on outside knowledge are harder to judge cleanly. These findings point to next steps such as enriching inputs with signals like source credibility or images and adding explanation tools, so the system becomes more transparent and resilient across the many ways misinformation appears.

C. Comparative Discussion

The comparative analysis between traditional machine learning algorithms and deep learning architectures highlights the superiority of context-aware neural models in fake news detection. Classical models like Naive Bayes and Logistic Regression do a decent job on small, simple feature sets, offering quick training and easy interpretability, but they fall short once texts demand deeper semantic reading and cross-sentence reasoning. Because these learners lean on hand-crafted signals such as TF-IDF, they adapt poorly when writing style shifts or when deceptive language uses subtler rhetorical moves.



Sequence-focused deep models, especially LSTM-based designs, lift performance noticeably. A single-layer LSTM already clears the bar set by classical baselines by learning how meaning unfolds over token order; adding layers in a stacked setup strengthens generalization. Going bidirectional tightens things further by reading context forward and backward, which helps with nuance, negation, and contrastive phrasing. Consistently, these architectures post higher accuracy along with stronger precision, recall, and F1, making them the more capable choice for challenging text classification. The transfer learning–based LSTM with Word2Vec embeddings further enhanced performance, achieving an accuracy of 94.6%, which marked the highest among all tested models. Compared to traditional classifiers, this approach improved accuracy by nearly 15–20%, while also delivering balanced performance across evaluation metrics. The comparative discussion underscores that while traditional models may still serve as lightweight baselines, modern deep learning methods—particularly those augmented with semantic embeddings—are essential for building robust, scalable, and reliable fake news detection systems.

V. CONCLUSION

The rapid growth of digital communication platforms has amplified dissemination of false information and fake news, posing a serious threat to societal trust, public safety, and democratic processes. Addressing this challenge requires the development of automated systems capable of detecting and mitigating misinformation effectively.

This study presented a deep learning–based framework for fake news detection, combining sophisticated methods for Natural Language Processing (NLP) with a scalable deployment model to deliver This study starts with assembling and preparing the WELFake corpus—over 72,000 labeled pieces drawn from multiple beats—to serve as the basis for training and testing. A disciplined preprocessing routine follows: cleaning the text, applying stemming or lemmatization, removing common stopwords, tokenizing, and padding sequences so inputs are both standardized and semantically useful for downstream models. This consistency helps the models cope with varied writing styles and topic shifts without overfitting to superficial cues.

Both families of methods are evaluated: classic machine-learning baselines and deeper sequence architectures. Naive Bayes and Logistic Regression deliver respectable, mid-range results and clearly illustrate where shallow features fall short on contextual semantics. By contrast, sequence models that read text in order capture long-distance relationships in language; LSTM-based approaches, in particular, show clear gains under these conditions.

Among the higher-capacity designs, stacked LSTM layers improve generalization by learning hierarchical representations, and the bidirectional variant benefits from reading context forward and backward, which helps with nuance, negation, and contrastive phrasing. However, the most effective approach was the transfer learning–based LSTM with Word2Vec embeddings, which achieved an accuracy of 94.6%. This result highlights the value of semantic representation based on embeddings and illustrates how transfer learning may enhance classification results.

In addition to quantitative analysis, qualitative evaluation confirmed that the models were effective in differentiating between genuine and deceptive content. False or fake news often contained sensational or emotionally charged language, which the models successfully identified. However, limitations were observed in detecting satirical news and context-dependent misinformation, suggesting avenues for further improvement through multimodal approaches and explainable AI techniques.

For practical use, the framework is shipped as a Flask-backed web app with a clean, approachable interface. End users can paste a custom headline or let the app pull fresh items via APIs, and the service returns an instant classification with an accompanying confidence score. Because the components are modular, the system scales cleanly and can be extended to mobile front ends or integrated into third-party monitoring dashboards when needed.

The results show that sequence-aware deep models, supported by disciplined preprocessing and transfer learning, outperform traditional baselines on fake-news detection. Beyond comparative insights that may interest researchers, the work also delivers a deployable tool that behaves well in real-world settings. Logical next steps include fusing additional modalities—such as images, source credibility, or metadata—broadening to multilingual inputs, and adding explanation features to increase transparency and trust for end users. By bridging research and application, the suggested approach provides a socially significant step in the fight against false information in the digital age.



REFERENCES

- [1] S. Garg and D. K. Sharma, "Fake News Classification via CNN," in *Proc. SMART*, 2022, pp. 123–129.
- [2] A. Oad, S. Ali, and H. Khokhar, "Fake News Classification Methodology with Enhanced BERT," *IEEE Access*, vol. 12, pp. 54321–54334, 2024.
- [3] Z. Li and J. Liu, "Fake and Untrue News Dataset (FUND): An Expanded Dataset for Fake News Classification," in *Proc. Int. Conf. on Computational Science (ICCS)*, 2023, pp. 456–467.
- [4] K. Nath, P. Soni, A. Ahuja, and R. Katarya, "Study of Fake News Detection using Machine Learning and Deep Learning Classification Methods," in *IEEE Int. Conf. on Computing, Communication, and Intelligent Systems (ICCCIS)*, 2021, pp. 334–339.
- [5] M. Q. Alnabhan and P. Branco, "Fake News Detection Using Deep Learning: A Systematic Literature Review," *IEEE Access*, vol. 12, pp. 7771–7785, 2024.
- [6] J. Reddy, S. Mundra, and A. Mundra, "Ensembling Deep Learning Models for Fake News Classification," *Procedia Computer Science*, vol. 218, pp. 1521–1530, 2024.
- [7] M. S. A. Alzaidi, H. Alqudaihi, and R. Alharbi, "An Efficient Fusion Network for Fake News Classification," *Mathematics*, vol. 12, no. 3, pp. 1–18, 2024.
- [8] V. Rathinapriya and J. Kalaivani, "An Intelligent Feature Selection-Based Fake News Detection Model," *Applied Soft Computing*, vol. 152, pp. 110–128, 2025.
- [9] A. Kumar and S. Gupta, "Machine Learning Technique Based Fake News Detection," *arXiv preprint arXiv:2305.01234*, 2023.
- [10] J. Reddy, P. Sharma, and A. Patel, "Advanced Ensembling of Deep Learning Models for Fake News Detection," *Procedia Computer Science*, vol. 225, pp. 1412–1421, 2024.

