

Implementation of Bidirectional Long Short-Term Memory and Convolutional Neural Network in Detecting Hoax Content on Social Media

Oktavian A. Lantang¹, Raphael E. C. Sendow^{2*}, Feisy Diane Kambey³

Abstract—The advancement of internet technology has facilitated the spread of information, including false information or fake news. The dissemination of hoaxes on social media, such as Twitter, can cause confusion and negatively impact society. This study aims to implement a hybrid model that combines Bidirectional Long Short-Term Memory (Bi-LSTM) and Convolutional Neural Network (CNN) for hoax detection. The dataset used consists of English tweets containing both real and fake news, collected between 2020 and 2022, as provided by the TruthSeeker dataset. The model utilizes an embedding layer with word2vec, a Conv1D layer, and a BiLSTM layer to effectively capture temporal and spatial patterns in text data. Additionally, experiments were conducted by varying the number of BiLSTM units and CNN filters to analyze their impact on model performance. After conducting parameter experiments, the best results were achieved using a Conv1D layer with 64 filters and a BiLSTM layer with 64 neurons/units. The evaluation results on the test data indicate an accuracy of 96.14%, a precision of 96%, a recall of 96.25%, and an F1-score of 96%. These results demonstrate the model's high capability in accurately detecting hoaxes, which is significant for combating misinformation on social media. With its strong performance, the model has potential applications in real-time content moderation systems, early hoax detection tools, and digital literacy platforms to help reduce the spread of false information.

Index Terms—Hoaxes, bi-LSTM, CNN, word2vec, social media.

I. INTRODUCTION

Advancements in internet technology have made it easier to

access information through online platforms, including websites and social media like Twitter (now X), which has 528.3 million active users worldwide [1] and approximately 500 million tweets daily [2]. Despite its benefits, social media can also be used to spread hoaxes and false information that can be misleading [3]. In Indonesia, the Ministry of Communication and Information (KOMINFO) recorded over 11,000 hoaxes related to politics, health, and natural disasters between August 2018 and March 2023 [4]. This phenomenon also occurs in different countries, highlighting the urgent need to identify and address hoaxes more effectively due to their widespread impact on public trust, social stability, and economic loss. Therefore, conducting this research is crucial for developing automated solutions that can support early hoax detection. The use of bidirectional long short-term memory (Bi-LSTM) is particularly important because it captures context from both past and future word sequences, making it well-suited for understanding the structure and intent of hoax-related content in social media texts [5].

Several approaches have been applied in hoax detection research, including deep learning techniques such as CNN, LSTM, Bi-LSTM, and their hybrid models. For instance, Bi-LSTM with dropout achieved a maximum accuracy of 96.60% [6], while the LSTM-CNN hybrid model utilizing word2vec reached an accuracy of 79.71% [11].

Research also shows that the Bi-LSTM-CNN hybrid model with word2vec achieves better accuracy than standalone models such as CNN and Bi-LSTM, outperforming the Multi-Layer Perceptron (MLP) model [7].

Some previous studies have shown different results in hoax detection. These differences are due to factors like the type of model used, data set characteristics, amount of data, and parameters used. Thus, this study uses the Truth Seeker dataset to focus on hoax detection on Twitter. This dataset is obtained by crawling using keywords from the PolitiFact dataset.

The Bi-LSTM model has been shown to achieve high accuracy in text classification with specific configurations

Received: 3 March 2025; Revised: 26 May 2025; Accepted: 31 May 2025.

*Corresponding author

¹Oktavian A. Lantang, Dept. Informatics Engineering, Sam Ratulangi University, Manado, Indonesia (e-mail: oktavian_lantang@unsrat.ac.id).

²Raphael E. C. Sendow, Dept. Informatics Engineering, Sam Ratulangi University, Manado, Indonesia (e-mail: raphaelsendow026@student.unsrat.ac.id).

³Feisy Diane Kambey, Dept. Informatics Engineering, Sam Ratulangi University, Manado, Indonesia (e-mail: feisykambey@unsrat.ac.id).

based on previous studies. Therefore, this research aims to implement a hybrid Bi-LSTM and CNN model with word2vec embedding to enhance hoax detection accuracy on social media. The main contributions of this study include: (1) evaluating different configurations of CNN and Bi-LSTM to determine the optimal structure, and (2) demonstrating the model's effectiveness in detecting hoaxes using a large-scale, real-world Twitter dataset [7].

II. RELATED WORK

This chapter describes the literature review that supports research, including theory, methods/techniques, models, algorithms, metrics or performance measurements as a reference in conducting research and must refer to journal articles.

A. Convolutional Neural Network (CNN)

A convolutional neural network (CNN) is a type of neural network architecture that is widely utilized in image recognition tasks, but it has also shown strong performance in natural language processing (NLP), particularly for text classification tasks [16]. CNN consists of three key components: the convolutional layer, which extracts significant features from the input through convolution operations; the pooling layer, which decreases data dimensionality and mitigates overfitting by selecting the most relevant features; and the fully connected layer [28], which links all neurons to generate the final classification output, typically using a Sigmoid activation function for binary classification tasks [17].

B. Bidirectional Long Short-Term Memory (Bi-LSTM)

Bidirectional long short-term memory (Bi-LSTM) is a variant of LSTM that utilizes two layers processing input in opposite directions: one moving forward and the other backward. This dual-directional flow helps the model understand the full context of each word within a sequence by considering both preceding and succeeding words. As a result, Bi-LSTM is capable of generating richer and more precise representations of text, making it highly suitable for various natural language processing applications [7].

C. Related Research

Reference [11] proposed a hybrid LSTM-CNN model using word2vec for detecting hoaxes related to COVID-19 on Twitter. Their model achieved an accuracy of 79.71%, outperforming the standalone LSTM and CNN architectures. This indicates that combining sequential and spatial feature extraction improves performance, especially in short social media texts.

Reference [10] implemented a CNN model with word2vec (skip-gram) for hoax news detection. Their model achieved 91% accuracy, precision, and recall, highlighting CNN's effectiveness in capturing local textual patterns. However, its performance heavily relied on the choice of filter size and embedding dimensions, which may limit generalization.

Reference [14] proposed a CNN-BiLSTM model with GloVe embeddings, achieving 96% accuracy. This architecture benefitted from CNN's feature extraction and BiLSTM's

contextual understanding, performing well particularly on large datasets.

Reference [26] employed XGBoost combined with sentiment and source-based features for hoax detection using the TruthSeeker dataset. The model achieved 93.35% accuracy. Its advantage lies in interpretability and integration of metadata, though it lacks the deep semantic understanding offered by neural networks.

All studies focus on hoax/fake news detection using different model combinations. The key similarity is the use of word embeddings to capture semantic meaning. However, the models differ in architecture and features. By ref. [11] and [14], they used hybrid models (LSTM/CNN), but only [14] included bidirectionality, which likely contributed to higher accuracy. Meanwhile, [10] relied solely on CNN, resulting in strong local feature capture but limited contextual comprehension. Also, [26] utilized traditional machine learning (XGBoost) with engineered features, offering better explainability but slightly lower accuracy.

Similarly model [11] was better than standalone models but had lower accuracy. Moreover, CNN model [10] effectively extracted features but lacked sequence context. Another CNN-BiLSTM model [14] achieved high accuracy by combining contextual and spatial features, although it required a large dataset. XGBoost model [26] was interpretable and used metadata but lacked deep semantic capture.

While previous studies have demonstrated the effectiveness of deep learning architectures such as LSTM, CNN, and their hybrid combinations, many of them either lacked bidirectional context modeling (as in [11]) or did not fully exploit spatial-temporal relationships in textual data (as in [10]). Moreover, some models were tested on limited datasets or lacked fine-tuned hyperparameter optimization, which may affect generalizability. In contrast, our study builds upon these works by integrating both CNN and Bi-LSTM in a complementary manner, using a large-scale, real-world dataset (TruthSeeker), and applying comprehensive experimentation with various filter and unit configurations to achieve more robust and accurate results. This allows our model to better capture both contextual dependencies and textual patterns, which is essential for detecting nuanced hoax content on social media platforms.

III. RESEARCH METHOD

This study adopts the CRISP-DM (Cross Industry Standard Process for Data Mining) framework to analyze and evaluate the performance of a hybrid model that combines Bi-LSTM and CNN for detecting hoax content on social media platforms. The framework consists of six main stages: business understanding, data understanding, data preparation, modeling, evaluation, and deployment. In the context of this study, these stages are represented by processes such as data collection, preprocessing, word embedding (as data transformation), data splitting, model development, and performance evaluation. Each stage is designed to ensure a systematic and structured approach to extracting useful knowledge from text data. The complete

workflow of this research, adapted from the CRISP-DM methodology, is illustrated in Fig. 1.

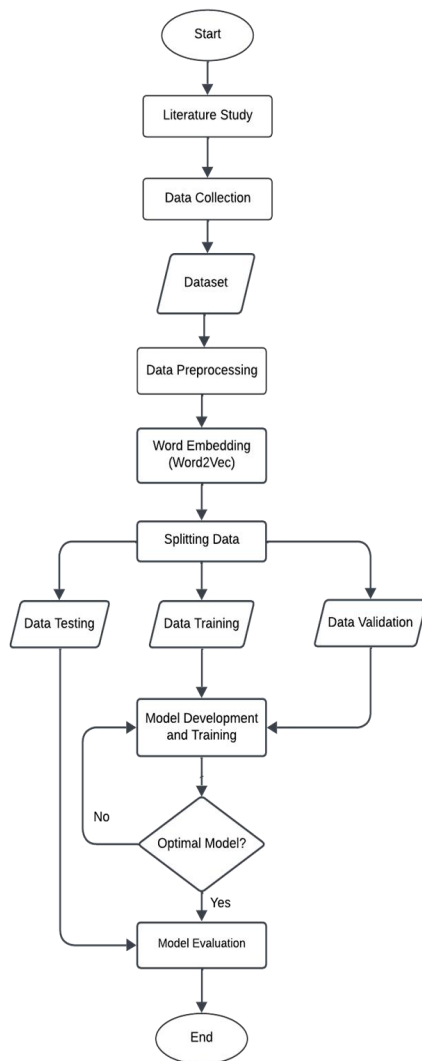


Fig. 1. Research stages.

A. Data Collection

The dataset used in this study is the Truth Seeker Dataset by Dadkhah et al., which consists of more than 134,000 tweets associated with 579 verified news stories and 479 false ones. Recognized as one of the most extensive datasets for analyzing fake news on social media, the data was gathered through the Twitter API using keywords sourced from the PolitiFact dataset [18].

B. Data Preprocessing

The preprocessing stage is needed to cleanse tweets and news texts, which often do not conform to the EYDs and contain symbols, numbers, or abbreviations. This stage aims to facilitate data processing in model building [19]. Data cleaning eliminates irrelevant elements, including URLs, usernames,

email addresses, hashtags, punctuation marks, new line characters, the word "amp," and repeated unnecessary characters, ensuring the text is cleaner and ready for analysis. Case folding converts all text to lowercase to prevent discrepancies between uppercase and lowercase letters. Stopword removal eliminates frequently used words that carry little meaningful information, such as "the," "and," and "is" [20]. Lemmatization reduces words to their base form (lemma), for example, changing "running" to "run" and "better" to "good" [21]. Tokenization breaks text into smaller units (tokens), such as words or subwords, to simplify further processing and analysis [22].

C. Word Embeddings (Word2Vec)

Word embeddings convert words into numerical vectors, enabling the model to capture the meaning and relatedness of words, where words with similar meanings are positioned near each other within the vector space. One widely used method for creating word embeddings is word2vec, which employs neural network models to produce these vector representations [23]. Word2vec includes two main architectures: Continuous Bag of Words (CBOW) aims to estimate a target word using the context provided by neighboring words, while the Skip-gram model functions oppositely by using a target word to predict its contextual surroundings [24]. In this study, the word2vec model is implemented using the CBOW technique via the Gensim library, generating vectors with a vocabulary of 53,236 terms and an embedding size of 100 dimensions.

D. Splitting Data

The dataset is divided into three parts: training, validation, and testing sets. The training set takes up the largest portion, while the validation and test sets share the remaining data equally. This research adopts a split ratio of 70% for training, 15% for validation, and 15% for testing to provide a reliable evaluation of the model's performance.

E. Model Development

At this stage, a hoax detection model is constructed by integrating Convolutional Neural Network and Bidirectional Long Short-Term Memory. The detailed structure of the proposed model is depicted in Fig. 2.

To optimize the model's performance, several hyperparameters were carefully chosen based on empirical experiments and insights from previous studies. A batch size of 32 was selected to balance training stability and computational efficiency. The learning rate was set to 0.001, as it provided steady convergence without overshooting during training. A dropout rate of 0.3 was applied to mitigate overfitting, while L2 regularization was included to further enhance generalization. The embedding dimension was fixed at 100 to ensure meaningful word representations without excessive computational cost.

Furthermore, the preprocessing techniques applied—such

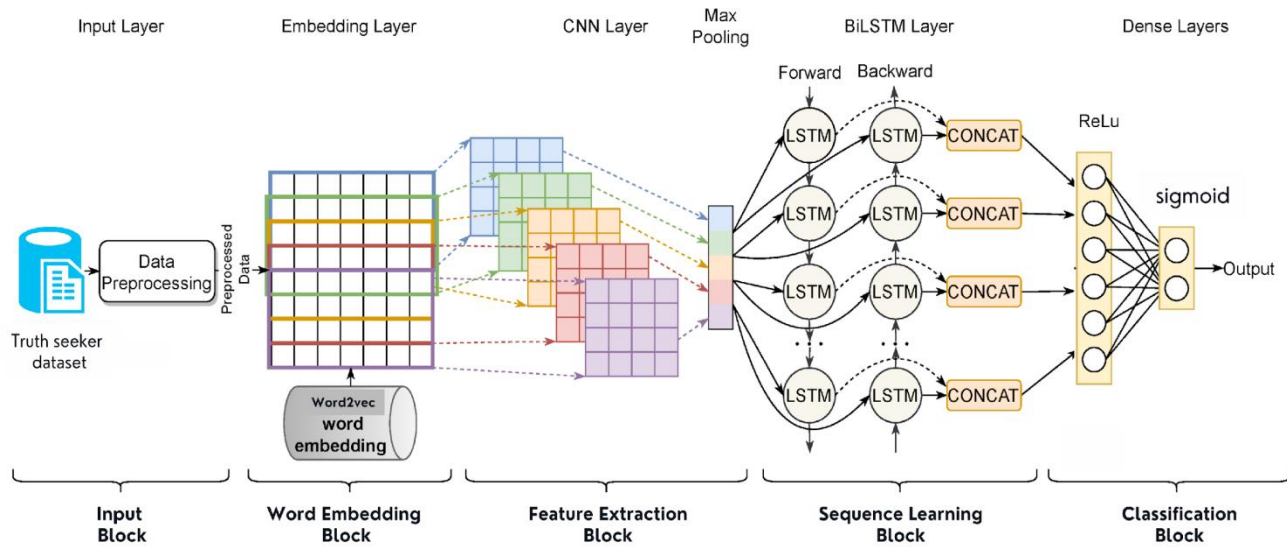


Fig. 2. The proposed model architecture.

as stopword removal, lemmatization, and tokenization—had a significant impact on the model's performance. Removing stopwords helped reduce noise in the input text, while lemmatization ensured that words were standardized to their base forms, allowing the model to better recognize patterns across similar words. These techniques collectively contributed to improving the model's ability to generalize and accurately classify hoax and non-hoax content. A detailed list of the hyperparameters used is provided in Table 1.

Table 1.

Hyperparameter Setting	
Hyperparameter	Value
Embedding Size	100
Word embedding	CBOW
Dropout rate	0.3
Batch size	32
Learning rate	1×10^{-3}
Optimizer	Adam
Regularizer	L2

F. Model Testing

This study assesses the model's performance by comparing four different combinations of Bi-LSTM unit counts and CNN filters, namely CNN 32-BiLSTM 32, CNN 32-BiLSTM 64, CNN 64-BiLSTM 32, and CNN 64-BiLSTM 64, to determine the best configuration for detecting hoax content on social media.

Model performance was assessed using a Confusion Matrix, which provides a comprehensive summary of prediction outcomes, including true positives (TP), true negatives (TN), false positives (FP), and false negatives (FN). Based on this matrix, several key evaluation metrics were derived, including:

- 1) *Accuracy*: Indicates the percentage of correct predictions out of the total number of predictions.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

- 2) *Precision* reflects how accurately the model identifies the

positive class.

$$Precision = \frac{TP}{TP + FP} \quad (2)$$

- 3) *Recall*: Reflects the model's capability to correctly identify all relevant instances of the positive class.

$$Recall = \frac{TP}{TP + FN} \quad (3)$$

- 4) *F1 score*: Combines both Precision and Recall into a single metric, offering a balanced evaluation of their trade-off [27].

$$F1 - Score = 2 \times \frac{Recall \times Precision}{Recall + Precision} \quad (4)$$

These metrics are used to evaluate and compare the performance of different model configurations in handling binary classification problems.

IV. RESULT

This chapter explains the findings considering what was already known about the issues being investigated.

A. Dataset

After going through the preprocessing stage, a dataset consisting of 98,019 text data was obtained, which has been cleaned and processed. This dataset is divided into 68,611 data for training, 14,702 data for validation, and 14,703 data for testing. Each text has a maximum length of 47 words, which is set based on the distribution of text length in the dataset.

The dataset is labeled into two categories: hoax (1) and non-hoax (0), making it suitable for binary classification in detecting hoax content on social media. With a large dataset, the model is expected to identify more complex patterns in both hoax and non-hoax texts. Table 2 presents the dataset distribution after preprocessing.

Table 2.
Detail Dataset

Total	Training	Validation	Testing	Label	Length (Max)
98.019	68.611	14.702	14.703	2	47

B. Comparison of CNN and BiLSTM Configurations

In these stages, combinations of the number of BiLSTM units and the number of CNN filters are evaluated to determine the best model configuration for detecting hoax content on social media. The model was tested using four main parameter combinations: CNN 32-BiLSTM 32, CNN 32 -BiLSTM 64, CNN 64-BiLSTM 32, and CNN 64-BiLSTM 64. The evaluation is conducted by comparing the accuracy of the test data (testing accuracy), with the results presented in Table 3.

Table 3.
Comparison of Text Accuracy for Different Configuration

Model Configuration	Testing Accuracy (%)
CNN 32 - BiLSTM 32	95.54%
CNN 32 - BiLSTM 64	95.51%
CNN 64 - BiLSTM 32	95.48%
CNN 64 - BiLSTM 64	96.14%

The test results show that the CNN 64-BiLSTM 64 configuration provides the highest accuracy of 96.14% compared to other configurations. Meanwhile, the CNN 32-BiLSTM 32 configuration obtained an accuracy of 95.54%, followed by CNN 32-BiLSTM 64 with 95.51% and CNN 64-BiLSTM 32 with 95.48%.

These results suggest that increasing the number of CNN filters and BiLSTM units enhances the model's performance. Models with 64 CNN filters and 64 BiLSTM units can capture more features from the text, thus increasing the accuracy in detecting hoax content. Therefore, the CNN 64-BiLSTM 64 configuration was selected as the best model for further evaluation.

The superior performance of the CNN 64-BiLSTM 64 configuration can be attributed to the increased model capacity, which allows the architecture to extract more complex spatial and sequential features from the text data. The larger number of CNN filters (64) enhances the model's ability to detect more nuanced patterns and phrases, while the 64 BiLSTM units improve its capability to capture longer and more meaningful contextual dependencies from both directions in a tweet. This synergy between convolutional and recurrent layers likely contributed to the observed performance gains.

However, this improvement comes at the cost of higher computational requirements and longer training time, which may present challenges for real-time deployment or environments with limited hardware resources. These trade-offs should be considered when adapting the model for practical applications.

C. Performance Evaluation of the Best Model

Once the best-performing model is identified from the previous experiments, a more in-depth evaluation is carried out to analyze its effectiveness in the classification task. This involves plotting the training accuracy and loss to monitor the model's learning progress, along with assessing its performance on the test dataset through key evaluation metrics such as accuracy, precision, recall, F1-score, and loss. A comprehensive summary of the model's performance is provided in Table 4.

Table 4.
Performance Evaluation Metrics

	Precision	Recall	F1-Score	Support
Non hoaks	0.95	0.96	0.95	5713
Hoaks	0.97	0.97	0.97	8989
Macro avg	0.96	0.96	0.96	14702
Weighted avg	0.96	0.96	0.96	14702
Testing Accuracy	0.9614			
Testing Loss	0.1201			

According to the evaluation results in Table 4, the model achieves a test accuracy of 96.14%, which indicates that the model can detect hoax and non-hoax content with a high level of accuracy. Additionally, the test loss is 0.1201, indicating a low prediction error rate. This suggests that the model is well-trained and does not suffer from overfitting.

Regarding precision, recall, and F1-score, the model demonstrates a balanced performance across both classes, with a precision value of 0.95 for non-hoaxes and 0.97 for hoaxes. This result indicates that the model has a high ability to identify hoax content correctly, as well as minimal errors in classifying non-hoax news as hoaxes.

A recall of 0.96 for non-hoax and 0.97 for hoax indicates that the model effectively identifies nearly all samples in their respective classes. This is further supported by an F1-score of 0.95 for non-hoax and 0.97 for hoax, demonstrating a well-balanced performance between precision and recall in detecting both categories.

At the overall level, the macro and weighted averages of precision, recall, and F1-score are approximately 0.96, indicating that the model maintains consistent performance across both classes without bias.

Figures 3 and 4 present the accuracy and loss graphs, providing insights into the model's convergence and stability during training. The graph shows a gradual increase in training accuracy from the beginning to the end of training. In the first epoch, the accuracy started at 81.31% and continued to rise, reaching 97.37% in the final epoch before early stopping. Meanwhile, the training loss significantly decreased from 0.5583 to 0.0827, indicating that the model effectively learned patterns in the data.

The validation curve also shows a trend in line with the

training curve. Val_accuracy increased from 90.67% in the first epoch to reach 95.99%, while val_loss decreased from 0.2599 to 0.1231. Although there is a slight fluctuation in val_loss in some epochs, the difference between training loss and validation loss remains relatively small. This indicates that the model does not suffer from significant overfitting. In addition, the use of ReduceLROnPlateau helps adjust the learning rate when the loss does not decrease significantly. This can be seen from the decrease in learning rate from 0.001 to 0.0000625 in the last epoch, which helps the model achieve a more stable convergence.

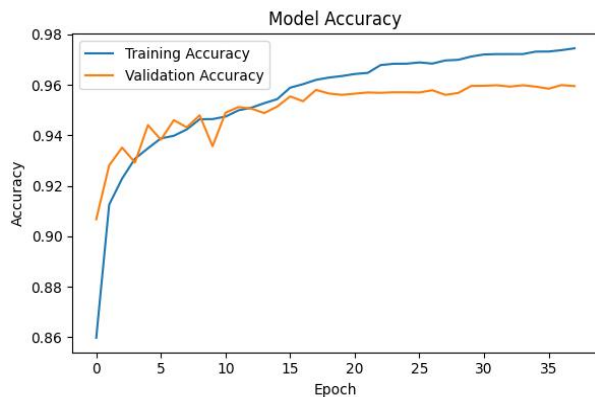


Fig. 3. Model accuracy.

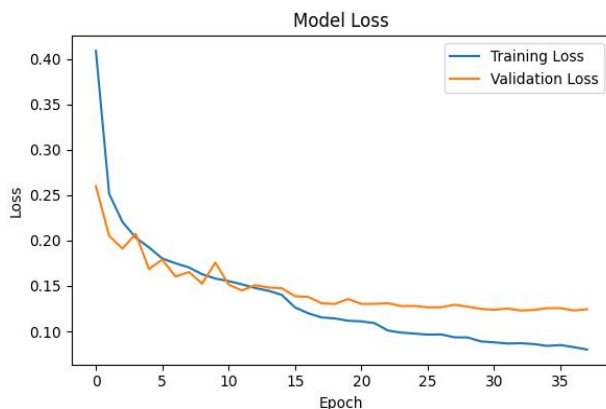


Fig. 4. Model loss.

To further analyze the model's classification effectiveness, a Confusion Matrix is employed, as depicted in Figure 5. This matrix visualizes how the model classifies instances into hoax (1) and non-hoax (0) categories, displaying the counts of TP, TN, FP, and FN.

The confusion matrix reveals that the model has a low rate of misclassification, with only a few instances of false positives (FP) and false negatives (FN). This suggests that the model effectively identifies hoax content while minimizing errors in classifying non-hoax news as hoaxes or vice versa.

Based on these evaluation results, it can be concluded that the CNN-BiLSTM model in this study demonstrates competitive performance in detecting hoax content on social media. This model serves as a solid foundation for developing more accurate hoax detection systems.

D. Comparison of Previous Studies

To measure how effective the proposed approach is, we compared it with previous studies that employed different architectures for hoax detection. This comparison aims to evaluate the performance improvements achieved with the CNN-BiLSTM model. Table 5 presents the accuracy comparison between the proposed method and prior research.

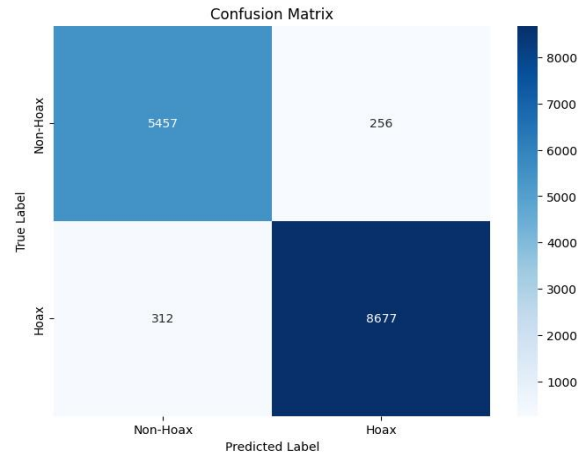


Fig. 5. Confusion matrix.

Table 5.
Comparison Results

Method	Architecture	Accuracy
Ref. [11]	LSTM-CNN	79.71%
Ref. [10]	CNN	91%
Ref. [14]	BiLSTM-CNN	96%
Ref. [26]	XGBoost	93.35%
Proposed Method	CNN-BiLSTM	96.14%

Based on Table 5, the proposed CNN-BiLSTM method achieves the highest accuracy of 96.14%, surpassing other architectures used in previous studies. Compared to [11] LSTM-CNN model, which achieved 79.71%, and [10] CNN model with 91% accuracy, the proposed method shows a significant improvement. In addition, it outperforms the BiLSTM-CNN model of [14] which achieved 96%, indicating the effectiveness of the proposed improvements.

Notably, the XGBoost model [26], which achieved 93.35% accuracy, used the same dataset as this study. While XGBoost effectively uses decision tree-based learning with feature engineering, the CNN-BiLSTM model surpasses it by capturing spatial and sequential text patterns more comprehensively. CNN captures essential textual features, while BiLSTM enhances contextual comprehension by processing information in both forward and backward directions.

Despite achieving higher accuracy, the comparison should be interpreted with caution due to possible differences in preprocessing steps, model complexity, and evaluation protocols. For instance, models like XGBoost relied heavily on handcrafted features such as sentiment scores and source credibility, while the proposed method leverages deep neural networks that automatically learn feature representations.

Moreover, unlike some prior studies, this work did not incorporate ensemble techniques, attention mechanisms, or multi-modal inputs, which may further enhance performance but also increase system complexity. These methodological differences may partly explain the performance gaps observed.

The higher accuracy achieved in this study can be attributed to several key factors. First, the hybrid architecture of CNN-BiLSTM allows the model to capture both local features through convolutional layers and long-range dependencies via bidirectional recurrent processing. This dual capability provides a more comprehensive understanding of the text structure compared to models that rely solely on either CNN or LSTM. Second, the use of the word2vec embedding trained with a large vocabulary and 100-dimensional vectors contributes to capturing semantic similarity between words more effectively than traditional sparse representations. Additionally, careful hyperparameter tuning—including the number of CNN filters and BiLSTM units—enabled the model to achieve an optimal balance between complexity and generalization. The large and balanced dataset also helped improve the model's robustness and minimize overfitting, further supporting its superior performance.

These findings highlight that the CNN-BiLSTM architecture significantly enhances hoax detection accuracy on social media, emphasizing its potential as a strong model for classifying fake news.

V. CONCLUSION

This study proposed and assessed a hybrid architecture that integrates Bidirectional Long Short-Term Memory (Bi-LSTM) with Convolutional Neural Network (CNN) for detecting hoax content on social media. The findings reveal that the combination of these models successfully captures both the contextual sequence and spatial characteristics of textual data, contributing to strong classification outcomes. Based on performance indicators, including accuracy, precision, recall, F1-score, and the confusion matrix, the model demonstrated high effectiveness, achieving 96.14% accuracy, 96% precision, 96.25% recall, and a 96% F1-score on the testing dataset. These results confirm the model's capability to accurately differentiate between hoax and legitimate content, with only a few misclassifications.

Among all configurations tested, the CNN 64-BiLSTM 64 model achieved the highest accuracy and the most stable performance. The model's effectiveness also improved with larger training datasets and increased epochs, highlighting its strong generalization capability for hoax detection tasks. These findings suggest that the BiLSTM+CNN architecture—particularly with the CNN 64-BiLSTM 64 setup—is a promising solution for misinformation detection and provides a solid foundation for future improvements and research in social media analysis.

This study has both practical and theoretical implications.

The CNN-BiLSTM model, with high accuracy and balanced performance, can be used in real-time applications for automated hoax detection in social media monitoring, fact-checking platforms, and content moderation tools, particularly during critical events.

From a theoretical perspective, the findings reinforce the value of hybrid deep learning architectures in natural language processing tasks, especially in handling complex textual patterns that require both spatial and sequential analysis. The study also contributes to the growing body of literature supporting the effectiveness of word2vec embeddings in text classification and highlights how parameter tuning (e.g., CNN filters and BiLSTM units) can significantly impact model performance. Nonetheless, the study has certain limitations that could be addressed in future work. Further research could involve the use of Indonesian news data collected directly from social media platforms such as Twitter, allowing for broader topic coverage. Additionally, exploring different word embedding techniques such as FastText, GloVe, TF-IDF, or BERT, along with more thorough hyperparameter tuning, may further enhance the model's performance. It is also recommended to develop a publicly accessible hoax detection system to help mitigate the widespread dissemination of false information.

REFERENCES

- [1] A. U. Hussna, M. G. R. Alam, B. F. Alkhamees, M. M. Hassan, and M. Z. Uddin, "Dissecting the infodemic: An in-depth analysis of COVID-19 misinformation detection on X (formerly twitter) utilizing machine learning and deep learning techniques," *Heliyon*, vol. 10, no. 18, Art. no. e37760, Sep. 2024, doi: 10.1016/j.heliyon.2024.e37760.
- [2] L. Samaras, E. Garcia-Barriocanal, and M.-A. Sicilia, "Sentiment analysis of COVID-19 cases in greece using twitter data," *Expert Systems with Applications*, vol. 230, Art. no. 120577, Jun. 2023, doi: 10.1016/j.eswa.2023.120577.
- [3] H. B. Aji and E. B. Setiawan, "Mendeteksi Konten hoax di media sosial menggunakan Bi-LSTM dan RNN," *Building of Informatics, Technology and Science (BITS)*, vol. 5, no. 1, pp. 114–125, 2023.
- [4] T. C. Praha, Widodo and M. Nugraheni, "Indonesian fake news classification using transfer learning in CNN and LSTM," *JOIV: International Journal on Informatics Visualization*, vol. 8, no. 3, pp. 1213–1221, 2024, doi: 10.62527/joiv.8.3.2126.
- [5] E. Gabarron, S. O. Oyeyemi and R. Wynn, "COVID-19-related misinformation on social media: A systematic review," *Bulletin of the World Health Organization*, vol. 99, no. 6, pp. 455–463A, 2021.
- [6] B. P. Nayoga, R. Adipradana, R. Suryadi and D. Suhartono, "Hoax Analyzer for Indonesian News Using Deep Learning," *Procedia Computer Science*, vol. 179, p. 704–712, 2021.
- [7] B. Jang, M. Kim, G. Harerimana, S.-u. Kang and J. W. Kim, "Bi-LSTM model to increase accuracy in text classification: Combining word2vec CNN and attention mechanism," *Applied Sciences*, vol. 10, no. 17, Art. no. 5841, 2020, doi: 10.3390/app10175841.
- [8] H. Bangyal, R. Qasim, N. Rehman, Z. Ahmad, H. Dar, L. Rukhsar, Z. Aman and J. Ahmad, "Detection of fake news text classification on COVID-19 using deep learning approaches," *Computational and Mathematical Methods in Medicine*, vol. 2021, pp. 1–14, 2021.
- [9] A. R. I. Fauzy and E. B. Setiawan, "Detecting Fake news on social media combined with the CNN methods," *Jurnal RESTI (Rekayasa Sistem Dan Teknologi Informasi)*, vol. 7, no. 2, pp. 271–277, 2023.
- [10] C. Yoviananda and T. M. Fahrudin, "Implementation of deep learning to detect indonesian hoax news with convolutional neural network method,"

- IJEEIT International Journal of Electrical Engineering and Information Technology*, vol. 4, no. 2, pp. 86–93, 2022.
- [11] P. N. Anggreyani and W. Maharani, “Hoax detection tweets of the COVID-19 on twitter using LSTMCNN with word2vec,” *Jurnal Media Informatika Budidarma*, vol. 6, no. 4, pp. 2432–2437, 2022.
- [12] R. Rhanoui, M. Mikram, S. Yousfi and S. Barzali, “A CNN-BiLSTM model for document-level sentiment analysis,” *Machine Learning and Knowledge Extraction*, vol. 1, no. 3, pp. 832–847, 2019.
- [13] L. Lestandy and A. Abdurrahim, “Effect of word2vec Weighting with CNN-BiLSTM Model on Emotion Classification,” *Jurnal Nasional Pendidikan Teknik Informatika: JANAPATI*, vol. 12, no. 1, pp. 99–107, 2023.
- [14] W. K. Sari, I. S. B. Azhar, Z. Yamani and Y. Florensia, “Fake news detection using optimized convolutional neural network and bidirectional long short-term memory,” *Computer Engineering and Applications Journal*, vol. 13, no. 3, pp. 25–33, 2024.
- [15] I. Segura-Bedmar and S. Alonso-Bartolome, “Multimodal fake news detection,” *Information*, vol. 13, no. 6, Art. no. 284, 2022, doi: 10.3390/info13060284.
- [16] Y. Kim, “Convolutional neural networks for sentence classification,” *Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP)*, pp. 1746–1751, 2014.
- [17] P. A. Aritonang, M. E. Johan and I. Prasetyawan, “Aspect-based sentiment analysis on application review using CNN,” *Ultima Infosys: Jurnal Ilmu Sistem Informasi*, vol. 13, no. 1, pp. 54–61, 2022.
- [18] S. Dadkhah, X. Zhang, A. G. Weismann, A. Firouzi, and A. A. Ghorbani, “The largest social media Ground-Truth dataset for Real/Fake content: TruthSeeker,” *IEEE Transactions on Computational Social Systems*, vol. 11, no. 3, pp. 3376–3390, Oct. 2023, doi: 10.1109/tcss.2023.3322303.
- [19] W. Pamungkas and S. Suryani, “Deteksi hoax untuk berita hoax covid 19 Indonesia Menggunakan CNN,” B. S. thesis, Universitas Telkom, 2021. [Online]. Available: <https://repositori.telkomuniversity.ac.id/pustaka/172363/deteksi-hoax-untuk-berita-hoax-covid-19-indonesia-menggunakan-cnn.html>
- [20] A. Nurkholis, S. Styawati, S. Alim, H. Saputra and A. Ferriyan, “Sentiment analysis of COVID-19 booster vaccines on twitter using multi-class support vector machine,” *Applied Information System and Management (AISM)*, vol. 8, no. 1, pp. 29–36, 2025, doi: 10.15408/aism.v8i1.42911.
- [21] H. P. Fitriani, I. Ruslianto and R. Hidayati, “Implementasi metode naive bayes classifier untuk aplikasi filtering email spam dengan lemmatization berbasis web, coding jurnal komputer dan aplikasi,” *Coding: Jurnal Komputer dan Aplikasi*, vol. 6, no. 2, pp. 13–24, 2018, doi: 10.26418/coding.v6i2.25487.
- [22] D. Jurafsky and C. Manning, *Natural Language Processing*, Standford, 2012.
- [23] D. I. Afidah, Dairoh, S. F. Handayani and R. W. Pratiwi, “Pengaruh parameter word2vec terhadap performa deep learning pada klasifikasi sentimen,” *Jurnal Informatika: Jurnal Pengembangan IT*, vol. 6, no. 3, pp. 156–161, 2021.
- [24] T. Mikolov, K. Chen, G. S. Corrado, and J. Dean, “Efficient estimation of word representations in vector space,” *ArXiv (Cornell University)*, Jan. 2013, doi: 10.48550/arxiv.1301.3781.
- [25] A. K. Tyagi and V. K. Reddy, “Performance analysis of under-sampling and over-sampling techniques for solving class imbalance problem,” *International Conference on Sustainable Computing in Science, Technology & Management (SUSCOM-2019)*, pp. 1305–1315, 2019. ArXiv preprint, vol. 3781, 2013.
- [26] D. Gupta, “Detection of fake news on twitter using machine learning: An xgboost-based,” *International Journal of Advanced Research*, vol. 12, no. 08, pp. 948–964, 2024.
- [27] C. Rahmad, N. Nurfaidah, S. Adhisuwarnjo and M. Hani’ah, “Mask detection app uses haar cascade and convolutional neural network to alert comply with health protocols,” *Applied Information System and Management (AISM)*, vol. 6, no. 2, pp. 77–82, 2023, doi: 10.15408/aism.v6i2.31396.
- [28] M. B. Tamam, H. Hozairi, M. Walid and J. F. A. Bernardo, “Classification of sign language in real time using convolutional neural network,” *Applied Information System and Management (AISM)*, vol. 6, no. 1, pp. 39–46, 2023, doi:10.15408/aism.v6i1.29820..