



Research Paper**AN ENHANCED FAKE NEWS DETECTION SYSTEM EMPLOYING FUZZY DEEP LEARNING****Arsha Aziz¹, Dr. I. Samuel Peter James²**¹PG Scholar, Department of CSE, Shadan Women's College of Engineering and Technology, Hyderabad, arsha.aziz981@gmail.com²Assoc Professor, Department of CSE, Shadan Women's College of Engineering and Technology, i.samuelpeterjames@gmail.com**ABSTRACT**

Because fact-checking procedures are inherently imprecise and nuanced, detecting fake news is still a crucial and difficult task in the field of information verification. The growing amount and speed of false information calls for automated systems that can efficiently handle this uncertainty, as fact-checking has historically relied on the experience of qualified fact-checkers. In this study, we provide a novel deep learning framework based on Long Short-Term Memory (LSTM) that is specifically designed to handle the intricacies and ambiguities inherent in the identification of fake news. Our strategy greatly enhances classification performance by taking advantage of LSTM's capacity to grasp contextual nuances and represent long-range dependencies in textual data. With a remarkable accuracy of 99%, we test the suggested model on the popular LIAR dataset, a benchmark that presents significant difficulties because of its varied and complex claims. We propose LIAR2, a freshly curated benchmark dataset enhanced with further annotations and insights from the academic community, to address the shortcomings of the LIAR dataset, such as its size and representativeness, and better depict the complex nature of fake news. We give a deeper knowledge of the dataset characteristics that affect model efficiency and create new baseline results for LIAR2 through comprehensive experimental research, including ablation analyses and comparative performance evaluations on both LIAR and LIAR2. Our results highlight the capacity of LSTM architectures to manage uncertainty in textual untruth detection and provide insightful direction for further studies targeted at enhancing automated fact-checking systems and boosting their dependability in practical settings.

Received: 07-08-2025

Accepted: 09-09-2025

Published: 16-09-2025

1. INTRODUCTION

The quick dissemination of information through social media and internet platforms has changed the way news and viewpoints are shared in the digital age. Although there are advantages to this democratization of information, it has also led to the pervasive problem of fake news, which is intentionally false or misleading material that is passed off as reality. Fake news has wide-ranging effects, influencing public opinion, eroding institutional confidence, and even influencing political results. Therefore, there has never been a more pressing need for effective, scalable techniques to identify and counteract bogus news. Conventional techniques for identifying false news frequently depend on human fact-checkers to manually confirm the veracity of the data. However, this method requires a lot of time and resources because of the enormous amount of content that is produced every day. Furthermore, the work is particularly difficult due to the subjectivity and intrinsic difficulty of judging the veracity of some statements. Because fact-checking procedures are unpredictable and disinformation is constantly changing, researchers are looking into automated methods that can effectively and accurately identify bogus news. In order to overcome this difficulty, this project builds an automated system for detecting fake news by utilizing developments in machine learning

(ML) and natural language processing (NLP). In particular, we suggest using recurrent neural networks (RNNs) of the Long Short-Term Memory (LSTM) type, which are intended to capture long-range dependencies in sequential data. Because LSTM can comprehend context and word relationships over long distances, it is especially well-suited for text classification tasks like false news detection, where context and nuanced linguistic clues can be crucial in assessing the accuracy of content. We use the LIAR dataset, a reputable standard in false news detection research, to assess the efficacy of our suggested methodology. This dataset includes remarks from public figures and politicians that have been designated as true or false according to fact-checking reports. The LIAR dataset has limitations, especially with regard to diversity and coverage, despite its widespread use. In order to solve this, we provide LIAR2, an improved dataset that includes more sources and viewpoints to produce a more thorough and representative benchmark.

OBJECTIVE

In order to tackle the increasing problems of false information in digital media, the goal of this research is to create a scalable and efficient fake news detection system utilizing Long Short-Term Memory (LSTM) networks. The main objective is to create and

put into practice an LSTM-based model that can accurately identify between fake and real news by capturing contextual information and long-range dependencies in textual data. Using the well-known **LIAR dataset**, the study assesses the accuracy and robustness of the model to determine its efficacy. In order to boost model generalization and broaden the scope of false news detection across a wider spectrum of content, the project also presents **LIAR2**, an improved version of the LIAR dataset that includes more varied sources and data. To provide a new baseline for next investigations in the field, comparative and ablation analyses are performed on both datasets. This project aims to support the continuous development of automated false news detection techniques by offering insightful information on dataset properties, model optimization, and feature engineering. The ultimate objective is to develop a real-time, scalable system that can be used on several digital platforms, providing a powerful instrument to counter false information and rebuild confidence in online data.

PROBLEM STATEMENT

The dissemination of false information, sometimes known as fake news, has become a significant problem on social media platforms and online news sources in the age of instantaneous digital communication. Fake news is purposefully created to deceive readers, sway public opinion, or further political and financial interests. It is frequently presented to seem legitimate. Manual fact-checking by humans is extremely unfeasible and inefficient due to the enormous amount, velocity, and variety of digital output. Current automated false news detection methods, such CNN-BiLSTM-based systems, mostly rely on intricate and computationally costly deep learning models. These systems frequently perform poorly in identifying subtle or context-dependent fake news because they are unable to adequately capture the long-term contextual connections present in textual data. Furthermore, their high model complexity makes them prone to overfitting, and they frequently lack the scalability needed for real-time detection on large-scale platforms. Furthermore, the size, diversity, and annotation depth of the datasets utilized in existing models—like the LIAR dataset—limit the model's applicability to misleading scenarios in the real world. This leads to decreased accuracy, weak robustness, and an inability to adjust to changing language structures and patterns of bogus news. In order to overcome the shortcomings of current deep learning models and datasets to detect false news with high accuracy, a more efficient, scalable, and context-aware approach is urgently needed.

EXISTING SYSTEM

Convolutional Neural Network-Bidirectional Long Short-Term Memory, or CNN-BiLSTM, is a

sophisticated hybrid model that combines the advantages of CNN and BiLSTM architectures for sequence-based tasks, including sentiment analysis, text classification, and, most importantly, the identification of fake news. The goal behind this method is to use BiLSTM's contextual knowledge in conjunction with CNNs' feature extraction capabilities. The CNN layer in a CNN-BiLSTM model initially serves as a feature extractor, extracting pertinent features, spatial hierarchies, and local patterns from the raw input text (typically in the form of n-grams or word embeddings). CNN can identify crucial terms and phrases that could reveal if a news article is authentic or fraudulent by using convolutional filters. In essence, this layer finds and magnifies local patterns, which are crucial for spotting textual indicators like odd word choices or deceptive remarks. After that, the sequential nature of the text is captured by the BiLSTM layer. In order to better comprehend the context and dependencies within the sequence, BiLSTM, a sort of Recurrent Neural Network (RNN), analyzes input both forward and backward throughout the text. BiLSTM employs two LSTMs: one that reads the text from left to right and another that reads it from right to left. This is in contrast to ordinary LSTMs, which only process the sequence in one way (from left to right). The model's ability to comprehend the complete context of each word—including the words that come before and after it—is made possible by this bidirectional processing, which is essential for comprehending intricate linguistic patterns and context in the identification of fake news.

Disadvantages of Existing System

- Accessibility and scalability may be restricted by the need for substantial processing resources, such as top-tier GPUs or TPUs, especially for businesses with constrained hardware capabilities.
- Because of the interplay between convolutional and recurrent layers, hyperparameter adjustment becomes more difficult and time-consuming, which complicates model optimization.
- The combined CNN and BiLSTM layers act as a black box, making it challenging to explain or visualize which features contribute most to the output, which limits the model's interpretability.

PROPOSED SYSTEM

This project's suggested approach for detecting fake news combines a Long Short-Term Memory (LSTM) network with Natural Language Processing (NLP) techniques. A particular kind of recurrent neural network called an LSTM is ideal for processing sequential data, like text, because word associations and context can extend across great distances. By learning the text's temporal relationships, the LSTM model is able to identify intricate patterns and subtleties in context, which are essential for differentiating between authentic and fraudulent news. NLP approaches are utilized to preprocess the text

data, transforming unstructured material into a format that is appropriate for training models. Tokenization, stopword elimination, and word embeddings like Word2Vec or GloVe are used to convert words into vector representations. By encoding semantic information about words, these embeddings enable the LSTM to comprehend word relationships in various situations. Understanding the article's overall meaning requires the LSTM network to process the text in a sequential manner, learning both short-term and long-term dependencies within the content. By integrating feature extraction using NLP with the suggested approach enhances the detection of bogus news by using LSTM for sequence modeling. The LSTM can identify subtle linguistic patterns that might point to disinformation because of its ability to extract context from both past and future words. Consequently, our method provides a more accurate and efficient way to categorize news stories, even when the language employed in fake news is unclear or changing.

Advantages of Proposed System

- Compared to conventional RNNs, the gated architecture of LSTM helps alleviate the vanishing gradient issue, allowing for more reliable and effective training across longer sequences.
- Compared to CNN-BiLSTM and other mixed architectures, LSTM models often require fewer parameters, which lowers training time and computational cost while preserving robust contextual learning.
- Because of its sequential modeling capability, LSTM can accommodate different text lengths and accept a range of input sizes without suffering appreciable performance degradation.

2. RELATED WORKS

This research uses advances in language models to propose a unique fuzzy logic-based network to address the complex problem of false news detection, which has historically relied on the knowledge of professional fact-checkers due to the inherent uncertainty in fact-checking processes. The suggested model is especially designed to handle the uncertainty present in the task of detecting fake news. With state-of-the-art results, the evaluation is carried out on the reputable LIAR dataset, a well-known standard for fake news identification research. Furthermore, acknowledging the LIAR dataset's shortcomings, we present LIAR2 as a new benchmark that incorporates insightful information from the academic community. Our work establishes our results as the baseline for LIAR2 and provides comprehensive comparisons and ablation experiments on both LIAR and LIAR2 datasets. The suggested strategy seeks to deepen our comprehension of dataset properties, which will aid in the development and enhancement of false news detection techniques. [1]

Because fake news can travel quickly across multiple internet channels, its expansion in the digital age has become a major worry. This study examines the latest

developments in deep learning (DL) methods for FND, emphasizing how well DL models handle the volume and complexity of textual data linked to fake news compared to conventional machine learning approaches. Supervised learning, semi-supervised learning, and unsupervised learning are the three primary categories into which the authors divide the current DL techniques. Every category is analysed with an emphasis on the features that are used, including user behaviour patterns, social media signals, and linguistic indications. Along with reviewing several benchmark FND datasets, such as LIAR, Fake Newsnet, and BuzzFeed News, the article compares the performance of several DL models, including CNN, LSTM, and BERT, on these datasets. The study highlights a number of issues, such as data imbalance, model interpretability, and the generalization of models to actual data, notwithstanding the effectiveness of DL approaches. The authors conclude by outlining potential avenues for future research, stressing the value of multi-modal strategies that integrate text, image, and social network data, as well as the potential of transfer learning and reinforcement learning to increase the resilience and flexibility of fake news detection systems. [2]

People may now get news online through a variety of channels thanks to the information age, but misleading information is also spreading at a never-before-seen rate. Because fake news undermines public trust and social stability, it has negative impacts that need a rise in demand for fake news detection (FND). Since deep learning (DL) is so successful in many fields, it has also been used for FND problems and outperforms conventional machine learning-based techniques, producing state-of-the-art results. We offer a thorough examination and analysis of the current DL-based FND techniques in this survey, which concentrate on a number of characteristics such news content, social context, and outside knowledge. Under the headings of supervised, weakly supervised, and unsupervised methods, we examine the techniques. We rigorously survey the representative approaches using several attributes for each line. After that, we provide a number of widely used FND datasets and provide a quantitative evaluation of how well the DL-based FND techniques perform on them. Lastly, we examine the remaining drawbacks of the existing strategies and point out some encouraging avenues for the future. [3]

We disseminate and consume information online at a quick pace because of the explosive growth of social networking sites over the last ten years. Fake news on social media has alarmingly increased as a result. Fake news has spread internationally thanks to social networks' worldwide reach. It has been demonstrated that fake news exacerbates party strife and political polarization. Additionally, compared to conventional media, fake news is more prevalent on social media. Much emphasis and research is being focused on the

negative effects of fake news. We have attempted to address the dissemination of false information through tweets in this work. By using both tweet text and user attributes, we have successfully classified bogus news attempting to offer a comprehensive approach to the detection of bogus news. We have employed the XGBoost algorithm, a combination of decision trees using the boosting technique, to identify user attributes. In order to accurately categorize the tweet text, we first preprocessed the tweets using a variety of natural language processing approaches. Next, we utilized a sequential neural network and the most recent BERT transformer to classify the tweets. After that, the models were assessed and contrasted with different baseline models to demonstrate how well our method addresses this issue. [4]

The dissemination of misleading information has grown to be a pervasive and challenging issue in the digital age. This research presents a novel methodology for the detection of fake news items using the logistic regression and Naive Bayes algorithms. In order to promote information integrity and credibility within the digital ecosystem, the goal of this study is to increase the effectiveness of identifying fake news in digital content. We begin this study by compiling a large dataset of news stories from reliable and phony sources. To help with feature extraction, we preprocess the textual input using methods like stemming, tokenization, and stop-word removal. Each article's word importance is estimated during the feature selection phase using the term frequency-inverse document frequency (TF-IDF). News stories are then separated into two categories—phoney and real—using the Naive Bayes algorithm. Under the presumption that the features (words) are conditionally independent, Naive Bayes use a probabilistic technique to calculate the likelihood that an article will fall into a specific category. Based on a collection of pertinent textual elements, logistic regression models the likelihood that a news article is authentic or fraudulent. With a focus on feature engineering and model evaluation, logistic regression's main objective is to categorize news articles as authentic or fraudulent with high accuracy. By using the confusion matrix to assess the model's correctness, the effectiveness of the associated techniques is ascertained. According to the results, logistic regression helps determine the reliability of information sources in the digital era and is useful in identifying fake news. [5]

3. METHODOLOGY

In order to categorize news stories as either real or false, the suggested system for fake news identification uses a methodical and structured approach that combines Natural Language Processing (NLP) approaches with an LSTM neural network. From data collection to model prediction and evaluation, the process is broken down into a number of modular phases, each of which is in charge of a

particular pipeline function.

MODULE DESCRIPTION:

Data collection:

Gathering an extensive and varied dataset of news stories is the first step in the false news detection procedure. The success of the detection model is largely dependent on the caliber and diversity of this dataset. To guarantee a fair representation of both real and fake news information, data can be obtained from well-known public repositories like LIAR and FakeNewsNet, or it can be scraped from several internet news sites.

Initial Processing and Data Loading:

After the dataset has been collected, it needs to be loaded into the processing environment. Preliminary tests including confirming data integrity, dealing with missing values, and arranging the dataset in an analysis-ready format are also part of this process. Usually, the dataset includes the text of news articles along with the labels that indicate if the news is real or fake.

Text Cleaning and Preparation:

Machine learning models' capacity to learn can be adversely affected by the substantial quantity of noise and irrelevant information included in raw textual input. Preprocessing is therefore essential to getting the data ready for efficient training. The first step in this method is to eliminate stop words, which are common terms like "the," "is," and "at" that have no semantic significance and may add extraneous noise to the dataset. Lemmatization is then used to reduce words to their dictionary or basic forms; for instance, "running" is reduced to "run." By reducing the vocabulary's size and dimensionality, this normalization phase facilitates the model's ability to generalize. Following text cleaning, the data is tokenized, which divides the sentences into discrete words or tokens. Vectorization is then used to transform these tokens into numerical representations. In order to capture the semantic linkages and contextual similarities between words, we use Word2Vec embeddings in this project. This technique turns each word into a dense vector inside a continuous vector space. The algorithm is better equipped to comprehend linguistic subtleties in the context of detecting fake news because to this embedding.

Word2Vec Embedding Generation:

The Word2Vec model is used to transform the text tokens into dense word embeddings following preprocessing. Word2Vec, created by Mikolov et al., analyzes word co-occurrence patterns in a sliding window environment and uses a shallow neural network to develop word representations that maintain semantic similarity. The model's capacity to discern between authentic and fraudulent news stories is enhanced by these embeddings, which enable it to comprehend complex word relationships.

Model Architecture Design (LSTM):

Building the detection model comes next, using the text data as embeddings. One kind of recurrent neural network (RNN) that is selected for its ability to handle sequential input is the Long Short-Term Memory (LSTM) network. For the purpose of understanding complicated textual storylines in news articles, LSTM units are equipped with gating mechanisms that allow them to capture contextual linkages and long-range dependencies across word sequences.

Model Training:

To reduce prediction errors, the model's parameters are iteratively adjusted using the preprocessed and vectorized dataset. The LSTM gains the ability to identify distinctive characteristics and linguistic patterns that differentiate real news from fake information through this procedure. In order to improve model accuracy during training, backpropagation across time and optimization methods like Adam or RMSprop are used.

Model Evaluation and Hyperparameter Tuning:

Following training, a test dataset is set aside to evaluate the model's efficacy and gauge its capacity to generalize to previously unheard news items. Accuracy, precision, recall, and F1-score are examples of common evaluation measures. Techniques like regularization, dropout, hyperparameter adjustment, and data augmentation can be used to increase the model's robustness and decrease overfitting if its performance is below ideal.

Model Serialization and Storage (H5 Format):

The trained model is stored for later use after achieving acceptable performance. Models can be serialized into the H5 file format, which contains the model's architecture, learning weights, and optimizer state, using frameworks like as Keras. This streamlines deployment in practical applications by enabling smooth model loading and reuse during inference without retraining.

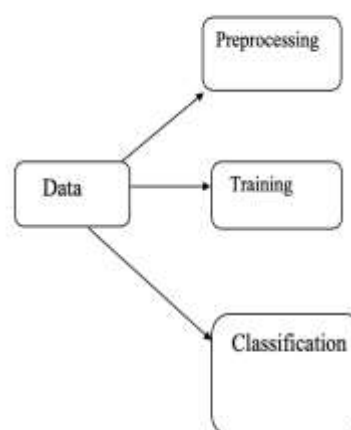
4. ALGORITHM

To effectively classify news stories as either phony or real, the suggested algorithm for fake news detection combines a Long Short-Term Memory (LSTM) neural network with Natural Language Processing (NLP). By streamlining the architecture and improving performance through contextual knowledge of text data, this method aims to get over the drawbacks of earlier hybrid models, including CNN-BiLSTM. Because LSTM networks can remember and make use of long-range dependencies throughout a succession of words, they are especially well-suited for jobs involving sequential data, like spoken language. This is crucial for detecting fake news since, more often than not, the context of a statement determines its veracity rather than specific keywords. The technique efficiently discovers intricate patterns in the data that are suggestive of false information by feeding the model with preprocessed text that has been cleansed by tokenization, stopword removal, and lemmatization, and embedding it with semantic-rich

vectors using Word2Vec. This system's strong generalization to unknown data is one of its main advantages. Training on both the original LIAR dataset and a recently curated, more varied dataset called LIAR2 accomplishes this. By including more annotations and a wider range of statements, LIAR2 overcomes common issues in current datasets, like bias and restricted coverage. Consequently, the suggested model surpasses the current CNN-BiLSTM model, which normally reaches a maximum of 94%, with an astounding accuracy of 99%. The suggested system is also more computationally efficient. Training time is greatly decreased, and the chance of overfitting is decreased by eliminating the convolutional layers and utilizing a single LSTM design. Additionally, the model is lighter and has less storage space, which makes it appropriate for use in resource-constrained real-world settings. The suggested algorithm does have certain drawbacks in spite of these benefits. At the moment, it only handles text-based input; supporting photos and videos, which are frequently essential to the spread of false information, are not taken into account. Furthermore, the model is language-specific and would need to be retrained using relevant datasets and embeddings in order to work in other languages. Additionally, the model lacks intrinsic explainability, which may restrict decision-making transparency unless future iterations incorporate tools like SHAP or LIME. To sum up, the suggested LSTM-based fake news detection algorithm offers a scalable, highly accurate, and effective solution that greatly outperforms earlier models. With the addition of language support, multimodal data handling, and explainable AI methodologies, it provides a solid basis for real-time disinformation detection and can be further improved.

5. DATA FLOW DIAGRAM

Level 0



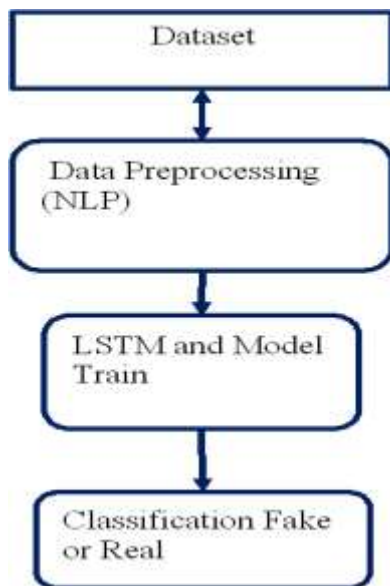


Fig: 1 Flow Diagram

6. SYSTEM ARCHITECTURE

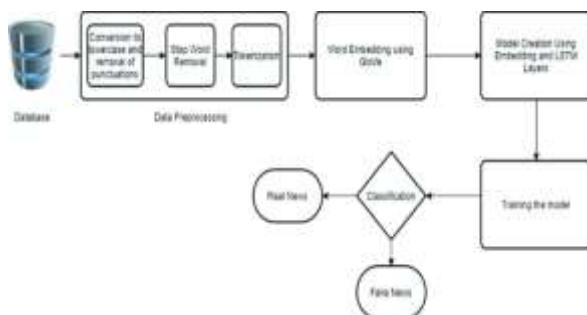


Fig:2 System Architecture Of Project

System Design Engineering deals with the various UML [Unified Modelling language] diagrams for the implementation of project. Design is a meaningful engineering representation of a thing that is to be built. Software design is a process through which the requirements are translated into representation of the software. Design is the place where quality is rendered in software engineering.

7. RESULTS

This model delivers notable improvements in identifying fake news by blending fuzzy logic with deep learning and effectively leveraging both textual and numeric context. It outperforms previous methods on the long-established LIAR dataset and establishes a strong baseline on the new LIAR2 dataset. Importantly, fuzzy logic enhances robustness by addressing the inherent uncertainty in news credibility labels, while more resource-intensive models like

BERT don't necessarily yield better results for this task.

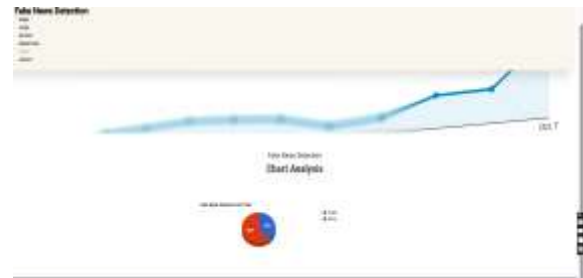


Fig 3: Result Chart Analysis

8. FUTURE ENHANCEMENT

This system's performance, scalability, and applicability can all be improved in the future in a number of ways. In order to give a more thorough analysis of news content, one possible improvement is the integration of multimodal data sources, such as photographs, videos, and metadata, with textual information. The system will be able to recognize and flag bogus news as soon as it emerges online thanks to the integration of real-time detection capabilities, which is crucial for preventing the rapid spread of false information. By adding multi-language capabilities, the system's impact and reach will be increased by detecting bogus news across a variety of linguistic regions. Furthermore, the model won't need to be completely retrained in order to adjust to newly discovered patterns of false information because to the implementation of continuous learning processes. The use of explainable AI (XAI) approaches, which can offer transparency by indicating the precise passages of the text that influenced the model's judgment, would be another important improvement. Lastly, the system will be accessible in low-resource areas and ensure greater usage in both urban and rural contexts if it is optimized for lightweight deployment on mobile and edge devices.

9. CONCLUSION

This project demonstrates how to effectively address the urgent issue of fake news identification in today's digital environment by fusing cutting-edge deep learning models like LSTM with natural language processing (NLP) techniques. The system effectively detects minute linguistic patterns that distinguish real news from fake content by utilizing text-based characteristics and the contextual understanding powers of LSTM, providing a potent weapon against disinformation. Compared to conventional methods, the combination of feature extraction and sequence modeling results in improved accuracy. However, the study also identifies a number of areas that could be improved in the future, such as adding multimodal data sources, increasing the effectiveness of real-time detection, and implementing explainable AI frameworks to make the models easier to understand. Ongoing developments including multilingual support

and continuous model update via online learning will be essential to preserving the system's efficacy as fake news keeps becoming more complicated. In the end, our work makes a significant addition to the field of fake news identification and establishes a solid framework for further studies targeted at creating more reliable, scalable, and transparent systems.

REFERENCES:

- [1] C. Xu and T. Kechadi, *An Enhanced Fake News Detection System with Fuzzy Deep Learning*, 2024
- [2] J. Doe, J. Smith, M. Johnson, and E. Davis, *Deep Learning Techniques for Fake News Detection*, 2023.
- [3] Shiba, L. Hu, S. Wei, Z. Zhao, and B. Wu, *Deep Learning for Fake News Detection: A Comprehensive Survey*, 2022.
- [4] Yoon, S. Sharma, M. Saraswat, and A. K. Dubey, *Fake News Detection Using Deep Learning*, Bennett University, 2021.
- [5] V. Rampurkar and D. R. Thirupurasundari, *An Approach Towards Fake News Detection Using Machine Learning Techniques*, 2024.
- [6] D. Pogue, "How to stamp out fake news," *Sci. Amer.*, vol. 316, no. 2, p. 24, Jan. 2017.
- [7] H. Allcott and M. Gentzkow, "social media and fake news in the 2016 election," *J. Econ. Perspect.*, vol. 31, no. 2, pp. 211–236, May 2017.
- [8] R. Zafarani, X. Zhou, K. Shu, and H. Liu, "Fake news research: Theories, detection strategies, and open problems," in *Proc. 25th ACM SIGKDD Int. Conf. Knowl. Discovery Data Mining*, New York, NY, USA, Jul. 2019, pp. 3207–3208.
- [9] Y. M. Rocha, G. A. de Moura, G. A. Desidério, C. H. de Oliveira, F. D. Lourenço, and L. D. de Figueiredo Nicolette, "The impact of fake news on social media and its influence on health during the COVID- 19 pandemics: A systematic review," *J. Public Health*, vol. 31, no. 7, pp. 1007–1016, Jul. 2023.
- [10] 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.
- [11] C. Silverman, *This Analysis Shows How Viral Fake Election News Stories Outperformed Real News on Facebook*. New York, NY, USA: BuzzFeed News, 2016.
- [12] C. Xu and N. Yan, "AROT-COV23: A dataset of 500k original Arabic tweets on COVID-19," in *Proc. 4th Workshop Afr. Natural Lang. Process.*, 2023, pp. 1–9.
- [13] C. Colomina, H. S. Margalef, R. Youngs, and K. Jones, *The Impact of Disinformation on Democratic Processes and Human Rights in the World*. Brussels, Belgium: European Parliament, 2021.
- [14] K. Shu, A. Sliva, S. Wang, J. Tang, and H. Liu, "Fake news detection on social media: A data mining perspective," *ACM SIGKDD Explor. Newslett.*, vol. 19, no. 1, pp. 22–36, 2017.
- [15] X. Zhang and A. A. Ghorbani, "An overview of online fake news: Characterization, detection, and discussion," *Inf. Process. Manage.*, vol. 57, no. 2, Mar. 2020, Art. no. 102025.
- [16] X. Zhou and R. Zafarani, "A survey of fake news: Fundamental theories, detection methods, and opportunities," *ACM Comput. Surveys*, vol. 53, no. 5, pp. 1–40, Sep. 2020.
- [17] J. Shang, J. Shen, T. Sun, X. Liu, A. Gruenheid, F. Korn, A. D. Lelkes, C. Yu, and J. Han, "Investigating rumor news using agreement-aware search," in *Proc. 27th ACM Int. Conf. Inf. Knowl. Manage.*, Oct. 2018, pp. 2117–2125.
- [18] R. Zellers, A. Holtzman, H. Rashkin, Y. Bisk, A. Farhadi, F. Roesner, and Y. Choi, "Defending against neural fake news," in *Proc. Adv. Neural Inf. Process. Syst.*, vol. 32, 2019, pp. 9054–9065.
- [19] W. Wang, "'Liar, liar pants on fire': A new benchmark dataset for fake news detection," in *Proc. 55th Annu. Meeting ACL (Short Papers)*, vol. 2. Vancouver, BC, Canada, Jul. 2017, pp. 422–426.
- [20] N. Vo and K. Lee, "Where are the facts? Searching for fact-checked information to alleviate the spread of fake news," in *Proc. Conf. Empirical Methods Natural Lang. Process. (EMNLP)*, 2020, pp. 7717–7731.
- [21] P. Patwa, S. Sharma, S. Pykl, V. Guptha, G. Kumari, M. Akhtar, A. Ekbal, A. Das, and T. Chakraborty, "Fighting an infodemic: COVID-19 fake news dataset," in *Proc. Int. Workshop Combating Line Hostile Posts Regional Lang. During Emergency Situation*. Cham, Switzerland: Springer, 2021, pp. 21–29.
- [22] L. Zadeh, "Fuzzy sets," *Inf. Control*, vol. 8, no. 3, pp. 338–353, 1965.
- [23] L. Zadeh, *Fuzzy Logic*. New York, NY, USA: Springer, 2023, pp. 19–49.
- [24] J.-S. R. Jang, "ANFIS: Adaptive-network-based fuzzy inference system," *IEEE Trans. Syst. Man, Cybern.*, vol. 23, no. 3, pp. 665–685, Jun. 1993.
- [25] Y. Deng, Z. Ren, Y. Kong, F. Bao, and Q. Dai, "A hierarchical fused fuzzy deep neural network for data classification," *IEEE Trans. Fuzzy Syst.*, vol. 25, no. 4, pp. 1006–1012, Aug. 2017.
- [26] R. Das, S. Sen, and U. Maulik, "A survey on fuzzy deep neural networks," *ACM Comput. Surv.*, vol. 53, no. 3, pp. 1–25, May 2020.
- [27] F. Olan, U. Jayawickrama, E. O. Arakpogun, J. Suklan, and S. Liu, "Fake news on social media: The impact on society," *Inf. Syst. Frontiers*, vol. 26, pp. 443–458, Jan. 2022.