

Detection of hoax content using support vector machine, Naive Bayes, and random forest algorithms

Rama Sahtyawan¹, Anton Yudhana²

¹Department of Information Technology, Faculty of Engineering and Information Technology, Jenderal Achmad Yani University, Yogyakarta, Indonesia

²Department of Electrical Engineering, Faculty of Industrial Technology, University Ahmad Dahlan, Yogyakarta, Indonesia

Article Info

Article history:

Received Sep 8, 2025

Revised Apr 14, 2026

Accepted Jul 9, 2026

Keywords:

Hoax detection

Machine learning

Naive Bayes

Random forest

Support vector machine

Term frequency-inverse
document frequency

ABSTRACT

The rapid dissemination of hoax information through digital media poses a major danger to public perception and societal order. Automated detection systems are crucial to combat this issue. This research providing a side-by-side evaluation of three machine learning (ML) models: support vector machine (SVM), Naive Bayes (NB), and random forest (RF) for classifying news text as hoax or non-hoax. A total of 4,009 articles were first cleaned, broken into tokens, stripped of common stopwords, and lemmatized. Next, the term frequency-inverse document frequency (TF-IDF) approach was employed to pull out the key features. Evaluation of the models relied on metrics such as accuracy, precision, recall, and F1-score. The top results came from recorded by SVM, which reached 97% accuracy and a balanced F1-score of 96%, followed by RF at 95% accuracy. Meanwhile, NB only achieved 87% accuracy and showed a stronger tendency toward false positives. These results lead to the conclusion that SVM is the optimal algorithm for hoax news detection, thanks to its superior and well-balanced classification capability. What sets this study apart is its use of both comparative performance evaluation and statistical testing on an English-language dataset, providing a robust methodological foundation for future development and adaptation of hoax detection systems in other linguistic contexts, including Indonesian.

This is an open access article under the [CC BY-SA](https://creativecommons.org/licenses/by-sa/4.0/) license.



Corresponding Author:

Rama Sahtyawan

Department of Information Technology, Faculty of Engineering and Information Technology

Jenderal Achmad Yani University

Siliwangi Street, West Ring Road, Banyuraden, Sleman, Yogyakarta, Indonesia

Email: ramasahtyawan@gmail.com

1. INTRODUCTION

The digital era has radically changed how information circulates, giving everyone immediate access to news through the internet and social networks. But this ease of access carries a major downside: hoaxes, false information created intentionally to mislead the public for political, financial, or social purposes are now spreading uncontrollably [1]. In the real world, this issue leads to serious fallout, such as shaping what people believe, undermining social harmony, and sparking acts of aggression [2]. For instance, a hoax regarding child abductions in Indonesia triggered public panic and tragic vigilante actions before being debunked by authorities [3]. This highlights just how urgent and essential it is to have reliable, automated tools that can identify hoax content quickly and on a large scale.

Within content-based detection, several algorithms have been explored. The Naive Bayes (NB) algorithm is frequently employed as a baseline due to its computational efficiency and simplicity in text classification [4]. The application of the NB classifier has been established as a classic and efficient baseline

method for automated fake news detection in digital text classification [5]. Ensemble learning methods in machine learning (ML) have been extensively utilized to improve the robustness and overall performance of systems designed to detect deceptive information [6]. But because it heavily relies on the idea that all features are unrelated to one another, this approach frequently struggles when applied to complicated, real-life text data [7]. Support vector machines (SVM) are well-known for performing exceptionally well when working with many dimensions, which makes them particularly effective when handling text that has been converted through techniques like term frequency-inverse document frequency (TF-IDF) [8]. Their ability to find an optimal hyperplane for separation has led to strong results in previous hoax detection studies [9]. Many studies compare models using different preprocessing pipelines or evaluation metrics, making it difficult to draw concrete conclusions about their relative strengths and weaknesses in a controlled environment [10].

Apart from their theoretical value, systems designed to spot hoaxes also serve an essential real-world function: preserving the trustworthiness of online information environments. With these tools, social media companies can catch and label deceptive or untrue posts before they go viral [11], while also supporting fact-checking efforts by integrating with external knowledge bases for rapid claim verification [12]. The presence of such systems enhances public trust and strengthens the credibility of digital media platforms [13]. Furthermore, hoax detection has pedagogical value in engineering and computer science education, such as demonstrating the role of NLP in secure communication or showcasing the challenges of deploying artificial intelligence (AI) models on resource-constrained embedded systems or internet of things (IoT) devices [14], [15]. These interdisciplinary applications highlight the broad impact and relevance of automated misinformation detection technologies.

This suggested methodology follows a structured sequence that includes getting the data ready, applying thorough text preprocessing steps (such as cleaning, breaking into tokens, discarding common words, and lemmatization), and finally pulling out features via TF-IDF [16], model training, and a multi-faceted evaluation. The novelty and value of this research are threefold: i) it provides a clear, empirical benchmark of these three algorithms on an identical, preprocessed dataset using a consistent evaluation framework [17], ii) it moves beyond mere accuracy metrics by employing a detailed diagnostic analysis—including precision [18], reveal more than just the top-performing model—they also explain the reasoning behind it by looking at each algorithm's pattern of generating false alarms (false positives) or missing actual cases (false negatives) [19], and iii) the findings offer practical guidance for developers and researchers: SVM is recommended for balanced accuracy [20], random forest (RF) for minimizing missed hoaxes [21], and NB for computational speed where some accuracy can be sacrificed [22]. The remainder of this article is laid out in the manner described below [23]. Tree-based classifiers, specifically RF and J48 decision trees, provide a structured approach to analyzing and identifying patterns within misleading news articles [24]. The integration of SVM with TF-IDF continues to serve as a powerful and highly reliable methodology for unveiling the truth in online content [25].

2. METHOD

A quantitative strategy was adopted for this study, relying on text analysis to identify hoax news. The workflow consists of multiple key phases: gathering data, preprocessing the text, extracting features, and finally running classification through ML algorithms. This setup aims to automatically and precisely uncover textual patterns that separate hoax articles from legitimate ones. An overview of this research process: start, dataset fake news, preprocessing, feature extraction TF-IDF, classification (SVM, NB, and RF), evaluation, and end.

2.1. Data collecting

The research draws on a publicly available dataset called 'Fake News Detection,' which comes from Kaggle and can be accessed via the following link: <http://kaggle.com/datasets/jruvika/fake-news-detection>. This collection includes news pieces that have been categorized either as 'hoax' or 'non-hoax.' In total, there are 4,009 articles in the dataset, with a fairly even split between the two classes. Four key fields are present in every record in Table 1.

Table 1. Dataset attributes

Feature	Description
URLs	Links contains news articles
Headline	Featured news headlines
Body	The main content or content of a news article
Label	Article classification: hoax (1) and non-hoax (0)

For classification purposes, only the body feature is used as the main source of text, while labels are used as a marker for news categories, whether they are hoax or not. Figure 1 shows the distribution of data by label. Label 0 represents non-hoax with a total of 2,137 articles, while label 1 represents hoax (fake news) with a total of 1,872 articles. This relatively balanced distribution is important so that the ML algorithm can be not biased against any of the classes.

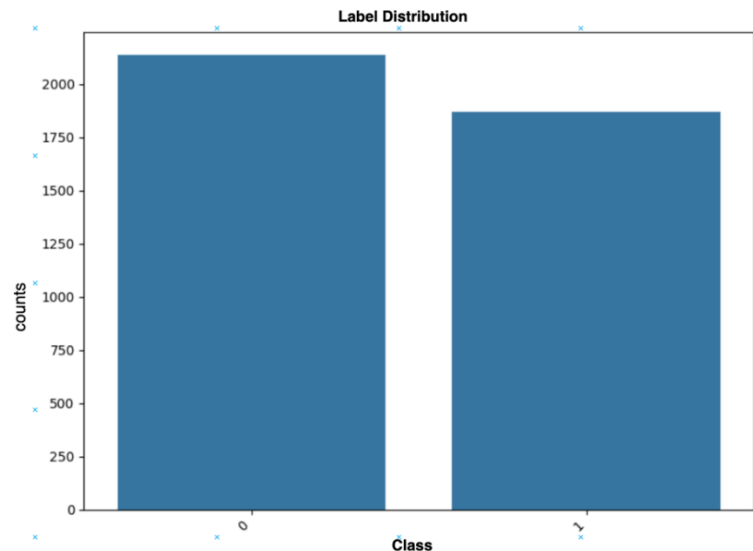


Figure 1. Label distribution

2.2. Data preprocessing

The preprocessing step gets text data ready so ML models can process it effectively. Raw text often comes with unwanted elements like punctuation, special characters, web links, or irregular formatting that can hurt how accurately the model performs. The preprocessing pipeline consists of several sequential steps, illustrated in Figure 2.



Figure 2. Preprocessing stages

The preprocessing steps include: cleaning (removing numbers, punctuation, URLs, emojis, and extra spaces); case folding (lowercasing all text for uniformity); normalization (fixing slang and variants, e.g., "u"→"you", "covid19"→"covid"); tokenization (splitting text into words); stopwords removal (eliminating low-meaning words like "is," "the"); and lemmatization (reducing words to dictionary root form, e.g., "running"→"run"). At the end of the pipeline, in a single document compared to how often they show up across all documents, making sure that words carrying more useful information get more weight when the model learns.

2.3. Feature extraction term frequency–inverse document frequency

Once the raw text has passed through preprocessing and been stripped of unnecessary components, the following phase is feature extraction. Its job is to convert the text-based information into numeric form so that ML models can work with it. For this research, we applied the TF-IDF approach, a highly popular choice for representing text when the goal is document classification.

Which shows how important that word is within that specific piece of text. IDF, on the other hand, measures how distinctive or meaningful a word is across the whole set of documents (the corpus). Words that pop up often across many documents end up with low IDF scores because they don't help tell documents apart very well. Meanwhile, words that are rare get high IDF scores, pointing to their potential usefulness for distinguishing between different classes of documents.

2.4. Classification model

Here, the classification step takes place to tell apart hoax articles from legitimate ones, using the features that were pulled out earlier via the TF-IDF technique. The three ML algorithms applied in this phase: SVM is being a classifier that relies on margins, SVM searches for the best possible boundary known as a hyperplane that splits two categories while maximizing the gap on both sides. NB relies on probability principles and stems from Bayes' theorem, operating under the assumption that all features are unrelated to one another. When handling text data, a specific version called Multinomial NB is applied. RF belongs to the ensemble family of algorithms. It brings together multiple decision trees.

3. RESULTS AND DISCUSSION

This research aims to identify the algorithm with the highest level of accuracy in detecting text-based hoax news. The entire process is carried out using Google Collaboratory as a cloud-based programming environment, as well as the Windows 10 Pro operating system to support the documentation and reporting process of research results.

3.1. Data collection

At the start, the data collection consisted of 4,009 articles, a mix of hoaxes and real news. A preliminary check found 1,125 duplicates (caused by the same news appearing across different sources) and 21 blank entries in the body column, which were unusable. After removing these, 2,863 unique and valid articles remained. This cleanup was essential to prevent duplication bias and ensure a cleaner, more representative dataset for training.

3.2. Data preprocessing

The data preparation step was performed in order to clean and standardize news text before it was used in model training. This process included text cleaning, case folding, word normalization, tokenization, stopword removal, and lemmatization. After the preprocessing stage, the data gets split, with 80% allocated for training and the remaining 20% reserved for testing. Out of a total of 2.863 data, 2.290 training data, and 573 test data were obtained, as shown in Table 2.

Table 2. Dataset splitting

Data	Amount of data
Data training	2.290
Data testing	573
Total data	2.863

3.3. Term frequency-inverse document frequency feature extraction

Once preprocessing finished, TF-IDF was applied to extract features, where a word's weight rises with its frequency output is displayed in Figure 3. In addition, a WordCloud visualization was performed to display the words that appear most frequently in the corpus, as shown in Figure 4.

	visible	visibly	vision	visionary	visit	visitation	visitor
0	0.0	0.0	0.0	0.0	0.015158	0.0	0.000000
1	0.0	0.0	0.0	0.0	0.000000	0.0	0.000000
2	0.0	0.0	0.0	0.0	0.032118	0.0	0.000000
3	0.0	0.0	0.0	0.0	0.000000	0.0	0.000000
4	0.0	0.0	0.0	0.0	0.000000	0.0	0.000000
...	...	...	...	...	...	...	...
2285	0.0	0.0	0.0	0.0	0.000000	0.0	0.000000
2286	0.0	0.0	0.0	0.0	0.000000	0.0	0.000000
2287	0.0	0.0	0.0	0.0	0.018683	0.0	0.023908
2288	0.0	0.0	0.0	0.0	0.015531	0.0	0.000000
2289	0.0	0.0	0.0	0.0	0.000000	0.0	0.000000

Figure 3. TF-IDF implementation results

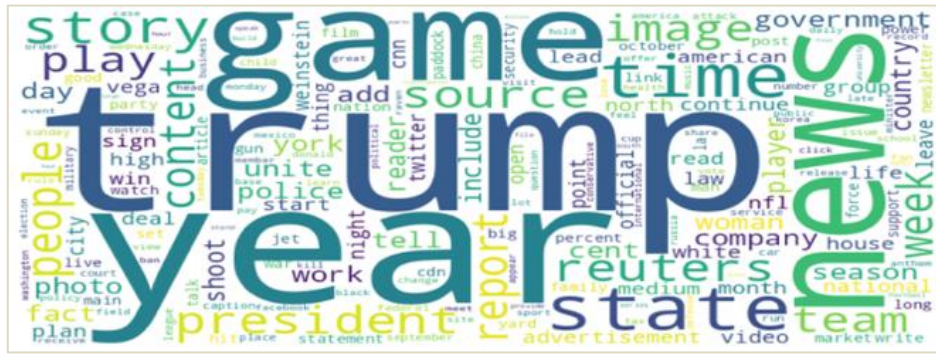


Figure 4. WordCloud visualization

3.4. Machine learning model training

For model training, this research employed three classifiers: SVM, NB, and RF. Each algorithm was trained on data that had already been preprocessed and transformed into numerical features via the TF-IDF method. The SVM model was run with a linear kernel and its standard default settings.

3.5. Evaluation results

To assess the classification capability of each algorithm, the models were evaluated using several standard performance metrics: accuracy, precision, recall, and F1-score. These four measures collectively give a fair assessment of how effectively each algorithm performs in distinguishing hoax and real news. The comparative results are presented in Table 3.

Table 3. Classification model performance comparison

Model	Accuracy	Precision	Recall	F1-score
SVM	0.97	0.97	0.96	0.96
NB	0.87	0.90	0.85	0.86
RF	0.95	0.96	0.94	0.95

SVM achieved the best results (97% accuracy and balanced precision-recall), proving robust at distinguishing hoax from real news. RF (95% accuracy) minimized false negatives well, though its false positives were slightly above SVM's. NB was fast to train but less precise on real news, likely due to its word-frequency sensitivity. Figures 5–7 show the confusion matrices for all three models.

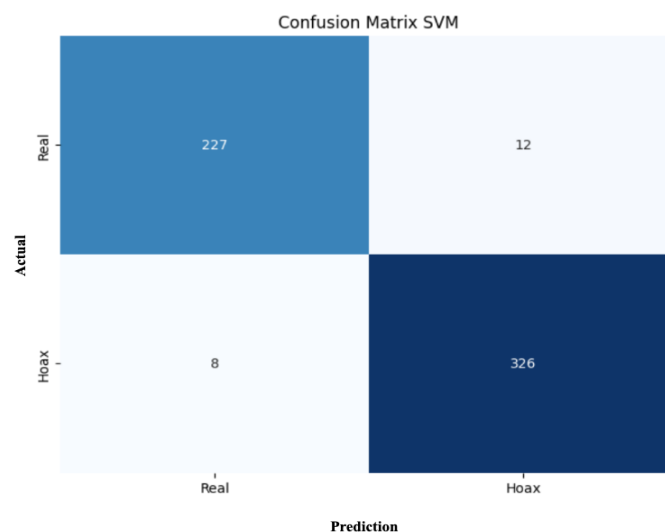


Figure 5. Confusion matrix SVM

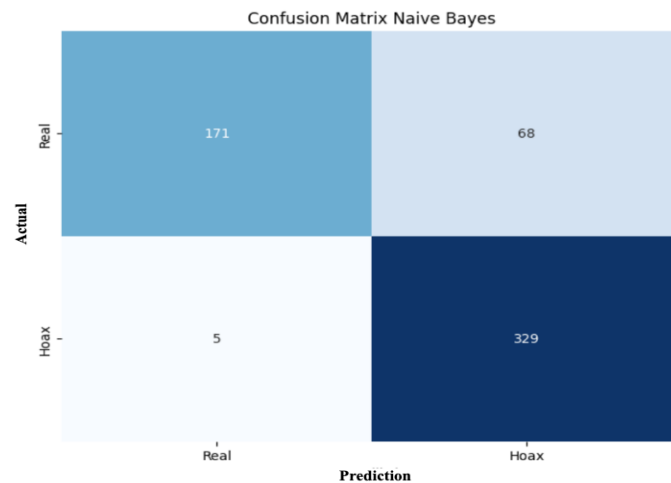


Figure 6. Confusion matrix NB

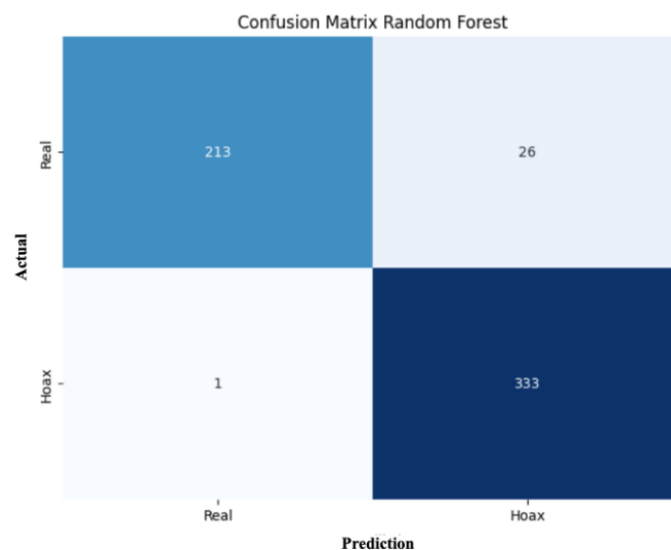


Figure 7. Confusion matrix RF

SVM showed strong and stable performance with 227 real and 326 hoax correct, plus only 12 false positives, and 8 false negatives. NB, however, had 171 reals, and 329 hoax correct but suffered from 68 false positives, far higher than other models, along with just 5 false negatives, revealing its over-sensitivity in labeling news as hoaxes. RF achieved 213 real and 333 hoax correct, with 26 false positives and only 1 false negative, it excellent at detecting hoaxes despite having slightly more false positives than SVM.

3.6. Receiver operating characteristic curve

Receiver operating characteristic (ROC) curves measure hoax vs. real classification performance by showing the balance between the rate of correctly identified positives and falsely flagged positives across different cutoff points. A larger area under the curve (AUC) (up to 1.0) indicates superior model performance. Figure 8 compares ROC curves for SVM, NB, and RF.

The ROC curve shows that the three algorithms have excellent classification performance in the context of detecting hoax news. SVM is the model with the most optimal performance probabilistic based on the value $AUC=1.00$. Meanwhile, the NB and RF models obtained AUC values of 0.98 and 0.99, respectively. Evaluations using ROC and AUC curves provide more in-depth information than just accuracy or F1-scores, as they take into account all classification thresholds.

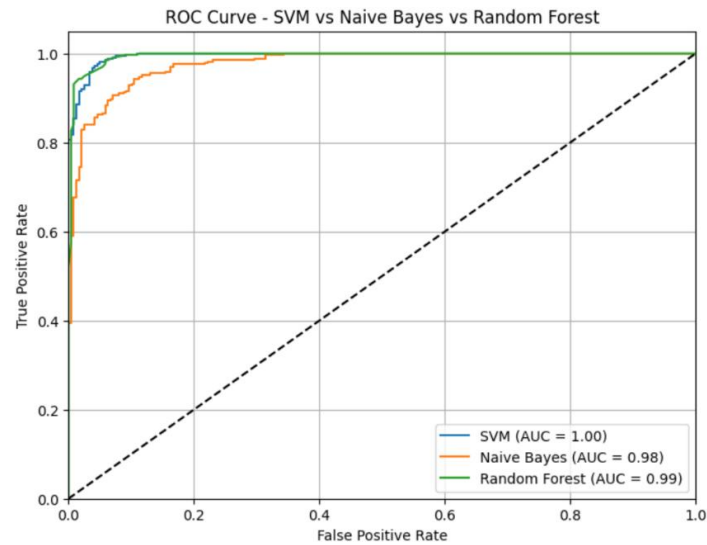


Figure 8. Curve ROC

3.7. Cross-validation performance and statistical significance

A 5-fold cross-validation (CV) was carried out to assess how well the models performed, using three classifiers: SVM, NB, and RF. Table 4 shows the accuracy scores recorded for each fold. To test if model performance differences were significant, a one-way ANOVA was applied to their 5-fold CV accuracies. The F-statistic (94.55) and p-value (0.0000) indicate genuine differences, not random noise. Figure 9 boxplot visualizes accuracy variation and stability across folds.

Table 4. 5-Fold CV results

Model	Fold 1	Fold 2	Fold 3	Fold 4	Fold 5
SVM	0.9803	0.9738	0.9607	0.9738	0.9519
NB	0.8646	0.8318	0.8296	0.8362	0.8231
RF	0.9651	0.9607	0.9214	0.9520	0.9279

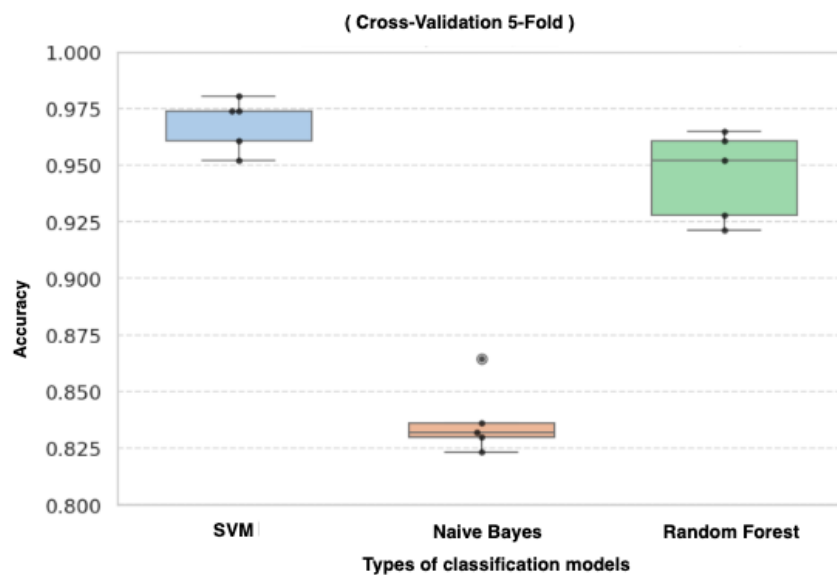


Figure 9. Results of one-way ANOVA on model accuracy (5-fold CV)

From the boxplot, SVM shows the least fluctuation across folds, pointing to consistent behavior, whereas NB displays more noticeable variation, suggesting greater sensitivity to data distribution. RF shows moderate variance and stable behavior. These findings reinforce that SVM not only achieves superior accuracy but also demonstrates higher consistency.

4. CONCLUSION

This study provides a comprehensive empirical evaluation of ML algorithms—specifically SVM, NB, and RF—for automated hoax news detection using TF-IDF feature extraction. The primary contribution of this research lies in establishing a rigorous, controlled benchmarking pipeline on a preprocessed text dataset, validated through both standard evaluation metrics and 5-fold CV paired with statistical significance testing (one-way ANOVA, $F=94.55$ and $p<0.0001$). The experimental results demonstrate that SVM achieves the superior overall classification capability with 97% accuracy, an F1-score of 0.96, an AUC of 1.00, and exceptionally low error rates across folds, establishing it as the most effective baseline model among those tested.

From a practical perspective, these findings offer concrete design insights for practitioners, software developers, and digital platform administrators implementing automated content moderation systems. Depending on operational priorities, deployment choices can be tailored: SVM provides optimal, balanced performance for general automated filtering; RF is highly effective for applications prioritizing minimal false negatives (missing only 1 hoax in testing); while NB offers a lightweight alternative suited for resource-constrained or real-time environments where computational speed is prioritized over slight trade-offs in precision. The proposed comparative evaluation contributes significantly to the broader field of hoax detection by moving beyond surface-level accuracy to reveal granular algorithm-specific behaviors, such as error trade-offs and stability across CV folds. Furthermore, by establishing a standardized, statistically validated framework on textual data, this benchmark serves as a reliable methodological reference for adapting and scaling ML-based misinformation detection tools across diverse domain environments and multilingual contexts, including Indonesian text classification.

Despite these contributions, a few limitations of the study should be noted. The evaluation relied primarily on classic TF-IDF representations and surface-level textual features from a single dataset, which may not fully capture complex contextual semantics, subtle sarcasm, or rapidly evolving online misinformation patterns. To address these limitations, future research directions should focus on: advanced model architectures: incorporating deep learning and pre-trained transformer-based language models (e.g., BERT, RoBERTa, or contextual embeddings) to better capture semantic nuance, multilingual and cross-lingual adaptation: extending the framework to non-English corpora, particularly Indonesian news datasets, to evaluate cross-linguistic robustness.

ACKNOWLEDGMENTS

The authors would like to thank LLDIKTI Region 5 yogyakarta for the 2025 publication Assistance program grant, as well as the support of the Study Program Information Technology, Faculty of Engineering and Information Technology, Jenderal Achmad Yani University, Yogyakarta, Indonesia.

FUNDING INFORMATION

Authors state no funding involved.

AUTHOR CONTRIBUTIONS STATEMENT

This journal uses the Contributor Roles Taxonomy (CRediT) to recognize individual author contributions, reduce authorship disputes, and facilitate collaboration.

Name of Author	C	M	So	Va	Fo	I	R	D	O	E	Vi	Su	P	Fu
Rama Sahtyawan	✓	✓	✓	✓	✓	✓	✓	✓	✓	✓			✓	✓
Anton Yudhana		✓				✓		✓	✓	✓	✓	✓		

- C : Conceptualization
- M : Methodology
- So : Software
- Va : Validation
- Fo : Formal analysis
- I : Investigation
- R : Resources
- D : Data Curation
- O : Writing - Original Draft
- E : Writing - Review & Editing
- Vi : Visualization
- Su : Supervision
- P : Project administration
- Fu : Funding acquisition

CONFLICT OF INTEREST STATEMENT

Authors state no conflict of interest.

DATA AVAILABILITY

The dataset underlying this study's findings <http://kaggle.com/datasets/jruvika/fake-news-detection>. This dataset includes all parameters used for model training and validation.

REFERENCES

- [1] X. Zhou and R. Zafarani, "A Survey of Fake News," *ACM Computing Surveys*, vol. 53, no. 5, pp. 1–40, Sep. 2021, doi: 10.1145/3395046.
- [2] A. S. Matemilola and S. Aliyu, "Development of an enhanced naive bayes algorithm for fake news classification," *Science World Journal*, vol. 19, no. 2, pp. 512–517, Jul. 2024, doi: 10.4314/swj.v19i2.28.
- [3] S. Ahmed, A. Qaiser, J. Shahzad, M. Ali, and N. Khan, "Detecting and Analyzing Deceptive Information in News Articles: A Study Using a Dataset of Misleading Content," *International Journal of Emerging Engineering and Technology*, vol. 2, no. 1, pp. 52–56, Jul. 2023, doi: 10.57041/ijeet.v2i1.930.
- [4] X. Zhang and A. A. Ghorbani, "An overview of online fake news: Characterization, detection, and discussion," *Information Processing & Management*, vol. 57, no. 2, Mar. 2020, doi: 10.1016/j.ipm.2019.03.004.
- [5] M. Granik and V. Mesyura, "Fake news detection using naive Bayes classifier," in *2017 IEEE 1st Ukraine Conference on Electrical and Computer Engineering (UKRCON)*, May 2017, pp. 900–903, doi: 10.1109/UKRCON.2017.8100379.
- [6] I. Ahmad, M. Yousaf, S. Yousaf, and M. O. Ahmad, "Fake News Detection Using Machine Learning Ensemble Methods," *Complexity*, vol. 2020, pp. 1–11, Oct. 2020, doi: 10.1155/2020/8885861.
- [7] R. K. Kaliyar, A. Goswami, and P. Narang, "FakeBERT: Fake news detection in social media with a BERT-based deep learning approach," *Multimedia Tools and Applications*, vol. 80, no. 8, pp. 11765–11788, Mar. 2021, doi: 10.1007/s11042-020-10183-2.
- [8] V. Pandey, V. Sharma, S. Vats, and S. P. Yadav, "From Misinformation to Facts: A Detailed Review of Fake News Detection Methods," in *2024 International Conference on Communication, Computing and Energy Efficient Technologies (I3CEET)*, Sep. 2024, pp. 449–454, doi: 10.1109/I3CEET61722.2024.10994139.
- [9] L. Breiman, "Random Forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, Oct. 2001, doi: 10.1023/A:1010933404324.
- [10] S. K. Maakar, S. V. Singh, and A. Zia, "Harnessing NLP and AI for Reliable News Verification in the Digital Era," in *2025 International Conference on Pervasive Computational Technologies (ICPCT)*, Feb. 2025, pp. 111–115, doi: 10.1109/ICPCT64145.2025.10939116.
- [11] E. Stewart, "Detecting Fake News: Two Problems for Content Moderation," *Philosophy & Technology*, vol. 34, no. 4, pp. 923–940, Dec. 2021, doi: 10.1007/s13347-021-00442-x.
- [12] S. Hangloo and B. Arora, "Evidence-Aware Fake News Detection: A Review," in *2023 International Conference on Advanced Computing & Communication Technologies (ICACCTech)*, Dec. 2023, pp. 81–86, doi: 10.1109/ICACCTech61146.2023.00022.
- [13] B. S. A. Nasser and S. S. Abu-Naser, "Leveraging AI for Effective Fake News Detection and Verification," *Arab Media & Society*, no. 37, Dec. 2024, doi: 10.70090/BSAN24FN.
- [14] S. Debnath, M. Senbagavalli, K. Ramalakshmi, S. Naik, and N. Begum, "Security of Social Media Content for AI-Embedded Systems," in *Embedded Artificial Intelligence*, Boca Raton: Chapman and Hall/CRC, 2025, pp. 293–308.
- [15] A. N. Boruah, M. Goswami, M. Kumar, and O. Loyola-González, *Embedded Artificial Intelligence*. Boca Raton: Chapman and Hall/CRC, 2025.
- [16] J. Ramos, "Using TF-IDF to Determine Word Relevance in Document Queries," in *Proceedings of the First Instructional Conference on Machine Learning*, vol. 242, no. 1, pp. 29–48, 2003.
- [17] V. K. Singh, I. Ghosh, and D. Sonagara, "Detecting fake news stories via multimodal analysis," *Journal of the Association for Information Science and Technology*, vol. 72, no. 1, pp. 3–17, Jan. 2021, doi: 10.1002/asi.24359.
- [18] J. Devlin, M.-W. Chang, K. Lee, and K. Toutanova, "BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding," in *Proceedings of the 2019 Conference of the North*, 2019, pp. 4171–4186, doi: 10.18653/v1/N19-1423.
- [19] A. Qazi, R. H. Goudar, R. Patil, G. S. Hukkeri, and D. Kulkarni, "Leveraging BERT, DistilBERT, and TinyBERT for Rumor Detection," *IEEE Access*, vol. 13, pp. 72918–72929, 2025, doi: 10.1109/ACCESS.2025.3563301.
- [20] D. Sharma, B. Singh, S. Agarwal, H. Kim, and R. Sharma, "FakedBits- Detecting Fake Information on Social Platforms using Multi-Modal Features," *KSI Transactions on Internet and Information Systems*, vol. 17, no. 1, pp. 51–73, Jan. 2023, doi: 10.3837/tiis.2023.01.004.
- [21] K. Shu, A. Sliva, S. Wang, J. Tang, and H. Liu, "Fake News Detection on Social Media," *ACM SIGKDD Explorations Newsletter*, vol. 19, no. 1, pp. 22–36, Sep. 2017, doi: 10.1145/3137597.3137600.
- [22] W. Y. Wang, "'Liar, Liar Pants on Fire': A New Benchmark Dataset for Fake News Detection," in *Proceedings of the 55th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, 2017, pp. 422–426, doi: 10.18653/v1/P17-2067.
- [23] M. Del Vicario, W. Quattrociocchi, A. Scala, and F. Zollo, "Polarization and Fake News," *ACM Transactions on the Web*, vol. 13, no. 2, pp. 1–22, May 2019, doi: 10.1145/3316809.
- [24] R. Jehad and S. A. Yousif, "Fake News Classification Using Random Forest and Decision Tree (J48)," *Al-Nahrain Journal of Science*, vol. 23, no. 4, pp. 49–55, Dec. 2020, doi: 10.22401/ANJS.23.4.09.
- [25] R. Rizal, A. Faturahman, A. Impron, I. Darmawan, E. Haerani, and A. Rahmatulloh, "Unveiling the Truth: Detecting Fake News Using SVM and TF-IDF," in *2025 International Conference on Advancement in Data Science, E-learning and Information System (ICADEIS)*, Feb. 2025, pp. 1–6, doi: 10.1109/ICADEIS65852.2025.10933324.

BIOGRAPHIES OF AUTHORS

Rama Sahtyawan     received his Master of Computer Science in Faculty of Mathematics and Natural Sciences, Universitas Gadjah Mada, Indonesia. He is currently lecturing with the information technology at Jenderal Achmad Yani University, Yogyakarta, Indonesia. His research areas are machine learning, cyber security, and internet of things. He can be contacted at email: ramasahtyawan@gmail.com.



Anton Yudhana     is a professor at the Department of Electrical Engineering Science at Ahmad Dahlan University, Indonesia. He holds a Doctoral degree in Electrical Engineering, Universiti Teknologi Malaysia, Malaysia, Expertise: data communication and intelligent network and computer vision application in intelligent control. Achievement as Gold Winner of the 2024 DIKTISAINTEK Academic Leader Award. He can be contacted at email: eyudhana@ee.uad.ac.id.