

Augmented Fake News Detection Model Using Machine Learning

Arisha Farha¹ and Afsaruddin²

¹PG Student, Department of Computer Science & Engineering, Integral University, Lucknow, U.P., INDIA

²Assistant Professor, Department of Computer Science & Engineering, Integral University, Lucknow, U.P., INDIA

¹Corresponding Author: arishaaz@student.iul.ac.in

Received: 28-05-2023

Revised: 15-06-2023

Accepted: 30-06-2023

ABSTRACT

In today's time, fake news has become like a virus for any social media platform, which destroys the uniqueness of that platform itself. Because a fake news is sent to hurt the sentiments of any person, society or religion. That's why today we need a computer artificial intelligent based model that can detect any fake news before it is posted. All social media platforms have worked in this direction, but somewhere it seems that their model is insufficient to catch such fake news. Because some social media companies have tried to decide whether the news is fake or not on the basis of some predefined datasets. And some companies have searched only on the keywords of the news that the news is fake. This proves that we need a model that is based on the old dataset, and the current news dataset and keywords. Along with this, it is also important to pay attention to the timing, place and type of news, while these things are not taken care of in the existing models. So I would like to include all these parameters in my model to help detect fake news. If we recognize the Fake News at the right time, then we can take the right steps at the right time. Computer based models are not always accurate, so the model should also have the facility to compare with real news. If news is compared with current news then 76% of fake news can be detected at the same time. Therefore, the model should also have the facility of comparative review.

Keywords— Fake News, Genuine, Highlights, Models

I. INTRODUCTION

The emergence of social media marks the emergence of a new media platform for the 21st century. Social media platforms provide a platform people can express their views freely. This is a platform where they can tell the world about themselves, about family, about society, about religion, about their customs. Through social media, they can also tell about any injustice done to them or they can also tell about the injustice done to someone else. And for this it is not necessary that they should contact any reporter of any media house, they can directly connect with the public through this platform. Because those who run news media houses, they run on the basis of their profit and loss and many times it is inspired by news politics or some business. But no one can suppress your voice on social media, because your news directly reaches the public and there is no cut or

editing in it. This means that a common citizen can also be a reporter and social media can act as a virtual platform. But as we all know that a coin does not have only one side, if something has its good, then it also has some drawback. And it is the same with social media as well, there is some good in it and there is some bad. Since Key is an independent platform and editing here is not editing of the news posted by you, that's why some people also take wrong advantage of it. They use social media to spread their wrong feelings in the society, which spreads confusion in the society. And they send wrong facts to prove their wrong point. And similar news comes in the category of Fake News[1,2]. If any fake news is posted on social media, then it hurts the sentiments of any society, person, or religion. And sometimes this fake news takes a very frightening form in the new society and sometimes it provokes riots, stampede. Many times through fake news, politicians try to turn the tide of elections by keeping wrong facts in the public. Due to this, the innocent people are not able to take the right decision of right and wrong and due to which a wrong candidate wins the election. That is why if the public is blindly believing on social media, then it is necessary that social media should maintain its credibility. And this is possible only when social media companies use filters on their platform so that any fake news can be identified before it is transmitted. And we could stop it before it was pressed[3,4,5]. For this, many social media platforms have developed many artificial models, but in many places these models seem to be spreading. Because these models ignore many important facts. Just like it is not said that every criminal leaves behind some evidence, similarly every fake news has its own identity. If any news is fake[6,7], then its place, time, or its relation must have been tampered with. And if any news is true then it must be present on some website somewhere in this vast world of internet. That is why if any fake news is posted, it is necessary to first analyze its category, time, and place. And also it is necessary to see whether this news is present somewhere else or not and on the basis of these parameters we can decide that the news is fake. And for this it is necessary that we set a dataset of real news from Previo. And this dataset is categorized on the basis of timing, location, category, and title of the news. Along with this, we prepare a dataset of current news on the basis of these parameters so that we can compare the posted news with

old facts and current news. And in today's time, there is no better place than Twitter to get current news. We can use Twipy tool for this work, through Twipy we can get word wise news. Because every big personality, politician, media houses must tweet or retweet the news on Twitter. And in this way we can easily get current evidence which we can compare with fake news. In this way we will have two types of datasets one is the data set of old news categorized on the basis of time, place, category. And the second dataset that we have collected through Tweepy[8,9]. After this it is necessary that we filter this dataset because it may also contain some news which is not Grammarly correct. That's why in the next step we can merge both the complete datasets then using NLP will remove the Grammarly Unfit News. Then after this we can classify the dataset using classification algorithms so that we can prepare Trained and Test dataset. Because machine learning model plays a huge role in prediction time trained and test dataset. Here it is also necessary that the trained and test datasets are taken in the correct ratio. Mostly such analysis should have a ratio of 8:2 or 7:3 in which 80% dataset should be trained and 20% dataset should be test. The ratio of the Trained dataset should be high so that the machine can easily decide what the news is like. Both True and False news should be included when creating Trained and Test datasets. Because if we carry only true news then machine will not experience fake news that's why we should include both types of news. After this, we should prepare the model in which the title, time, place, category and content of the news have been made the basis[10]. And on the basis of this model the machine decides whether the news is fake or not. And your model is right that wrong depends on its accuracy and the accuracy of such system should be above 90%. Because any fake news is very harmful for any society, person or religion and fake news ruins their image. If such a model is used on social media, then the machine first reads the news. It then breaks the news into keywords and categorizes the news into titles, times, places, categories, and sources. The news is then filtered through NLP and then classified using the classification algorithm. Then on the base of the Predefine model, the machine predicts whether the new is fake or not and if the news is fake then the machine does not allow it to be posted. If seen carefully, China has done a lot of work in this matter, their tools remove the unfit news according to them very fast[11].

II. METHODOLOGY

Following learning algorithms in conjunction with our proposed methodology to evaluate the performance of fake news detection classifiers. Logistic Regression As we are classifying text on the basis of a wide feature set, with a binary output (true/false or true article/fake article), a logistic regression (LR) model is used, since it provides the intuitive equation to classify

problems into binary or multiple classes [27]. We performed hyperparameters tuning to get the best result for all individual datasets, while multiple parameters are tested before acquiring the maximum accuracies from LR model. Mathematically, the logistic regression hypothesis function can be defined as follows

$$H_{\theta}(x) = \frac{1}{1 + e^{-}} (\beta_0 + B_1 x)$$

Logistic regression uses a sigmoid function to transform the output to a probability value; the objective is to minimize the cost function to achieve an optimal probability. The cost function is calculated as shown in

$$\text{cost}h\theta(x,y) = \begin{cases} \log h \\ -\log(1-h\theta(x)) \end{cases} \text{ where } y = 1^{\theta(x)}$$

DECISION TREE ALGORITHM

Decision tree as the name recommends it is a stream like a tree structure that chips away at the rule of conditions. It is proficient and has solid calculations utilized for prescient investigation. It has predominantly describes that incorporate interior hubs, branches and a terminal hub.

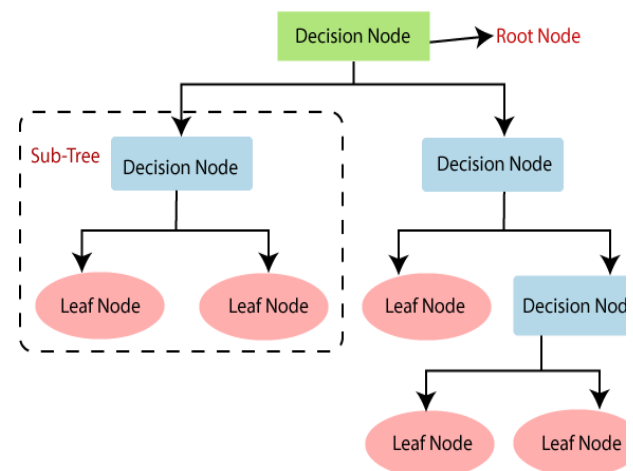


Figure 1: Structure of decision tree algorithm

To aid in decision-making, the multilayer perceptron (MLP) has been employed, which is a widely used technique for classification and regression tasks. The MLP classifier utilizes the probabilistic outputs $p(c)$ of the preceding layer as new features for training. The suggested MLP has four layers: an input layer, two hidden layers, and an output layer with different numbers of neurons, primarily for multi-class classification tasks. The activation of the neurons is accomplished using ReLu functions, while SoftMax functions are employed for classification. The stacked layer of the MLP is trained utilizing the features

extracted from the input features, specifically those for the MLP classifier. The prediction of a specific class, such as fake news, was denoted by $p(c)$ and can be calculated as follows. Let w_i denotes to the neuron i weight, θ represents the weights of the deep learners that were trained in the prior phase, and x_i denotes the corresponding output of the previous layer. Then $p(c)$ which represents the score of correctly forecasting a particular class (such as fake news) can be calculated as follows [12].

$$P(\text{class} - \text{label}) = c = \sum_{i=1}^n w c_i * x_i + \theta$$

Then, the logistic function is calculated for each predicted class as follows.

$$\log_{it} (P_{\text{class label} = c} x_c) = \log \cdot \left(\frac{P_{\text{class label} = c}}{1 - P_{\text{class label} = c}} \right)$$

The x_c logit is the logistic function of the predicted class which maps the values from the range $(-\infty, +\infty)$ into $[0, 1]$. Let $x \rightarrow \{x_1, x_2, x_3, \dots, x_c\}$ vectors contain the scores (x_c) predicted by the MLP for class c , then the final predicted class label is calculated based on the SoftMax function as follows[13].

$$\forall x_c \in \vec{X} \text{ calculate } \frac{e^{x_c}}{\sum_{c=0}^k e^{x_c}} \text{ append } \vec{y}$$

where $V^{\rightarrow}$ is a vector that contains the probability of the input features belonging to the class by finding the maximum probability, the predicted class label is determined.

$$\text{predicted class} = \text{index_of_max}(V^{\rightarrow})$$

The complete process can be better understood using the below algorithm:

- **Branches** - Division of the entire tree is called branches.
- **Root Node** - Represent the entire example that is additionally separated.
- **Parting** - Division of hubs is called parting.
- **Terminal Node** - Node that doesn't part additionally is known as a terminal hub.
- **Decision Node** - It is a hub that additionally gets additionally separated into various sub-hubs being a sub hub.
- **Pruning** - Removal of subnodes from a choice hub.
- **Parent and Child Node** - When a hub gets separated further then that hub is named as parent hub while the partitioned hubs or the

sub-hubs are named as a youngster hub of the parent hub.

NAIVE BAYES

Naive Bayes is a basic directed AI calculation that utilizes the Bayes' hypothesis with solid autonomy suspicions between the highlights to obtain results. That implies that the calculation simply expects that each information variable is free. It truly is an innocent supposition to make about certifiable information. For instance, in the event that you utilize Naive Bayes for assumption examination, given the sentence 'I like Harry Potter', the calculation will take a gander at the individual words and not the full sentence [14].

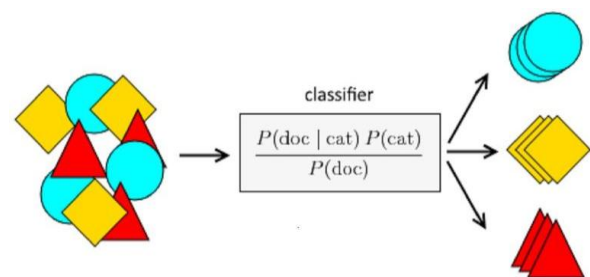


Figure 2: Naive Bayes Classifier

In a sentence, words that remain close to one another impact the significance of one another, and the situation of words in a sentence is additionally significant. Be that as it may, for the calculation, phrases like 'I like Harry Potter', 'Harry Potter like I', and 'Potter I like Harry' are something very similar. Turns out that the calculation can viably take care of numerous perplexing issues. For instance, fabricating a book classifier with Naive Bayes is a lot simpler than with more advertised calculations like neural organizations. The model functions admirably even with lacking or mislabeled information, so you don't need to 'take care of' it a huge number of models before you can receive something sensible in return. Regardless of whether Naive Bayes can take as much as 50 lines, it is extremely powerful. This allows us to analyze the likelihood of an occasion dependent on the earlier information on any occasion that identified with the previous occasion. So for instance, the likelihood that cost of a house is high, can be better evaluated on the off chance that we know the offices around it, contrasted with the appraisal made without the information on the spot of the house. Bayes' hypothesis does precisely that [15,16].

$$P(A | B) = \frac{P(B | A) P(A)}{P(B)}$$

Above condition gives the essential portrayal of the Bayes' hypothesis. Here A and B are two occasions and, $P(A|B)$: the contingent likelihood that occasion A

happens, given that B has happened. This is otherwise called the back likelihood. $P(A)$ and $P(B)$: likelihood of A and B without respect of one another. $P(B|A)$: the contingent likelihood that occasion B happens, given that A has happened. Presently, how about we perceive how this suits well to the reason for AI. Take a basic AI issue, where we need to take in our model from a given arrangement of attributes (in preparing models) and afterward structure a theory or a connection to a reaction variable. Then, at that point we utilize this connection to anticipate a reaction, given ascribes of another occasion [25,26,27]. Utilizing the Bayes' hypothesis, it's conceivable to assemble a student that predicts the likelihood of the reaction variable having a place with some class, given another arrangement of traits. Consider the past condition once more. Presently, accept that A is the reaction variable and B is the information quality. So as per the condition, we have $P(A|B)$: contingent likelihood of reaction variable having a place with a specific worth, given the info ascribes. This is otherwise called the back likelihood. $P(A)$: The earlier likelihood of the reaction variable. $P(B)$: The likelihood of preparing information or the proof. $P(B|A)$: This is known as the probability of the preparation information. Hence, the above condition can be modified as [17]

$$\text{posterior} = \frac{\text{prior} \times \text{likelihood}}{\text{evidence}}$$

How about we take an issue, where the quantity of characteristics is equivalent to n and the reaction is a boolean worth, for example it very well may be in one of the two classes. Likewise, the qualities are categorical (2 classes for our case). Presently, to prepare the classifier, we should compute $P(B|A)$, for every one of the qualities in the case and reaction space. This implies, we should ascertain $2 \times (2^n - 1)$, boundaries for learning this model. This is obviously ridiculous in most pragmatic learning spaces. For instance, on the off chance that there are 30 boolean traits, we should assess multiple billion boundaries [18].

Naive Bayes Algorithm

The intricacy of the above Bayesian classifier should be diminished, for it to be commonsense. The gullible Bayes calculation does that by making a supposition of contingent autonomy over the preparation dataset. This definitely lessens the intricacy of previously mentioned issue to simply $2n$. The suspicion of contingent freedom expresses that, given arbitrary factors X , Y and Z , we say X is restrictively autonomous of Y given Z , if and just if the likelihood circulation overseeing X is autonomous of the worth of Y given Z . At the end of the day, X and Y are restrictively autonomous given Z if and just if, given information that Z happens, information on whether X happens gives no data on the probability of Y happening, and information on whether Y happens

gives no data on the probability of X happening. This suspicion makes the Bayes calculation, innocent. Given, n distinctive characteristic qualities, the probability currently can be composed as [19].

$$P(X_1 \dots X_n | Y) = \prod_{i=1}^n P(X_i | Y)$$

Here, X addresses the characteristics or highlights, and Y is the reaction variable. Presently, $P(X|Y)$ becomes equivalent to the results of, likelihood dispersion of each property X given Y [20].

PASSIVE AGGRESSIVE CLASSIFIER

Passive-Aggressive algorithms are generally used for large-scale learning. It is one of the few 'online-learning algorithms'. In online machine learning algorithms, the input data comes in sequential order and the machine learning model is updated step-by-step, as opposed to batch learning, where the entire training dataset is used at once. This is very useful in situations where there is a huge amount of data and it is computationally infeasible to train the entire dataset because of the sheer size of the data. We can simply say that an online-learning algorithm will get a training example, update the classifier, and then throw away the example [21,22,23].

Passive Aggressive Algorithm remains passive for a correct classification outcome, and turns aggressive in the event of a miscalculation. Its purpose is to make updates that correct the loss, causing very little in the norm of weight vector.

GENERATING NEWS FEATURE VECTOR

A. COUNT VECTORIZER

Count Vectorizer is a representation of text that describes the occurrence of words within a document. It involves two things:

- A vocabulary of known words.
- A measure of the presence of known words.

B. TF-IDF VECTORIZER

TF-IDF is a statistical measure that evaluates how relevant a word is to a document in a collection of documents. This is done by multiplying two metrics: how many times a word appears in a document, and the inverse document frequency of the word across a set of documents [24,25,26].

Hence, obtaining the accuracy and printing the true and false positives and negatives using confusion matrix. Confusion matrix is nothing but a table used to describe the performance of classification model or classifier on set of data set.

$$\text{Accuracy} = \frac{TP+TN}{TP+FP+TN+FN}$$

Confusion Matrix: Basically this metrics represent how many results are correctly predicted and how many results are not correctly predicted [27].

Actual Class	Predicted Class		
		Class = Yes	Class = No
	Class = Yes	True Positive	False Negative
	Class = No	False Positive	True Negative

Figure 3: Model of confusion matrix

III. IMPLEMENTATION STEPS

- [1] Download Dataset from Kaggle website.
- [2] Libraries required for the project.
 - **Numpy:** NumPy is a library for the Python programming language, adding support for large, multi-dimensional arrays and matrices, along with a large collection of high-level mathematical functions to operate on these arrays.
 - **Pandas:** Pandas is a software library written for the Python programming language for data manipulation and analysis. In particular, it offers data structures and operations for manipulating numerical tables and time series.
 - **Sklearn:** Scikit-learn is a free software machine learning library for the Python programming language. It features various classification, regression and clustering algorithms.
 - **Nltk:** The Natural Language Toolkit, or more commonly NLTK, is a suite of libraries and programs for symbolic and statistical natural language processing for English written in the Python programming language.
- [3] Remove all non words from news column
- [4] Remove all stop words using nltk
- [5] Then create features using TfidfVectorizer
- [6] Then create X as input variable having feature and Y as output variable having label 0 for fake and 1 for real.
- [7] Then split data in training and testing.
- [8] Then pass training data to Passive Aggressive Classifier algorithm. Passive Aggressive Classifier belongs to the category of online learning algorithms in machine learning. It works by responding as passive for correct classifications and responding as aggressive for any miscalculation.
 - **Passive:** If the prediction is correct, keep the model and do not make any changes. i.e., the data in the example is not enough to cause any changes in the model.
 - **Aggressive:** If the prediction is incorrect, make changes to the model. i.e., some change to the model may correct it.
- [9] Then we will check accuracy of model and deploy it on flask web framework.

[10] Our web application will take a news as input and will classify it as Fake or Real according to model.

IV. CONCLUSION

Thus it is proved that Fake News is like a curse for social media platforms. And the only way to get rid of it is with a good strong Fakes News Detection Tool. A good fake news detection model must include all the parameters of fake news. Like the timing of fake news, here the timing is important because political fakenews is transmitted only at the time of elections. This is why the place is important because many times fake news is sent to stop the work of any religious place or any development. Along with this, source of fake news means what is the source of news from where the news is being transmitted. And the content of the news also matters. In the fake news model, it is necessary that both true and false news should be included in the dataset and they have a correct ratio. as well as the ratio of the Trined and Test datasets.

REFERENCES

- [1] Abadi, M., Barham, P., Chen, J., Chen, Z., Davis, A., Dean, J., Devin, M., Ghemawat, S., Irving, G., Isard, M., Kudlur, M., Levenberg, J., Monga, R., Moore, S., Murray, D.G., Steiner, B., Tucker, P.A., Vasudevan, V., Warden, P., Wicke, M., Yu, Y., Zhang, X.: Tensorflow: A System for Large-Scale Machine Learning. CoRR abs/1605.08695 (2016), <http://arxiv.org/abs/1605.08695>
- [2] Aghakhani, H., Machiry, A., Nilizadeh, S., Kruegel, C., Vigna, G.: Detecting Deceptive Reviews using Generative Adversarial Networks. CoRR abs/1805.10364 (2018), <http://arxiv.org/abs/1805.10364>
- [3] Srivastava, S., Haroon, M., & Bajaj, A. (2013, September). Web document information extraction using class attribute approach. In *2013 4th International Conference on Computer and Communication Technology (ICCCCT)* (pp. 17-22). IEEE.
- [4] Costin BUSIOC, Stefan RUSSETI, Mihai DASCALU, " A Literature Review of NLP Approaches to Fake News Detection and Their Applicability to Romanian- Language News Analysis", 2020, Romanian Ministry of Education and Research, CNCS - UEFISCDI, project number PN-III- P1-1.1-TE-2019-1794, within PNCDI III.
- [5] Khan, R., Haroon, M., & Husain, M. S. (2015, April). Different technique of load balancing in distributed system: A review paper. In *2015 Global Conference on Communication Technologies (GCCT)* (pp. 371-375). IEEE.
- [6] Khan, W., & Haroon, M. (2022). An

- unsupervised deep learning ensemble model for anomaly detection in static attributed social networks. *International Journal of Cognitive Computing in Engineering*, 3, 153-160.
- [7] Alim Al Ayub Ahmed, Ayman Aljarboub, Praveen Kumar Donepudi, "Detecting Fake News using Machine Learning: A Systematic Literature Review", 2020, IEEE conference.
- [8] Haroon, M., & Husain, M. (2015, March). Interest Attentive Dynamic Load Balancing in distributed systems. In *2015 2nd International Conference on Computing for Sustainable Global Development (INDIACom)* (pp. 1116-1120). IEEE.
- [9] Khan, W., & Haroon, M. (2022). An efficient framework for anomaly detection in attributed social networks. *International Journal of Information Technology*, 14(6), 3069-3076.
- [10] Husain, M. S., & Haroon, D. M. (2020). An enriched information security framework from various attacks in the IoT. *International Journal of Innovative Research in Computer Science & Technology (IJIRCST)* ISSN, 2347-5552.
- [11] Haroon, M., & Husain, M. (2013). Analysis of a Dynamic Load Balancing in Multiprocessor System. *International Journal of Computer Science engineering and Information Technology Research*, 3(1).
- [12] Razan Masood and Ahmet Aker, "The Fake News Challenge: Stance Detection using Traditional Machine Learning Approaches", In *Proceedings of the 10th International Joint Conference on Knowledge Discovery, Knowledge Engineering and Knowledge Management (KMIS 2018)*, pages 128-135
- [13] Khan, W. (2021). An exhaustive review on state-of-the-art techniques for anomaly detection on attributed networks. *Turkish Journal of Computer and Mathematics Education (TURCOMAT)*, 12(10), 6707-6722.
- [14] Husain, M. S. (2020). A review of information security from consumer's perspective especially in online transactions. *International Journal of Engineering and Management Research*, 10.
- [15] Haroon, M., & Husain, M. (2013). Different Types of Systems Model For Dynamic Load Balancing. *IJERT*, 2(3).
- [16] Khan, N., & Haroon, M. (2022). Comparative Study of Various Crowd Detection and Classification Methods for Safety Control System. Available at SSRN 4146666.
- [17] Siddiqui, Z. A., & Haroon, M. (2023). Research on significant factors affecting adoption of blockchain technology for enterprise distributed applications based on integrated MCDM FCEM-MULTIMOORA-FG method. *Engineering Applications of Artificial Intelligence*, 118, 105699.
- [18] Sohan De Sarkar, Fan Yang, "Attending Sentences to detect Satirical Fake News", *Proceedings of the 27th International Conference on Computational Linguistics*, pages 3371-3380 Santa Fe, New Mexico, USA, August 20- 26, 2018.
- [19] Abdullah-All-Tanvir, Ehesas Mia Mahir, Saima Akhter "Detecting Fake News using Machine Learning and Deep Learning Algorithms", 2019 7th International Conference on Smart Computing & Communications (ICSCC).
- [20] Hadeer Ahmed, Issa Traore, and Sherif Saad. Detection of online fake news using n-gram analysis and machine learning techniques. In *International Conference on Intelligent, Secure, and Dependable Systems in Distributed and Cloud Environments*, pages 127-138. Springer, 2017.
- [21] Hunt Allcott and Matthew Gentzkow. Social media and fake news in the 2016 election. *Journal of Economic Perspectives*, 31(2):211-36, 2017.
- [22] Peter Bourgonje, Julian Moreno Schneider, and Georg Rehm. From clickbait to fake news detection: an approach based on detecting the stance of headlines to articles. In *Proceedings of the 2017 EMNLP Workshop: Natural Language Processing meets Journalism*, pages 84-89, 2017. 12
- [23] Yimin Chen, Niall J Conroy, and Victoria L Rubin. Misleading online content: Recognizing clickbait as false news.
- [24] Ali, A. M., Ghaleb, F. A., Mohammed, M. S., Alsolami, F. J., & Khan, A. I. (2023). Web-Informed-Augmented Fake News Detection Model Using Stacked Layers of Convolutional Neural Network and Deep Autoencoder. *Mathematics*, 11(9), 1992.
- [25] Ahmad, I., Yousaf, M., Yousaf, S., & Ahmad, M. O. (2020). Fake news detection using machine learning ensemble methods. *Complexity*, 2020, 1-11.
- [26] Haroon, M., Tripathi, M. M., & Ahmad, F. (2020). Application of Machine Learning In Forensic Science. In *Critical Concepts, Standards, and Techniques in Cyber Forensics* (pp. 228-239). IGI Global.
- [27] Hiramath, C. K., & Deshpande, G. C. (2019, July). Fake news detection using deep learning techniques. In *2019 1st International Conference on Advances in Information Technology (ICAIT)* (pp. 411-415). IEEE.