

A Smart Approach for Precise Fake News Detection Using Bi-Directional LSTM and Self-Attention Mechanism

O. Jeba Singh^{1,*}, R. Remya², G. Dhivya³, M. Manikandan⁴, S. Rubin Bose⁵

¹Center for Academic Research, Alliance University, Bengaluru, Karnataka, India.

²Department of Electronics and Communication Engineering, Sri Krishna College of Engineering and Technology, Coimbatore, Tamil Nadu, India.

³Department of Computer Science and Engineering, Sai Vidya Institute of Technology, Bengaluru, Karnataka, India.

⁴Department of Electronics and Communication Engineering, Presidency University, Bengaluru, Karnataka, India.

⁵School of Computer Science and Engineering, SRM Institute of Science and Technology, Ramapuram, Chennai, Tamil Nadu, India.

jeba.singh@alliance.edu.in¹, remiamernath@gmail.com², dhivyaasrigopal@gmail.com³, email2mani86@gmail.com⁴, rubinbos@srmist.edu.in⁵

Abstract: The legitimacy of information and the trust of the public are both in jeopardy because of the swift rise in the amount of false information that is disseminated through social media and other computerised platforms. Statistical characteristics are the primary foundation for most conventional techniques that are used to detect false information in the news. However, these characteristics often fail to account for the intricate variations in semantics and context found in news stories. This paper introduces an enhanced fake news detection model that blends Bi-directional Long Short-Term Memory (Bi-LSTM) networks with a self-attention mechanism to address this restriction. Bi-LSTM networks analyse text in both forward and backward directions, enabling them to capture long-range contextual correlations within the content. The self-attention mechanism is an additional feature that is incorporated into the model to improve its performance. It does this by dynamically assigning distinct words and phrases in the text to different levels of priority. The Kaggle fake news dataset, also known as the ISOT public dataset, is used to evaluate the model's performance.

Keywords: Fake News Detection; Bi-LSTM Networks; Self-Attention; Deep Learning; Data Ecosystem; Low Detection; Nuanced Cases; Long Short-Term Memory (LSTM); Manipulated News.

Received on: 25/12/2024, **Revised on:** 02/03/2025, **Accepted on:** 14/04/2025, **Published on:** 22/11/2025

Journal Homepage: <https://www.fmdbpublish.com/user/journals/details/FTSCL>

DOI: <https://doi.org/10.69888/FTSCL.2025.000484>

Cite as: O. J. Singh, R. Remya, G. Dhivya, M. Manikandan, and S. R. Bose, "A Smart Approach for Precise Fake News Detection Using Bi-Directional LSTM and Self-Attention Mechanism," *FMDB Transactions on Sustainable Computer Letters*, vol. 3, no. 4, pp. 173–182, 2025.

Copyright © 2025 O. J. Singh *et al.*, licensed to Fernando Martins De Bulhão (FMDB) Publishing Company. This is an open access article distributed under [CC BY-NC-SA 4.0](https://creativecommons.org/licenses/by-nc-sa/4.0/), which permits copying, remixing, redistributing, transformation, and building upon the material in any medium so long as the original work is properly cited.

1. Introduction

The proliferation of fake news on digital platforms, particularly on social media, has become a significant concern in the modern data ecosystem. Nowadays, the spread of misleading, fabricated, or manipulated news has significantly impacted public

*Corresponding author.

perception, shaped political opinions, and even altered the course of major global events. The rise of online platforms, such as Facebook, Instagram, and Twitter, has made it easy to share information with large groups of people. But it finds it more difficult for individuals to discriminate between trustworthy news and falsehoods [1]. Therefore, in situations such as elections, public health crises, and environmental disasters. During this period, misinformation can create far-reaching consequences. In light of the need for an automated system to efficiently detect and mitigate the spread of fake news, a proposal is made. Various methods for identifying fake news have largely relied on linguistic and statistical features extracted from news articles. These methods typically focus only on identifying certain keywords, sentence structures, or patterns within the text. While these approaches can be useful in some cases, they often fall short in capturing the deeper contextual and semantic meaning of the content [2]. Fake news can be subtly written, incorporating elements of truth to mislead the reader, making it difficult for traditional models to detect inconsistencies effectively. Furthermore, these models tend to rely on predefined rules and may not generalise well to diverse datasets, resulting in low detection accuracy, especially in complex or nuanced cases [3]. To overcome these obstacles, new advances in deep learning have been made, driven by more enlightened approaches. In particular, neural networks are a promising method, especially with the implementation of the Long Short-Term Memory (LSTM) algorithm.

It is more suitable for handling consecutive information in text form. LSTM models are designed to capture long-range dependencies within the text, making them well-suited for understanding complex narratives. However, traditional LSTM models process text in a unidirectional manner, meaning they can analyse data from start to finish or from finish to start, but not both. This limitation can limit the model's ability to fully understand the news article's context, which is crucial for accurate fake news detection [4]. Bi-LSTM networks address this issue by processing the text in both forward and backward directions, enabling the model to know the whole context of a sentence or article. This bidirectional processing enables Bi-LSTM networks to retain more comprehensive contextual information, making them better equipped to detect subtle inconsistencies and misleading narratives common in fake news. This selective focus enables the system to prioritise the most relevant parts of an article while disregarding less significant information, thereby enhancing both the accuracy and efficiency of the detection process. Unlike traditional models that treat all words equally, the self-attention mechanism allows for focusing on critical elements, making it more effective for analysing lengthy and complex news articles [5]. Evaluated on publicly available ISOT fake news dataset. The proposed model demonstrates superior performance compared to conventional fake news detection systems. The model's scalability and reliability make it a robust solution for real-world applications, offering a promising method to counteract incorrect data and enhance the credibility of online data.

This work highlights the advanced deep learning techniques in the ongoing battle against fake news and misinformation on digital platforms. In addition to Bi-LSTM, incorporating a self-attention mechanism further enhances the model's capabilities. Self-attention mechanisms enable the model to assign varying levels of significance to different words and phrases within the text. By focusing on the most relevant words, the model can ignore less significant details, improving its efficiency and accuracy in detecting fake news. This dynamic focus enables the model to handle long, complex articles more effectively, where certain pieces of information may be more critical than others in determining the content's veracity. This paper proposes an enhanced fake news identification model that combines the strengths of Bi-LSTM and self-attention procedures. By leveraging deep learning techniques, the model can better understand context, prioritise important information, and detect subtle forms of misinformation that other models might overlook. The proposed model is evaluated using publicly available datasets, demonstrating its superior performance compared to traditional methods for detecting fake news. This approach offers a promising solution to the growing problem of misinformation, providing a reliable and scalable tool for recognising fake news in real-time across various digital platforms [6].

1.1. Key Contributions of this Study

- Development of a Self-Attention and Bi-LSTM (SA- Bi-LSTM) model for precise fake news identification.
- Comprehensive investigation of LSTM, Bi-LSTM, and Bi-LSTM with self-attention mechanisms.
- Performance evaluation using various metrics to assess the model's effectiveness. Demonstration of the proposed model's improvements over traditional approaches.
- The integration of self-attention is pivotal for capturing long-range dependencies in textual data, while Bi-LSTM models effectively retain context from both forward and backward sequences.

Together, these methods enable the model to analyse complex linguistic patterns within fake news articles. By using these techniques, the proposed model aims to surpass the performance of traditional fake news identification systems.

2. Literature Review

Identifying fake news has emerged as a prominent research focus in recent years. Early approaches primarily relied on traditional machine learning techniques such as Support Vector Machines (SVM), Decision Trees, and Naive Bayes. These models focused on handcrafted features, such as word frequency, sentiment analysis, and readability metrics. While useful,

these methods often struggled with high-dimensional text data and the nuances of language, resulting in limited accuracy. As fake news detection tasks became more complex, there was a need for more advanced models capable of handling such intricacies. The emergence of deep learning revolutionised the approach to text classification tasks, including fake news detection. Models such as Convolutional Neural Networks (CNNs) and Recurrent Neural Networks (RNNs) have demonstrated considerable success in processing textual data. CNNs have been particularly effective at extracting local features from text, especially Long Short-Term Memory (LSTM) networks, which are adept at modeling temporal dependencies in sequential data. According to Sastrawan et al. [7], a fake news detection method is highlighted using CNN, Bi-LSTM, and ResNet architecture. This method is trained using four different datasets. Upon comparing the outputs, the bi-LSTM performs better than the CNN and ResNet results.

Moreover, Matheven and Kumar [8] discussed a deep learning-based approach for detecting fake news in natural language. In this scheme, a word-to-Vee model is integrated with an LSTM model, thereby exposing the characteristics of both models. The model is trained and tested on a dataset comprising both real news and fake news. Independent variables, including the number of training cycles, data diversity, and vector size, significantly influence the system's accuracy. Furthermore, Almandouh et al. [9] described a technique to improve the performance of fake news detection. In this approach, various techniques, including machine learning, ensemble learning, and deep learning, are applied. Furthermore, to optimise the performance, the hyperparameters of transformer-based models, including BERT, XLNet, and RoBERTa, are tuned. The two datasets, AFND and ARABICFAKETWEETS, have been processed. Furthermore, four deep learning models—CNN+LSTM, RNN+LSTM, CNN+RNN, and Bi-GRU+LSTM—are discussed to analyse their superior performance in terms of accuracy, F1-score, and loss metrics. Mohawesh et al. [10] demonstrated a framework for fake news detection using multilingual deep learning with a capsule neural network.

Accordingly, a semantic approach is presented to identify fake news based on relational parameters, such as entities, facts, and sentiments, which are derived from the text. The performance of this method is evaluated using the Tallip fake news dataset, specifically in the conversion from English to English, Hindi, Indonesian, Vietnamese, and Swahili. Another study, conducted by Rasul [11], could be applied to fake news identification, yielding significant improvements over traditional machine learning models. However, these models were limited in their ability to capture the full context of a news article, as they considered only the information, not both the information and its context simultaneously. The introduction of self-attention mechanisms, particularly with the Transformer architecture, has opened up possibilities for natural language processing tasks. Recent studies have demonstrated that self-attention mechanisms substantially enhance the performance of text classification models by facilitating a deeper understanding of a document's nuances [12].

This capability is particularly beneficial for detection, where certain phrases or patterns may be more likely to contain fabricated content than others. Bi-LSTM models have further enhanced the ability to understand text by processing it in both forward and backward directions. This dual perspective allows Bi-LSTMs to capture a complete understanding of the context within a sentence [13]. Due to their ability to consider context, Bi-LSTMs outperform standard LSTMs across various text classification tasks, including sentiment analysis and fake news detection. Combining Bi-LSTM with self-attention mechanisms, it has emerged as an approach to address the limitations of earlier models. Recent studies presented by Zhou et al. [14] proposed a model that integrates Bi-LSTM with self-attention for fake news detection, demonstrating an improved accuracy. This hybrid model leverages self-attention to focus on relevant parts of text while utilizing Bi-LSTM to maintain context from both directions. This combination provides a solution to the challenges of recognising fake news.

3. Proposed Method

3.1. Dataset

The dataset utilised in this work is the ISOT dataset, which was carefully curated to provide a balanced representation of real and fake news articles [15]. An esteemed and widely trusted news outlet, while the fake news samples were derived from various unreliable sources, including Wikipedia and other dubious sites identified by PolitiFact, a well-known fact-checking organisation. The dataset spans multiple domains, with a particular focus on political and international news, reflecting the kinds of stories where misinformation is often prevalent. The dataset is split into two CSV files: one containing real news articles labelled '1' and another containing fake news articles labelled '0'. The real news file comprises 21,417 items, and the fake news file comprises 23,481 items, for a total of 44,898 data points. The dataset includes key fields for each article, including title, text, label (real or fake), and publication date. Articles were predominantly sourced from topics relevant to the contemporary issues surrounding fake news during that period. The raw data was cleaned and processed, though the original punctuation errors in fake news articles were retained to preserve authenticity. The ISOT dataset serves as a critical resource for training and testing models for fake news identification, offering real-world challenges that mirror the complexity of distinguishing genuine from fake news (Figure 1).

author	published	title	text	language	site_url	main_img_url	type	label	title_without_stopwords	text_without_stopwords	hasImage
Barracuda Brigade	2016-10-26T21:41:00.000+03:00	muslims busted they stole millions in govt ben...	print they should pay all the back all the mon...	english	100percentfedup.com	http://bb4sp.com/wp-content/uploads/2016/10/Fu...	bias	Real	muslims busted stole millions govt benefits	print pay back money plus interest entire fam...	1.0
reasoning with facts	2016-10-29T08:47:11.259+03:00	re why did attorney general loretta lynch plea...	why did attorney general loretta lynch plead t...	english	100percentfedup.com	http://bb4sp.com/wp-content/uploads/2016/10/Fu...	bias	Real	attorney general loretta lynch plead fifth	attorney general loretta lynch plead fifth bar...	1.0
Barracuda Brigade	2016-10-31T01:41:49.479+02:00	breaking weiner cooperating with fbi on hillar...	red state infox news sunday reported this mor...	english	100percentfedup.com	http://bb4sp.com/wp-content/uploads/2016/10/Fu...	bias	Real	breaking weiner cooperating fbi hillary email ...	red state fox news sunday reported morning ant...	1.0
Fed Up	2016-11-01T05:22:00.000+02:00	pin drop speech by father of daughter kidnappe...	email kayla musler was a prisoner and torture...	english	100percentfedup.com	http://100percentfedup.com/wp-content/uploads/...	bias	Real	pin drop speech father daughter kidnapped kill...	email kayla musler prisoner tortured isis cha...	1.0
Fed Up	2016-11-01T21:56:00.000+02:00	fantastic trumps point plan to reform healthc...	email healthcare reform to make america great ...	english	100percentfedup.com	http://100percentfedup.com/wp-content/uploads/...	bias	Real	fantastic trumps point plan reform healthcare ...	email healthcare reform make america great sin...	1.0

Figure 1: Sample dataset

3.2. Bidirectional LSTM

While traditional LSTMs are highly effective at learning from sequential data, they have a limitation: they can only consider past information when making predictions (Figure 2). This unidirectional processing means the model may miss valuable information that could be derived from future data points within the sequence. Bidirectional LSTMs (Bi-LSTMs) were introduced to address this limitation by allowing the model to examine input data in both forward and backward directions [16].

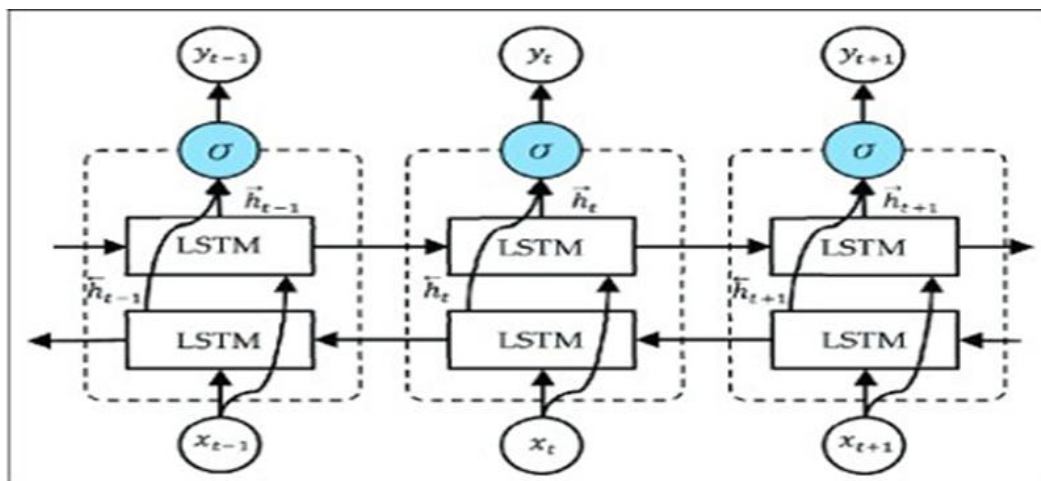


Figure 2: Architecture of bidirectional LSTM

In a Bi-LSTM, two LSTM layers are utilised — one examines the input from the beginning to the end (forward LSTM), and the other from the end to the beginning (backward LSTM). At each time, the hidden states from both the forward and backward LSTMs are combined to provide a more comprehensive representation of the information. This bidirectional methodology enables the arrangement to gather data from both past and future contexts, making it particularly useful for tasks such as fake news identification, where both prior and subsequent words in an article can provide important clues about the news's authenticity.

3.3. Bidirectional LSTM with Self-Attention

Despite improved performance, Bi-LSTMs still have limitations in their ability to focus on specific parts of the input. In huge natural language processing (NLP) tasks, certain words or phrases are more important than others in the determination of the final output [17]. This is also true in fake news identification, where specific keywords or terms may indicate whether a news item is genuine. To address this, a self-attention mechanism can be integrated into the Bi-LSTM model (Figure 3).

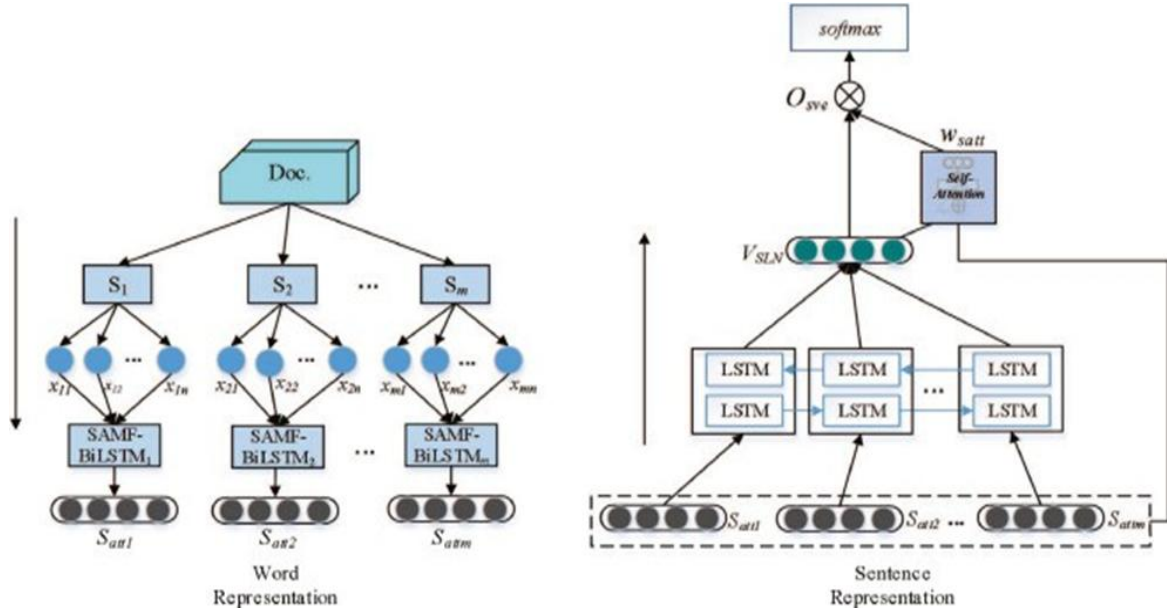


Figure 3: Bidirectional LSTM with self-attention

The self-attention mechanism assigns higher attention scores to words that are more relevant to the task at hand, thereby weighting the significance of words in the input sequences. In the proposed model, the output is passed through an attention layer that computes attention weights for each word in the input sequence. These attention weights are learned during training and used to assign greater importance to words more likely to influence the final classification (real or fake news) [18]. By incorporating self-attention, the model can focus on the most crucial words in the input sequence, thereby improving its ability to distinguish genuine from fake news articles accurately. To prevent overfitting, randomly set a portion of the inputs to '0' during training. The final output of the setup is obtained using a dense layer with a sigmoid activation function, which produces a binary prediction indicating whether the news article is real or fake.

3.4. Self-Attention Mechanism

The Self-attention mechanism is widely used in computer vision, excelling at image classification, object recognition, and more by enabling models to focus on key points in input data selectively. In this study, a self-attention layer is incorporated into the neural network architecture [19]. The process begins by transforming the query, key, and value. Using convolutional layers with a kernel size of 1, these matrices are reshaped into smaller dimensions. The dot product of the query and key matrices generates an energy matrix, which is then passed through a Softmax function to produce an attention weight, highlighting the most important features. This map is combined with the original input feature map and then passed through the remaining network layers. The self-attention mechanism enables the network to dynamically adjust its focus, improving performance across various vision tasks. Mathematically, the attention process is defined by:

$$\text{Attention} = \text{softmax}(\text{query} \times \text{key}) \quad (1)$$

$$\text{Out} = \text{attention} \times \text{value} \quad (2)$$

This enhances the model's capacity to prioritise essential features for the task at hand.

4. Result and Discussion

The setup was evaluated using both real and fake news dataset images, which were grouped into training and test sets. After training for a specific epoch count, its performance is assessed using multiple metrics. Accuracy examined the overall performance of the model; precision evaluated how accurately the model identified fake news; recall assessed the ability of the setup to identify fake news; and the F1-score provided a single performance metric that combined precision and recall. Additionally, a confusion matrix was used to provide a detailed breakdown of the classification results, displaying true positives and true negatives, as well as false positives and false negatives. These offered insights into the ability to identify fake news accurately and highlighted areas where it should struggle.

4.1. Input Dataset (Header)

The Author column lists the names of the individuals who wrote the news articles. This information helps identify the writer's credibility and background. In cases where the author is not specified, the entry is marked as "No Author," which may indicate content aggregated from multiple sources or anonymously published material. Knowing the author can also help analyse potential biases or expertise in the subject matter. The Published column provides the exact date and time when each news article was made publicly available. The timestamp format includes both the day and the precise time, allowing readers or analysts to track the timeliness and relevance of the content. This is particularly useful for monitoring breaking news or comparing articles published on the same topic at different times (Figure 4).

author	published	title	text	language	site_url	main_img_url	type	label	title_without_stopwords	text_without_stopwords
Alex Ansary	2016-11-05T00:33:12.210+02:00	bill clinton and hillary lolita express pedoph...	opinion hillary is the whore of babylon and is...	english	amtvmedia.com	http://www.amtvmedia.com/wp-content/uploads/20...	3	Fake	bill clinton hillary lolita express pedophilia...	opinion hillary whore babylon human tracy twym...
Dr. Patrick Slattery	2016-11-16T00:55:00.000+02:00	why our survival depends on the defeat of jews...	share ndr duke and pastor dankof quote jews b...	english	davidduke.com	http://davidduke.com/wp-content/uploads/2016/1...	5	Real	survival depends defeat jewish power	share dr duke pastor dankof quote jews boatin...
Phyllis Bentley	2016-11-27T23:11:00.000+03:00	why you should drink carrot juice daily how to...	keywords better tasting food gmos homestead ...	english	naturalnews.com	http://10667-presscdn-0-56.pagely.netdna-cdn.c...	6	Fake	sugar industry funding research sugarcoat truth	back pain pain side important treat properly r...
No Author	2016-10-26T23:47:00.000+03:00	top aide hillary still not perfect in her head...	email \n\nhillary clinton personally ordered a...	english	awdnews.com	http://awdnews.com/images/14774968271.jpg	1	Fake	top aide hillary still perfect head wikileaks	email hillary clinton personally ordered consu...
Lily Dane	2016-11-02T17:45:27.945+02:00	former us state dept official us intelligence ...	posted on october by davidswanson \npicture ...	english	thedailysheepie.com	No Image URL	3	Fake	assange donald trump wont allowed win clinton ...	illegal immigrants flooding across americas so...

Figure 4: Input dataset (Header)

The Title column displays the headline or title of each article. Although the titles may sometimes appear truncated due to display constraints, they are designed to give a concise summary of the article's main subject. Headlines often employ attention-grabbing words to attract readers, making them a key factor in assessing the article's focus and potential bias. The Text column contains selected snippets or excerpts from the main content of the news articles. These summaries highlight the central ideas or key points, enabling readers to quickly comprehend the content without needing to read the full text. This column is useful for tasks such as sentiment analysis, keyword extraction, and automated summarisation. The Language column specifies the language in which the article is written. For the provided dataset, all entries are in English, ensuring uniformity for natural language processing and analysis tasks. Knowing the language is critical for text-based classification, translation, and cross-lingual studies. The Site URL column provides the web address where the original article was published.

This allows users to verify the source, access the full article, and check the context in which the content appeared. It also plays a role in evaluating the article's trustworthiness, as reputable sites are generally more reliable. The Main Image URL column contains the direct link to the news article's featured image. These images are often used to summarise the story or to attract the reader's attention visually. Images can also be analysed for content, relevance, or signs of manipulation when assessing the authenticity of news. The Type column contains a numerical value that may represent different categories, classifications, or source types related to the article. Examples of values include 1, 3, or 5, although the precise meaning of each number is not explicitly defined in the dataset. This column can be important for organising articles by topic, format, or source, and may support automated filtering or categorisation. The Label column indicates whether the article is classified as Fake or Real. This is crucial for research and applications involving the detection of fake news, content verification, and automated fact-checking. The label provides ground truth for training machine learning models, helping them learn patterns that distinguish trustworthy from misleading news content.

4.1.1. Key words for Classification

The image you provided is a word cloud of keywords used for classification, likely from a text analysis or machine learning model. The word cloud in the article represents the frequency, with larger words representing higher frequency or importance in the corpus. In this particular word cloud, prominent keywords include "Trump," "said," "people," "one," and "Clinton." These terms are likely to feature prominently in discussions of political news, given the presence of prominent figures such as Donald Trump and Hillary Clinton. Other significant words, including "time," "election," "country," and "American," suggest themes related to political campaigns, governance, and public discourse. Smaller yet notable words, such as the "right," "need,"

The Bias (green) category includes articles that are heavily opinionated or one-sided. These articles often reflect the author's personal or organisational viewpoints rather than objective reporting. Bias can subtly influence readers' perceptions, making it important to identify and analyse such content, especially when building models to detect misinformation or politically slanted news. The Hate (red) category encompasses articles that promote hate speech, discriminatory language, or content designed to demean or attack individuals or groups. This category signals malicious intent and raises ethical concerns, underscoring the need for careful monitoring and detection to prevent the dissemination of harmful or inflammatory material. The Satire (purple) category comprises articles that humorously or sarcastically critique politics, society, or current events. While satirical content is not intended to mislead, it can sometimes be misinterpreted as factual, presenting a unique challenge in news classification and automated detection tasks.

The State (brown) category represents articles published by state-controlled or government-influenced media outlets. These articles may reflect official narratives, propaganda, or selective reporting, making it crucial to differentiate between independent reporting and content shaped by institutional interests. The Junk Science (pink) category includes articles that promote pseudoscientific claims or misinformation about science and technology. These articles can mislead readers about scientific facts, public health, or technological developments, emphasising the importance of verifying sources and evidence. The Fake (grey) category covers a smaller portion of the dataset but includes articles with entirely fabricated content. These articles are intentionally deceptive, aiming to misinform readers with false stories or claims, and are a key focus for fake news detection research.

4.1.3. Output

The image distinguishes between websites that publish Real News and those that publish Fake News (Figure 7).

REAL NEWS OUTPUT

```
Websites publishing real news:['100percentfedup.com', 'addictinginfo.org', 'dailywire.com', 'davidduke.com', 'fromthetrenchesworldreport.com', 'frontpagemag.com', 'newstarget.com', 'politicususa.com']
```

FAKE NEWS OUTPUT

```
Websites publishing fake news:['21stcenturywire.com', 'abcnews.com.co', 'abedanger.net', 'abovetopsecret.com', 'activistpost.com', 'adobochronicles.com', 'ahtribune.com', 'allnewspipeline.com']
```

Figure 7: Outputs

Real News Output (Top Section): This section lists reliable sources of real news. Some of these websites include:

- 100percentfedup.com
- addictinginfo.org
- dailywire.com
- davidduke.com
- newstarget.com

These sites are categorised as authentic news sources and are recognised for delivering factual, accurate content.

Fake News Output (Bottom Section): This part highlights websites known for publishing fake or misleading news. Some websites in this category include 21stCenturyWire.com, abcnews.com.co, and AboveTopSecret.com. These sources are flagged as unreliable, with a tendency to publish fabricated or exaggerated content, often intended to mislead the audience. This diagram emphasises the importance of distinguishing between credible and non-credible news sources, guiding readers on where to find trustworthy information and where to approach with scepticism. The provided image compares Real News and Fake News articles based on their authors, source URLs, labels, and main image URLs (Figure 8).

- **Real News (Left Panel):** This section lists articles labelled as “Real” from the website 100percentfedup.com. All entries are authored by “Fed Up.” An image accompanies each article. The uniformity across the author and source sites suggests that these articles were classified as real and came from a consistent, trustworthy source.
- **Fake News (Right Panel):** This section contains articles labelled as “Fake” from 21stcenturywire.com. Various authors are listed, including “No Author,” “Shawn Helton,” and “Mike Rivero,” indicating that some articles lack

clear attribution, a potential red flag for fake news. The image URLs, which display diverse and potentially sensational visuals, are linked to articles that were classified as fake, underscoring their misleading nature.

	site_url	label	main_img_url
author			
Fed Up	100percentfedup.com	Real	
Fed Up	100percentfedup.com	Real	
Fed Up	100percentfedup.com	Real	
Fed Up	100percentfedup.com	Real	

	site_url	label	main_img_url
author			
No Author	21stcenturywire.com	Fake	
No Author	21stcenturywire.com	Fake	
Shawn Helton	21stcenturywire.com	Fake	
Mike Rivers	21stcenturywire.com	Fake	
No Author	21stcenturywire.com	Fake	
Shawn Helton	21stcenturywire.com	Fake	

Figure 8: Real and fake news

The comparison between these two panels illustrates the differences between real and fake news, particularly emphasising the importance of clear authorship, reliable sources, and factual reporting in distinguishing real news from fake or misleading information.

5. Conclusion

The integration of the self-attention mechanism with Bidirectional Long Short-Term Memory (Bi-LSTM) offers a powerful and accurate approach for identifying fake news in digital media. Bi-LSTM captures the sequential dependencies of the text by processing the data in both forward and backward directions, which are essential for understanding context. However, relevant parts of the text identify key elements that are crucial for discriminating between real and fake news. This combination enables the model to handle complex and nuanced content more effectively than simpler approaches, resulting in higher accuracy in identifying fake news. Furthermore, the setup's flexibility enables it to be applied across various types of data and languages, making it a versatile tool in the global war against misinformation. Although computationally intensive, advancements in machine learning and hardware are making these models more accessible. The self-attention with Bi-LSTM approach represents a significant advancement in identifying fake news by leveraging deep contextual understanding and dynamic focus, contributing to more reliable and intelligent information for processing systems.

Acknowledgment: The authors express their sincere gratitude to Alliance University, Sri Krishna College of Engineering and Technology, Sai Vidya Institute of Technology, Presidency University, and SRM Institute of Science and Technology for their constant support and academic guidance. We also thank all faculty and contributors from these institutions for their insights and encouragement, which made this work possible.

Data Availability Statement: The datasets generated and analyzed for this study, Smart Approach for Precise Fake News Detection Using Bi-Directional LSTM and Self-Attention Mechanism, are available with the author team and can be provided upon reasonable request. All data shared will comply with institutional guidelines and participant privacy requirements.

Funding Statement: This research and manuscript preparation were carried out solely through the collective efforts of the authors. No external financial support, grants, or funding agencies were involved in any stage of the work.

Conflicts of Interest Statement: The authors hereby declare that there are no financial, institutional, or personal conflicts of interest that could have influenced the outcomes of this study. All referenced sources and citations are properly acknowledged.

Ethics and Consent Statement: Ethical procedures were strictly followed throughout the study. Necessary permissions were obtained from the relevant organizations, and informed consent was secured from all participants involved in the data collection process.

References

1. H. Allcott and M. Gentzkow, "Social Media and Fake News in the 2016 Election," *Journal of Economic Perspectives*, vol. 31, no. 2, pp. 211–236, 2017.
2. S. Ghosh and T. Vealee, "Fracking Sarcasm using Neural Network Models," *Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, California, United States of America, 2016.
3. T. Mikolov, K. Chen, G. Corrado, and J. Dean, "Efficient Estimation of Word Representations in Vector Space," *arXiv preprint arXiv:1301.3781*, 2013. Available: <https://arxiv.org/abs/1301.3781> [Accessed by 04/10/2024].
4. J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," in *Proc. NAACL-HLT*, Minnesota, United States of America, 2019.
5. H. Rashkin, E. Chooi, J. Y. Jang, S. Volkova, and Y. Choi, "Truth of Varying Shades: Analyzing Language in Fake News and Political Fact-Checking," in *Proc. Conf. Empirical Methods in Natural Language Processing (EMNLP)*, Suzhou, China, 2017.
6. F. Farrhangian, R. M. O. Cruz, and G. D. C. Cavalcanti, "Fake news detection: Taxonomy and comparative study," *Information Fusion*, vol. 103, no. 3, p. 102140, 2024.
7. I. K. Sastrawan, I. P. A. Bayupati, and D. M. S. Arsa, "Detection of fake news using deep learning CNN–RNN based methods," *ICT Express*, vol. 8, no. 3, pp. 396–408, 2022.
8. A. Matheven and B. V. D. Kumar, "Fake News Detection Using Deep Learning and Natural Language Processing," in *Proc. 9th Int. Conf. Soft Computing and Machine Intelligence (ISCMI)*, Ontario, Canada, 2022.
9. M. E. Almandouh, M. F. Alrahmawy, M. Eisa, M. Elhoseny, and A. S. Tolba, "Ensemble based high performance deep learning models for fake news detection," *Scientific Reports*, vol. 14, no. 11, p. 26591, 2024.
10. R. Mohawesh, S. Maqsood, and Q. Althebyan, "Multilingual deep learning framework for fake news detection using capsule neural network," *Journal of Intelligent Information Systems*, vol. 60, no. 3, pp. 655–671, 2023.
11. P. Rasul, "Fake News Detection Using Machine Learning," *Indonesian Journal of Computer Science*, vol. 12, no. 4, pp. 1602–1609, 2023.
12. N. F. Baarir and A. Djeflal, "Fake News detection Using Machine Learning," in *Proc. 2nd Int. Workshop Human-Centric Smart Environments for Health and Well-being (IHSH)*, Boumerdes, Algeria, 2021.
13. X. Zhang and A. A. Ghorbani, "An overview of online fake news: Characterization, detection, and discussion," *Information Processing and Management*, vol. 57, no. 2, p. 102025, 2020.
14. P. Zhou, W. Shi, J. Tiann, Z. Qi, B. Li, H. Hao, and B. Xu, "Attention-Based Bidirectional Long Short-Term Memory Networks for Relation Classification," in *Proc. 54th Annu. Meeting Assoc. for Computational Linguistics*, Berlin, Germany, 2016.
15. Kaggle, "ISOT Fake News Dataset," 2021. <https://www.kaggle.com/datasets/csmalarkodi/isot-fake-news-dataset> [Accessed by 14/10/2024].
16. J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of deep bidirectional transformers for language understanding," *arXiv preprint arXiv:1810.04805*, 2018. <https://arxiv.org/abs/1810.04805> [Accessed by 04/10/2024].
17. J. Jiang, P. Li, A. U. Haq, A. Saboor, and A. Ali, "A novel stacking approach for accurate detection of fake news," *IEEE Access*, vol. 9, no. 2, pp. 22626–22639, 2021.
18. M. Hakak, S. Alazab, S. Khan, T. R. Gadekallu, P. K. R. Maddikunta, and W. Z. Khan, "An ensemble machine learning approach through effective feature extraction to classify fake news," *Future Generation Computer Systems*, vol. 117, no. 4, pp. 47–58, 2021.
19. Y. Khan, M. T. I. Khondaker, S. Afroz, G. Uddin, and A. Iqbal, "A benchmark study of machine learning models for online fake news detection," *Machine Learning with Applications*, vol. 4, no. 2, p. 100032, 2021.