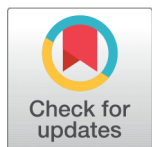


RESEARCH ARTICLE

 OPEN ACCESS

Received: 28-02-2024

Accepted: 18-06-2024

Published: 06-07-2024

Citation: Alharthi W, Alzahrani ME (2024) Toward an Efficient Hybrid Clustering and Classification Approach for Fake News Detection in Social Network. Indian Journal of Science and Technology 17(26): 2754-2762. <https://doi.org/10.17485/IJST/V17I26.579>

* Corresponding author.

wafaalharthi10@gmail.com

Funding: None

Competing Interests: None

Copyright: © 2024 Alharthi & Alzahrani. This is an open access article distributed under the terms of the [Creative Commons Attribution License](#), which permits unrestricted use, distribution, and reproduction in any medium, provided the original author and source are credited.

Published By Indian Society for Education and Environment ([iSee](#))

ISSN

Print: 0974-6846

Electronic: 0974-5645

Toward an Efficient Hybrid Clustering and Classification Approach for Fake News Detection in Social Network

Wafa Alharthi^{1*}, Mohammad Eid Alzahrani¹

¹ Department of Computer Science, Faculty of Computing & Information, Al-Baha University, Al-Baha, Saudi Arabia

Abstract

Objectives: The spread of fake news on social media has become a pressing issue in recent times. Despite various organizations' efforts to address this problem, it continues to persist, necessitating finding more effective solutions. This study implements a machine learning-based approach for identifying fake news on social media with improved accuracy. **Methods:** The study's methodology utilizes a combination of Convolutional Neural Network (CNN) and Long Short-Term Memory (LSTM) to extract semantic features from news articles. Then, a hybrid clustering and classification approach is used, which combines K-means clustering and Artificial Neural Network (ANN) classifiers. Data used is a secondary data consisting of 23481 world news articles obtained from the Reuters.com website and 21417 unreliable news from polifact.com. During the training process, the number of units and layers on the model was tuned to optimize model performance. The model was compared to other baseline models such as KNN, SVM, Decision tree, and Boosted decision tree to establish the best performing model. **Findings:** The results from the two algorithms were weighted to final classifications using Bayesian probability theory. The proposed approach achieved an accuracy of 99.78%, a sensitivity of 100%, and specificity equal to 99.73%. The model's precision is 99.74%, indicating its ability to identify fake news. The F-score of the approach is 99.87%, indicating that the model strikes a good balance between correctly classifying fake news articles and reliable news articles. The approach outperformed other machine learning classifiers, including KNN, SVM, Decision Tree, and Boosted Decision Tree. **Novelty :** The study applies a hybrid approach with a classification and clustering algorithm to improve detection of fake news on social media, the approach is tested with varied real-world datasets to establish its robustness under different vocabularies and vocabulary sizes.

Keywords: Artificial Intelligence; Machine Learning; Deep Learning; Hybrid clustering approach; Classification; CNN; LSTM; Boosted Decision Tree; Social platforms; K-means

1 Introduction

In modern society, technological innovations have changed the way things are done at a very high speed, with social media being a significant transformational force. News content is a fundamental aspect of the global population's day-to-day interactions at personal and societal levels. The proliferation of fake news has attracted a lot of research around the topic, with the majority of literature seeking to come up with an advanced approach which detects fake news quickly. Such approaches could save news users time and other related resources in finding credible content on social media⁽¹⁾.

Moreover, Fake news is a multidisciplinary issue that requires joint professional efforts and a comprehensive framework. Additionally, there exist multiple strategies and theories targeting different fake news perspectives, such as propagation methods, textual style, source, and existing knowledge. All of these need to be pieced together into a framework. This highlights the high demand for research literature in this field to contribute to curbing the spread of fake news⁽²⁾.

With the advancements in Artificial Intelligence (AI) technology and innovative applications, there is potential to develop classification approaches that can effectively distinguish fake news content from reliable news sources. This avenue has been increasingly explored by researchers in recent years. For instance, in 2020, Facebook launched SimSearchNet++ AI, which it uses to match images with their manipulated versions, making it easier to detect practices such as image blur and cropping. They consider it important for visuals-first platforms such as Instagram. During the March 2020 US elections, the AI was able to label 180 million pieces of content on Facebook, which were later debunked by third-party fact-checkers. In the following month, the AI labeled 50 million pieces of content related to COVID-19⁽³⁾. The literature expressed difficulties with hate speech, where two pieces of misinformation can be expressed differently by rephrasing one, using a different image, or even changing the content from graphic to text. This highlights the potential for evading detection by simply changing vocabulary. Additionally, with the rapid change in events and vocabulary, AI can be highly challenged. According to Facebook, there is still a long way to go; actually, the accuracy of the AI was 80.5% in 2019 and 94.7% in 2020, leaving room for improvement⁽⁴⁾.

Twitter, on the other hand, acquired Fabula AI in 2019⁽⁵⁾ to detect misinformation. However, as late as 2021, an exploratory study on misinformation by⁽⁶⁾ still shows a significant level of misinformation on Twitter.

Ozbay et al.⁽⁷⁾ proposed a two-step fake news detection system. In the first step, they used NLP to preprocess data, followed by the application of 23 machine learning classifiers in the second step. Text mining techniques were utilized to preprocess news articles before converting them into vector representation using term frequency (TF). These vectorized texts served as features during the modeling process. Their literature review led them to the conclusion that linguistic cues, particularly satire, are highly effective in detecting deceptive information. However, the classifiers generally performed poorly. The best classifier achieved an accuracy of 65.5% and an F-score of 67.5%, indicating a poor balance between precision and recall. Furthermore, the best classifier achieved 74.7%, suggesting that it achieves high precision by predicting most incoming articles as fake.

Thepade et al. developed a hybrid fake news detection system that combines a Deep Structured Semantic Model (DSSM) and an Improved Recurrent Neural Network (RNN). The DSSM is specialized in feature extraction, while the RNN is specialized in learning long-range temporal features. This system outperformed traditional ML approaches with an accuracy of 99%. With this accuracy, the model leaves room for improvement in terms of accuracy.

Braşoveanu et al.⁽⁸⁾ developed a hybrid model that combines machine learning, semantics, and NLP. The semantic approach involved collecting metadata from third parties, including sentiments and named entities. The NLP approach involved extracting relations from a Knowledge Graph (KG) and training word embeddings, while the machine learning part included training classic machine and deep learning models. They concluded that their approach yields better scores on most classifiers and recommend extending this approach to fake news reviews. However, the model also achieved low performance. Out of all the experiments, the best result was on Liar data, where the overall accuracy was 54.9%.

Zhang et al.⁽⁹⁾ combined linguistic and grammar-based features for classification. They utilized gated recurrent units (GRU) to extract latent features and a gated diffusive unit (GDU) to combine latent characteristics regarding news articles, their topics, and their creators. They attained a fake news detection accuracy of 86.40%. The analysis was conducted by utilizing the followers' intersection and retweet graph, as well as classification approaches, to establish whether the information shared by a particular source was reliable or not. Additionally, other models such as machine learning and deep learning were employed in developing news detection models.

Nasir et al. used a hybrid approach combining CNN and RNN neural networks and achieved a higher accuracy rate than other classic machine learning methods, with 60% accuracy on a fake news dataset and 99% on an ISOT dataset. They suggest further research using more complex ANN layers. With the observed accuracy, there is room for improvement.

These reviewed literature reveals the intense use of NLP to create a representation that organizes unstructured natural language, either by identifying grammatical relationships among the constituents of the text (syntactic structure) or by understanding the meaning conveyed by the text (semantic structure). LSTM and Word Embedding models are among the common models in NLP. Application of classification techniques to NLP-based features extracted by such approaches showed

success in most of the literature from the hybrid approaches it is seen that combining different feature extraction approaches showed success in A research gap exists in the performance of the reviewed approaches additionally all the hybrid approaches focused in feature extraction with no effort to implement a hybrid approach on the classification phase.

This study seeks to contribute to the fight against the spread of fake news by introducing a fake news detection system that can detect fake news articles with high precision. The system is a hybrid one, combining classification and clustering approaches. At the base of the system are CNN and LSTM layers that learn features from news articles and pass the learned features to two paths. The first path has ANN layers that categorize news articles as fake or reliable. The second path has a K-means cluster model that clusters news articles into fake and reliable latent clusters. At the end of the system is a Naive Bayes classifier that merges the results from the two paths using Bayes theory. Additionally, the study adds new literature on how to implement hybrid fake news detectors and proposes future research areas that require to be explored in order to improve existing fake news solutions.

2 Methodology

The study makes use of secondary data with 23481 world news articles obtained from the Reuters.com website, and 21417 unreliable news articles obtained from polifact.com, during the years 2016 and 2017⁽¹⁰⁾. All articles had at least 200 characters.

2.1 Data cleaning and Pre-processing

The collected articles were pre-processed using NLP. Firstly, Unicode symbols and punctuations used in place of ASCII symbols and punctuations were translated back to ASCII using a lexical tools lookup table⁽¹¹⁾. A tokenization layer was attached to the data preprocessing pipeline to convert the articles into word and sentence tokens. A pretrained Global Vectors for Word Representation (GloVe) was used to map articles into a vector representation, with a vocabulary size equal to 400,000 words, each represented by a vector of length 100⁽¹²⁾. The chosen maximum token is 20000 words. This means each word token is represented by a vector with 100 values representing the relationship between the word token and other word tokens. It is from these vectors that the model learns the semantic patterns of both categories. This means that during model training the features are the semantic relationship between words. The maximum token means only the top 20,000 most common tokens will be considered, meaning rare tokens are considered negligible. The final step was to categorize 80% of the data into training sets, 15% into testing sets, and 5% into validation sets. Figure 1 is a representation of the data preparation and modeling part while Table 1 summarizes the details of the GloVe representation.

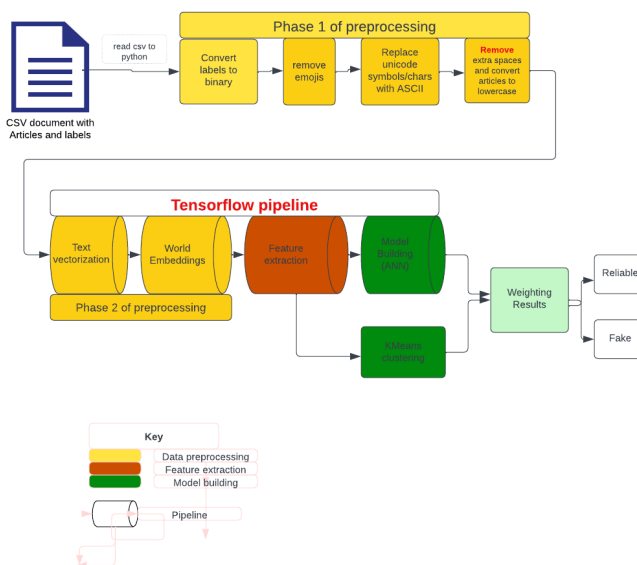


Fig 1. Flow chart representation of the data pre-processing phase

Table 1. Performance of the proposed approach across benchmark datasets

	Detail
Embedding Type	Glove
Source	https://nlp.stanford.edu/projects/glove/#discuss
Vocabulary Size	400000
Max Tokens	20000
Hits	18830
Misses	1170

2.2 Proposed approach

The study employed a hybrid system combining classification and clustering approaches. At the base of the system are CNN and LSTM layers that learn features from news articles and pass the learned features to two paths. The first path has ANN layers that categorize news articles as fake or reliable. The second path has a K-means cluster model that clusters news articles into fake and reliable latent clusters. At the end of the system is a Naive Bayes classifier that merges the results from the two paths using Bayes theory. Natural language Processing (NLP) is the component of the proposed system dedicated to preprocessing and extracting features from the articles, research works of literature about NLP were reviewed to acquire background information on NLP, evaluate its effectiveness on those systems, and assess its integrality with other approaches. The proposed approach yields two sets of prediction results: one generated from the ANN model and the other from the K-means clustering approach. To consolidate these predictions into a unified output, a Naïve Bayes classifier is employed.

The workings of CNN and LSTM are often considered a black box, making it difficult to provide feedback on why articles are classified as fake when using binary classification. For the sake of transparency, a Local Interpretable Model-agnostic Explanations (LIME) technique is proposed to help annotate the word relationships that lead to the article classification. For each article predicted, LIME makes a prediction and then generates samples by randomly altering the original text. Predictions on these samples and the similarity of the samples to the original text are used to train an interpretable model, which approximates the behavior of the original classifier in the local neighborhood of the original text instance⁽¹³⁾.

2.3 Benchmark datasets

This study utilized four benchmark datasets to establish the performance of the approach under different real-world stressors. The first is “Coronavirus tweets NLP data” which consists of about 36,447 tweets labeled as extremely positive, positive, extremely negative, negative, and neutral sentiment scores⁽¹⁴⁾. The second was the “Fake News dataset,” which is a small sample of 6,335 news articles. The motivation behind including the data in this study is to understand the ability of the proposed algorithm to learn from small samples and balanced samples⁽¹⁵⁾. The “Getting Real about Fake News Dataset” with 12,999 news articles retrieved from 244 websites using webhose.io⁽¹⁶⁾. The “Liar dataset” consists of 12,791 manually labeled articles from politifact.com⁽¹⁷⁾. The labels are a likert scale with levels: half-true, False, mostly-true, barely-true, true, and pants-fire.

2.4 Baseline models

The proposed approach was compared to currently existing classifiers, including decision trees, the K-nearest neighbor algorithm, Support Vector Machines, and boosted decision trees, in terms of efficiency and classification accuracy. These were implemented on the features extracted from the main data.

3 Results and Discussion

3.1 Proposed approach

The ANN classifier on the implemented system generally performed well. It correctly classified 99.78% of the articles in the testing set. Moreover, the model predicted 99.80% of the fake news articles and 99.75% of reliable news articles on the testing set correctly. The model had a high precision (99.77%) which means it rarely predicts fake news articles as reliable. The F1 –score for the model is 99.78% indicating a good balance between precision and recall. The model training history shows that the model was able to learn rapidly achieving minimal loss in the first few epochs. The model loss on the training and testing set remains close indicating no over fitting see Figure 2. The K- means path clustered the articles into two latent clusters 0 and 1. This being an unsupervised approach, these clusters are not analogous to the 0/1 labels of the ground truth labels. Each class

was first matched to the corresponding cluster. The model correctly clustered 99.76% of the articles on the testing set. It also clustered 99.74% of the fake news articles correctly and 99.78% of the reliable news articles. The precision of the model was 99.80% while the F1-score was 99.80% indicating a good balance between precision and recall. Cluster separation was good for the model as can be seen in Figure 3.

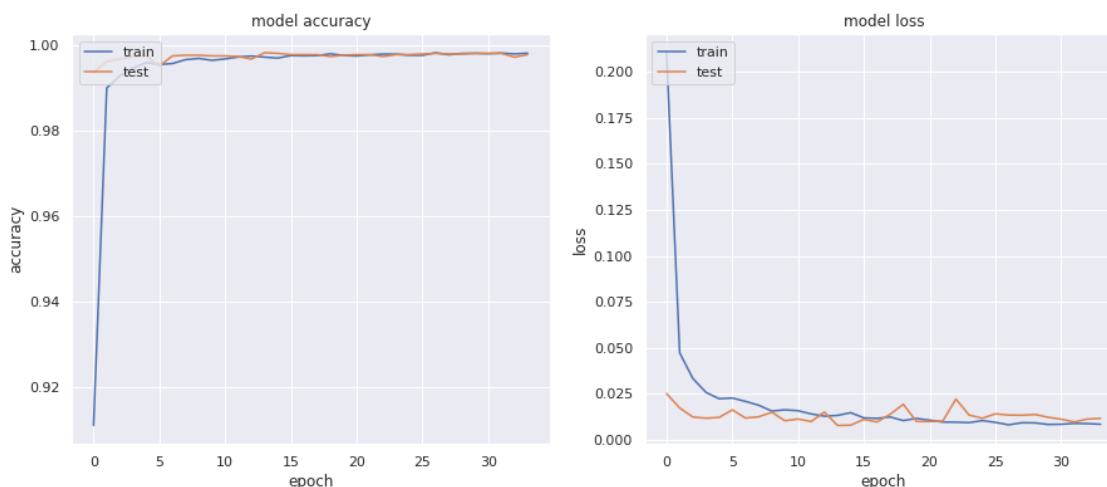


Fig 2. Model training accuracy and loss. It shows that the model was able to learn rapidly achieving minimal loss in the first few epochs. The model loss on the training and testing set remains close indicating no over fitting

A naïve Bayes classifier was trained on the two outputs to output a weighted prediction. It achieved 99.87% accuracy on the validation set and was able to identify 100% of the fake news articles on the validation set and 99.73% of reliable news articles on the validation set. The precision and F1 scores were 99.74% and 99.87% respectively.

The feature extraction part in this study had a high number of parameters, particularly on the word embedding layer where the number of units was set high to capture many word features in the articles. As a result, the models scored high in terms of accuracy with minimal overfitting. This is in tandem with the findings of Sekeroglu et al. (17) findings that with a large number of parameters that represent real-world scenarios, AI-based solutions provide more accurate results. The findings are also in agreement with (18) argument that NLP can be integrated with data processing methods for better classification results. In this study, NLP was properly interpreted with other TensorFlow pipeline preprocessing steps, resulting in a refined dataset.

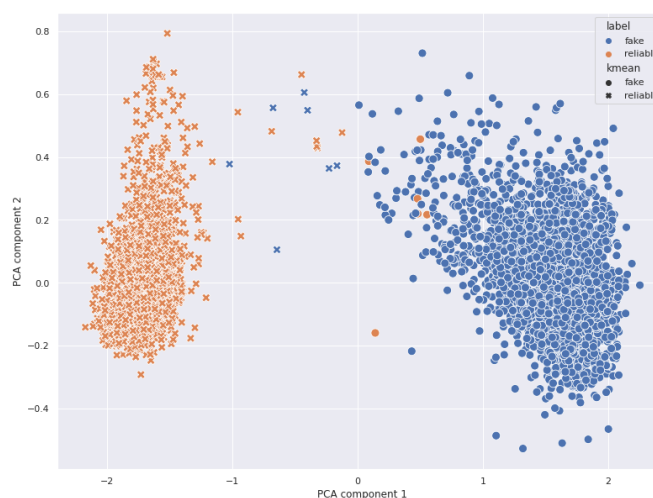


Fig 3. PCA-based visualization of the separation between the fake news cluster and the reliable news cluster

Figure 4 is a sample article predicted as Real. LIME explanations assign a score to word tokens, which are annotated on the article. The strength of the color coding indicates how each word or sentence token contributes to the classification. Blue color indicates words contributing to the probability of the article being Real, while Orange shows the words contributing to the probability of the article being Fake.

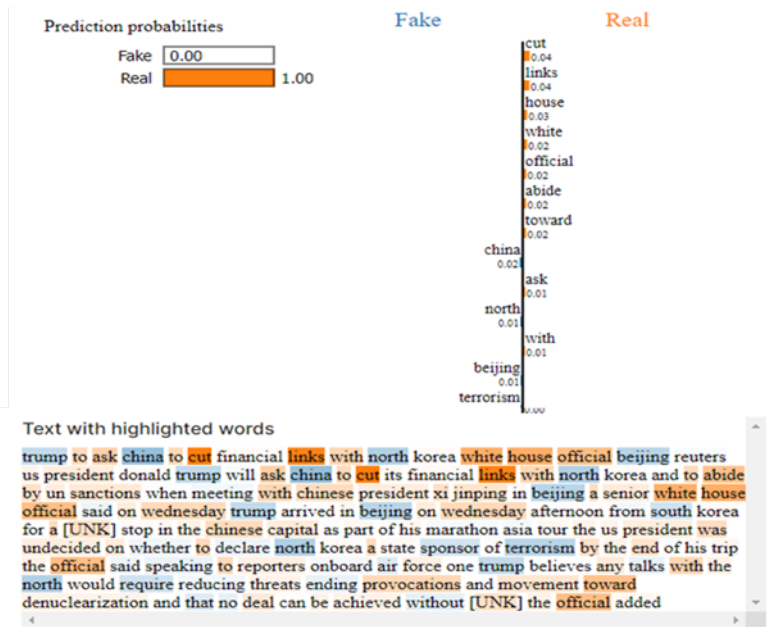


Fig 4. Visualization of Lime explanations for a sample nes article that was predicted as fake, blue color denotes words contributing to the probability of the article being fake while orange denotes words contributing to the probability of the article being Real

3.2 Benchmark datasets

By exposing the proposed approach to the “Fake news dataset”, which is a small dataset with fairly balanced outcome categories, the embedding layer found 18964 (95%) of the top 20000 tokens on the pre-trained GloVe embeddings. For the “Liar dataset” which had poorly balanced outcome categories 5507 words could not be found on the pre-trained embeddings, this accounts for 47.30% of the considered tokens which is slightly higher compared to the other datasets, this is indicative of a slightly different semantic structure. On the “Coronavirus tweets NLP - Text Classification” 4340 (27.71%) tokens were not discovered from the pre-trained embeddings, the dataset is characterized by fairly balanced outcome categories. For the “Getting Real about Fake News” which featured poorly balanced outcomes about 1484 (7%) words were not found on the pre-trained word tokens. Table 2 reports the performance of the proposed approach on these datasets.

The results of this benchmark are inconsistent with the conclusion of van de Schoot et al. ⁽¹⁹⁾ that machine learning models are efficient in extracting data for eligibility purposes and synthesis of results, such that they can compensate for data deficiencies and produce high accuracy with deficient data. Exposing the proposed system to different datasets showed that imbalanced datasets may lead to imbalanced sensitivity and specificities. The larger class also tends to influence the overall accuracy at a greater scale.

Table 2. Performance of the proposed approach across benchmark datasets

Model	Metric	Getting about News	Real Fake	Fake Dataset	news	Original Dataset	Liar Dataset	Corona virus tweets NLP
ANN	Accuracy	0.9041		0.8968		0.9978	0.7309	0.8173
	Sensitivity	0.3913		0.8415		0.9980	0.0000	0.8194
	Specificity	0.9727		0.9563		0.9975	1.0000	0.8151
	Precision	0.6569		0.9539		0.9977	0.0000	0.8244

Continued on next page

Table 2 continued

K MEANS	F-score	0.9041	0.8968	0.9978	0.7309	0.8173
	Accuracy	0.5646	0.9063	0.9976	0.5967	0.8023
	Sensitivity	0.9174	0.8740	0.9974	0.6569	0.8870
	Specificity	0.5174	0.9410	0.9978	0.5745	0.7125
	Precision	0.2027	0.9409	0.9980	0.3625	0.7658
	F-score	0.3320	0.9062	0.9977	0.4671	0.8219
Weighted	Accuracy	0.9041	0.9084	0.9978	0.7309	0.8173
	Sensitivity	0.3913	0.8780	0.9980	0.0000	0.8193
	Specificity	0.9727	0.9410	0.9975	1.0000	0.8151
	Precision	0.6569	0.9412	0.9978	0.0000	0.8244
	F-score	0.4905	0.9085	0.9979	0.0000	0.8219
	Accuracy	0.9077	0.5016	0.9987	0.7048	0.5415
Validation set	Sensitivity	0.3889	0.9936	1.0000	0.0000	0.1235
	Specificity	0.9723	0.0248	0.9973	1.0000	0.9780
	Precision	0.6364	0.4968	0.9974	0.0000	0.8542
	F-score	0.4828	0.6624	0.9987	0.0000	0.2158

¹ The ANN, K-Means, and weighted results are based on the testing set.

² The validation set results are the weighted results of ANN and K Means on the validation set.

3.3 Baseline models

The proposed approach was compared to four other ML methods; the four were fit on the same data, by taping the extracted features from the CNN and LSTM layer. Considering the fact that the naïve Bayes weighting was done on the testing set. The comparison between baseline models was done on the validations set, not to give the proposed approach a head start because it has seen the testing set, the model parameters were first tuned using grid search to ensure optimal performance. For accuracy, the proposed approach achieved the highest (99.74%) followed by the support vector classifier with 90.77%. For specificity and precision, the proposed approach was the best with 100% and 99.73%. The proposed approach was also the best in terms of sensitivity and F score. Table 3

Table 3. Performance of the proposed approach across benchmark datasets

Model	Metric	ANN K-means	and KNN	SVM	Decision tree	Boosted decision tree
Testing set	Accuracy	0.9978	0.9026	0.9062	0.9026	0.9062
	Sensitivity	0.9980	0.4652	0.4522	0.4783	0.4739
	Specificity	0.9975	0.9610	0.9669	0.9593	0.9640
	Precision	0.9978	0.6149	0.6460	0.6111	0.6374
	F-score	0.9979	0.5297	0.5320	0.5366	0.5436
	Accuracy	0.9987	0.9015	0.9077	0.9015	0.9031
Validation set	Sensitivity	1.0000	0.4444	0.4583	0.4583	0.4861
	Specificity	0.9973	0.9585	0.9637	0.9567	0.9550
	Precision	0.9974	0.5714	0.6111	0.5690	0.5738
	F-score	0.9987	0.5000	0.5238	0.5077	0.5263

3.4 Comparative analysis

The study posted higher performance results compared to the reviewed literature. With an accuracy equal to 99.78%, precision of 99.77%, and recall of 99.80%, the model outperforms Facebook's SimSearchNet++, which reported an accuracy of 94.7% by 2020. Similarly, the model surpasses all 23 algorithms in Ozbay et al. ⁽⁷⁾ With the highest achieving 65.5% accuracy and 74.7% precision. The model also outperforms Thepade et al. Deep Structured Semantic Model (DSSM) and an Improved Recurrent Neural Network (RNN) approach, which achieved 99% accuracy. For Braşoveanu et al. ⁽⁸⁾, the accuracy of 54.9% achieved with the Knowledge Graph feature extraction approach is much lower than the results of this study. Additionally, the study achieved better results than Zhang et al. ⁽⁹⁾, which had 86.40% accuracy using gated recurrent units (GRU) and gated diffusive

units (GDU) to extract and combine latent features regarding news articles, topics, and creators. Furthermore, the approach outperforms Nasir et al., which achieved 99% accuracy using a hybrid approach combining CNN and RNN layers.

4 Conclusion

The study has successfully developed an AI-integrated approach for evaluating social media news articles, which can be used by consumers and social media platforms to identify fake news. The approach combines deep learning classification using artificial neural networks and K-means clustering. The study has also identified typical characteristics of fake and credible news, including slightly lower spam rates and common vocabulary in fake news articles, which can serve as early indicators of false information. The study has benchmarked the approach on real-life datasets of varying complexities and found that large, well-balanced datasets result in a good balance between sensitivity and specificity and faster model convergence. The study has also compared different classifiers and found that the projected methodology does well in terms of accuracy, specificity, and precision, with boosted decision trees achieving a good balance between specificity and sensitivity.

A major limitation of this study is the correct definition of “fake news”. Molina et al.⁽²⁰⁾ Believe that the true definition of fake news has become highly debatable and has even been used to undermine reliable news. Many concepts with overlapping meanings such as “Misinformation”, “disinformation”, “fake news”, “post-truth”, “bullshit”, “information pollution”, or “information disorder” exist to define different elements and phases of fake news⁽²¹⁾. Additionally, vocabulary use like the use of satire can lead to news content being classified as fake⁽²²⁾. This means that some articles flagged as fake news propagators might actually be credible sources just expressed in differently. Furthermore, some domains not blacklisted might also promote fake news expressed with different vocabulary choices. The secondary data used in this study considers news articles crawled from Reuters as reliable and also relies on Polifact.com fact check for fake news. With the human involvement in fact checking process for both Reuters and Polifact the human bias pointed out by Molina et al.⁽²⁰⁾ is eminent considering the complex elements of fake news enshrined in different vocabularies. These can cause wrong labeling and consequently having models learn wrong patterns. Such models will lead to wrong predictions which could affect the validity of the conclusions drawn from this study.

References

- 1) Collins B, Hoang DT, Nguyen NT, Hwang D, Chen BC, Wu Z, et al. Efficient Object Embedding for Spliced Image Retrieval. In: 2021 IEEE/CVF Conference on Computer Vision and Pattern Recognition (CVPR).. 2021;p. 14960–14970. Available from: <http://dx.doi.org/10.1109/cvpr46437.2021.01472>.
- 2) Lee T. The global rise of “fake news” and the threat to democratic elections in the USA. *Public Administration and Policy*. 2019;22(1):15–24. Available from: [http://dx.doi.org/10.1108/pap-04-2019{\char"2013relax}0008](http://dx.doi.org/10.1108/pap-04-2019{\char).
- 3) Metaverse. Here's how we're using AI to help detect misinformation [Internet]. . 2020. Available from: <https://ai.facebook.com/blog/heres-how-were-using-ai-to-help-detect-misinformation/>.
- 4) Wiggers K. Facebook's improved AI isn't preventing harmful content from spreading. 2020. Available from: <https://venturebeat.com/ai/facebook-s-improved-ai-isnt-preventing-harmful-content-from-spreading/>.
- 5) Agrawal P. Twitter acquires Fabula AI to strengthen its machine learning expertise.. 2019. Available from: https://blog.twitter.com/en_us/topics/company/2019/Twitter-acquires-Fabula-AI.
- 6) Shahi GK, Dirkson A, Majchrzak TA. An exploratory study of COVID-19 misinformation on Twitter. *Online social networks and media*. 2021;22:100104. Available from: <https://doi.org/10.1016/j.osnem.2020.100104>.
- 7) Ozbay FA, Alatas B. Fake news detection within online social media using supervised artificial intelligence algorithms. *Physica A: Statistical Mechanics and its Applications*. 2020;540. Available from: <http://dx.doi.org/10.1016/j.physa.2019.123174>.
- 8) Braşoveanu AM, Andonie R. Semantic Fake News Detection: A Machine Learning Perspective. *Advances in Computational Intelligence*. 2019;p. 656–667. Available from: http://dx.doi.org/10.1007/978-3-030-20521-8_54.
- 9) Zhang J, Dong B, Yu PS. FakeDetector: Effective Fake News Detection with Deep Diffusive Neural Network. In: 2020 IEEE 36th International Conference on Data Engineering (ICDE). 2020. Available from: <http://dx.doi.org/10.1109/icde48307.2020.00180>.
- 10) Siddik AB. Fake and true news dataset. 2020. Available from: https://figshare.com/articles/dataset/Fake_and_True_News_Dataset/13325198.
- 11) Lexical Tools . . . Available from: <https://lhncbc.nlm.nih.gov/LSG/Projects/lvg/current/docs/designDoc/UDF/unicode/NormOperations/mapSymbolToAscii.html>.
- 12) Pyrovolakis K, Tzouveli P, Stamou G. Multi-Modal Song Mood Detection with Deep Learning. *Sensors*. 2022;22(3):1065–1065. Available from: <https://doi.org/10.3390/s22031065>.
- 13) Mardaoui D, Garreau D. An Analysis of LIME for Text Data . . Available from: <https://arxiv.org/abs/2010.12487>.
- 14) Nazar S, Bustam MR. Artificial Intelligence and New Level of Fake News. In: and others, editor. IOP Conference Series: Materials Science and Engineering;vol. 879 of 1. 2020;p. 2–5. Available from: <http://dx.doi.org/10.1088/1757-899x/879/1/012006>.
- 15) Zhou X, Zafarani R. Fundamental Theories, Detection Methods, and Opportunities. *ACM Computing Surveys*. 2020;53(5):1–40. Available from: <http://dx.doi.org/10.1145/3395046>.
- 16) Risdal M. Getting Real about Fake News.. . Available from: <https://www.kaggle.com/datasets/mrisdal/fake-news>.
- 17) Sekeroglu B, Abiyev R, Ilhan A, Arslan M, Idoko JB. Systematic Literature Review on Machine Learning and Student Performance Prediction: Critical Gaps and Possible Remedies. *Applied Sciences*. 2021;11(22):10907–10907. Available from: <http://dx.doi.org/10.3390/app112210907>.

- 18) Khurana D, Koli A, Khatter K, Singh S. Natural language processing: state of the art, current trends and challenges. *Multimedia tools and applications*. 2022;82(3):3713–3757. Available from: <https://doi.org/10.1007/s11042-022-13428-4>.
- 19) Van De Schoot R, De Bruin J, Schram R, Zahedi P, De Boer J, Weijdemans F, et al. An open source machine learning framework for efficient and transparent systematic reviews. *Nature Machine Intelligence*. 2021;3(2):125–133. Available from: <http://dx.doi.org/10.1038/s42256-020-00287-1>.
- 20) Molina MD, Sundar SS, Le T, Lee D. “Fake News” Is Not Simply False Information: A Concept Explication and Taxonomy of Online Content. *American Behavioral Scientist*. 2021;65(2):180–212. Available from: <http://dx.doi.org/10.1177/0002764219878224>.
- 21) Baptista J, Gradim A. A working definition of fake news. *Encyclopedia*. 2022;2(1):632–677. Available from: <https://doi.org/10.3390/encyclopedia2010043>.
- 22) De Paor S, Heravi B. Information literacy and fake news: How the field of librarianship can help combat the epidemic of fake news. *The Journal of academic librarianship*. 2020;46:102218–102218. Available from: <https://doi.org/10.1016/j.acalib.2020.102218>.