

Cross comparison of COVID-19 fake news detection machine learning models

Muhammad Qadees ¹ and Abdul Hannan ²

¹Affiliation not available

²National university of Science and technology EME Rawalpindi

December 16, 2023

Abstract

The quick advancement of technology in internet communication and social media platforms eased several problems during the COVID-19 outbreak. It was, however, used to spread untruths and misinformation regarding the illness and the immunization. In this study, it is examined whether machine-learning algorithms (Naive Bayesian, Random Forest, Logistic Regression, Decision Tree, and Support Vector Machine, as well as Gradient Boost, Bagging, AdaBoost, Stochastic Gradient Descent, and Multi-layer Perceptron) can automatically classify and point out fake news text about the COVID-19 pandemic posted on social media platforms. The “COVID19-FNIR DATASET” was used to train, test, and fine-tune machine learning models in order to predict the sentiment class of each fake news item on COVID-19. The results were assessed using a variety of evaluation metrics (confusion matrix, classification rate, true positives rate, etc.). The findings collected demonstrate an extremely high level of accuracy when compared to other models.

Cross comparison of COVID-19 fake news detection machine learning models

Muhammad Qadees, Abdul Hannan
CEME, NUST
ISLAMABAD, PAKISTAN

Abstract— The quick advancement of technology in internet communication and social media platforms eased several problems during the COVID-19 outbreak. It was, however, used to spread untruths and misinformation regarding the illness and the immunization. In this study, it is examined whether machine-learning algorithms (Naive Bayesian, Random Forest, Logistic Regression, Decision Tree, and Support Vector Machine, as well as Gradient Boost, Bagging, AdaBoost, Stochastic Gradient Descent, and Multi-layer Perceptron) can automatically classify and point out fake news text about the COVID-19 pandemic posted on social media platforms. The "COVID19-FNIR DATASET" was used to train, test, and fine-tune machine learning models in order to predict the sentiment class of each fake news item on COVID-19. The results were assessed using a variety of evaluation metrics (confusion matrix, classification rate, true positives rate, etc.). The findings collected demonstrate an extremely high level of accuracy when compared to other models.

Index Terms—Fake News Detection, COVID-19, Social Media, Machine Learning, Fine Tuning ML algorithms, Text Classification

I. INTRODUCTION

For instance, Facebook, TikTok, and Twitter prioritize news sharing, communication, engagement, and cooperation. This is used by businesses to advertise their goods and draw people in as well as for personal sharing. Fortunately, these platforms are now approachable and user-friendly thanks to the development and accessibility of mobile applications. Controlling the dissemination of false information on social media, which touches on a variety of subjects including politics, the environment, economy, and health, is one of the major issues. False information and news stories Publishers may be motivated by a variety of factors, including amusement, shaping the public's perception of an issue, boosting website traffic, advancing an unbalanced viewpoint, etc. They are either dishonest or generally unlawful [1].

Fake news is information that has been purposefully spread and contains incorrect facts in an effort to sway public opinion in favor of a certain agenda for political, social, or economic advantage, or just for fun. Fake news raises more eyebrows in the public than reliable sources of information, which is worrying [2]. In addition, compared to authentic news, false news disseminates far more quickly and influences people's perceptions much more deeply. Because of this, most individuals accept and disseminate such news information without hesitation or knowledge [1]. The propagation of false

Information has a variety of negative effects, including major negative effects including poor decision-making, cyberbullying, animosity in society, and violence. The present COVID-19 epidemic has brought to light the worrying effects of the dissemination of misleading information. The continuing epidemic is one of the most catastrophic public health problems now occurring, with a wide range of effects on millions of individuals [3]. Fake news has been spreading on internet channels throughout the epidemic, which has alarmed the people. The economy of the nation was harmed, citizens' faith in their governments was diminished, certain items were promoted in order to generate enormous profits, and inaccurate preventative and treatment advice was spread [4].

To identify false news, one must assess its reliability and categorize it as "Fake or Real" [5]. Content-based approaches and social context-based techniques make up the two primary groups of false news detecting methods [1]. In order to classify news as true or fake, the components used in content-based detection techniques are extracted directly from news material (such as headlines, body text, images/videos, and news sources). For social context-based detection to work, users and the news must share information [6].

The main part of this paper is the proposal of a fake news detection system that uses a variety of machine learning algorithms (Naive Bayesian, Random Forest, Logistic Regression, Decision Tree, and Support Vector Machine, Gradient Boost, Bagging, AdaBoost, Stochastic Gradient Descent, Multilayer Perceptron) to automatically identify and detect fake news data for the COVID-19 pandemic [7].

Users or AI developers can establish standards or weight the metrics in accordance with their significance. using the study's analysis of model performance using a variety of evaluation metrics. This allows them to select the best model that satisfies the requirements or to classify the models based on more than one metric. However, false news detection may be handled via multi-criteria decision-making, in which users rank deep learning models in accordance with their preferences, such as the quickest training time and the most accurate model with the best F1 score. The most accurate model that takes very long time to update or retrain on a fresh dataset may not be selected over models that can be updated fast and with good accuracy [8]. We assess the models used in this study based on how well they predict outcomes.

The related research on predicting bogus news is examined in Section 2.

I. LITERATURE REVIEW

Several research on the detection of false news will be covered in this section.

Umer et al. created a hybrid deep neural network architecture in 2020 [9] for identifying and identifying fake news on social media networks. The CNN and LSTM models are combined in the hybrid model. They decreased the feature vectors' dimensions using Chi-Square and Principal Component Analysis before presenting them to the classifier (PCA). The Fake News Challenges (FNC) dataset, which comprises of news items classified into four categories, was used to test the algorithm (Agree, Discuss, Disagree, and Unrelated). Feeding non-linear data into the PCA and chi-square results in the production of more interpretative features for the task of fake news identification. The findings showed that PCA outperformed Chi-square with an accuracy of 97.8%.

Traditional natural language processing methods are not appropriate for identifying attitudes in huge data, according to Khanam et al(2021) 's [10] proposal. As a result, they used computational categorization, machine learning algorithms, and self-organization maps to propose a model. They also conducted a comparison between the proposed technique and cutting-edge approaches. Six machine learning models were applied. Classification accuracy for each model is 75%, 74%, 74.5, 71.2%, and 70%, respectively. The outcomes demonstrate an intriguing phenomena whereby the performance of the proposed model becomes better as data amount increases. They arrived to the conclusion that the proposed technique had the highest accuracy (75%) after researching four algorithms (XGBoost).

Kim et al. (2019) [11] claim that CNN has obtained favorable attention and talk over in the area of emotion classification. They used CNN for sentiment-based data classification and further incorporated consecutive convolutional layers for this purpose. Then, they compared this model to other cutting-edge deep learning and machine learning techniques using three distinct datasets. For the CNN design, they suggest an embedding layer, two convolutional layers, a pooling layer, and a fully connected layer. Numerous works on machine learning concentrate on two or more sentiment labels. The proposed CNN model is compared to the NB, DT, SVM, and RF models. The results obviously show that sequential layer CNN leads when using datasets from movie reviews, customer reviews, and the Stanford emotion treebank, with accuracy rates of 81.06%, 78.3%, and 68.3%, respectively. Their model was also evaluated for ternary classification. before applying it to the MR Dataset. With a 68.3% accuracy rate, their model is the most accurate ML and DL model available.

The accuracy of sentiment analysis work is completely relying on the precision of a domain-specific vocabulary, claim Xia et al. (2020) [12]. Instead of using the complete review, they proposed a solution that makes use of the conditional random field technique and emotional characteristics from review fragments (CRF). The feature words are then asymmetrically weighted and SVM is used for classification. They used two distinct sources to gather their data. Audi review website in China given one dataset.

A4 automobile, while the information regarding the Samsung S7 phone is from the Amazon website.com. They ran three different experiments using (CRF+ asymmetric weighting+ SVM), (TDIDF+SVM), and (CRF + TDIDF+SVM) on the first, second, and third of these two datasets, respectively. The results clearly show that the conditional random field approach with asymmetric weighting improved the average accuracy of the Chinese dataset to 90% and the average accuracy of the English dataset to 91%. Many studies on sentiment analysis using machine learning, according to Ullah et al. (2020) [13], rely primarily on text, emoticons, or pictures. Emoticon-heavy text has never been taken seriously. In order to locate SA using both text and emoticons, they devised a model and algorithm. Both alone and together, they analyzed plain text and text that had emoticons. From Twitter, they acquired data on airline reviews. The vocabulary they built, which they utilized in their study, covers the most frequently used emoticons among all Twitter users. They also invented the emoticons. They investigated machine learning and deep learning techniques using SVM, NB, LR, Random Forest, LSTM, and CNN. LSTM and CNN outperform all other algorithms, with accuracy values of 0.89%, 0.81%, 0.88%, and 0.79% on (text + emoticons) and text, respectively. This proves categorically that machine learning techniques fall short of deep learning algorithms. The novel strategy raised text SA accuracy from 57% to 78% and text and emoticon accuracy from 65% to 89 when they differentiate their offered model to existing models. Medical and health reviews are not given much thought by NLP and DM researchers, according to Basiri et al. (2020) [14]. The 3W1DT and 3W3DT fusion models were released. The first fusion model combines a deep model with a traditional learning strategy (GRU, CNN, and 3CRNN with NB, DT, RF, and KNN). The second fusion model is composed of three deep models and one conventional model. They used a dataset of 215063 drug reviews that had been divided into three categories: neutral, unfavorable, and positive. After the initial test, NB outperforms all other algorithms when tests are conducted on a dataset. The 3CRNN-NB fared better than the others in the second trial using the 3W1DT. When the third experiment was undertaken using 3W3DT, NB outperformed the second fusion model with the highest level of accuracy. The suggested model exceeds the existing model with an accuracy of 88.36% when they compare it to their best model, 3W3DT-NB, which they then compare to.

According to Awwalu et al. (2019) [15], political parties should consider using twitter data on politics since it enables them to forecast the opinions of their followers based on their tweets. They put out a classification model that uses NB and a two-gram hybrid approach. By resolving the "zero count problem," this model enhances the precision and recall accuracy of -gram models. The two rounds of sentiment analysis in the proposed technique employ the least-order and highest-order -gram models. They made use of the Obama-McCain dataset. The tests demonstrate that the recommended method outperforms all prior research on the identical dataset. This model improves the unigram model's accuracy to 76.14%, the -gram model's accuracy to 67%, and the hybridized model's accuracy to 80%. It shows how merging

the unigram and -gram models

may improve sentiment prediction.

Aspect-level analysis is essential, according to Chen et al(2019) 's [16] proposal, however labelled data in relation to aspect-level analysis is the key roadblock in this field of study. They thus put out a design known as the transfer capsule network (TransCap). In essence, this is how knowledge is moved from the document level to the aspect level. They assess their methods using two datasets from SemEval 2014 task 4: evaluations of restaurants and laptops. They transferred knowledge from papers using reviews on Yelp, Amazon, and Twitter. Out of all investigated methodologies, the suggested methodology, on the restaurant and laptop datasets, respectively, achieves 79.5% and 73.87% accuracy.

Chatsiou et al. completed a sentence-level categorization task in 2020 [17]. They conducted many tests using CNN knowledgeable on top of pre-trained word vectors for sentence-level classification tasks. In addition to CNN, they experiment with Word2Vec+CNN, GloVe+CNN, ELMo+CNN, and BERT+CNN. The results unequivocally show that BERT+CNN beats the remaining pairings across two datasets —the manifestos project corpus for training the model and the coronavirus (COVID-19) news briefing corpus for evaluating the model's performance. An F1 score of 64.58% and 68.65% accuracy are achieved using BERT+CNN.

Human supervision and the identification of fraudulent stories, according to Koirala et al(2020) 's [18] hypothesis, are essentially impossible jobs. By giving computers the duty of identifying patterns, processing techniques have advanced to the point where ML models, DL models, and user interaction may be removed; nonetheless, a sizable dataset of both legitimate and fraudulent news is needed. Between January 15 and February 15, 2020, he acquired news from all around the world, but the information was unlabeled. After unneeded information was removed and the news items were labelled, the dataset now contains 2426 articles with the label "true" and 1646 articles with the label "false." Following classification studies, LR attained an accuracy of 75.65%, dense layer embedding an accuracy of 86.93%, LSTM layer embedding an accuracy of 86.9%, and bi-LSTM model an accuracy of 72.31%. The overview and findings of the literature review are shown in Table 1.

Paper	Algorithm	Accuracy (%)
Umer et al., (2020) [9]	PCA	97.8
Khanam et al., (2021) [10]	XGBOOST	75
Kim et al., (2019) [11]	CNN	81.06
Xia et al., (2020) [12]	SVM	91
Ullah et al., (2020) [13]	LSTM	89
Basiri et al., (2020) [14]	NB	88.36
Awwalu et al., (2019) [15]	NB	80
Chen et al., (2019) [16]	TransCap	79.5
Chatsiou et al., (2020) [17]	BertCNN	68.65
Koirala et al., (2020) [18]	LSTM	86.9

TABLE I: Literature Review

II. MATERIAL AND METHODS

A. Dataset

The COVID-19 Fake News Infodemic Research (CoVID19-FNIR) Dataset, provided by Diksha Shukla [19], is used in this study.

1) *Dataset summary:* The COVID-19-FNIR dataset is a compilation of verified news articles and posts about CoVID-19 that were gathered from Poynter, Twitter, and various online sources in India, the United States, and Europe. The dataset includes both genuine news from verified news publishers and fake news that has been fact-checked. The data was collected between February and June 2020 and has undergone processing to remove special characters and unnecessary information.

2) *Data Format and File Structure:* The complete data set is in the *CoVID19-FNIR.zip* folder. The folder includes two files; (1) *fakeNews.csv*, and (2) *trueNews.csv*.

fakeNews.csv: The file *fakeNews.csv* is organized as follows. It contains the columns and the corresponding information as listed below. The last column, label, shows the classification label for the corresponding news item. Each row is one news item.

- Date: The date on which the article was published
- Link: Article's Poynter Link
- Text: Article text.
- Region: The region where the article is from.
- Country: The country where the article is from
- Explanation: Explanation for the article- why it was false
- Origin: The website where the article is taken from
- Origin_URL: The URL for the article's website origin.
- Fact_checked_by: Name was given of who fact-checked the article
- Poynter_Label: The multi-class classification label given by Poynter
- Label: The binary classification label we provided of 0 for false

trueNews.csv: The file *trueNews.csv* contains the following columns with last column label being the classification label for the corresponding news item. In this file all news items come from the twitter handles of trusted news sources and were assigned a classification label as 'True'.

- Date: The date on which the tweet was posted
- Link: Link to the article in tweets
- Text: The tweet text.
- Region: Region of the tweet.
- Username: Username/ Twitter handles of the authors of the new tweets
- Publisher: New publication organization official name.
- Label: Classification labels i-e True or False

Twitter Handles: The verified Twitter usernames of news sources where the *trueNews.csv*'s news samples were gathered.

- India: NDTV, The Hindu, India Today
- The United States of America:: CDC, The Washington PostThe New York Times,
- Europe: BBC News UK, Guardian News, Reuters UK

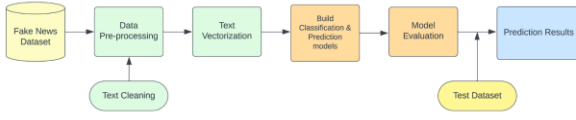


Fig. 1: Research Methodology

B. Pre-processing

The COVID19-FNIR dataset includes two files: trueNews.csv and fakeNews.csv. The true news file has 7 columns, while the fake news file has 9 columns. For the purposes of classification, we are only using the text and label columns. The text column contains the news headlines, and the label column indicates whether the information in the text column is true or fake. The text field for true news items includes links to websites and Twitter images where the articles are posted, while the text field for fake news items does not include any URLs. To ensure consistency between the true and fake news headlines, we applied several text-cleaning methods to the d[20].

- **Regex Substitution:** This is used to remove various elements from the text, such as punctuation, digits, HTML tags, and newline characters.
- **String Replacement:** This is used to remove specific types of strings, such as URLs, hashtags, and Twitter usernames, from the text.
- **Lower Casing:** This is used to standardize the text to lowercase.
- **Removal of Stop Words:** Stop words are typical terms from any language that don't significantly further the meaning of the information (Source: [Name of Source]). Examples of English stop words include 'is,' 'am,' 'are,' 'of,' and 'the.' Removing stop words from the text can help to reduce the feature's dimensionality space and potentially improve the performance of a classification model [21].
- **Removal of Spaces:** All the unwanted spaces were removed from the dataset.
- **Removal of Duplicated Rows:** All the duplicated rows were removed from the dataset. After removing the 234 duplicated rows from fakeNews.csv, the updated count can be seen in Figure 1.

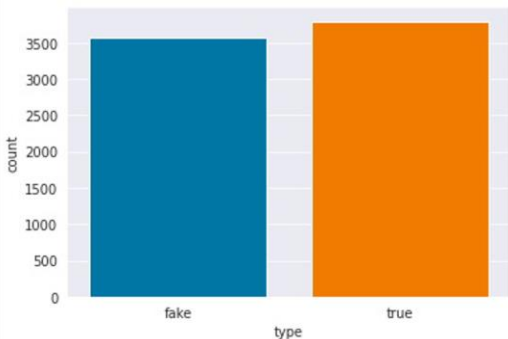


Fig. 2: Classification Results

C. Data Vectorization

We had to structure the data by transforming it such that the models could interpret it in order to run machine learning models on our textual data. We used the Bag of Words (BoW) model, which be text as numerical vectors, to do this. We choose the BoW model because it can be used with traditional machine learning techniques and is straightforward and efficient [22]. We also utilized the Term Frequency-Inverse Document Frequency (TF-IDF) to weight the BoW model's characteristics [23]. The TF-IDF takes into account the utility of a characteristic both in the overall dataset and within individual documents. A characteristic that appears frequently in a document is given more weight, whereas a feature that appears frequently is given less weight.

D. Classifiers

Single-label categorization is the process of learning from a collection of cases that share a single label. Binary classification is the name of the learning job when there are just two labels. The dataset used in this study is a binary classification problem. For this research, 10 **Machine Learning** algorithms were employed, fined tuned and cross-compared.

Naive Bayesian Algorithm: Based on the Bayes theorem, the A supervised learning approach for handling classification issues is the naive Bayes method. Eq.1. With a large training dataset, it's typically utilized for text classification [24]. A rapid and efficient classification method for producing machine learning models with accurate predictions is the Naive Bayes Classifier. Because it is a probabilistic classifier, it forecasts outcomes.

Gradient Boost Algorithm: Gradient boosting is a machine learning technique that builds models sequentially, with each new model aiming to cut down on the flaws in the preceding one [25]. Usually, this is accomplished by basing a new model on the residuals or mistakes of the old one. A gradient-boosting repressor is used when the goal column is continuous, while a gradient-boosting classifier is used when the target column is a categorical variable. The objective of gradient boosting is to minimize a loss function using gradient descent, and the specific loss function used will depend on the type of problem being solved (e.g., mean squared error for regression or log-likelihood for classification). Overall, gradient boosting is a powerful machine-learning technique that has been successful in a variety of applications.

Support Vector Machine: SVM is a supervised machine-learning technique for classifying and predicting data (Support Vector Machine). However, classification problems are where it is most commonly applied [26]. Each piece of data is represented as a point in n-dimensional space, and the value of each feature is the value that the SVM algorithm assigned to a certain place. The categorization is then completed by selecting the hyper-plane that most usefully discern the two classes. The support vectors are computed using the individual observation's coordinates. The SVM classifier acts as a divider between the two classes.

3) **Decision Tree Algorithm:** Decision trees falls in supervised machine-learning algorithms that are often employed for classification and regression problems [27]. They construct a tree-like model of decision-making based on data properties, with the purpose of producing predictions based on these features. In a decision tree, the root node mean the overall decision that the tree is trying to make, and the leaves of the tree represent the final prediction or classification. The internal nodes of the tree represent the different features or characteristics that the tree uses to make decisions. Decision trees are easy to interpret and visualize, which makes them useful for understanding and explaining the decision-making process of a model. They are also robust to noise and can handle missing values.

4) **Random Forest Algorithm:** random forest classifier is an assembling method technique which works on averaging the results of Decision Trees [28]. Higher the number of decisions trees, the better the results of the random forest classifier

5) **Bagging:** Bagging is a machine learning technique that involves generating diverse samples of training data using a bootstrapping sampling method [29]. This means that data

Points are selected at random from the training dataset with Replacement, which means that the same instance can be selected multiple times in a single sample. These samples are then used to train multiple weak or base learners independently and in parallel finally, the predictions of these individual learners are merged to get a more precise approximation. In the case of regression, this is accomplished by averaging the predictions. In classification difficulties, the class with the most votes is selected. Bagging is frequently used to boost the effectability of machine learning models [9] by decreasing overfitting and enhancing the model ensemble's diversity.

7) **AdaBoost:** Adaboost, also known as adaptive boosting, is a common machine learning boosting technique that is used to increase the prediction accuracy of "lazy" learning systems [30]. Boosting algorithms function by training several weak or "lazy" learners and then combining them to form a single, strong learner. Adaboost works by iteratively modifying training instance weights based on classification results. Each iteration increases the weights of misclassified cases while decreasing the weights of successfully classified examples, resulting in a better overall classifier. Adaboost is frequently used in classification problems and is well-known for its excellent prediction accuracy and ability to operate effectively with a wide range of poor learners.

8) **Stochastic Gradient Descent:** Stochastic Gradient Descent (SGD) is an efficient optimization algorithm for fitting linear classifiers and repressors to data under convex loss functions, such as those used in Support Vector Machines and Logistic Regression [25]. It has seen widespread use in the machine learning community for many years, but has earned particular notice in immediate years due to its ability to scale to large and sparse datasets. SGD has been effectively used to a wide range of large-scale projects machine learning tasks, particularly in the areas of text classification and natural language processing. It is a simple but powerful approach that

Algorithms	Accuracy %	F1-Score
Naïve Bayes	89.2	0.89
Gradient Boost	85.1	0.86
SVM	85.1	0.86
Decision Tree	85.8	0.86
Random Forest	91.6	0.92
Bagging	87.2	0.87
AdaBoost	84.8	0.85
SGD	91.5	0.92
Logistic Regression	90.6	0.91
MLP	90.4	0.90

TABLE II: Cross Comparison of the Score of Machine Learning Models

is capable of handling large datasets efficiently, making it a popular choice for many machine learning practitioners.

9) **Logistic Regression:** When the child variable is dichotomous, the best regression approach to use is logistic regression (binary). Logistic regression is a predictive research, just like other regression studies. Utilizing logistic regression is one method for illustrating and explaining the connection between one dependent binary variable and one or more independent nominal, ordinal, interval, or ratio-level variables. [31].

10) **Multi-layer Perceptron:** A popular machine learning approach for binary classification tasks is the perceptron. It performs well when categorizing data that can be divided into lines, but it might have trouble with more challenging data sets that don't fit this pattern, like the XOR problem. The XOR issue serves as an example that for some classification jobs, a linear boundary cannot be found that can accurately separate all of the data points. It is important to employ more sophisticated algorithms, such as the Multilayer Perceptron (MLP) [32], to handle these more complicated data sets. By employing a more complicated and adaptable architecture that can develop more intricate regression and classification models, MLPs are able to get around the perceptron's drawbacks. They have shown to be successful at a diversity of machine learning work and are frequently utilized for difficult data sets that are not linearly separable.

Using the random forest with hyper parameter optimization by Randomized- SearchCv, an accuracy score of 91.6% was attained. SGD scored the second-highest accuracy, coming in at 91.5%. The model that scored the lowest accuracy 84.8% was AdaBoost. Table II provides the Accuracy and F1 score for each of the 10 machine learning models.

III. RESULTS

Machine learning-based binary classifiers were applied and compared in this work. With an accuracy of 91.6% and an F1 score of 92%, Random Forest including hyper parameter optimization using RandomizedSearchCv outperformed all other models, while the quoted study [10] had an accuracy of 75% on XGBoost Classifier. A score of 75% indicates that, although being a significant improvement over the mentioned work, there is still room for development. Figure 2 displays the F1 scores for each of the algorithms used in this study.

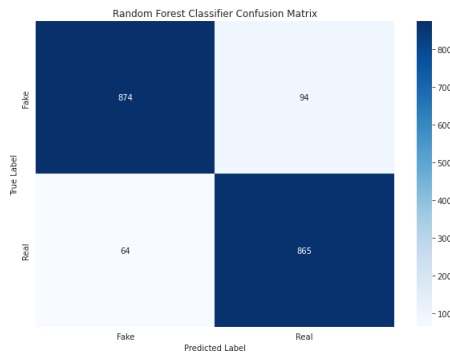


Fig. 3: Random Forest Confusion Matrix

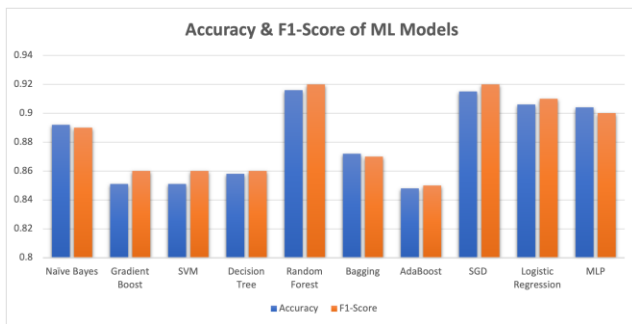


Fig. 4: Classification Results

IV. CONCLUSION AND FUTURE WORK

In this study, ten machine learning algorithms such as Naive Bayes, Gradient Boost, Support Vector Machine, Decision Tree, Random Forest, Bagging, Adaboost, Stochastic Gradient Descent, Logistic Regression, and Multi-Layer Perceptron were employed to detect the fake and true News on the COVID19-FNIR dataset. For the purpose of determining the best classifier and method that can be used to discern between the Fake and the Real News about the COVID-19 virus, many methodologies and tests were run on this dataset. The algorithms' Precision, Recall, and F1 Score were examined. We want to use multiple deep learning models which includes Bi-LSTM, LSTM, BERT, and CNN, in further studies. In the future, it may be possible to employ deep learning models in combination with several embedding methods, such as Word2Vec, RoBERTa, and Word embedding, to determine which approach performs best on the dataset applied in the study. The machine and deep learning technique that can improve the task may be used to a larger dataset and a dataset that is composed of various languages.

REFERENCES

- [1] S. Vosoughi, D. Roy, and S. Aral, "The spread of true and false news online," *science*, vol. 359, no. 6380, pp. 1146–1151, 2018.
- [2] S. Zaryan, "Truth and trust: How audiences are making sense of fake news," 2017.
- [3] M. Nicola, Z. Alsafi, C. Sohrabi, A. Kerwan, A. Al-Jabir, C. Iosifidis, M. Agha, and R. Agha, "The socio-economic implications of the coronavirus pandemic (covid-19): A review," *International journal of surgery*, vol. 78, pp. 185–193, 2020.
- [4] B. Al-Ahmad, A. Al-Zoubi, R. Abu Khurma, and I. Aljarah, "An evolutionary fake news detection method for covid-19 pandemic information," *Symmetry*, vol. 13, no. 6, p. 1091, 2021.
- [5] M. K. Elhadad, K. F. Li, and F. Gebali, "Detecting misleading information on covid-19," *Ieee Access*, vol. 8, pp. 165 201–165 215, 2020.
- [6] R. K. Kaliyar, A. Goswami, and P. Narang, "Fakebert: Fake news detection in social media with a bert-based deep learning approach," *Multimedia tools and applications*, vol. 80, no. 8, pp. 11 765–11 788, 2021.
- [7] H. Alhakami, W. Alhakami, A. Baz, M. Faizan, M. W. Khan, and A. Agrawal, "Evaluating intelligent methods for detecting covid-19 fake news on social media platforms," *Electronics*, vol. 11, no. 15, p. 2417, 2022.
- [8] Y. Tashtoush, B. Alrababah, O. Darwish, M. Maabreh, and N. Alsaedi, "A deep learning framework for detection of covid-19 fake news on social media platforms," *Data*, vol. 7, no. 5, p. 65, 2022.
- [9] M. Umer, Z. Imtiaz, S. Ullah, A. Mehmood, G. S. Choi, and B.-W. On, "Fake news stance detection using deep learning architecture (cnn-lstm)," *IEEE Access*, vol. 8, pp. 156 695–156 706, 2020.
- [10] Z. Khanam, B. Alwasel, H. Sirafi, and M. Rashid, "Fake news detection using machine learning approaches," in *IOP Conference Series: Materials Science and Engineering*, vol. 1099, no. 1. IOP Publishing, 2021, p. 012040.
- [11] H. Kim and Y.-S. Jeong, "Sentiment classification using convolutional neural networks," *Applied Sciences*, vol. 9, no. 11, p. 2347, 2019.
- [12] H. Xia, Y. Yang, X. Pan, Z. Zhang, and W. An, "Sentiment analysis for online reviews using conditional random fields and support vector machines," *Electronic Commerce Research*, vol. 20, no. 2, pp. 343–360, 2020.
- [13] M. A. Ullah, S. M. Marium, S. A. Begum, and N. S. Dipa, "An algorithm and method for sentiment analysis using the text and emoticon," *ICT Express*, vol. 6, no. 4, pp. 357–360, 2020.
- [14] M. E. Basiri, M. Abdar, M. A. Cifci, S. Nemati, and U. R. Acharya, "A novel method for sentiment classification of drug reviews using fusion of deep and machine learning techniques," *Knowledge-Based Systems*, vol. 198, p. 105949, 2020.
- [15] J. Awwalu, A. A. Bakar, and M. R. Yaakub, "Hybrid n-gram model using naïve bayes for classification of political sentiments on twitter," *Neural Computing and Applications*, vol. 31, no. 12, pp. 9207–9220, 2019.
- [16] Z. Chen and T. Qian, "Transfer capsule network for aspect level sentiment classification," in *Proceedings of the 57th annual meeting of the association for computational linguistics*, 2019, pp. 547–556.
- [17] K. Chatsiou, "Text classification of covid-19 press briefings using bert and convolutional neural networks," 2020.
- [18] A. Koirala, "Covid-19 fake news classification with deep learning," *Preprint*, 2020.
- [19] J. A. Saenz, S. R. Kalathur Gopal, and D. Shukla, "Covid-19 fake news infodemic research dataset (covid19-fnir dataset)," 2021. [Online]. Available: <https://dx.doi.org/10.21227/b5bt-5244>
- [20] S. A. Alasadi and W. S. Bhaya, "Review of data preprocessing techniques in data mining," *Journal of Engineering and Applied Sciences*, vol. 12, no. 16, pp. 4102–4107, 2017.
- [21] N. Hardeniya, J. Perkins, D. Chopra, N. Joshi, and I. Mathur, *Natural language processing: python and NLTK*. Packt Publishing Ltd, 2016.
- [22] R. Zhao and K. Mao, "Fuzzy bag-of-words model for document representation," *IEEE transactions on fuzzy systems*, vol. 26, no. 2, pp. 794–804, 2017.
- [23] F. Sebastiani, "Machine learning in automated text categorization," *ACM computing surveys (CSUR)*, vol. 34, no. 1, pp. 1–47, 2002.
- [24] T. Hastie, R. Tibshirani, J. H. Friedman, and J. H. Friedman, *The elements of statistical learning: data mining, inference, and prediction*. Springer, 2009, vol. 2.
- [25] F. A. Ozbay and B. Alatas, "Fake news detection within online social media using supervised artificial intelligence algorithms," *Physica A: Statistical Mechanics and its Applications*, vol. 540, p. 123174, 2020.
- [26] T. Joachims, "Text categorization with support vector machines: Learning with many relevant features," in *European conference on machine learning*. Springer, 1998, pp. 137–142.
- [27] S. Hakak, M. Alazab, S. Khan, T. R. Gadekallu, P. K. R. Maddikunta, and W. Z. Khan, "An ensemble machine learning approach through effective feature extraction to classify fake news," *Future Generation Computer Systems*, vol. 117, pp. 47–58, 2021.
- [28] P. Bharadwaj and Z. Shao, "Fake news detection with semantic features and text mining," *International Journal on Natural Language Computing (IJNLC) Vol*, vol. 8, 2019.

- [29] J. Del Ser, M. N. Bilbao, I. Laña, K. Muhammad, and D. Camacho, "Efficient fake news detection using bagging ensembles of bidirectional echo state networks," in *2022 International Joint Conference on Neural Networks (IJCNN)*. IEEE, 2022, pp. 1–7.
- [30] M. Amjad, G. Sidorov, A. Zhila, H. Gómez-Adorno, I. Voronkov, and A. Gelbukh, ""bend the truth": Benchmark dataset for fake news detection in urdu language and its evaluation," *Journal of Intelligent & Fuzzy Systems*, vol. 39, no. 2, pp. 2457–2469, 2020.
- [31] S. Kaur, P. Kumar, and P. Kumaraguru, "Automating fake news detection system using multi-level voting model," *Soft Computing*, vol. 24, no. 12, pp. 9049–9069, 2020.
- [32] R. Kumari and A. Ekbal, "Amfb: attention based multimodal factorized bilinear pooling for multimodal fake news detection," *Expert Systems with Applications*, vol. 184, p. 115412, 2021.