



A Comparative Study of Machine Learning Algorithm for Fake News Detection

Sakshi Jagdish Dave¹, Ms. Deepali Gavhane²

Student MCA-II, SVIMS, Savitribai Phule Pune University¹

Assistant Professor, SVIMS, Savitribai Phule Pune University²

Abstract: In today's digital era, the widespread use of social media and online platforms has enabled rapid dissemination of information—but also facilitated the spread of misinformation, commonly known as fake news. Such false information can distort public opinion, disrupt political processes, and cause widespread confusion. This research addresses the growing challenge of fake news detection by developing a binary classification system using machine learning techniques. The study compares the performance of two supervised algorithms—Naïve Bayes and Logistic Regression—applied to the Constraint@AAAI 2021 shared task dataset on COVID-19 fake news. The dataset underwent rigorous preprocessing, including text normalization, noise removal, stopword elimination, and TF-IDF feature extraction. Experimental results demonstrate that both models perform effectively in classifying real and fake news, with Naïve Bayes achieving an accuracy of 92.37% and Logistic Regression slightly outperforming it with 93.85%. These findings highlight the potential of lightweight machine learning models for reliable and efficient fake news detection, contributing to the fight against online misinformation and promoting trustworthy digital communication.

Keywords: Fake news, Machine Learning, Naïve bayes, Logistic Regression.

I. INTRODUCTION

In the present digital era, the ability to instantly create and distribute information is available to anyone with internet access. While this has greatly improved the speed and accessibility of communication, it has also enabled the rapid spread of inaccurate or misleading content. Such false information, often shared through social media and messaging platforms, can lead to public misunderstanding, fear, and even harmful real-world outcomes. These risks make it necessary to design smart detection systems that can identify and limit the spread of fabricated news before it causes damage [1].

To address this issue, several technological solutions have been explored, including Machine Learning (ML), Deep Learning (DL), Natural Language Processing (NLP), and metadata-based approaches. Although these techniques have achieved encouraging results, none are entirely free of shortcomings. Each method brings its own strengths and weaknesses, making it essential to identify the most effective option for fake news detection [2].

This research focuses on two supervised ML algorithms: Naive Bayes and Logistic Regression. Naive Bayes applies Bayes' theorem under the assumption that features are independent, allowing it to process text data efficiently and classify news quickly. Logistic Regression, on the other hand, is a well-established linear model that identifies the relationship between variables and binary outcomes, offering interpretability and consistent performance in distinguishing between genuine and fabricated news articles [3].

The aim of this study is to evaluate these models using measures such as accuracy, precision, recall, and F1-score. It is expected that Naive Bayes will perform efficiently with larger datasets, delivering faster training times and reasonable accuracy, while Logistic Regression may provide stronger precision when detecting news with subtle linguistic differences.

The need for this research arises from the serious consequences of misinformation, which in some cases have contributed to social unrest and violent incidents. By comparing these two algorithms, this study seeks to support the creation of more reliable and efficient systems for identifying false information early [4].

Another important factor in fake news detection is the quality of the training data. Variations in language, writing style, and subject matter can significantly influence the accuracy of detection models. For this reason, the study will also investigate how different datasets affect performance. Drawing inspiration from the work of Villela et al. [5], this research intends to add valuable insights to ongoing efforts against online misinformation. Since the success of any algorithm is



tied to both its intended use and the nature of the data it is trained on, this project will highlight the combinations of models and datasets that produce the most effective results in detecting fake news.

II. OBJECTIVES

- 1.To design and implement a machine learning-based framework for the detection of fake news by employing supervised algorithms such as Naïve Bayes and Logistic Regression to classify news articles as real or fake.
- 2.To perform comprehensive text preprocessing and feature extraction, including normalization, noise and stopword removal, stemming, lemmatization, and Term Frequency–Inverse Document Frequency (TF-IDF) transformation, to enhance the quality and representativeness of input data for model training.
- 3.To conduct a comparative performance analysis of the Naïve Bayes and Logistic Regression models using standard evaluation metrics such as accuracy, precision, recall, and F1-score, thereby identifying the algorithm that provides optimal performance for fake news classification.
- 4.To contribute to the ongoing efforts against online misinformation by demonstrating the applicability of lightweight, interpretable, and computationally efficient machine learning approaches for reliable and scalable fake news detection in digital environments.

III. LITERATURE REVIEW

The rise of misinformation on digital platforms has motivated researchers to adopt machine learning techniques for the automatic identification of fake news. Among the numerous algorithms available, Naïve Bayes (NB) and Logistic Regression (LR) remain two of the most widely used supervised models for text classification. Their popularity in natural language processing tasks comes from their simplicity, interpretability, and consistent performance on textual datasets.

Naïve Bayes is a probabilistic classifier that applies Bayes' theorem under the assumption of feature independence. This assumption allows it to be computationally efficient and effective, particularly when dealing with high-dimensional data such as news text. Despite its straightforward design, it has proven to be reliable in practice. For example, Granik and Mesyura (2020) applied the model to Facebook news posts and obtained an accuracy of about 74%, showing its value in detecting misinformation online. Similarly, Sutradhar et al. (2022) assessed several algorithms—including NB, LR on a dataset of around 1,876 news articles. Their findings showed NB delivering the highest performance among the tested models, although accuracy was limited to 56%, largely due to the small dataset size.

Logistic Regression, in contrast, estimates the likelihood of outcomes through a logistic function and is particularly effective when applied to frequency-based text features such as TF-IDF or n-grams. Gilda (2020) demonstrated this by applying LR with bi-gram features on a dataset of over 11,000 articles, achieving competitive results compared to other classifiers. More recently, Sudhakar and Kaliyamurthie (2022) tested both LR and NB on a political dataset of more than 44,000 records. Their experiments confirmed the strength of LR, which reached an accuracy of 98.7%, outperforming NB at 94.8%.

Other comparative studies further highlight the robustness of these approaches. Mykytiuk et al. (2023) evaluated six machine learning models—including LR, NB, Decision Tree, Random Forest, KNN, and Multilayer Perceptron—and reported near-perfect results, with LR and KNN achieving almost 99% accuracy. While these outcomes are striking, the absence of detailed dataset characteristics suggests the possibility of overfitting. In addition, survey papers by Merryton and Augasta (2020) and Pavan et al. (2020) emphasized that while deep learning architectures such as CNNs, RNNs, and LSTMs perform strongly on massive, unstructured datasets, NB and LR remain indispensable baselines because they are lightweight, transparent, and computationally efficient. Huang (2020) also illustrated their practicality by showing that fake news detection can be approached similarly to spam filtering; using a Kaggle dataset, the study demonstrated that even simple binary classifiers can achieve reliable outcomes with proper preprocessing.

Taken together, the literature demonstrates that Naïve Bayes and Logistic Regression remain effective and practical tools for fake news detection. While complex deep learning models can achieve high accuracy when abundant data and resources are available, NB and LR consistently offer dependable performance across a variety of datasets. Their combination of interpretability, speed, and reliability makes them suitable for benchmarking studies. Building on this evidence, the present research focuses specifically on these two models, aiming to provide a comparative evaluation of their strengths and weaknesses in the context of fake news classification.



IV. DATASET AND DATA PREPROCESSING

• Dataset

This study employs the Constraint@AAAI 2021 shared task dataset on COVID-19 fake news detection, which is publicly accessible via Kaggle [1]. The dataset contains thousands of news articles and social media posts, each labeled as either *real* or *fake*. It was chosen due to its reliable annotations, balanced class distribution, and relevance to COVID-19-related misinformation, making it a strong benchmark for evaluating machine learning models in this domain. To ensure unbiased evaluation, the dataset was randomly divided into 80% for training and 20% for testing, with stratified sampling to maintain class balance.

• Data Preprocessing

To prepare the dataset for model training, a rigorous preprocessing pipeline was applied:

1. **Text Normalization:** All text was converted to lowercase to maintain uniformity.
2. **Noise Removal:** Digits, punctuation, URLs, and special characters were removed.
3. **Stopword Removal:** Common stopwords such as “is,” “the,” and “and” were filtered out to reduce noise.
4. **Stemming and Lemmatization:** Words were reduced to their base forms to ensure consistency.
5. **Feature Extraction:** The cleaned text was transformed into numerical features using Term Frequency–Inverse Document Frequency (TF-IDF). Both unigrams and bigrams were considered, capturing single words and short word sequences for richer contextual representation.

• Model Implementation

Two supervised learning algorithms were implemented for classification:

- **Naïve Bayes:** A probabilistic classifier based on Bayes’ theorem with the assumption of conditional independence among features. It is computationally efficient and performs well with high-dimensional text data.
- **Logistic Regression:** A linear classifier that estimates probabilities for binary outcomes using the logistic function. When combined with TF-IDF features, it provides strong predictive performance and interpretability.

Both models were trained on the preprocessed training set and tested on the reserved test set.

• Evaluation Metrics

The effectiveness of both models was measured using widely adopted metrics:

- **Accuracy** – the proportion of correctly classified samples.
- **Precision** – the proportion of correctly predicted fake news among all predicted fake instances.
- **Recall** – the ability of the model to identify all fake news samples.
- **F1-Score** – the harmonic mean of precision and recall, balancing both.

In addition to these numerical measures, confusion matrices and classification report pie charts were generated to visualize classification outcomes and provide a clearer comparison of model performance.

V. METHODOLOGY

This research utilizes a supervised machine learning framework to tackle the issue of fake news detection, with a focus on comparing the performance of Naïve Bayes and Logistic Regression classifiers. The dataset employed in this study is sourced from the Constraint@AAAI 2021 shared task on COVID-19 fake news detection, accessible via Kaggle. It consists of thousands of news articles and social media posts, each containing a text field and a binary label indicating whether the content is real or fake. This dataset was chosen due to its reliable annotations, balanced class distribution, and COVID-19-specific content, making it well-suited for benchmarking machine learning models in this domain [1].

To ensure robust training and unbiased evaluation, the dataset was randomly split into 80% for training and 20% for testing.

Prior to model training, the textual data underwent a comprehensive preprocessing pipeline to standardize and clean the content. All text was converted to lowercase for consistency, while non-textual elements such as punctuation, digits, URLs, and special characters were removed. Common stop words, including terms like “is,” “the,” and “and,” were filtered out to minimize noise. Additionally, stemming and lemmatization techniques were applied to reduce words to their base forms, ensuring uniformity and improving the quality of features extracted from the text.

The cleaned text was then transformed into numerical features using Term Frequency–Inverse Document Frequency (TF-IDF) vectorization. This method emphasizes terms that are more informative while down-weighting frequently occurring but less meaningful words. Both unigrams and bigrams were extracted to capture individual words as well as short sequences of words, thereby enhancing the contextual information available to the classifiers. The study implemented two



supervised learning algorithms on the prepared features. The Naïve Bayes classifier, which applies Bayes' theorem with the assumption of conditional independence among features, was chosen for its efficiency and ability to handle high-dimensional text data. Logistic Regression was also applied, leveraging the logistic function to estimate probabilities for binary classification, offering both interpretability and strong predictive performance when combined with TF-IDF features.

Both models were trained on the preprocessed training set and evaluated using the reserved test set. Their performance was assessed using standard metrics: accuracy, precision, recall, and F1-score. Accuracy represents the proportion of correctly classified records, precision evaluates the correctness of fake news predictions, recall measures the ability to identify fake news comprehensively, and F1-score provides a harmonic mean of precision and recall. These metrics allowed a systematic comparison of the two classifiers' effectiveness.

This methodology facilitates a detailed analysis of the comparative strengths and limitations of Naïve Bayes and Logistic Regression, illustrating their applicability as practical tools to mitigate the spread of misinformation.



VI. RESEARCH MODEL

In this research, we implemented below models to evaluate the veracity of news articles. In this below there are research models.

1. "Naïve Bayes (Using Count vectorizer Features)": The bias theorem is the foundation of the NBC. We have used scikit-learn to get naive bayes classifiers.
2. "Logistic regression (Using word level tf-idf Features)": This is a supervised learning model. This means something where we use labeled data. It is also a classification model. It uses a sigmoid function. We got this model from scikit-learn.

VII. RESULTS OVERVIEW

The experimental results indicate that both models achieved strong performance, with Logistic Regression slightly outperforming Naïve Bayes in overall accuracy and F1-score.

Table I
Performance Metrics for Naïve Bayes Model

Metric	REAL	FAKE
Precision	0.90	0.95
Recall	0.95	0.90
F1-Score	0.92	0.93

**Overall Performance:**

Metric	Score
Accuracy	92.37%
Micro Average F1	0.92
Weighted Average F1	0.92

Table II
Performance Metrics for Logistic Regression Model

Metric	REAL	FAKE
Precision	0.93	0.95
Recall	0.95	0.93
F1-Score	0.94	0.94

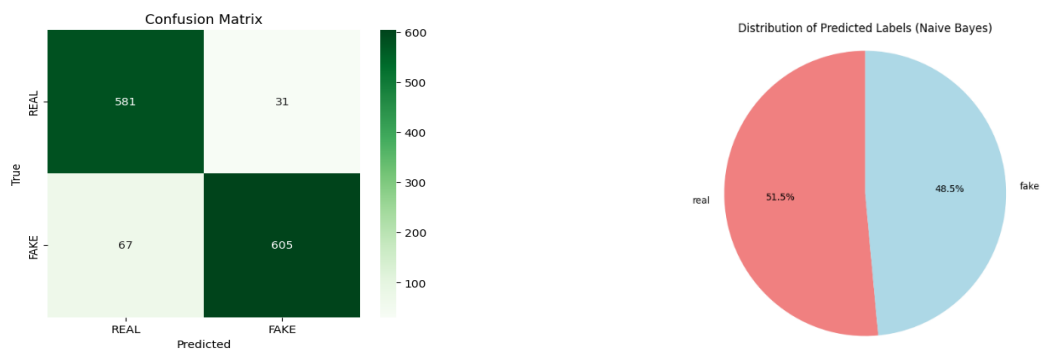
Overall Performance

Metric	Score
Accuracy	93.85%
Micro Average F1	0.94
Weighted Average F1	0.94

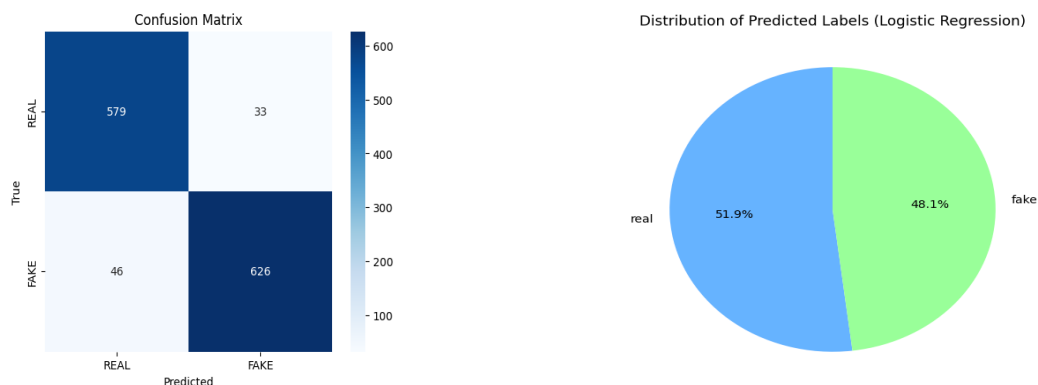
- Visualizations:**

The complement the numerical results, the following graphs were generated:

- ✓ **Confusion Matrix and Pie Chart for Naïve Bayes Model:**

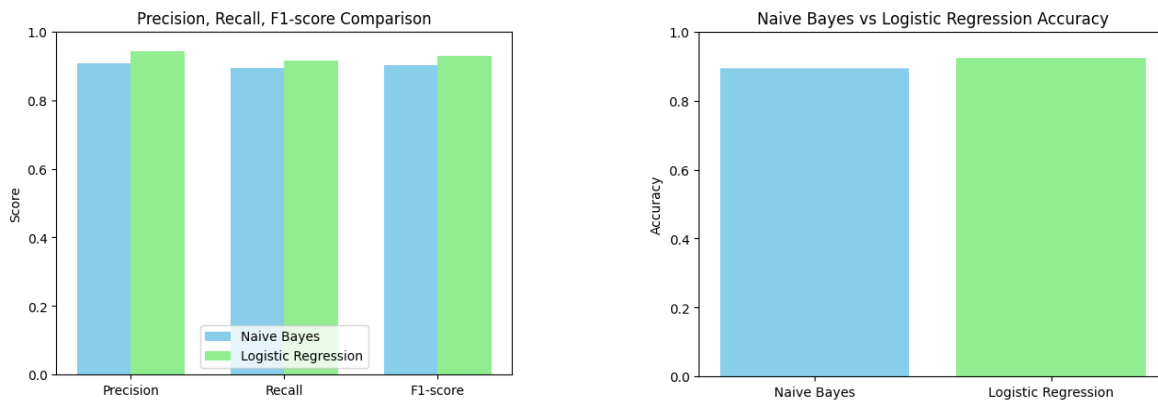


- ✓ **Confusion Matrix and Pie Chart for Logistic Regression Model:**





✓ Precision, Recall, F1-score and Accuracy Comparison (Bar Chart)



These figures provide visual confirmation of the numerical metrics, illustrating the distribution of correct and incorrect predictions across both models.

VIII. DISCUSSION

The outcomes of this study demonstrate that both Naïve Bayes and Logistic Regression are reliable approaches for identifying false and genuine news items. Despite their comparable effectiveness, Logistic Regression achieved marginally higher accuracy and F1-score. This variation arises from the distinct mechanisms by which the two models analyze data. The Naïve Bayes classifier treats all input terms as independent of one another, which simplifies the process but prevents it from detecting interactions among words. Conversely, Logistic Regression employs a weighted learning strategy that recognizes relationships between features, enabling it to detect deeper textual patterns and linguistic context. Another factor contributing to Logistic Regression's advantage is its capacity to fine-tune weights using iterative optimization, allowing greater adaptability to diverse language structures and writing styles. While this method delivers slightly better predictive capability, it also demands more processing time compared with the faster but less flexible Naïve Bayes algorithm. It is important to acknowledge that this investigation focuses solely on textual properties of news data. However, misinformation on social media frequently combines text with other forms of content, such as photographs, video clips, and user interactions. Since these multimodal features were not included, the model's contextual understanding remains limited. Future studies could enhance the framework by integrating multiple data formats or by applying advanced neural architectures, such as transformer-based systems, to interpret context and meaning across several information sources.

IX. CONCLUSION

This study evaluated Naïve Bayes and Logistic Regression for COVID-19 fake news detection using the Constraint@AAAI 2021 / Kaggle dataset. After rigorous preprocessing—including text normalization, stopwords removal, stemming, lemmatization, and TF-IDF vectorization—both models demonstrated strong performance. Naïve Bayes achieved 92.37% accuracy (macro F1-score 0.92), while Logistic Regression slightly outperformed it with 93.85% accuracy (macro F1-score 0.94).

The results highlight the importance of high-quality preprocessing and dataset selection. Naïve Bayes offers computational efficiency for large datasets, whereas Logistic Regression provides slightly higher precision and recall, making it suitable for applications where accurate identification of fake news is critical. These findings support the development of reliable machine learning systems to mitigate misinformation. Future work may explore hybrid or ensemble models to further enhance detection performance.

X. FUTURE SCOPE

Although the current study demonstrates the effectiveness of machine learning models such as Naïve Bayes and Logistic Regression for fake news detection, there remain several avenues for future research and enhancement. First, larger and more diverse datasets covering multiple domains (e.g., political news, health misinformation, financial fraud) can be incorporated to improve generalizability across contexts. Second, advanced deep learning models such as BERT, LSTM, or hybrid transformer architectures could be explored to capture semantic nuances and contextual information more effectively. Third, the integration of multimodal data—including images, videos, and social media metadata—may



significantly improve classification accuracy, since fake news often relies on both textual and visual cues. Furthermore, real-time detection systems can be developed for deployment in social media platforms to flag misinformation instantly and prevent its spread. Finally, ethical considerations such as explainability, fairness, and user privacy must be addressed to ensure that detection systems remain transparent, unbiased, and socially responsible.

REFERENCES

- [1]. J. Y. Khan, M. T. I. Khondaker, S. Afroz, G. Uddin, and A. Iqbal, *Mach. Learn. Appl.* **4**, 100032 (2021). <https://doi.org/10.1016/j.mlwa.2021.100032>
- [2]. J. Jouhar, A. Pratap, N. Tijo, and M. Mony, *Procedia Comput. Sci.* **233**, 763 (2024). <https://doi.org/10.1016/j.procs.2024.03.265>
- [3]. H. F. Villela, F. Corrêa, J. S. A. N. Ribeiro, A. Rabelo, and D. B. F. Carvalho, *J. Interact. Syst.* **14**, 1 (2023). <https://doi.org/10.5753/jis.2023.3020>
- [4]. A. Jain, H. Khatter, A. K. Gupta, and H. Khatter, in *Proc. 2nd Int. Conf. Intell. Commun. Comput. Tech. (ICICT)*, pp. 189–193 (IEEE, 2019). <https://doi.org/10.1109/ICICT46931.2019.8977659>
- [5]. U. Sharma, S. Saran, and S. M. Patil, *Int. J. Eng. Res. Technol. (IJERT)* **9**, 509 (2021). <https://www.ijert.org/fake-news-detection-using-machine-learning-algorithms>
- [6]. P. Kulkarni, S. Karwande, R. Keskar, P. Kale, and S. Iyer, *ITM Web Conf.* **40**, 03003 (2021). <https://doi.org/10.1051/itmconf/20214003003>
- [7]. L. Waikhom and R. S. Goswami, in *Proc. Int. Conf. Adv. Comput. Manage. (ICACM-2019)*, pp. 680–684 (SSRN, 2019). <https://ssrn.com/abstract=3462938>
- [8]. S. Mishra, P. Shukla, and R. Agarwal, *Wirel. Commun. Mob. Comput.* **2022**, 1575365 (2022). <https://doi.org/10.1155/2022/1575365>
- [9]. B. K. Sutradhar, M. Zonaid, N. J. Ria, and S. R. H. Noori, *Daffodil Int. Univ., Bangladesh* (n.d.). <https://doi.org/10.1063/5.0070019>
- [10]. A. Mykytiuk, V. Vysotska, O. Markiv, L. Chyrun, and Y. Pelek, in *CEUR Workshop Proc. (COLINS-2023)* (2023). <https://ceur-ws.org/Vol-3381/paper35.pdf>
- [11]. X. Zhou, K. Shu, and H. Liu, in *Proc. 12th ACM Conf. Web Search Data Min. (WSDM)* (ACM, 2019). <https://doi.org/10.1145/3289600.3291382>
- [12]. A. Altheneyan and A. Alhadlaq, *IEEE Access* **11**, 29447 (2023). <https://doi.org/10.1109/ACCESS.2023.3260763>
- [13]. A. M. Rajasenah and M. G. Augusta, *Test Eng. Manage.* **83**, 11572 (2020). <https://www.researchgate.net/publication/341902603>
- [14]. J. Huang, *J. Phys. Conf. Ser.* **1693**, 012158 (2020). <https://doi.org/10.1088/1742-6596/1693/1/012158>
- [15]. M. N. Pavan, P. R. Prasad, T. Gowda, T. S. Vibhakar, and S. Shidnal, *Int. Res. J. Eng. Technol. (IRJET)* **7**, 2151 (2020). <https://www.irjet.net/archives/V7/i12/IRJET-V7I1243.pdf>
- [16]. M. Sudhakar and K. P. Kaliyamurthi, *Meas. Sens.* **24**, 100495 (2022). <https://doi.org/10.1016/j.measen.2022.100495>
- [17]. M. Hisham, R. Hasan, and S. Hussain, *Sir Syed Univ. Res. J. Eng. Technol.* **13**, 118 (2023). <https://doi.org/10.33317/ssurj.565>
- [18]. Y. Mahmud, N. S. Shaeali, and S. Mutalib, *Int. J. Adv. Comput. Sci. Appl. (IJACSA)* **12**, 658 (2021). <https://doi.org/10.14569/IJACSA.2021.0121076>
- [19]. M. M. M. Hlaing and N. S. M. Kham, in *Proc. 11th Int. Workshop Comput. Sci. Eng. (WCSE 2021)*, pp. 59–65 (2021). <https://doi.org/10.18178/wcse.2021.02.010>
- [20]. S. Hangloo and B. Arora, *Int. J. Comput. Sci. Inf. Technol. (IJCSIT)* (2021). [Central Univ. Jammu]
- [21]. A. D. Radhi, H. A. H. Al Naffakh, A. I. Fuqdan, B. A. Hakim, and B. Al-Attar, *BIO Web Conf.* **97**, 00123 (2024). <https://doi.org/10.1051/bioconf/20249700123>
- [22]. J. Singh, F. Liu, H. Xu, B. C. Ng, and W. Zhang, *IEEE J. LaTeX Class Files* **1**, 1 (2024). <https://doi.org/10.48550/arXiv.2405.04165>
- [23]. H. L. Gururaj, H. Lakshmi, B. C. Soundarya, F. Flammmini, and V. Janhavi, *J. ICT Stand.* **10**, 509 (2022). <https://doi.org/10.13052/jicts2245-800X.1042>
- [24]. S. A. Al-Obaidi and T. Çağlıkantar, *Iraqi J. Comput. Sci. Math.* **5**, 12 (2024). <https://doi.org/10.52866/2788-7421.1200>
- [25]. S. Mishra, P. Shukla, and R. Agarwal, *Wirel. Commun. Mob. Comput.* **2022**, 1575365 (2022). <https://doi.org/10.1155/2022/1575365>