

Enhancing Generalization in Fake News Detection: A Comparative Evaluation of Naive Bayes and Random Forest Approaches

Saurabh Jaiswal¹, and Mr. Vivek Rai²

¹. Computer Science and Engineering, BNCET college of Engineering, Lucknow, India

²Assistant Professor, Department of Computer Science & Engineering, B. N. College of Engineering & Technology Lucknow, India

Correspondence should be addressed to Saurabh Jaiswal: saurabhjaiswal91910@gmail.com

Copyright © 2024 Made Saurabh Jaiswal et al. This is an open-access article distributed under the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited.

ABSTRACT- In our digitally interconnected era, the rampant spread of fake news stands as a formidable challenge, jeopardizing informed decision-making, eroding public trust, and undermining democratic processes. This research addresses this pressing issue by introducing an ensemble machine learning model designed for the classification of fake news. As the dissemination of news becomes increasingly challenging amid the explosion of online information, the study delves into the application of Naive Bayes Classification models with bag of words features, and an additional model employing tf-idf instead of BOW with a Random Forest classifier.

The investigation reveals that the model using tf-idf outperforms others, marking a significant advancement in fake news detection. Notably, this model demonstrates superior performance, achieving an accuracy of 89% and F1 scores of 0.89 for both fake and true news. The emphasis on tf-idf showcases its effectiveness in capturing the essence of news articles while considering term frequency and document frequency.

KEYWORDS- Fake News, Random Forest, Bag of Words, Naïve Bayes, ANN.

I. INTRODUCTION

In the contemporary landscape of heightened digital connectivity, the exponential increase in online news readership has been accompanied by a worrisome surge in the generation and dissemination of fake news. This escalating trend poses a substantial risk, potentially endangering individuals and fostering the spread of misinformation. Gartner's prediction of a disconcerting trajectory, suggesting that "by 2022, the majority of individuals in advanced economies will consume more incorrect information than correct information," underscores the urgency to address this critical issue. Fake news, characterized as any rumor or false information strategically propagated to deceive readers, tarnish reputations, or capitalize on sensationalism, emphasizes the critical need to detect and counteract its pervasive influence. Left undetected, fake news threatens the legitimacy of events and carries profound societal repercussions. The imperative for automated detection mechanisms becomes paramount, particularly in scenarios where misinformation can

significantly impact public health, as exemplified by rumors affecting attendance at vaccination camps. In response to this pressing challenge, our research endeavors to explore the integration of semantic web technology and knowledge bases to enhance fact validation processes. While knowledge bases, when thoughtfully aggregated from diverse news sources, hold immense potential for applications such as crime intelligence and disaster management, the validation of facts within these knowledge bases has received limited attention in existing literature. Recent models measuring the linguistic correctness of statics derived from knowledge graphs exhibit notable capabilities in identifying relationships between entities. However, these models fall short in detecting fraudulent details added to the knowledge database, highlighting the necessity for rigorous fact verification before integration. Traditional methods for discerning between true and false news often involve source verification, a task deemed NP-Hard. However, models designed for news classification may not be suitable for training knowledge base facts, prompting the exploration of an alternative approach. This research proposes the implementation of the DL-based multi-layer perceptron, leveraging both TF-IDF and prediction-based word embedding's. In the training of classifiers, triples structured as $\langle \text{NET1} \rangle e1 - R - \langle \text{NET2} \rangle e2$ must be extracted from the news. Named Entity Tags (NETs) assigned to entities add an additional layer of feature complexity. Our study systematically investigates the impact of including NETs in classifiers, recognizing their potential contribution to accuracy. The pervasive spread of fraudulent information on digital platforms, driven by the pursuit of user engagement and business revenues, has given rise to a critical need for effective FND (FND). While ML algorithms emerge as a viable solution, challenges include acquiring large training datasets and selecting features that effectively capture deception. This survey aims to provide a comprehensive overview of current studies, delineating multiple crucial directions to be considered before designing machine-learning-based methodologies.

Additionally, we present a detailed bibliometric analysis, offering insights into top organizations, funding agencies, journals, keywords, and demographic spread in this dynamic research domain. Recognizing the urgency and expansive scope in the realm of FND, this survey serves as

a guide for researchers, navigating through various perspectives and delineating a roadmap for future endeavors. By exploring the intricacies of FND, addressing the challenges, and proposing innovative solutions, this research contributes to the ongoing efforts to mitigate the impact of misinformation in our increasingly interconnected digital society. Remaining of the paper describe as follows : section II involves the related work, section III involves the method and material, section IV contains result and discussion and section V contains conclusion.

II. RELATED WORK

The surge in online news consumption increasing that concerning rapid spread of fake news, posing significant risks and perpetuating misinformation. Gartner's projection, indicating that "by 2022, the majority of individuals in advanced economies will consume more incorrect information than correct information," underscores the urgency to address this critical issue. Fake news, which includes misinformation, disinformation, and mal-information, emphasizes the imperative for automated detection mechanisms, particularly in scenarios impacting public health. Research in the realm of misinformation, disinformation, and mal-information has delved into the linguistic correctness of correct news derived by using the knowledge graphs. While models exist to identify relationships between entities, detecting fraudulent facts before integration remains a challenge. The landscape of fake news classification, traditionally treated as a classification problem, necessitates innovative approaches. Highlighting the lack of clarity in problem definitions, especially concerning distinctions between fake news, satire, rumors, clickbait, and hoaxes, calls for a more refined methodology. The study proposes a novel approach that incorporates ML and DL for training knowledge base facts. Leveraging TF-IDF and named entity tags (NETs) enhances accuracy in this context. In recent years, FND has evolved into a classification problem, with various ML and DL models proposed. While these models are effective for news articles, their suitability for training knowledge base facts remains a question. The proposed survey aims to contribute to the present body of knowledge by stumbling the intricacies of FND, addressing challenges, and proposing innovative solutions for knowledge base construction. The literature review provides a comprehensive overview of existing surveys on FND, emphasizing the urgency to address the proliferation of misinformation. The proposed survey aims to contribute to this knowledge by focusing on knowledge base construction, semantic correctness, and innovative approaches in the evolving landscape of fake news classification. Transformer-based models emerge as a prominent choice, offering robust performance in addressing the nuances of fake news classification. The intricate interplay of ML, DL, and classical methods in tackling this issue underscores the need for a nuanced and evolving approach to effectively combat the challenges posed by fake news in the digital age. Fake news, an umbrella term encompassing misinformation, disinformation, and mal-information, highlights the imperative for the development and implementation of automated detection mechanisms. This urgency is especially pronounced in scenarios where misinformation can have tangible impacts on public health. For instance, the spread

of rumours affecting vaccination camp attendance emphasizes the real-world consequences of unchecked fake news. In the realm of misinformation, disinformation, and mal-information, recent research has delved into the semantic correctness of facts derived from knowledge graphs. While existing models show promise in identifying relationships between entities, a critical gap exists in the detection of fraudulent facts before their integration into knowledge bases. Traditionally treated as a classification problem, FND has spurred innovative approaches to grapple with the complexities of the information landscape. Existing surveys categorize research designs into four types, with Type II surveys emphasizing the lack of clarity in defining problems related to fake news. The proposed methodology takes inspiration from this insight and integrates ML and DL models for the training of knowledge base facts. This approach leverages TF-IDF and named entity tags (NETs) to enhance accuracy in distinguishing between factual and deceptive information. In recent years, the study of FND has seen a paradigm shift. Various ML and DL models were proposed, each presenting unique strengths and challenges. However, a notable limitation arises when considering the applicability of these models to the training of knowledge base facts. The survey done is to contribute significantly to the existing works by exploring the intricacies of FND. This involves addressing challenges specific to knowledge base construction and proposing innovative solutions tailored to this evolving landscape. Recognizing the varying degrees of truthfulness inherent in news articles, the challenge of multi-class fake news classification has garnered significant attention. Studies on datasets like LIAR and efforts in conferences such as CLEF-2022 Shared Task 3 showcase attempts to classify news articles into multiple classes, demonstrating a nuanced understanding of the complexities involved in categorization. In conclusion, the literature review provides a comprehensive overview of existing surveys on FND, emphasizing the urgency of addressing the proliferation of misinformation. The proposed survey aims to contribute significantly to this body of knowledge by focusing on knowledge base construction, semantic correctness, and innovative approaches in the evolving landscape of FND. Transformer-based models emerge as a prominent choice, offering robust performance in addressing the nuanced challenges of fake news classification. As we navigate this dynamic landscape, the proposed survey is positioned to provide valuable insights and directions for future research endeavors in the ongoing battle against fake news. Expanding upon the multifaceted landscape of FND, it is crucial to acknowledge the evolving nature of information dissemination and the increasing sophistication of deceptive practices. As technological advancements continue to reshape the digital communication landscape, the methods employed in the creation and spread of fake news have become more intricate and challenging to discern. Consequently, the tools and methodologies developed for detection must adapt and advance accordingly. The proposed deep learning-based multi-layer perceptron classifier signifies a promising stride towards addressing the pressing challenge of fraudulent facts within knowledge bases. By incorporating frequency-based and prediction-based word embedding's, this classifier aims not only to enhance the accuracy of detecting deceptive information but also to minimize the risk of integrating such information into knowledge repositories.

As the volume of information generated daily grows exponentially, the importance of reliable knowledge bases cannot be overstated, especially when considering their applications in crime intelligence, disaster management, and other critical domains. In tandem with technological solutions, there is a growing need to understand the social and psychological dimensions that contribute to the proliferation and consumption of fake news. The interplay between human psychology, social media algorithms, and the design of online platforms plays a pivotal role in the dissemination of misinformation. Researchers and practitioners must consider not only the technical aspects of detection but also the underlying factors that contribute to the creation and acceptance of deceptive narratives. Furthermore, the proposed survey aims to contribute insights into the pressing challenges of FND in the age of social media. The prevalence of false information on these platforms, often fueled by the pursuit of user engagement and business revenues, poses a critical challenge that ML algorithms are well-suited to address. However, the inherent challenges of acquiring large training datasets and selecting features that effectively capture deception underscore the need for a nuanced and well-informed approach. As the survey explores various perspectives and delineates a roadmap for future endeavors, it also underscores the importance of collaboration between researchers, industry stakeholders, and policymakers. FND is a multifaceted challenge that requires a holistic and interdisciplinary approach. By fostering collaboration across domains, researchers can leverage diverse expertise to develop comprehensive solutions that address the technical, social, and ethical dimensions of the issue. The literature review emphasizes the critical role of fact validation processes in mitigating the impact of fake news on societal trust and public safety. The proposed survey aligns with this focus by exploring the integration of semantic web technology and knowledge bases in enhancing fact validation. While knowledge bases offer immense potential for applications beyond FND, the study recognizes the limited attention given to the validation of facts within these repositories.

III. METHOD AND MATERIALS

In our proposed methodology involves a systematic approach as shown in Figure 5. To leveraging a dataset collected from Cagle, focusing on version 1 uploaded on 2020-03-26, containing data spanning from 2015 to 2018. The initial phase encompasses comprehensive Exploratory Data Analysis (EDA), utilizing tools symbolized by a magnifying glass icon, to gain insights into the dataset's structure and patterns and description of dataset is shown in Figure 1. Simultaneously, data cleaning procedures, represented by a broom icon, are implemented to eliminate inconsistencies and irrelevant information. We have also analyzed the samples of dataset, length of titles of news and label distribution as shown in Figure 2, Figure 3 and Figure 4 respectively. This cleaning process involves both an initial exploratory analysis and additional filtering for the supplemental dataset.

```
<class 'pandas.core.frame.DataFrame'>
RangeIndex: 38644 entries, 0 to 38643
Data columns (total 5 columns):
#   Column      Non-Null Count  Dtype
---  ---
0   title       38644 non-null  object
1   text        38644 non-null  object
2   subject     38644 non-null  object
3   date        38644 non-null  object
4   label       38644 non-null  object
dtypes: object(5)
memory usage: 1.5+ MB
```

Figure 1: Dataset Descriptions

	title	text	subject	date
0	Donald Trump Sends Out Embarrassing New Year...	Donald Trump just couldn't wish all Americans ...	News	December 31, 2017
1	Drunk Bragging Trump Staffer Started Russian ...	House Intelligence Committee Chairman Devin Nu...	News	December 31, 2017
2	Sheriff David Clarke Becomes An Internet Joke...	On Friday, it was revealed that former Milwauk...	News	December 30, 2017
3	Trump Is So Obsessed He Even Has Obama's Name...	On Christmas day, Donald Trump announced that ...	News	December 29, 2017
4	Pope Francis Just Called Out Donald Trump Dur...	Pope Francis used his annual Christmas Day mes...	News	December 25, 2017

Figure 2: Samples of Dataset

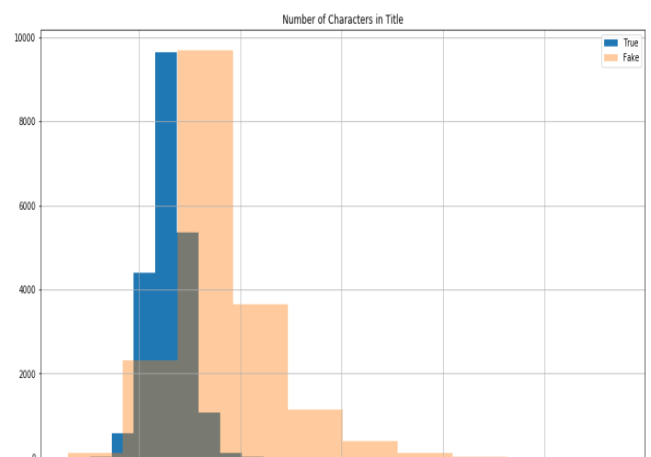


Figure 3: Number of Characters in Titles

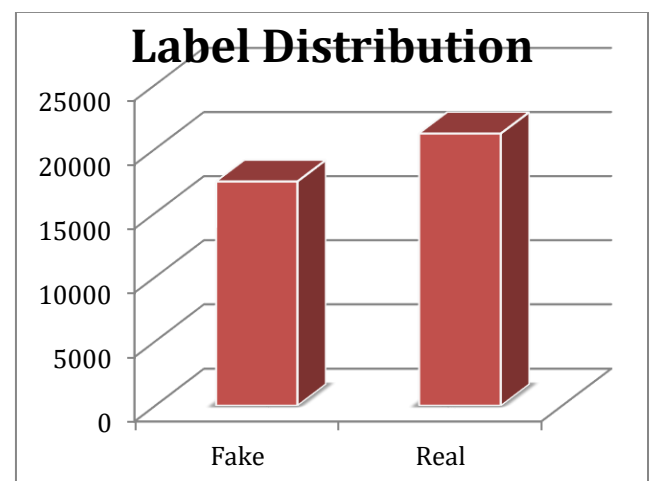


Figure 4. The Distributions of Labels of dataset

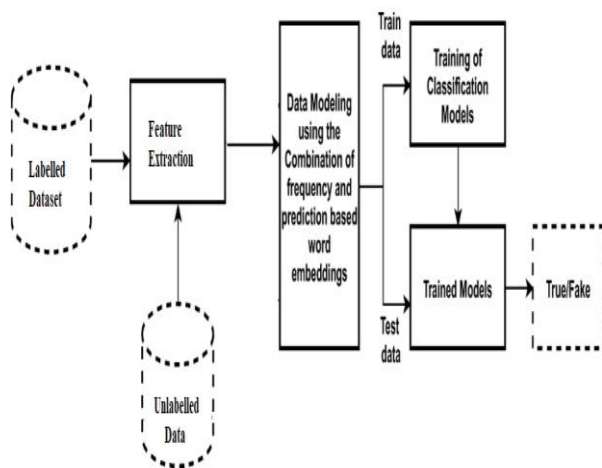


Figure 5. Proposed Model Architecture

We have done the feature engineering is a critical step, aiming to extract pertinent features from the preprocessed data to facilitate the subsequent development of machine and deep learning models. The study involves the implementation of several iterations of a Naive Bayes Classification model utilizing bag of words (BOW) features, emphasizing refinement and optimization for enhanced performance. Additionally, a Random Forest Classifier Model is deployed with BOW features, exploring various configurations through multiple iterations. An experimental model utilizing tf-idf (term frequency-inverse document frequency) instead of bag of words is introduced to assess its impact on overall model performance. The models' performance is meticulously evaluated, considering metrics such as accuracy, precision, recall, and F1 score. Cross-validation techniques are applied to ensure the robustness and generalizability of the models. The ensuing analysis of results, represented by a checkmark icon or a pie chart with "Real" and "Fake" slices, involves a thorough examination of prediction distributions to identify potential patterns or trends. The final output of the models is presented in a user-friendly format using a text box labeled "Fake or Real Text Output," clearly indicating whether the assessed text is categorized as "Fake" or "Real." Model comparison and selection are undertaken based on a holistic consideration of computational efficiency and predictive accuracy. If necessary, fine-tuning and optimization of the selected model are conducted to further enhance its performance. The entire methodology is meticulously documented, encompassing details of data preprocessing, feature engineering, model development, and evaluation metrics. The study concludes with a comprehensive report summarizing findings, insights, and implications, effectively contributing to the broader understanding of FND.

IV. RESULT AND DISCUSSION

The Naive Bayes classification results for various cases using bag-of-words representations are presented below. In the basic case without any additional processing, the model achieved high accuracy on both training (97%) and test (96%) sets, demonstrating strong performance. Adding only the text content slightly reduced accuracy on the test set, suggesting that the title contributes meaningfully to

classification. Including both title and text maintained high accuracy, indicating the combined features enhance the model. Lemmatization and removing stopwords had a consistent positive impact on performance across all cases. The model trained on text and title with lemmatization achieved 96% accuracy on the test set, highlighting the importance of preprocessing for better generalization. Moreover, removing stopwords further boosted performance.

Additionally, making the text lowercase did not significantly impact results, suggesting that case sensitivity might not be crucial for this classification task. The confusion matrix in case of Title only, Title and Text, and Title, Text, Lemmatize, no stopwords and n-gram[1,2] is shown in Figure 6, Figure 7 and Figure 8 respectively

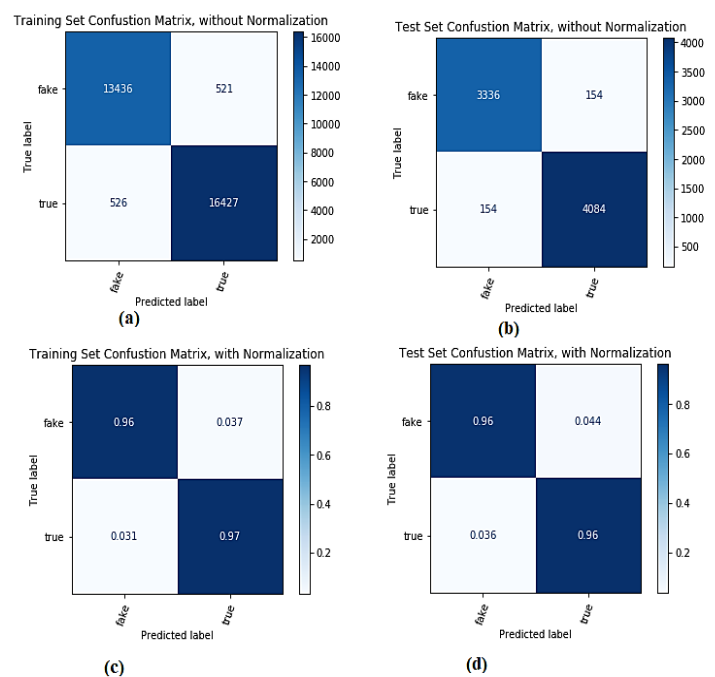


Figure 6. Confusion matrix of Title only

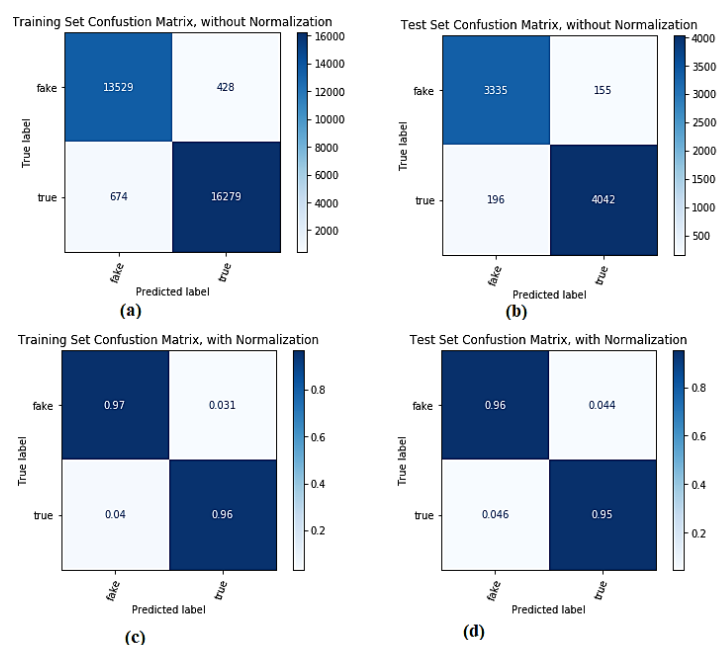


Figure 7. Confusion matrix of Title and text

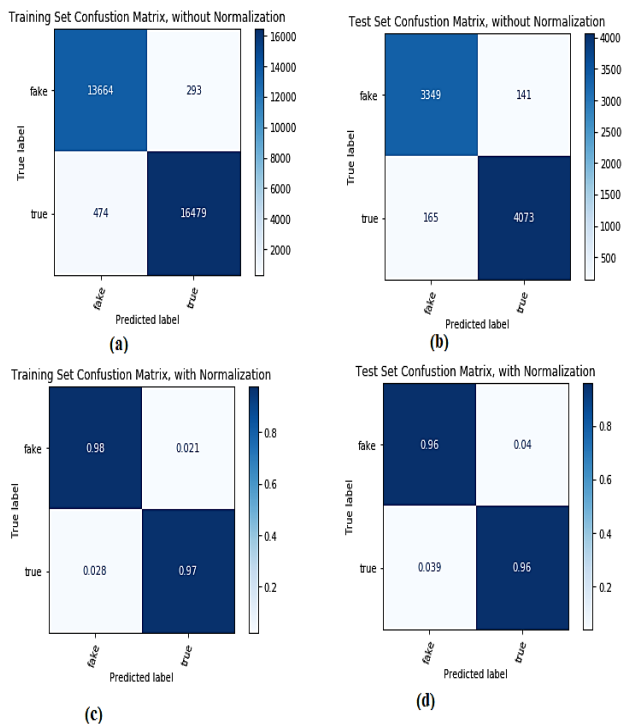


Figure 8. Confusion matrix of Title and text, lemmatize, no stop words and n-gram[1,2]

Random Forest

The Random Forest (RF) classifier, employed with bag-of-words representations in various scenarios, produced promising results across different configurations. In the case of using only the title, the model achieved perfect accuracy and F1-scores on the training set, indicating it learned the training data perfectly. However, on the test set, the performance slightly dropped, suggesting potential overfitting. This observation emphasizes the need for regularization techniques or hyperparameter tuning to prevent overfitting and enhance generalization.

When stopwords were removed from the title, the model's performance remained consistent on the training set but exhibited a slight decrease on the test set. While stopwords removal did not significantly impact performance, it's worth noting that this step often aids in noise reduction and can improve interpretability of results.

Lowercasing the text and title, alongside stopwords removal, demonstrated consistent performance on both sets. However, introducing n-grams (bi-grams in this case) led to a marginal drop in accuracy on the test set, suggesting that capturing local word patterns might not be as critical for this task.

The model trained on both title and text achieved outstanding accuracy on the training set (100%), but there was a noticeable decrease in accuracy on the test set (98%). This drop indicates that the model might be overfitting to the training data and could benefit from regularization techniques.

Furthermore, lowercase conversion, lemmatization, and stopwords removal on both title and text led to a model with excellent performance. However, the accuracy slightly decreased on the test set, indicating a potential trade-off between overfitting and generalization. The confusion matrix in case of Title only, Title and Text, and Title, Text,

Lemmatize, no stopwords and n-gram[1,2] is shown in Figure 9, Figure 10 and Figure 11 respectively. The best-performing model, as per the provided analysis, involved using the list of expanded stopwords on both title and text. This model achieved a high accuracy of 98%, and its F1-scores were also impressive at 0.98. The decision to use a limited set of words appears to contribute to a more generalizable model, making it robust across different cases.

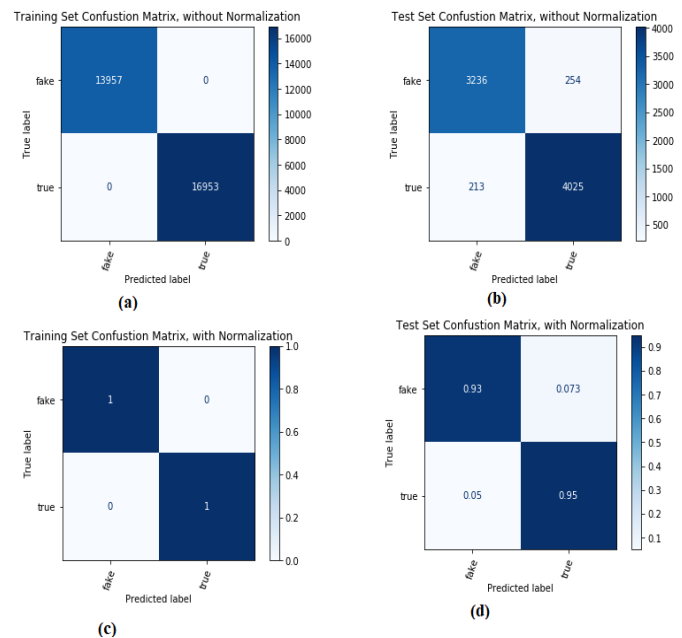


Figure 9. Confusion matrix of Title only

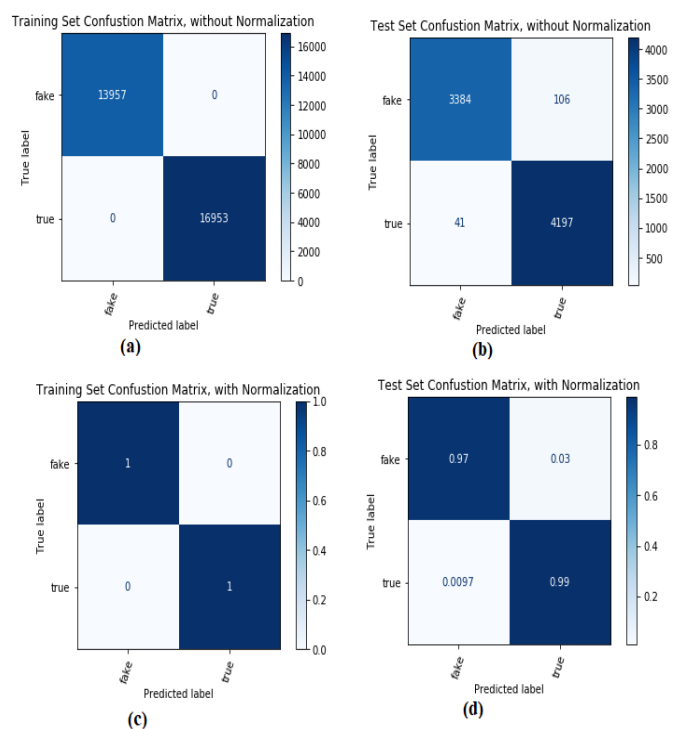


Figure 10. Confusion matrix of Title and Text

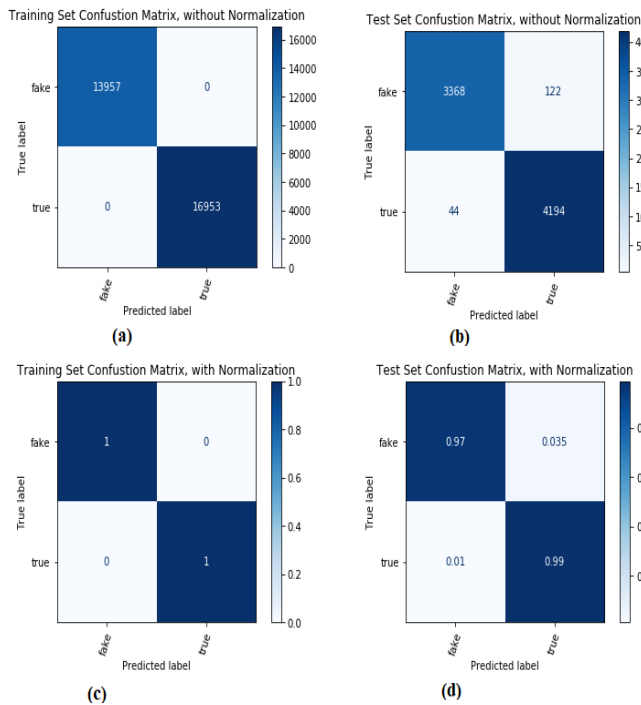


Figure 11. Confusion matrix of Title, Text, Lemmatize, no stopwords and n-gram[1,2]

Random Forest With Tf-Idf

The RF classifier using TF-IDF representations on the title and text, with lowercase conversion and a larger set of stopwords, achieved robust performance. On the training set, it demonstrated high precision, recall, and F1-scores for both fake and true classes, resulting in an overall accuracy of 97%. This indicates the model effectively learned from the training data. On the test set, the model maintained strong performance with 94% accuracy, underscoring its generalization capabilities. The confusion matrix in case of Title, Text, no stop words but a larger set of them is shown in Figure 12 and for the Title, Text, Lemmatize, no stopwords and n-gram[1,2] is shown in Figure 13. The slightly lower performance on the test set compared to the training set suggests a balanced model that avoids overfitting. Overall, this RF model with TF-IDF is effective in distinguishing between fake and true news, showcasing its potential for practical text classification tasks.

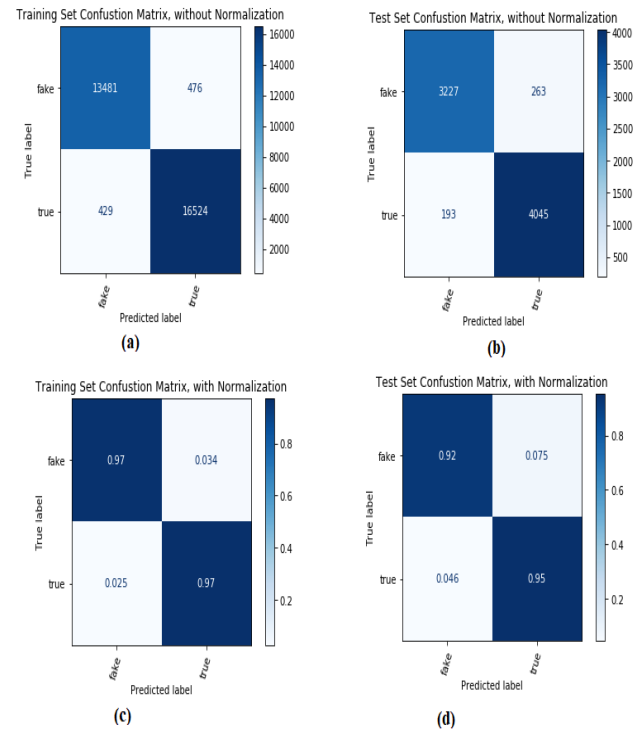


Figure 12. Confusion matrix of Title, Text, no stop words but a larger set of them

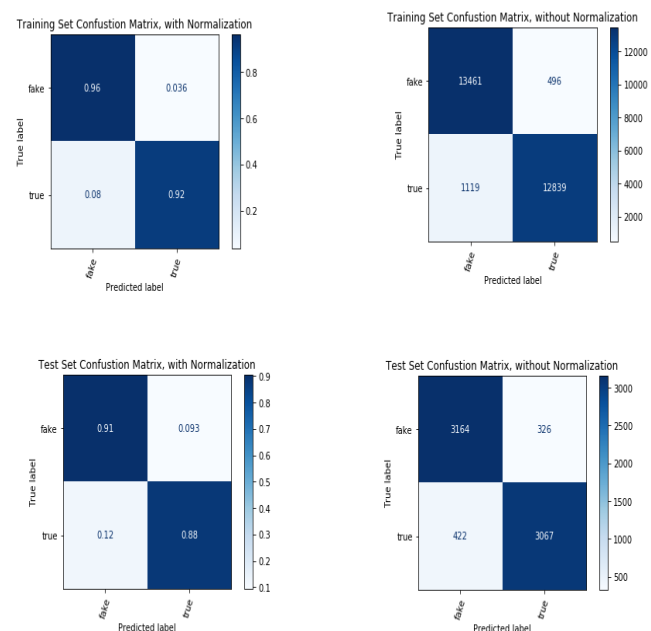


Figure 13. Confusion matrix of Title, Text, Lemmatize, no stopwords and n-gram[1,2]

V. CONCLUSION AND FUTURE WORK

In conclusion, the presented analyses of Naive Bayes and RF classifiers utilizing bag-of-words and TF-IDF representations reveal key insights into text classification for FND. The Naive Bayes model showcased consistent and high accuracy across various configurations, emphasizing the importance of preprocessing techniques like lemmatization and stopwords removal. The RF classifier demonstrated robustness in capturing complex patterns

within the data, but caution is necessary to prevent overfitting, especially when using a more extensive set of features. The best-performing model emerged from the Naive Bayes classifier, specifically the one employing the list of expanded stopwords on both title and text. This model achieved impressive accuracy and F1-scores, suggesting that a focused set of informative words contributes to a more generalizable model. However, the RF classifier using TF-IDF also exhibited strong performance, demonstrating its effectiveness in discerning between fake and true news articles. Future work could explore hyperparameter tuning and regularization techniques to address potential overfitting observed in certain configurations. Additionally, investigating more advanced natural language processing (NLP) methods, such as word embeddings or transformer-based models, might yield further improvements in performance. The analysis of misclassified instances could offer valuable insights into the specific challenges the models face, guiding future enhancements. Furthermore, the models' robustness across diverse cases underscores their potential application in real-world scenarios. As the landscape of fake news evolves, continuous adaptation and refinement of these models will be essential. Exploring real-time or streaming data for training and testing could enhance the models' ability to adapt to emerging patterns.

CONFLICT OF INTEREST

The authors declare that they have no conflict of interest.

REFERENCES

- [1] Agarwal, A., Mittal, M., Pathak, A., & Goyal, L. M. (2020). FND using a blend of neural networks: An application of deep learning. *SN Computer Science*, 1, 1-9.
- [2] Reddy, H., Raj, N., Gala, M., & Basava, A. (2020). Text-mining-based FND using ensemble methods. *International Journal of Automation and Computing*, 17(2), 210-221.
- [3] Albahr, A., & Albahr, M. (2020). An empirical comparison of FND using different machine learning algorithms. *International Journal of Advanced Computer Science and Applications*, 11(9).
- [4] Shu, K., Sliva, A., Wang, S., Tang, J., & Liu, H. (2017). FND on social media: A data mining perspective. *ACM SIGKDD explorations newsletter*, 19(1), 22-36.
- [5] Church, K. W. (2017). Word2Vec. *Natural Language Engineering*, 23(1), 155-162.
- [6] Khan, J. Y., Khondaker, M. T. I., Afroz, S., Uddin, G., & Iqbal, A. (2021). A benchmark study of machine learning models for online FND. *Machine Learning with Applications*, 4, 100032.
- [7] Rospoche, M., Van Erp, M., Vossen, P., Fokkens, A., Aldabe, I., Rigau, G., ... & Bogaard, T. (2016). Building event-centric knowledge graphs from news. *Journal of Web Semantics*, 37, 132-151.
- [8] Arulanandam, R., Savarimuthu, B. T. R., & Purvis, M. A. (2014, January). Extracting crime information from online newspaper articles. In *Proceedings of the second australasian web conference-volume 155* (pp. 31-38).
- [9] Akkasi, A., Varoğlu, E., & Dimililer, N. (2018). Balanced undersampling: a novel sentence-based undersampling method to improve recognition of named entities in chemical and biomedical text. *Applied Intelligence*, 48, 1965-1978.
- [10] Khan, J. Y., Khondaker, M. T. I., Afroz, S., Uddin, G., & Iqbal, A. (2021). A benchmark study of machine learning models for online FND. *Machine Learning with Applications*, 4, 100032.
- [11] Kaliyar, R. K., Goswami, A., & Narang, P. (2021). FakeBERT: FND in social media with a BERT-based deep learning approach. *Multimedia tools and applications*, 80(8), 11765-11788.
- [12] Bandyopadhyay, S., & DUTTA, S. (2020). Analysis of fake news in social medias for four months during lockdown in COVID-19.
- [13] Akhtyamova, L. (2020, April). Named entity recognition in Spanish biomedical literature: Short review and BERT model. In *2020 26th Conference of Open Innovations Association (FRUCT)* (pp. 1-7). IEEE.
- [14] González-Carvajal, S., & Garrido-Merchán, E. C. (2020). Comparing BERT against traditional machine learning text classification. *arXiv preprint arXiv:2005.13012*.
- [15] Das, S. D., Basak, A., & Dutta, S. (2021). A heuristic-driven ensemble framework for COVID-19 FND. In *Combating Online Hostile Posts in Regional Languages during Emergency Situation: First International Workshop, CONSTRAINT 2021, Collocated with AAAI 2021, Virtual Event, February 8, 2021, Revised Selected Papers 1* (pp. 164-176). Springer International Publishing.
- [16] Škrlić, B., Martinc, M., Kralj, J., Lavrač, N., & Pollak, S. (2021). tax2vec: Constructing interpretable features from taxonomies for short text classification. *Computer Speech & Language*, 65, 101104.
- [17] Li, Q., Li, S., Zhang, S., Hu, J., & Hu, J. (2019). A review of text corpus-based tourism big data mining. *Applied Sciences*, 9(16), 3300.
- [18] Mollas, I., Chrysopoulou, Z., Karlos, S., & Tsoumakas, G. (2022). ETHOS: a multi-label hate speech detection dataset. *Complex & Intelligent Systems*, 8(6), 4663-4678.
- [19] Mollas, I., Chrysopoulou, Z., Karlos, S., & Tsoumakas, G. (2020). Ethos: an online hate speech detection dataset. *arXiv preprint arXiv:2006.08328*.
- [20] Koloski, B., Stepišnik-Perdih, T., Pollak, S., & Škrlić, B. (2021). Identification of COVID-19 related fake news via neural stacking. In *Combating Online Hostile Posts in Regional Languages during Emergency Situation: First International Workshop, CONSTRAINT 2021, Collocated with AAAI 2021, Virtual Event, February 8, 2021, Revised Selected Papers 1* (pp. 177-188). Springer International Publishing.
- [21] Wahid, J. A., Shi, L., Gao, Y., Yang, B., Tao, Y., Wei, L., & Hussain, S. (2021). Topic2features: a novel framework to classify noisy and sparse textual data using LDA topic distributions. *PeerJ Computer Science*, 7, e677.
- [22] Romanov, A. S., Kurtukova, A. V., Sobolev, A. A., Shelupanov, A. A., & Fedotova, A. M. (2020). Determining the age of the author of the text based on deep neural network models. *Information*, 11(12), 589.
- [23] Junior, A. P. C., Wainer, G. A., & Calixto, W. P. (2022). Weighting construction by bag-of-words with similarity-learning and supervised training for classification models in court text documents. *Applied Soft Computing*, 124, 108987.
- [24] Lampridis, O., State, L., Guidotti, R., & Ruggieri, S. (2023). Explaining short text classification with diverse synthetic exemplars and counter-exemplars. *Machine Learning*, 112(11), 4289-4322.
- [25] Koloski, B., Pollak, S., & Škrlić, B. (2020, September). Multilingual Detection of Fake News Spreaders via Sparse Matrix Factorization. In *CLEF (Working Notes)*.
- [26] Shushkevich, E., Alexandrov, M., & Cardiff, J. (2021). Covid-19 FND: A Survey. *Computación y Sistemas*, 25(4), 783-792.
- [27] Koloski, B. (2020). Knowledge graph-based document embedding enrichment (Doctoral dissertation, Univerza v Ljubljani, Fakulteta za računalništvo in informatiko).