

MACHINE LEARNING FOR IDENTIFYING FAKE NEWS: TECHNIQUES AND APPLICATIONS

Olumide Emmanuel Adebayo

Department of Computer Science, Federal College of Agriculture Akure, Ondo State, Nigeria

Abstract: Fake news is a kind of propaganda in which false information is knowingly disseminated via news organizations and/or social media platforms. It is crucial to create methods of spotting false news information since its spread can have detrimental effects, such as influencing elections and widening political rifts. Today, the majority of people choose to obtain news online since it is convenient and affordable, ubiquitous, but this leads to the rapid spread of false information. News disseminated quickly among millions of users in a short amount of time as a result of the rise in the usage of social media platforms like Facebook, Twitter, etc. The effects of spreading false news are often much, from shaping public opinion to influencing election results in favour of particular politicians. With the use of the machine learning concept, this research seeks to perform binary classification of various news items that are available online. Additionally, it attempts to enable users to evaluate if a piece of news is true or false and to confirm the reliability of the website that is disseminating it.

Keywords: fake news, Machine Learning, machine learning, (ML Support Vector Machine (SVM), Facebook, twitter.

I. INTRODUCTION

In the last decade, Fake News phenomenon has experienced a very significant spread, favored by social networks. This fake news can be broadcasted for different purposes. Some are made only to increase the number of clicks and visitors on a site. Others, to influence public opinion on political decisions or on financial markets. Example of fake news is instance that occurred when the Facebook article "Pope Francis shocks world, endorses Donald Trump for president," received over 900,000 views. Despite the fact that this report was wholly untrue, it had an impact on the election. Another example of fake news was a post by Nnamdi Kanu speaking to IPOB supporters- "*The man you are looking at is not Buhari –His name is Jubril, he's from Sudan. After extensive surgery, they brought him back*", this post was viewed half a million times and prompted the President, standing for a second term, to make a public denial while on a visit to Poland. "This is the real me," he insisted.

The term "fake news" refers to information that has been deliberately and confirmed falsified in order to sway public impressions of actual events, facts, and viewpoints. The topic is blatantly fake news that the promoter is aware of, based on untrue statements, facts, or incidents. This is usually done in order to advance or enforce particular points of view, and it is typically achieved through the pursuit of political objectives.

Fake news is widely disseminated in society via social media platforms like Facebook and Twitter. Fake news is intentionally created content that mimics the style and tone of news media however does not conform to the editorial standards and procedures used by news organizations to verify the veracity and reliability of their reporting. Fake news is one of many types of deceit on social media that attracts people's attention, but it is a more dangerous form since it is produced with the intention of misleading others. When identifying fake news manually, several different methods and processes are used to verify the data. The methods can be categorized as manual or automated. The manual method might involve visiting fact-checking websites, crowdsourcing verified

news to contrast it with unverified news, and more. However, the amount of data that is gathered every day online is astounding. Manual fact-checking soon proves to be useless given how quickly information spreads on the internet. It is challenging to scale manual fact-verification due to the amount of data collected hence the need for automation.

The paper aim at designing a model that can classify news as fake or real using machine learning algorithm. With the aid of machine learning algorithms, users may categorize news as either true or fake and check the reliability of the website that originally published it. Machine learning algorithms is a subfield of artificial intelligence (AI) concerned with utilizing data and algorithms to simulate how people learn in order to increase accuracy over time. It's the study of algorithms for computers that can grow and learn on their own depending on information and experience.

Implementing this algorithm, support vector machine classifier will be used.

Types of Fake News

According to Merrimack College media professor Melissa Zimdars, there are four major forms of fake news.

Category 1: Websites that are shared on Facebook and other social media that are fake, deceptive, or frequently misleading. In order to increase likes, shares, and earnings, some of these websites may depend on "outrage" by employing skewed headlines and questionable or decontextualized material.

Category 2: Websites that may circulate misleading and/or potentially unreliable information

Category 3: Websites that sometimes use clickbait-y headlines and social media descriptions

Category 4: Satire/comedy sites, which can offer important critical commentary on politics and society, but have the potential to be shared as actual/literal news

No single topic falls under a single category - for example, false or misleading medical news may be entirely fabricated (Category 1), may intentionally misinterpret facts or misrepresent data (Category 2), may be accurate or partially accurate but use an alarmist title to get your attention (Category 3) or may be a critique on modern medical practice (Category 4). Some articles fall under more than one category.

Assessing the quality of the content is crucial to understanding whether what you are viewing is true or not. It is up to you to do the legwork to make sure your information is good.

II. LITERATURE REVIEW

In the work of Kaliyar *et al.* (2018) Fake News Detection Using A Deep Neural Network, Natural Language Processing, Machine Learning and Deep Learning Techniques was used to implement the model and the results were compared for accuracy. Nvidia DGX-1 supercomputer was also used for easy analysis of results and examine which model will precisely classify the given dataset into real and fake news. Naive Bayes, K nearest neighbors, Decision trees, Random forests, and Shallow Convolutional Neural Networks (CNN), Very Deep Convolutional Neural Networks (VDCNN), Long Short-Term

Memory Networks (LSTM), Gated Recurrent Unit Networks (GRU), as well as Combinations of Convolutional Neural Network with Long Short-Term Memory (CNN-LSTM), and Convolutional Neural Network with Gated Recurrent Unit Networks are machine learning models that was used for the model. The dataset used for testing and training the model as collected from <https://www.kaggle.com/c/fakenews/data>. The dataset contains combination of fake and real news: where (row*column) is (2080 *5) and (row*column) is (6335*3). When shallow CNN based model was used, accuracy of the model after 2 epochs is 91.3%, the depth of the network was increased by a few layers to check whether the performance is increased or not and the accuracy of the model was increased to 98.3%. Again, the model was implemented using the combination of CNN and LSTM. Accuracy

was reduced by a bit to 97.3% but precision and recall was effectively improved. Implementing Naïve Bayes model on the same dataset accuracy of 89% was obtain. Implementing Decision Tree Regressor on the data set, accuracy of (73%) was obtain. While random forest and K Nearest Neighbors produce accuracy of 71% and 53% accuracy respectively.

Hiramath *et al.* (2019) suggested a model for identifying fake news based on the following classification techniques Logistic regression (LR), Naive bayes (NB), Support vector machine (SVM), Random Forest (RF), and deep neural network (DNN), and the outcomes were evaluated. The memory analysis of various computations reveals that Logistic Regression requires less memory, while the time analysis of algorithms reveals that Deep Neural Networks take just 400 milliseconds, which is less time. Deep neural networks outperform the other four algorithms in terms of accuracy, with a score of 91%, according to the accuracy calculation. Regression, Support vector machines (SVM), Random Forest, and Deep Neural Network (DNN) classification approaches can therefore recognize fake news from a big array of datasets with ease and accuracy.

Machine learning-based categorization of news headlines was proposed by Tiwary *et al.* (2020). The study divides false news into the following classes: clickbait, propaganda, comment or opinion, and satire or humor. Logistic regression, decision trees, K-nearest neighbors, and random forests are some of the several classification techniques used for the project. After the features are extracted using the three vectorizers, the data is fed into each algorithm, and the results show that the usage of the Logistic Regression algorithm with the TF-IDF Vectorizer produces results that are more accurate than those of the other algorithms with the TF-IDF Vectorizer. About 73% of headline fake news was correctly identified and categorized, compared to up to 62% accuracy for Decision Tree and 66% accuracy for Random Forest, respectively. The model was able to classify HeadLines as fake or real.

A naïve Bayesian classifier is used by Granik *et al.* (2017) in the work Fake News Detection Using Naive Bayes Classifier to demonstrate a straightforward method for detecting fake news. A collection of data taken from Facebook news postings is used to evaluate this methodology. They assert that they can reach an accuracy of 74%. This model's rate is respectable but not the greatest because many other papers have used different classifiers to reach higher rates.

Ahmed *et al.* (2017) compared two distinct feature extraction strategies and six distinct classification techniques to develop a false news detection model that makes use of n-gram analysis and machine learning techniques. The results of the studies demonstrate that the so-called features extraction approach yields the greatest performances (TF-IDF). The classifier they utilized, the Linear Support Vector Machine (LSVM), has a 92% accuracy rate.

A model to detect news stories with comedy and satire was put out by Rubin *et al.* in 2016. 360 satirical news items from primarily four categories—civics, science, business, and what they labeled "soft news" ('entertainment/gossip pieces')— were investigated and scrutinized. They put up an SVM classification model based mostly on five features they generated after studying the satirical news. Absurdity, humor, grammar, negative affect, and punctuation are the five characteristics. Only three feature combinations—absurdity, grammar, and punctuation—were used to obtain their best precision of 90%.

In the paper A Framework for Real-Time Spam Detection in Twitter, Gupta *et al.* (2018) presents a framework based on various machine learning approaches that addresses a number of issues, such as accuracy shortage, time lag (BotMaker), and high processing time to handle thousands of tweets in a single second. First, they have gathered 400,000 tweets from the HSpam14 dataset. Then they go on to describe the 150,000 spam tweets and the 250,000 non-spam tweets in more detail. Along with the Top-30 terms from the Bag-of-Words model that are

yielding the largest information gain, they also deduced a few lightweight characteristics. Real-time spam detection is effectively handled by this method. Using their processed dataset, they also carried out a number of studies for detecting Twitter spam. They outperformed the previous result by almost 18% and were able to reach an accuracy of 91.65%.

Automatic Online Fake News Detection Combining Content and Social Signals by Vedova *et al.* (2018) first proposed a novel machine learning (ML) fake news detection method that outperforms other methods in the literature by combining news content and social context features, increasing its accuracy up to 78.8%. Second, they used their technique within a Facebook Messenger Chabot and tested it in a real-world setting, achieving an 81.7% accuracy in detecting fake news. They first discussed the datasets they utilized for their test, then showed the content-based approach they employed and the way they suggested to combine it with a social-based strategy that was previously published in the literature. Their objective was to categorize a news item as genuine or fake. The total dataset consists of 15,500 posts from 32 sites (14 conspiracy pages and 18 scientific pages), with more than 2,300,000 likes from more than 900,000 individuals. 6,577 (42.4%) of the postings are non-hoaxes, whereas 8,923 (57.6%) are hoaxes.

III. METHODOLOGY

The figure1 below illustrate the architecture of the proposed system. In the system, dataset (comprises of comments, likes etc. from social media article) for the model will be source online from Kaggle.com. The dataset will be split into training and testing data. The data for training the model will transform into digital form that can be processed by the computer, unwanted noise will be removed at the point of preprocessing. Features needed for the training the model will be extracted from the preprocessed data using TF-IDF. Support Vector Machine algorithm to build a decision model.

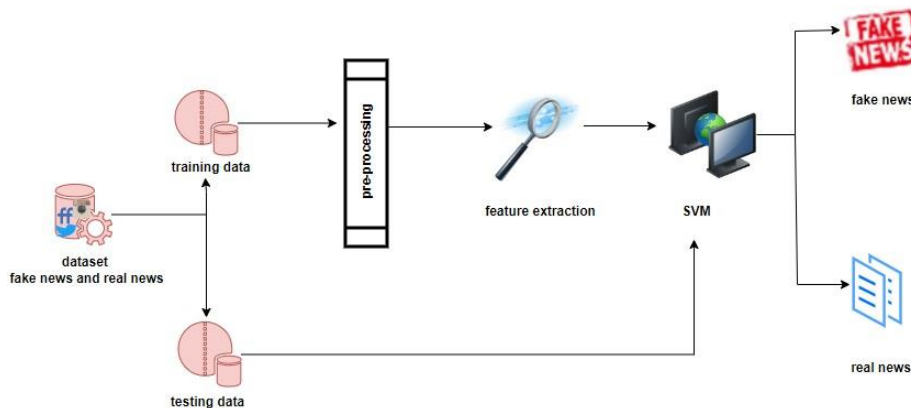


Figure 1: System Architecture Dataset

A dataset is a collection of data. In the case of tabular data, a data set corresponds to one or more database tables, where every column of a table represents a particular variable, and each row corresponds to a given record of the data set in question. The dataset lists values for each of the variables, such as for example title, id, author, label of an object, for each member of the data set. Dataset can also consist of a collection of documents or files. For the purpose of this research the dataset will be collected from <https://www.kaggle.com>. The dataset will contain combination of article of fake and real news.

Pre-Processing

The term Pre-processing the data is defined as the process of converting a data into an understandable format by cleaning it and preparing the text for classification. Texts from online contain usually lots of noise and uninformative parts such as scripts and advertisements. Pre-processing includes several steps such as online text cleaning, white space removal, expanding abbreviation. Stemming, stop words removal and feature Extraction. These will reduce the noise in the text which will help to boost up the performance of the classifier.

Stop Word Removal: Stop words are unimportant words that make noise when employed as text classification features. To link ideas or to aid with sentence structure, these terms are frequently employed in sentences. The following words are regarded as stop words: articles, prepositions, conjunctions, and some pronouns. Common terms like a, about, an, are, as, at, be, by, for, from, how, in, is, of, on, or, that, the, these, this, too, was, what, when, where, who, will, etc. will be eliminated. Each document will have these terms deleted, after which the processed papers will be kept and carried on to the next stage.

Stemming: Tokenizing the data is the first stage; the next is to standardize the tokens. Stemming is the process of restoring words to their original form, which also reduces the amount of word types or classes in the data. For instance, "run" will be used in place of phrases like "running," "Ran," and "runner." Stemming helps us to quickly and effectively classify data. In addition, Porter stemmer, the most widely used stemming algorithm because of its accuracy, will be employed.

Feature Extraction

The data for analysis must be translated into a machine-readable format in order for the computer to comprehend it. Machine learning algorithms have simplified certain human jobs by carrying them out in a manner similar to that of a person, yet they are still unable to comprehend spoken human language. Therefore, this language must be further transformed into numerical stuff that computers can readily understand. The Bag of Words (BoW) model is implemented for this job using machine learning methods. The main function of BoW is to extract characteristics from a document. The purpose of this approach is to create a vocabulary from the text document and track the frequency of each word inside it. The approach solely considers word occurrences, not their placement in the text. The Python Scikit-Learn package provides many methods, including Count vectorizer, Tf-idf vectorizer, and Hashing vectorizer, which vectorize the text with various ways, making the implementation of this model simple. For the purpose of this research Tf-idf vectorizer will be applied for the feature extraction.

Tf-idf Vectorizer: The TF-IDF algorithm weighs the importance of each word in a document. The term frequency (number of times a word appears in a document divided by the total number of words in the document) is first calculated by TF-IDF. However, TF-IDF devalues words that appear in more documents since common words like "the" and "or" typically exist in numerous papers. The only words left in our text analysis are those that are pertinent. The sum of all the TF-IDF values will be 1. Although this function may seem quite basic, it has benefits and influences Google searches. Researchers are able to ascertain the topics and keywords of each article by employing tfidf. The tfidf value formula is provided below. $tf_{i,j}$ stands for the quantity of times that i occurs in j . The total number of documents in the data set is denoted by the letter N . The term " df_i " stands for "documents that contain i ".

$$tf-idf_{i,j} = tf_{i,j} * \log(\frac{N}{df_i})$$

df_i

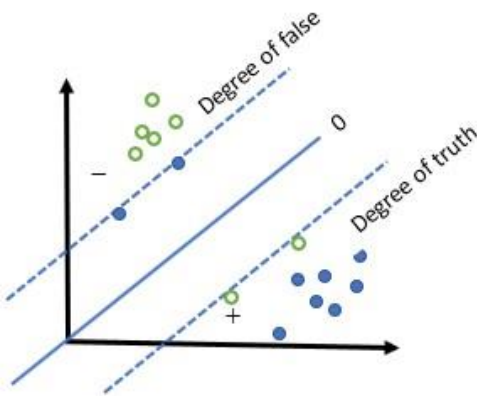
Table 1: Features from article

Feature Name	Description
Number of urls	The number of URLs included in this tweet
Number of tweets	The number of tweets this Twitter user sent
Number of userfavourites	The number of favorites this Twitter user received
account age	The age (days) of an account since its creation until the time of sending the most recent tweet
Number of follower	The number of followers of this Twitter user
Number of retweets	The number of retweets this tweet
Number of hashtag	The number of hashtags included in this tweet
Number of usermention	The number of user mentions included in this tweet

Machine Learning Algorithms

Artificial intelligence (AI) applications such as machine learning allow systems to automatically learn from their experiences and get better over time without having to be explicitly designed. The goal of machine learning is to create computer systems that can access data and utilize it to acquire knowledge on their own. While **TF-IDF** only collects pertinent characteristics from documents, determining if news is fake or true is a quite different challenge. Utilizing support vector machines is one method of doing this. Based on how frequently a feature appeared in the training data for each class, this classifier determines which class a document is most likely to belong to under the presumption that all features are unrelated. The classifier will next select the document's class by examining the portion with the highest likelihood.

Support Vector Machine: A support vector machine represents a data set as points in space divided into categories by a distinct gap or line that stretches as far as feasible. The additional data points are now mapped into the same area and categorized into one of many categories based on which side of the line or separation they land on.

**Figure 2: Graph showing hyperplane**

$$\begin{aligned}
 &Dec \\
 &X \ 100 \quad if \ Dec > 0 \\
 &P = \left\{ \begin{array}{l} \frac{Max_{dec}}{Min_{dec}} \\ \frac{Dec}{Min_{dec}} \end{array} \right. \quad (1)
 \end{aligned}$$

Where:

- Dec is the decision function value;
- Max_{dec} and Min_{dec} are the maximum and minimum values of the decision function;
- p is the percentage of truth or fake
- In SVM, we plot data points as points in an n-dimensional space (n being the number of features you have) with the value of each feature being the value of a particular coordinate.
- Support Vectors – Datapoints that are closest to the hyperplane is called support vectors. Separating line will be defined with the help of these data points.
- Hyperplane – As we can see in the in figure 2, it is a decision plane or space which is divided between a set of objects having different classes. **Evaluation Matrix**

The performance evaluation will be carried out based on the following metrics and is presented as:

i. Accuracy is percentage of the correctly classified records out of the overall total records

TP+TN

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} \quad (2)$$

TP+TN+FP+FN

ii. Sensitivity (recall) is the ratio of actually positive outcome predicted to the total number of the positive outcome that

could have been predicted.

TP

$$Sensitivity = \frac{TP}{TP+FN} \quad (3)$$

TP+FN

iii. Precision is the proportion of the actually positive classifications out of all those classified positive

$$Precision = \frac{TP}{TP+FP} \quad (4)$$

TP+FP Where:

False negative (FN): an outcome where the model incorrectly predicts negative class

True Negative (TN): an outcome where the model correctly predicts negative class

True Positive (TP): it means that the outcome where the model correctly predicts positive class

False Positive (FP): an outcome where the model incorrectly predicts positive class

IV. CONCLUSION

Social media is becoming more and more popular, and more people are choosing to read their news there instead of through traditional news outlets. Social media has also been used to distribute fake information, which has a detrimental effect on both individual users and society as a whole. fake is deliberately spread by individual or group of people with the intention of inciting violence and anger in our society. Youth of today often are those who face the psychological impacts. This paper aim at automating fake news detection to avert the circumstance that might arise as result of the spread of fake news.

V. ACKNOWLEDGEMENT

We are indeed grateful to God Almighty for giving us the wisdom to write this paper. We are thankful to all staffs and lecturers of computer science department Federal College of Agriculture Akure, Ondo State Nigeria (FECA). My sincere appreciate also goes to all the past and present Head of department of Computer science FECA Mr. O.O. Oyeneye, Dr. V.O Onibon and Dr. W.O. Adesanya and finally the provost of the college Dr. A.A. Fadiyimu.

REFERENCES

- Kaliyar, R. K. (2018). *Fake News Detection Using A Deep Neural Network. 2018 4th International Conference on Computing Communication and Automation (ICCCA)*. doi:10.1109/ccaa.2018.8777343
- Tiwari, V., Lennon, R. G., & Dowling, T. (2020). *Not Everything You Read Is True! Fake News Detection using Machine learning Algorithms. 2020 31st Irish Signals and Systems Conference (ISSC)*. doi:10.1109/issc49989.2020.
- 9180206
- Ahmed, H., Traore, I., & Saad, S. (2017). *Detection of Online Fake News Using N-Gram Analysis and Machine Learning Techniques. Intelligent, Secure, and Dependable Systems in Distributed and Cloud Environments*, 127–138. doi:10.1007/978-3-319-69155-8_9
- Granik, M., & Mesyura, V. (2017). *Fake news detection using naive Bayes classifier. 2017 IEEE First Ukraine Conference on Electrical and Computer Engineering (UKRCON)*. doi:10.1109/ukrcon.2017.8100379
- Gupta, H., Jamal, M. S., Madisetty, S., & Desarkar, M. S. (2018). *A framework for real-time spam detection in Twitter. 2018 10th International Conference on Communication Systems & Networks (COMSNETS)*. doi:10.1109/comsnets. 2018.8328222
- Manzoor, S. I., Singla, J., & Nikita. (2019). *Fake News Detection Using Machine Learning approaches: A systematic Review. 2019 3rd International Conference on Trends in Electronics and Informatics (ICOEI)*. doi:10.1109/icoei. 2019.8862770
- Aphiwongsophon, S., & Chongstitvatana, P. (2018). *Detecting Fake News with Machine Learning Method. 2018 15th International Conference on Electrical Engineering/Electronics, Computer, Telecommunications and Information Technology (ECTI-CON)*. doi:10.1109/ecticon.2018.8620051
- Balpande, V., Baswe, K., Somaiya, K., Dhande, A. & Mire, P.(2021). Fake News Detection Using Machine Learning. International Journal of Scientific Research in Computer Science, Engineering and Information Technology ISSN :
- 2456-3307 (www.ijsrcseit.com), Volume 7, Issue 3, Page Number: 533-542, doi : <https://doi.org/10.32628/CSEIT12173115533>

- Vedova, B. M. L., Tacchini, E., Moret, S., Ballarin, G., DiPierro, M., & de Alfaro, L. (2018). *Automatic Online Fake News Detection Combining Content and Social Signals*. 2018 22nd Conference of Open Innovations Association (FRUCT). doi:10.23919/fruct.2018.8468301
- Hiramath, C. K., & Deshpande, G. C. (2019). *Fake News Detection Using Deep Learning Techniques*. 2019 1st International Conference on Advances in Information Technology (ICAIT). doi:10.1109/icaity47043.2019.89872
- Rubin., Victoria, L., et al.: Fake news or truth? Using satirical cues to detect potentially misleading news. In: Proceedings of NAACL-HLT (2016)