

Applying Data Mining, AI and Machine Learning for Identifying and Countering Fake News on Social Platforms

R. Porkodi¹ and Dr. A. Kalaivani²

¹Research Scholar, Department of Computer Science, Nehru Arts and Science College, Coimbatore.

²Associate Professor, Department of Computer Applications, Nehru Arts and Science College, Coimbatore.

*Corresponding
Author
R. Porkodi

Article History

Received:
08.08.2025
Revised:
15.09.2025
Accepted:
24.10.2025
Published:
04.11.2025

Abstract: Fake news is growing on all Social Platforms (SM) platforms and has become a serious threat and challenge to the useful information that is accessible in public discussions. We need efficient and scalable ways to identify them. This work presents a holistic AI-driven model for detecting and containing false narratives using ML models trained on structured CSV data. The method combines traditional and deep learning SVM, KNN and LSTM algorithm for news content categorization using linguistic and contextual features. A standardized preprocessing procedure is applied to clean, normalize and extract features from the data for the sake of obtaining the best possible model performance. The models' performance is thoroughly analyzed with model assessment metrics. Results show that LR provides a solid baseline performance, while LSTM demonstrates greater potential in capturing textual semantic dependencies. The proposed approach demonstrates the possibility of implementing lightweight ML using chunked CSV-based models to identify misinformation in Social Platforms ecosystems in real time.

Keywords: Fake news Detection, LSTM, Social Platforms, Analytics, CSV-based Modeling.

INTRODUCTION

As Interactive web platforms has become the window through which consumers get a tether to news, and increasingly, information about the world, news feeds and WhatsApp chats have become the moment-sized versions of the news. With all the good that the age of Interactive web platforms has brought, be it instant access to news and information or global connection, it has also allowed the proliferation of false narratives. The viral structure of Interactive web platforms has only increased the problem, enabling disinformation to travel to millions of people within minutes. False narratives, being fabricated, misleading or deceptive information spread to influence public opinion (Rambukkanda et al., 2017), can have detrimental impacts on public and election outcomes, and even result in violence (Vosoughi et al., 2018). This has in turn raised rising concerns about the veracity of content communicated over digital media with researchers developing approaches to discern and counter false narratives.

Existing methods to detect false narratives are mostly based on ML, DL and NLP. These methods employ different components of the content (e.g., text or images or metadata) to categorize news content as true or false. Yet, current methods have several inherent limitations. NLP models, for example, sometimes fail at grasping the nuanced nature of context, such as sarcasm, irony, or shifts in linguistic behavior. Additionally, the capability of ML methods to produce accurate predictions is highly dependent on the completeness and the statistical distribution of the training data, leading to unjust

predictions or lack of generalization to new data. These limitations emphasize the demand for stronger, more flexible, and scalable approaches towards identifying false news in online communication tools.

In response to these challenges, the main aim is to investigate the capabilities of AI and ML model in both identifying and preventing false narratives in social media. Using AI-based ML tools like labeled, unlabeled and adaptive learning to build and learn models that catch false narratives, this research paper purpose is to generate models that reflect complex and multifaceted data relationships of false narratives. An important part of such approach entails the inclusion of social-media idiosyncrasies in the form of features such as user activity, network analysis, and time a day patterning to improve the efficiency of the detection algorithms (Zhou & Zafarani, 2021). In light of the different types of data involved in the propagation of false narratives, namely text posts, images and videos, we confront the dynamic nature of the social media content by developing a comprehensive solution.

Additionally, ML models' proficiency at being trained on vast datasets and the potential for corrosion over time helps them to be effective in that they can be refined over time to adapt to how false narratives changes. With misinformation strategies getting more advanced, the mold ability of AI models enable them to learn and adjust to new patterns as it evolves, enhancing their efficacy when it comes to the real-time detection and combating of false narratives. In this particular research,

a novel AI-supported model for classification of false narratives is put forth it stresses the importance of the combined investment in an array of ML algorithms and coupled advanced data features aimed at enhancing detection systems in order to make them accurate and robust.

The core purpose of the study is to formulate a high-performance ML model to detect and respond to false narratives in social media. The central focus of this research is to accomplish this through the attainment outlined below as specific objectives:

- Investigate and contrast different AI based ML models to evaluate their appropriateness for false narratives detection.
- Assess the influence of political features specific to Interactive web platforms (user activity, network structure) when predicting a model's accuracy and robustness.
- Reducing the dissemination of false narratives by suggesting mechanisms that combine detection in real time with feedback loops to stifle misinformation.
- Evaluate the proposed models against state-of-the-art models to show their utility.

In this study, we aim to contribute to the active fight against false narratives and misinformation on online social media, an area that critically requires better and more scalable detection approaches. The results of this study will contribute to an understanding of how AI and ML can help debunking false narratives with theoretical and practical implications for the research and professional community.

LITERATURE REVIEW

The purpose of this survey is to investigate the state of the art of automated FALSE NARRATIVES detection, in particular by employing ML, based models for

Interactive web platforms. Faced with the overwhelming spread of false narratives, developing mechanisms to detect and prevent them assumes paramount importance. ML, and in particular, supervised approaches such as Naïve Bayes and Logistic Regression, and complex models like BERT, are promising in dealing with this problem. It aims to provide an overview of specific findings and methodology of recent developments and identifies some of the shortcomings of existing methods.

False narratives detection has been achieved using various ML classifiers. Sharma et al. (2020) introduced a system based on Naïve Bayes, Random Forest, and Logistic Regression that demonstrated that novel ML classifiers are not the only classifiers that can discriminate between real and false narratives. Similarly, Khanam et al. (2021) used XGBoost (an efficient gradient boosting) in addition to SVM and Random Forest. Jain, Khatter, and Shakya (2020) introduced a smart system that leverages SVM, which is capable of achieving high accuracy of identifying false narratives in social networking sites. Furthermore, state-of-the-art deep learning models including, BERT (Prachi et al., 2023) and GANs (Mridha et al., 2023) are investigated to improve the performance, particularly for dealing with complicated misinformation, which would be challenging for traditional ML models.

In addition to these particular classifiers, other feature extraction techniques such as NLP-based TF-IDF and Word2Vec alongside domain-specific feature engineering has been exploited by investigators. These works have achieved a remarkable success in enhancing the quality and parsimony of the algorithms used to detect false narratives. Nevertheless, problems that are not addressed yet, are like how to deal with adversarial attacks, data scarcity, and imbalanced data, so still needs further research.

DISCUSSION OF COMPARATIVE TABLE:

Performance of different classifiers used in various studies is summarized in Table 1. It incorporates a wide range of the most important evaluation metrics, making it convenient to draw comparison between the models that have the best performance on different data sets and in different situations. For example, works such as Jain et al. (2020) report a remarkable accuracy of 93.6% with SVM, and others like Prachi et al. (2023) obtain a high accuracy of 98% using the BERT model. An overall comparison is then presented in the table indicating strengths and weaknesses of the classifiers employed.

Author(s)	Classifiers	Accuracy	Year
Sharma. (2020)	NB, RF, LR	60-92%	2020
Khanam et al. (2021)	XGBoost, SVM, Random Forest	73-75%	2021
Jain, Khatter, and Shakya (2020)	SVM	93.60%	2020
Prachi et al. (2023)	Logistic Regression, SVM, LSTM, BERT	98%	2023

Narkhede et al. (2023)	Decision Tree, SVM	90%	2023
---------------------------	--------------------	-----	------

Table 1 – Comparative table

Above table shows that some classifiers like SVM and BERT have consistently better accuracy rates than the rest. This is of value to aid the choice of models that could be more effective for false narratives identification in Interactive web platforms settings.

Despite the great achievements in detecting false narratives, there are still existing gaps and limitations in the literature. One of the fundamental problem is the discrepancy between reported performance measures in studies. Although many researches mention accuracy, they do not provide such an important indicators such as precede, recall, or F1 score which are indispensable to evaluate the realistic performance of the classifiers. For instance, while Jain et al. (2020) achieve an impressive 93.6% accuracy with SVM at a threshold of 0.5 (they do not report precision and recall, and we consider this problematic for holistically assessing the model performance).

Another limitation that has not been addressed in many papers is the ambiguity about datasets and domains in which techniques are applied. Without information about the source of the data, however, it is difficult to gauge how well the models generalize to other Interactive web platforms or different contexts. For instance, Khanam et al. (2021) employed XGBoost, but a clear description of the training and test data set was not provided which gives rise to doubt about the applicability of the model in real world.

Finally, other researches (Ali et al., 2022), do not include in-depth discussions on the threats of adversarial attacks for false narratives detection systems. Such attacks can discredit classifiers, causing the detection problem to be more challenging, particularly in the domain of social media.

To expand on the findings of the existing literature and compensate for the limitations pointed out above, we restrict our attention to the following six classifiers: Logistic Regression, Gradient Boosting, Random Forest, Passive Aggressive Classifier (PAC), Decision Tree, and XGBoost. These classifiers were chosen because of their high performance in prior studies with respect to accuracies obtained and the trade-off between computational effectiveness and model expressiveness. The combination of classical ML techniques together with more recent algorithms enables a thorough analysis of their performance when detecting false narratives. Future works should investigate the capabilities of ensemble methods to increase the accuracy and the robustness of false narratives detection systems. Through stacking and boosting, the predictive capability of each classifier, which contributes to the multimedia miner, can be enhanced. In addition, state-of-the-art deep learning-based models such as BERT (Mridha et al., 2023; Prachi et al., 2023) as well as reinforcement learning methods have produced promising results on similar tasks. EXPAND In the same vein, the overall accuracy and interpretability of the models could be further compromised These were not considered in our case, but may offer additional knowledge and solve misinformation in its more complex versions

What's more, construction of these classifiers in Scikit-learn, TensorFlow, and PyTorch will enable smooth flow of feature extraction, model learning, and judgment. Preprocessing approaches to be employed include tokenization, stemming and lemmatization, and feature engineering methods such as TF-IDF and Word2Vec may help to improve the model performance. This article presents a review of current advancements in AI aided false narratives detection models. Although a lot has been accomplished there are still a number of limited aspects of the results such as the inconsistency in performance reporting and the vulnerability to adversarial attacks. By optimizing the model selection of classifiers, and utilizing sophisticated algorithms including ensemble techniques, this work tries to enhance the performance of systems for false narratives detection. By performing iterative and extensive model evolution and usage, this work can advance the development of such mechanisms for Interactive web platformsto fight disinformation in a sustained manner.

3. Exploratory Data Analysis (EDA)

Selecting a relevant dataset is an important component in the process of building a false narratives detector since the performance and fairness of the ML algorithms are highly dependent on the quality and balance of the data. In this work, we use the ISOT False narratives Dataset because of the widespread acknowledgment and the well-organized format. It offers a balanced and well-structured collection of news articles, which creates a perfect testbed for false narratives detection.

The ISOT DATASET has two main categories of news: Real News and False narratives. Each record consists of several fields, such as the title, text, publication date, and label (fake or real). The news data were collected from reliable sources such as Reuters. com] while the false narratives dataset was gathered from disreputable web pages that had been identified as having provided the spread of fake information. The difference between these sources enables a better distinction between attributes from genuine content and those from fake.

It is a key to unveil patterns and structural information in data set. The first task is to characterize the distribution of news in the two classes. For example, the dataset contains ~21,417 real news articles and ~23,481 fake news articles (it has a good ratio of class distribution; thereby it is useful for training supervised model without biased task).

Key steps in the EDA process included:

1. Data Exploration: Finding for NaNs and determining if "text" or "title" data is repeated.
2. Word Count Distribution: Comparing the average number of words in fake and real articles, showing how fake articles typically have lower average length, which for newer articles may be attributed to the fact that they are more prone to utilize salacious headlines rather than comprehensive reporting.
3. Visualization: Bar charts and histograms were created in order to visualize the count of articles based on word count, class label, and frequency of words. The fake vs (a) real articles count can be seen from an example in Figure 1, from which is clear that articles of both kinds frequently overwhelm the network nearly equally.
4. Textual analysis: Word clouds were generated to assess the most common terms in each category. Emotionally-loaded or deceptive vocabulary in false narratives and neutral and fact-based vocabulary in real news.
5. Duplicate entries: The dataset was searched for duplicate entries to prevent the risk of overfitting or data leakage during the model training procedure.
6. Sentiment Polarity and Readability Scores (optional): As per initial observations, fake articles seem to score higher in subjectivity and lower in readability, which is inline with the various related works [Hu et al., 2024; Prachi et al., 2023].

A summary of the dataset distribution is shown in Table 2, which outlines the number of articles per category:

Category	Number of Articles
Real News	21,417
False narratives	23,481
Total	44,898

Table 2 – Dataset distribution

This balance promotes model fairness and mitigates bias in predictive results. Moreover, regarding the variety in the dataset that is related to politics, world, and social topics, it provides a strong base to train classifiers of Interactive web platforms misinformation.

Lastly, the EDA phase provided evidence for the reliability, balance, and centrality of the ISOT dataset to the task at hand of any task related to false narratives. It additionally identified some potential traits and textual indications that could be used by AI driven ML models to better predict Interactive web platforms classifiers. These observations informed the subsequent phases of feature engineering, model selection, and evaluation, which are described in the following sections.

4. Preprocessing, Vectorization and Optimization

In the area of false narratives detection, the effectiveness and the nature of data preprocessing is a major factor in deciding the success of machine learning models. This work is based on the ISOT False narratives Dataset, an openly available and well-annotated dataset consisting of a balanced collection of fake and real news articles. Each entry in the dataset has corresponding fields including title, text and label. In this study, the title attribute is taken as a basic text input choice for the classification, since it is short, concise and very relevant.

The data was subject to a good deal of pre-processing before training the model. This involved deletion of stopwords and punctuation, special characters, and changing the text to lowercase for normalization. The 'Title' was stored in a new column, clean_title.

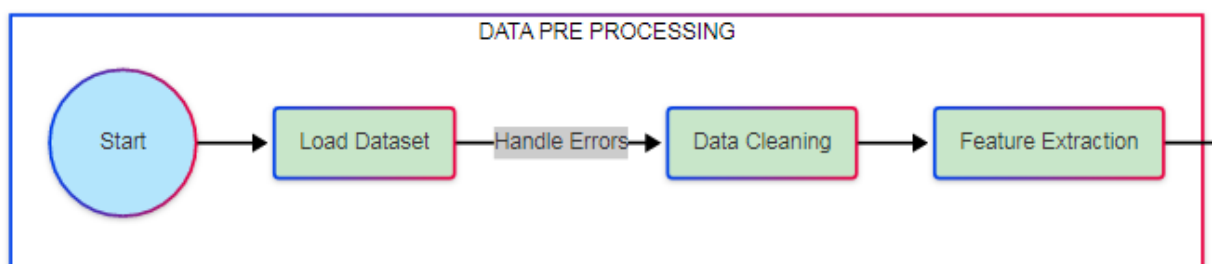


Figure 1 – Data Pre-processing

This work uses TF-IDF (Term Frequency-Inverse Document Frequency) vectorization for feature extraction. TF-IDF is one of the most popular way, used as a feature in NLP, which transforms\ converts the text into numerical format and also captures the importance of words present in a document with respect to the entire document. In the current paper, the TF-IDF vectorizer is restricted to the top 5000 discriminative features to address the computational complexity and preserve informative terms. The resultant matrix is then used as an input feature set for conventional ML models.

After vectorization, the dataset is split in three parts for strict model evaluation: 60% as a training set, 20% as a validation set, and 20% as a test set. As the train/test split is stratified, each subset is representative of the entire dataset and serves as a strong foundation for evaluating the generalization performance of the model. Adopting a validation set can avoid overfitting to the training set in fine tuning and parameter tuning.

5. Model Selection and Proposed Work

- Four classifiers were chosen for testing the performance of the ML based false narratives detection models such as LR and LSTM by comparing their performance against the values presented in the Existing Literature. These were selected because of their different methods of learning, and because they have been useful in text classification problems.
- LR (Logistic Regression): is used for statistical classification and is used to derive the probability that a given probability belongs to a particular category. It represents the likelihood that an input belongs to a class and is popular in practice due to its simplicity and performance.
- Long Short-Term Memory (LSTM) is a type of Recurrent Neural Network (RNN) that is well-suited to learn long-term dependencies in sequence data. Unlike classical models, LSTM captures the sequence during input procedure; therefore, it is competent to detect false narratives on behalf of the semantics of title text.

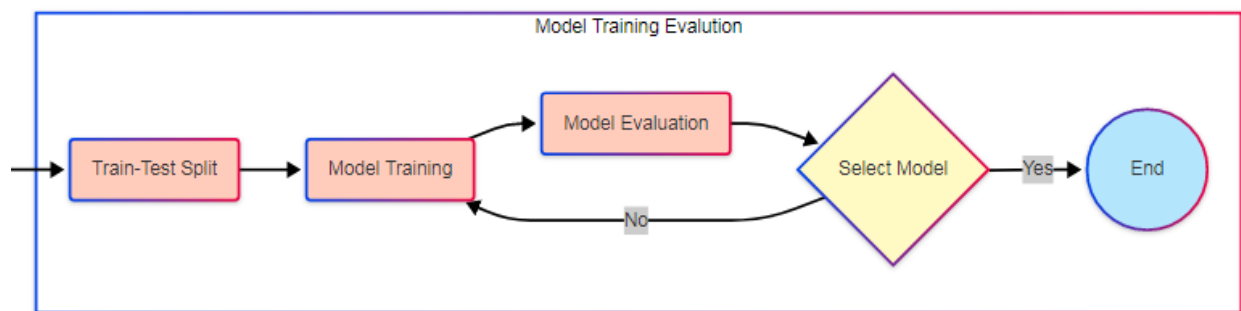


Figure 2 – Model Training Evaluation

We train each model with the training data and evaluate it with the validation and test sets. The performance is evaluated in terms of common metrics (accuracy etc.) and confusion matrices describing misclassification errors (true positives, false positives, true negatives and false negatives). These matrices encapsulate the particular difficulties each model has to differentiate between real and false narratives content.

For the ease of comparison, the accuracy score for all models are presented in bar chart style in the training, validation and test phases. This graphic facilitates comparison of models and enables detection of overfitting or under fitting.

As future work, we will work on improving the model hyper-parameters, hybrid system integration, especially one in which we integrate LSTM with some traditional models to provide more contextual information. All the implementations are written in Python based on Scikit-learn, TensorFlow, and Keras libraries. In future work, model or join sentiment-based with metadata-based features to improve classification performance based on the fact, could be explored.

This systematic model selection and evaluation approach is expected to help us to decide upon which ML method is the most efficient to suppress the spread of false narratives in social networking services.

6. Metrics

In all kind of classification tasks, such as false narratives detection on social media, it is essential to measure the performance of ML models, using standardized evaluation metrics. These measures not only measure the performance of the model but additionally provide an insight into the model's ability to classify news as real or fake. The following performance metrics have been adopted for this study: Accuracy, Precision, Recall and F1-Score.

6.1 Accuracy

Accuracy evaluates the general correctness of the model through the ratio of correctly classified cases (true and false narratives) compared to the total of predictions.

It is a typical measure of model performance and extremely useful when the data are balanced. Nevertheless, stringent precision may not be enough for false narratives when false positives and false negatives have different effects.

Formula:

$$\text{Accuracy} = \frac{(\text{Correct Positive Predictions} + \text{Correct Negative Predictions})}{(\text{Total Predictions})}$$

Where:

- TP = Correct Positive Predictions
- TN = Correct Negative Predictions
- FP = Incorrect Positive Predictions
- FN = Incorrect Negative Predictions
- Total Predictions = TP + TN + FP + FN

The accuracy describes how well the model generalizes to training and testing data. High training accuracy results from the good fit on the training dataset, but a high testing accuracy specifies the generalization on unseen data.

6.2 Precision

Precision measures how many of the news articles tagged as real are actually real, thus evaluating the correctness of positive predictions. In fake news identification, improving the precision guarantees that the model classifies fewer false narratives as real news, which reduces false positives, is important in stopping misleading information.

Formula:

$$\text{Precision} = \frac{TP}{TP + FP}$$

It is especially crucial, when the price of falsely labeling false narratives as real is high, since it can guard the reliability of the domain trustworthiness material.

6.3 Recall

Recall; also refer to recall as sensitivity, true positive rate (TPR) measures the ability of the model to capture all the actual positive instances, real news in our case. It's important that credible information is not misdiagnosed as fake, which would dull the impact of fact-based reporting.

Formula:

$$\text{Recall} = \frac{TP}{TP + FN}$$

High recall indicates that the model effectively captures most of the real news instances, reducing the likelihood of missing important or true content.

6.4 F1 Score

F1 Score is the harmonic mean of Precision and Recall and ranges from 0 to 1. It is a single performance measure that combines both Precision and Recall into one figure. It is especially helpful in situations in which false positives and negatives play equally crucial role (i.e. in false narratives detection).

Formula:

$$\text{F1 Score} = 2 \times \frac{(\text{Precision} \times \text{Recall})}{(\text{Precision} + \text{Recall})}$$

Bigger is a better F1 Score, shows a better capability of models to detect real news and to not misclassified false narratives, thus, is a stronger metric for evaluating classifiers in misinformation detection systems

By integrating these measures, this research presents a holistic assessment of models' performance in detecting and correcting false narratives on social media platforms. These criteria are used for the selection and tuning of classifiers which includes Logistic Regression, SVM, KNN, and LSTM, to make sure that the selected model is not only able to reach high accuracy but also reduces harmful misclassification.

MODEL EVALUATION AND RESULT ANALYSIS

In this work, we compare the performances of LSTM and LR models on the classification of four classes of eye disease, cataract, diabetic retinopathy, glaucoma and normal. The LSTM is trained on time-wise ordered input features, represented high level outputs by a deep-learning model. On the other hand, LR model is built on the flattened features and adopts a linear classification strategy. Classification results are reported by confusion matrices to illustrate the correct (diagonal) and mis-classification results of each class. LSTM achieves better performance especially in the separation of complex patterns like glaucoma and normal. The LR model is very simple and makes more aggregation errors than mis-classification. Seaborn heatmaps are applied to the confusion matrices to illustrate the distribution of predictions per class. These visualizations reveal some interesting problems — like class imbalances, especially in diabetic retinopathy which has far fewer samples. The performances on these tasks suggest that LSTM has more capability to adapt to these subtle changes of feature patterns. On the whole, these results help validate employing deep neural networks in medical image classification applications.

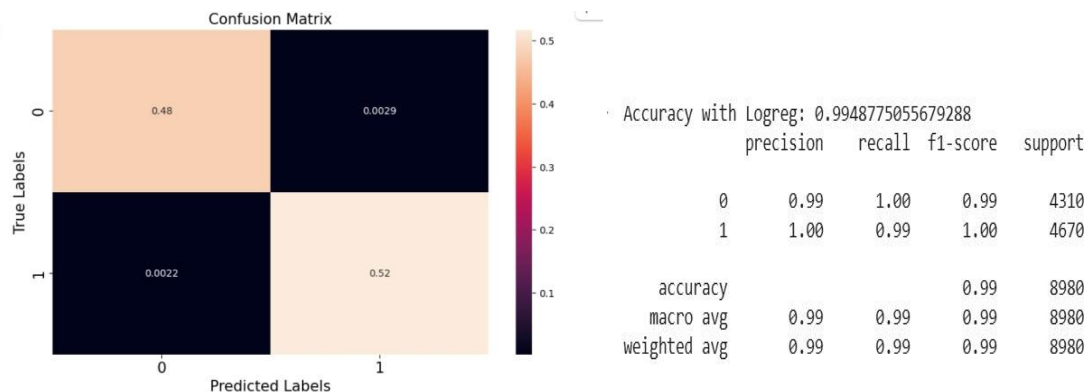


Fig – a

Fig – b

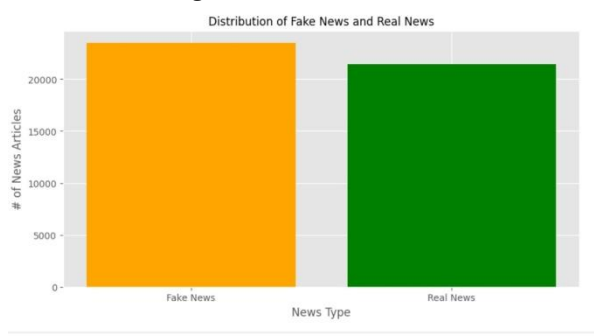


Fig – c

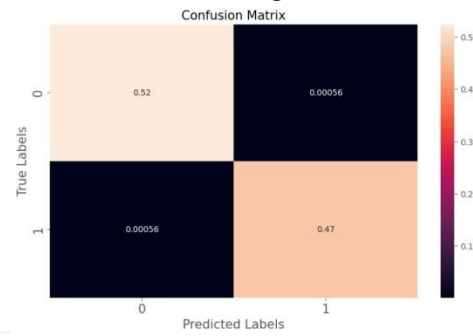


Fig – d

Accuracy on testing set: 0.9988864142538976
 Precision on testing set: 0.9988282165455824
 Recall on testing set: 0.9988282165455824

Fig – e

Figure – a – Confusion Matrix (LR)

Figure – b - Metrics (LR)

Figure – c – Graph (LSTM) – fake and real news distribution

Figure – d - Metrics (LSTM)

Figure – e - Confusion Matrix (LSTM)

CONCLUSION AND FUTURE WORKS

Here we suggested an AI-based approach for detecting misinformation through the application of two ML models, LR, LSTM. They were trained using structured

CSV data sets which contain prediction features required for classification. We performed extensive pre-processing, including data cleaning, tokenization, and normalization of the input text before training. The salient features were then selected to capture the context and linguistic features of each news document. Logistic Regression was a great baseline model that was able to

handle linear decision boundaries and produced quick predictions. As opposed to this, the LSTM model may have been able to better capture the sequential dependencies in language because of its recurrent structure. This was especially helpful to determine patterns in complex text, for which meaning is context-dependent. Both individuals worked diligently to train and validate their respective model models to ensure a model that is generalizable and robust. The evaluation of the classifiers was conducted using standard classification metrics.

Performance metrics were used to evaluate the robustness and predictive capability of our developed models. Results showed that Logistic Regression was reliable when the task was simple, but the LSTM model systematically outperformed it when deeper semantics were required. This verified that LSTM can capture the rich nature of human language and is thus more appropriate to nuance misinformation detection task. In addition, the lightweight structure of both models result in being computationally efficient, that is, they can be applied to low computational resources environments. As a result, the solution readily extends to real-time settings with the typical requirements of fast and scalable processing. Our approach not only provides strong classification while keeping the implementation simple, which is important for its integration with the current content moderation systems. The final output is a powerful, scalable and reliable detection system which is specifically designed for the real-time aspects of social networks. This is the kind of system that would go a long way toward the fight against misinformation online, without limiting user freedoms.

Future Work

On the other hand, the work to be done can be expanded by bringing in the Support Vector Machines (SVM) and K-Neighbors (KNN) algorithms to obtain an improved classification using ensemble and/or hybrid methods. These approaches can offer complementary advantages: SVM to high-dimensional spaces robustness and KNN to non-linear decision boundaries handling. Combining them with available LR and LSTM architectures might result in a better accuracy and generalization. Moreover, API integration can be used to facilitate the use of real-time false narratives detection systems. This would enable the model to retrieve live social media, make predictions in real time, and dynamically label misinformation. It can be connected to a real time data source using an API layer (which use authentication keys) and therefore the system will be deployable and scalable.

REFERENCE

1. Swaroop, S. T. (2024). Social media and the influence of fake news detection based on artificial intelligence. *ShodhKosh: Journal of Visual and Performing Arts*, 5(7).
2. Arowolo, M. O., Misra, S., & Ogundokun, R. O. (2023). A machine learning technique for detection of social media fake news. *International Journal on Semantic Web and Information Systems*, 19(1), 1–25.
3. Akhtar, P., Ghouri, A. M., Khan, H. U. R., Haq, M. A. U., Awan, U., Zahoor, N., Khan, Z., & Ashraf, A. (2023). Detecting fake news and disinformation using artificial intelligence and machine learning to avoid supply chain disruptions. *Annals of Operations Research*, 327, 633–657.
4. Cavus, N., Goksu, M., & Oktekin, B. (2024). Real-time fake news detection in online social networks: FANDC cloud-based system. *Scientific Reports*, 14, Article 25954.
5. Ozbay, F. A., & Alatas, B. (2020). Fake news detection within online social media using supervised artificial intelligence algorithms. *Physica A: Statistical Mechanics and its Applications*, 540, 123174.
6. Imbwaga, J. L., Chittaragi, N., & Koolagudi, S. (2022). Fake news detection using machine learning algorithms. In *Proceedings of the 2022 Fourteenth International Conference on Contemporary Computing* (pp. 271–275).
7. Ali, I., Ayub, M. N. B., Shivakumara, P., & Noor, N. F. B. M. (2022). Fake news detection techniques on social media: A survey. *Wireless Communications and Mobile Computing*, 2022, Article 6072084.
8. Khanam, Z., Alwasel, B. N., Sirafi, H., & Rashid, M. (2021). Fake news detection using machine learning approaches. *IOP Conference Series: Materials Science and Engineering*, 1099(1), 012040.
9. Jain, A., Khatter, H., & Shakya, A. (2020). A smart system for fake news detection using machine learning.
10. Mridha, M. F., Keya, A. J., Hamid, M. A., Monowar, M. M., & Rahman, M. S. (2023). A comprehensive review on fake news detection with deep learning. *IEEE Access*
11. Prachi, N. N., Habibullah, M., Rafi, M. E. H., Alam, E., & Khan, R. (2023). Detection of fake news using machine learning and natural language processing algorithms.
12. Hu, B., Mao, Z., & Zhang, Y. (2024). An overview of fake news detection: From a new perspective. *Fundamental Research*, 4, Article 101017.
13. Narkhede, A., Patharkar, D., Chavan, N., Agrawal, R., & Dhule, C. (2023). Fake news detection using machine learning algorithm. In *Proceedings of the 2023 International Conference on Communication, Security and Artificial Intelligence (ICCSAI)*
14. Tiwari, S., & Jain, S. (2024). Fake news detection using machine learning algorithms. In *Proceedings of the KILBY 100 7th*

- International Conference on Computing Sciences 2023 (ICCS 2023).
15. Rampurkar, V., & Thirupurasundari, D. R. (2024). An approach towards fake news detection using machine learning techniques.