

International Journal of
Engineering Research and Science & Technology



ISSN : 2319-5991

www.ijerst.com

Email: editor@ijerst.com or editor.ijerst@gmail.com

Twitter Data-Driven Depression Detection using Machine Learning

¹Mrs.Zunaira Begum ²Syed Abdul Rehman, ³ MD Shanawaaz, ⁴ Mir Murtuza Ali Hashmi.

¹Assistant Professor, Department of CSE-AIML, Lords Institute of Engineering & Technology.

²³⁴ Student Department of CSE-AIML, Lords Institute of Engineering & Technology.

Abstract—

These depressive symptoms are plaguing today's society, and the number of those afflicted is steadily rising. Unfortunately, not everyone who is struggling with depression is aware of it; in fact, some of them are able to recognize it. Conversely, many are using social media as a "diary" to document their emotional condition. Using machine learning algorithms, many types of study have been carried out to identify cases of depression in social media posts. Researchers can learn whether social media users are suffering from depression by analyzing the data that is publicly accessible. In order to distinguish between data that is depressed and data that is not, a machine learning algorithm can sort the data into the appropriate categories. The purpose of the proposed study is to identify user sadness using social media data. The Twitter data is then put into two distinct kinds of classifiers, which are Naïve Bayes and a hybrid model, NBTree. The findings will be evaluated based on the greatest accuracy value to select the best algorithm to diagnose depression. The findings demonstrate that both algorithms are equally effective, since they achieve the same degree of accuracy.

Keywords—Depression, Social Media, Twitter, Classification, Hybrid, NBTree, Naïve Bayes

I. INTRODUCTION

Mental health is being negatively impacted by the alarming rise of depression cases. More than 264 million individuals throughout the globe suffer from depression, making it another significant condition [1]. There are several reasons such as unexpected changes in surrounds, changes in the brain's neurotransmitter level due to some sensation assaults, or even hereditary traits [2]. Medication or therapy

sessions may alleviate depression. Due to a lack of understanding of depression, a large number of individuals continue to go untreated, despite the fact that many sufferers are receiving therapy. As a result, patients may have negative ideas including withdrawing from society, acting irrationally, contemplating suicide, and becoming too reliant on medicinal substances like antidepressants. A loss of motivation and interest in once-enjoyed activities is one way in which depression may impair everyday functioning. If depression is not identified, it might eventually cause damage to the person's body and brain. Thanks to technological advancements, almost 3.8 billion individuals are now active social media users [3]. Evolvement and trendsetting of social media among individuals welcome the users to communicate their feelings, ideas, and opinion instantaneously. For others, the ease and anonymity of social media has made it the perfect outlet for the repressed emotions and ideas that plague their brains. Users of social media applications are given a platform to express themselves verbally, whether through the sharing of opinions or brief notes pertaining to their feelings and emotions. Healthcare, entertainment, politics, perspective communication, and tourism are just a few of the real-world intelligent applications that might benefit from the content published on social media [4]. Users' material or points provide academics with useful data for studying mental states [5]. As social media apps continue to flourish, more and more individuals are choosing to share their mental health issues with the world on these platforms rather than keeping them private. It is possible to assess the user's mental health and identify areas for improvement by analyzing the data stored in their account. The most common definition of Twitter is a microblogging service that facilitates the posting of brief narratives (no more than 140 characters) using user-friendly interfaces. Nearly every user's tweets are fully

accessible via their personal Twitter API, making them totally public. Using Twitter's API, you may execute advanced searches, such as retrieving all tweets mentioning a certain subject [6]. Twitter helps us to comprehend and supports more than 35 languages throughout the globe. Twitter provides a place for individuals to analyze and comprehend global or local trends and happenings. The Twitter API allows Twitter to reap substantial benefits [7]. Consequently, researchers may tell whether someone is sad only by looking at their publicly accessible tweets. A text's positive, negative, or neutral tone may be determined by sentiment analysis. Each word in the text will be given a score by a sentiment analysis system according to the predetermined polarity. You can tell whether the person is happy or sad just by looking at their face. We assign a positive, negative, or neutral sentiment score to each tweet using the sentiment analysis approach. Twitter sentiment analysis uses sentiment scores to quantify the user's opinion of a tweet or remark. Following the application of sentiment analysis to tweets, a machine learning system can determine whether the tweets are depressed or not by analyzing the labels attached to the sentiment scores. Using the data collected from tweets and an algorithm trained on sentiment ratings, the researchers were able to identify cases of depression. How well the machine learning system detects is determined by how accurate it is. The objective of the article is to identify cases of depression using tweets via the use of two distinct machine learning methods, namely Naïve Bayes and a hybrid model called NBTtree. Here is the structure of the remaining portion of the paper: Section II summarizes the associated work. Section III lays forth the suggested structure and Section IV concludes the results, while section V summarizes them.

II. LITERATURE REVIEW

This section will explain what depression is, go over current depression detection systems that use different kinds of machine learning algorithms, and then identify where the current work is lacking and how the proposed work would fill those gaps. Multiple sections make up the section.

A. Overview of Depression

When people talk about common health conditions, depression often comes up since it is a mental health problem. Serious difficulties, including the danger to one's life, may arise from ignoring or not treating depression. A complex interplay of social, biological, and psychological variables initiates the early stages of depression. The reality of depression may lead to major and complex issues [8]. There are seven distinct forms of depression, with clinical depression and bipolar illness serving as the primary umbrella terms [9]. For around two weeks, the patient will experience the symptoms of major depressive illness, at which time they will feel

1. Irritability, worthlessness, and feelings of powerlessness
2. Either eat less and lose weight, or eat more and put on weight
3. Death contemplation and attempted suicide
4. Having trouble sleeping, waking up too early, or sleeping too much
5. Feeling down almost all the time
6. Decline of enthusiasm for once enjoyed pastimes

Mood fluctuations from manic or hypomanic episodes to depressive episodes are the driving force behind manic depression. Mild to severe mood fluctuations that last a few days to a few weeks or more are possible. Either extremes—too active or too depressed—may characterize the person. When someone is manic, they may exhibit improper social conduct, excessive ecstasy, lack of sleep, excessive talkativeness, racing thoughts, hyperenergy, and bad judgment.

B. Machine learning application on detect depression via social media

Through data exploration, machine learning is able to discover fascinating patterns and new information. Researchers in the past have used publicly accessible Facebook data to do depression analyses. Based on the linguistic style and emotional tone of the words

used, the researchers conducted the study [10]. In addition, researchers used the SVM method with various kernels to perform classifications, and the results demonstrate that the technique achieves superior accuracy [11]. In 2016, another researcher by the name of Nadeem played around with the idea of leveraging Twitter data to identify Major Depressive Disorder using the Naïve Bayes and SVM algorithms. According to the final findings, Naïve Bayes is more effective than SVM [12]. Depressive disorder detection using Twitter data also makes use of hybrid model machine learning. According to [13], the sentiment analysis job is successfully completed by the naïve Bayes-SVM hybrid model. From 2016 to 2020, the researchers in Table I summarized their work on a distinct algorithm and its limitations. In order to identify depressive symptoms, the majority of the researchers used only one machine learning algorithm. When trying to figure out which algorithm is best for a certain situation, researchers are keeping an eye on the accuracy number. The researchers also utilized a single one-size dataset for every method. Overall, it reveals that Naïve Bayes is outperformed in most of the study trials with remarkable accuracy. The aforementioned limitations of the papers are both a weakness and an opportunity to highlight the algorithm's performance-enhancing capabilities.

III.RESEARCH METHODOLOGY

The framework begins with data collection by using Twitter scrapper tools and is stored in a .csv file. The first step in data pre-processing is cleaning up the raw data. The data will be tokenized, stemming, and lemmatization as a part of normalizing the data. Next, the data will be analyzed by using sentiment analysis to obtain a score of words. The data are fed into two different classifiers which are Naïve Bayes and NBTree.

Research Paper	Algorithm Approach	Findings	Limitation
[12]	- Decision Tree - Linear SVM - Naïve Bayes - Logistic Regression	Naïve Bayes gives high accuracy	- Usage of the old dataset - Focus on user confession
[14]	- SVM - Naïve Bayes	Naïve Bayes gives high accuracy	The dataset is in Arabic
[5]	- Decision Tree - K-Nearest Neighbour - SVM	The decision tree gives the best results	The data focus on the comments
[13]	NB-SVM hybrid model	The accuracy of the model is 85%	Takes time to compare the long-short snippets.
[11]	- SVM - Decision Tree - Naïve Bayes	SVM gives high accuracy	- Usage of the different kernel in SVM - Cannot avoid overfit data

A train set and a test set are created from the data. In order to guarantee that the classifier learns, model development makes use of the training data. The model will be fed the test data once it has learnt about the evaluation data. We will compare the two classifiers' accuracy results to see which approach works better. The following flowchart, based on Fig. 1, illustrates the entire implementation structure.

A. Data Collection

The data is collected in two different volumes which are 1000 data and 3000 data via Tweepy. You may access the Twitter API with the help of the Tweepy Python package. In order for Tweepy to extract the data, the Twitter Developer must provide the necessary API requirements. In order to provide access to the API and token for data extraction, the Twitter developer need certain details. The data is then extracted and stored in the format of CSV as compatible for analysis. A single column in the dataset, "text," stores all of the tweets retrieved from Twitter.

B. Data Preprocessing

The data collection is imported into Jupyter Notebook, an open-source code editor that can be accessed via the Anaconda command line. Removed from the dataset were URLs, retweets, mentions, and stopwords.

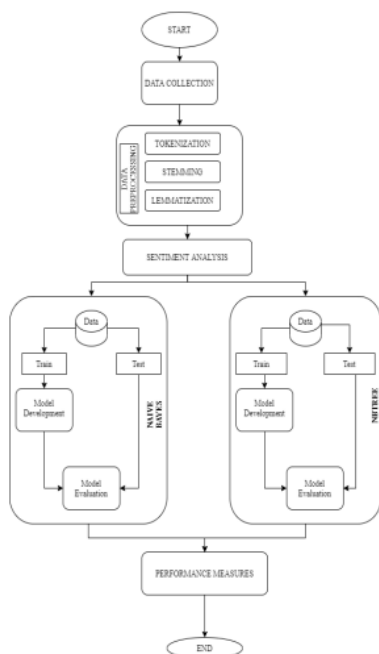


Fig.1. General flow of framework

After that, the dataset is tokenized by dividing the phrase into individual words or tokens and storing them in an array. Following tokenization, stemming and lemmatization are applied to the sentences. The process of stemming involves removing affixes from a word in order to recover its original form. The word "walking" will become "walk" as an example of a normalizing word. By using lemmatization, a word may be reduced to its root form. It uses the definition of a word to determine its lemma. Stemming differs from lemmatization in that it discards the last few letters, such as "s" or "-ing," which might cause misspellings and ambiguities in meaning. But lemmatization knows what the term is and gets it back to its useful basic form.

C. Sentiment Analysis

The next step is to transform the cleaned dataset into a string format. After the data is cleaned, TextBlob is

used to assess sentiment. Python users may easily count sentiment scores together with topic and polarity by importing the Textblob package. Using polarity and subjectivity, we rate every statement in the dataset. The polarity score is then assigned to new properties using the -1, 0 and 1 criteria. Labeled data is defined by the polarity score; data is considered negative if the score is less than 0. The score is considered neutral if it is precisely 0. A positive score is one that is more than zero. According to the polarity score, a phrase is not depressed if its score is greater than zero. A depressing phrase, on the other hand, has a score lower than zero.

D. Classification The scored dataset is fed into the WEKA.

The open-source graphical user interface (GUI) program WEKA has a library of machine learning algorithms for use in data mining and analysis. The experiment's train-test split ratio is 70:30. The proportion is used to divide the scored dataset. Data containing labels is used to train and evaluate the classifiers. For Weka, the Naïve Bayes and NBTree classifiers are used. The technique is done for the 3000 data dataset as well to discover any major differences. a) The Naïve Bayes algorithm: This statistical classification method begins with the Bayes Theorem and assumes that each characteristic inside a class is completely separate from all other features. According to the Naive Bayes classifier, each feature's impact on a class is considered to be separate from that of any other features. Naïve Bayes is great at predicting numerous classes and uses a probabilistic strategy to forecast the test data's class. However, the issue of Naïve Bayes is that assumes the values are independent among the characteristics which might not apply in the real-life scenario. The Naïve Bayes model is shown in Figure 2, and the formula of the Bayes is shown in Equation (1). Let note C is a class label and x_n are the feature inputs

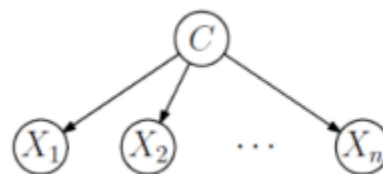


Fig. 2. Naïve Bayes Model [15]

$$P(c|x) = \frac{P(x|c)P(c)}{P(x)} \quad (1)$$

- $P(c|x)$ = the posterior probability of the class where C is the target and predictor x is attributes.
- $P(c)$ = the class's prior probability
- $P(x|c)$ = The probability of predictor per class the class known as the likelihood
- $P(x)$ = predictor's prior probability

Naive Bayes classifier calculates the probability by following some simple steps:

Step 1 - Determine the prior probability using the class labels provided.

Step 2 - Find the possibility of each attribute in the dataset for each class.

Step 3- Use the Bayes formula to get the posterior probability by plugging in values.

Step 4 - The higher probability class is the one that, given the input, has the best chance of being true. b) NBTree : It is a hybrid model made up of Naïve Bayes and Decision Tree.. An open-source program called WEKA uses the algorithm that Ron Khavi described. The NBTree method resembles a decision tree with Naïve Bayes characteristics at the node representing the leaves. The parent node is a property that separates the dataset value into two groups. Every route from the root to the leaf node has its own rule created for it. Each branching condition serves as an antecedent rule, and the leaf determines the consequent rule's class. Figure 3 displays the NBTree model. The oval form is the feature that divides the data into a more specific set, while the rectangular form is the leaf node that the Naïve Bayes classifier uses for classification. In this case, the "v" stands for the node's attribute value. The Naïve Bayes Model's insight is shown in Figure 4. The nominal value and characteristics are shown in the first column. This column keeps track of the numerical characteristics' class values and frequency counts or normal

distribution values. Finding the algorithm that is most effective at identifying depression is the last step. A measure of accuracy is the proportion of classes that were properly predicted out of all the forecasts. Each measure of classification performance undergoes the computation. Once the optimal method for depression detection has been determined, the flow will conclude. The confusion matrix, which serves as an assessment measure, is shown in Figure 5. One way to determine which algorithm produces the best classifier for a given collection of problems is to use a confusion matrix, which summarizes the method's performance. Using the confusion matrix, we may provide a method for determining an algorithm's optimal accuracy. The confusion matrix displays the actual class in the rows and the expected class in the columns.

The confusion matrix is divided into 4 parts which are True Positive, True Negative, False Positive, and False Negative.

- When a model accurately predicts a positive class, it is said to be a true positive.
- "False Positive": this happens when a model's predictions for the positive class are off.
- By "True Negative," we mean the case where the model accurately predicts a negative class.

In model training, a false negative occurs when the negative class predictions are incorrect. In order to determine the algorithms' accuracy, precision, and recall, these four components are crucial. You may use the confusion matrix for multi-class as well as binary classification.\

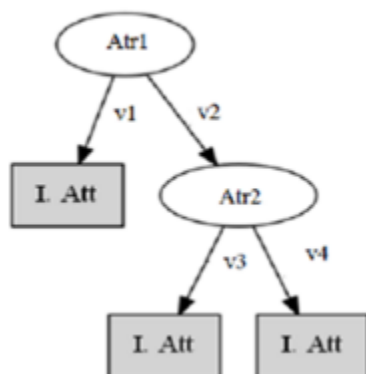


Fig. 3. NBTree approach [16]

Atr - i	Class = 1	Class = 2	...	Class = m
val-1	FC11	FC21		FCm1
val-2	FC12	FC22		FCm2
...				
val-n	FC1n	FC2n		FCmn

Fig. 4. Null Bayes leaf model [16]

		Actual Values	
		Positive (1)	Negative (0)
Predicted Values	Positive (1)	TP	FP
	Negative (0)	FN	TN

Fig. 5. Matrix of Perplexity[17]

The confusion matrix calculation formula table is shown in Table 1. Precision, Accuracy, and Recall may all be determined using the confusion matrix. The ratio of accurately anticipated observations to the total observations is known as accuracy, and it is a standard performance metric. Return is the percentage of positive observations that were accurately predicted relative to the total number of observations in the actual class, where the class is

yes/positive/1. One measure of accuracy is the proportion of positive observations that were accurately predicted relative to the total number of positive observations. The following table provides the formula for the computation.

IV. RESULTS AND DISCUSSIONS

In order to find the accuracy score, the section covered the outcomes of feeding cleaned Twitter datasets into the machine algorithm. The bar chart in Figure 6 displays the split of the 1000 tweets dataset according to their label score. The number of tweets is broken down as follows: 596 favorable, 303 neutral, and 98 negative. Figure 7 is a bar chart that displays the classification of the 3000 tweet dataset according to their label score. In all, there are 1,730 positive tweets, 915 neutral ones, and 345 negative ones. For the 1000 data file dataset, the test set for Naïve Bayes and NBTree is 250, while for the 3000 data file dataset, it is 750. Weka is used for loading datasets and implementing the framework.

Name	Formula
Accuracy	$\frac{TP + TN}{TP + TN + FP + FN}$
Recall	$\frac{TP}{TP + FN}$
Precision	$\frac{TP}{TP + FP}$

TABLE II. FORMULA FOR CONFUSION MATRIX [17]

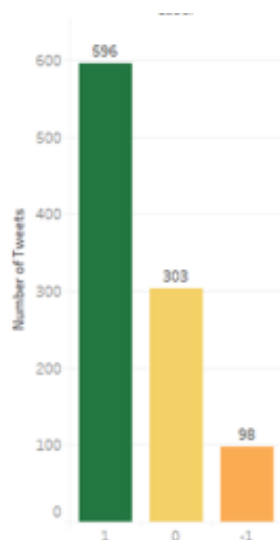


Fig. 6. Analyzed 1000 tweets according to their labels

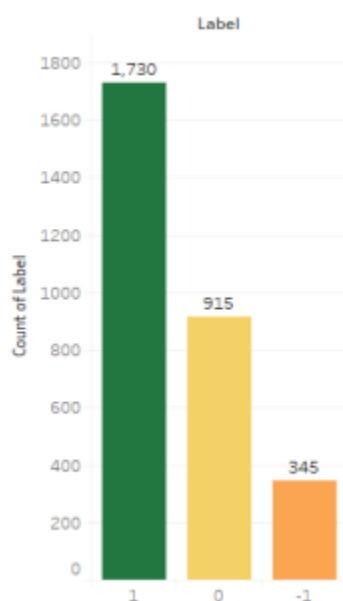


Fig. 7. Analyzed 1000 tweets according to their labels

Figure 8 displays the production outcomes acquired from the 1000 tweets dataset that is labeled using Naïve Bayes. Classification accuracy is 92.34%, according to the results. The findings for the 1000 tweets dataset categorized using NBTree are shown in Fig. 9. Classification accuracy is 92.34%, according to the results. In Figure 10, we can see the outcomes of applying Naïve Bayes classification on a dataset consisting of three thousand tweets. The result reveals that 97.31 percent of the data is appropriately labeled. Results for the 3000 tweets dataset identified using NBTree are shown in Fig. 11. The result reveals that 97.31 percent of the data is appropriately labeled.

Summary

Correctly Classified Instances	253	92.3358 %
Incorrectly Classified Instances	21	7.6642 %
Kappa statistic	0.8561	
Mean absolute error	0.1666	
Root mean squared error	0.2608	
Relative absolute error	44.9835 %	
Root relative squared error	59.8143 %	
Total Number of Instances	274	

Fig. 8. A thousand tweets' worth of output from Naïve Bayes

Summary

Correctly Classified Instances	253	92.3358 %
Incorrectly Classified Instances	21	7.6642 %
Kappa statistic	0.8561	
Mean absolute error	0.1666	
Root mean squared error	0.2608	
Relative absolute error	44.9835 %	
Root relative squared error	59.8143 %	
Total Number of Instances	274	

Fig. 9. The output of NBTree with one thousand tweets

Summary

Correctly Classified Instances	796	97.3105 %
Incorrectly Classified Instances	22	2.6895 %
Kappa statistic	0.951	
Mean absolute error	0.0885	
Root mean squared error	0.172	
Relative absolute error	23.425 %	
Root relative squared error	39.7121 %	
Total Number of Instances	818	

Fig. 10. The outcome of using Naïve Bayes with three thousand tweets

```

=== Summary ===

Correctly Classified Instances      796      97.3105 %
Incorrectly Classified Instances    22      2.6895 %
Kappa statistic                    0.951
Mean absolute error                 0.0885
Root mean squared error             0.172
Relative absolute error             23.425 %
Root relative squared error         39.7121 %
Total Number of Instances          818
  
```

Fig. 11The output of NBTree with three thousand tweets

Table III displays the findings derived from the categorization model. The accuracy value for the 1000 data file is the same for NBTree and Naïve Bayes. Both NBTree and Naïve Bayes provide the same accuracy value for the 3000 data set. Overall, the NBTree and Naïve Bayes are both equally have a larger accuracy value

TABLE III ACCURACY OF THE CLASSIFICATION

Algorithm	Accuracy	
	1000 data	3000 data
Naive Bayes	92.34	97.31
NBTree	92.34	97.31

V. CONCLUSION

Finally, mental health is being negatively impacted by the epidemic of depression. Aside from that, there is a deluge of data available on one's social media accounts due to the rapid expansion of social media. Twitter is an app that lets users express themselves concisely and immediately via the use of 140-character tweets. The researchers may use each tweet to extract and evaluate the information supplied inside. We assign a positive, negative, or neutral sentiment score to each tweet using the sentiment analysis approach. To classify the tweets into the appropriate categories, a machine learning algorithm is given the tagged tweets. The chosen machine learning algorithms, Naive Bayes and NBTree, were

tested on two twitter datasets of varying sizes to see how well they could distinguish between tweets that were depressed and those that were not. The NBTree algorithm offers an accuracy of 97.31% to categorize the depressed and non-depressive tweets on the 3000 tweets dataset and 92.34% on the 1000 tweets dataset. Conversely, when it comes to datasets with 3000 tweets, Naïve Bayes has a success rate of 97.31% and with 1000 tweets, it drops to 92.34%. In terms of experimental accuracy, Naïve Bayes and NBTree are on par with each other, indicating equivalent efficiency. But for now, we're only working on the wording for this project. Improving the work in the future may be achieved by analyzing the tweets of a certain person at a specific moment to identify if they are depressed or not.

REFERENCES

- [1] "Depression," World Health Organization. https://www.who.int/health-topics/depression#tab=tab_1 (accessed Sep. 09, 2020).
- [2] L. Goldman, "Depression: What it is, symptoms, causes, treatment, types, and more," 2019. <https://www.medicalnewstoday.com/articles/8933> (accessed Oct. 10, 2020).
- [3] S. Kemp, "Digital 2020: 3.8 billion people use social media - We Are Social UK - Global Socially-Led Creative Agency," Jan. 30, 2020. <https://wearesocial.com/blog/2020/01/digital-2020-3-8-billion-people-use-social-media> (accessed Oct. 10, 2020).
- [4] Y.-T. Wang, H.-H. Huang, and H.-H. Chen, "A Neural Network Approach to Early Risk Detection of Depression and Anorexia on Social Media Text."
- [5] M. R. Islam, M. A. Kabir, A. Ahmed, A. R. M. Kamal, H. Wang, and A. Ulhaq, "Depression detection from social network data using machine learning techniques," *Heal. Inf. Sci. Syst.*, vol. 6, no. 1, pp. 1–12, 2018, doi: 10.1007/s13755-018-0046-0.
- [6] A. Sistilli, "Twitter Data Mining: Analyzing Big Data Using Python | Toptal."

<https://www.toptal.com/python/twitter-datamining-using-python> (accessed Feb. 04, 2021).

[7] Fontanella Clint, "How to Get, Use, & Benefit From Twitter's API." <https://blog.hubspot.com/website/how-to-use-twitter-api> (accessed Mar. 05, 2021).

[8] "Depression." <https://www.who.int/news-room/factsheets/detail/depression> (accessed Oct. 10, 2020).

[9] N. Schimelpfening, "7 Most Common Types of Depression," Aug. 26, 2020. <https://www.verywellmind.com/common-types-of-depression-1067313> (accessed Sep. 09, 2020).

[10] M. R. Islam, A. R. M. Kamal, N. Sultana, R. Islam, M. A. Moni, and A. Ulhaq, "Detecting Depression Using K-Nearest Neighbors (KNN) Classification Technique," Int. Conf. Comput. Commun. Chem. Mater. Electron. Eng. IC4ME2 2018, pp. 4–7, 2018, doi: 10.1109/IC4ME2.2018.8465641.

[11] H. S. ALSAGRI and M. YKHLEF, "Machine Learning-Based Approach for Depression Detection in Twitter Using Content and Activity Features," IEICE Trans. Inf. Syst., vol. E103.D, no. 8, pp. 1825–1832, 2020, doi: 10.1587/transinf.2020edp7023.

[12] M. Nadeem, "Identifying Depression on Twitter," pp. 1–9, 2016.

[13] M. Gaikar, J. Chavan, K. Indore, and R. Shedge, "Depression Detection and Prevention System by Analysing Tweets," SSRN Electron. J., pp. 1–6, 2019, doi: 10.2139/ssrn.3358809.

[14] M. M. Aldarwish and H. F. Ahmad, "Predicting Depression Levels Using Social Media Posts," Proc. - 2017 IEEE 13th Int. Symp. Auton. Decentralized Syst. ISADS 2017, pp. 277–280, 2017, doi: 10.1109/ISADS.2017.41.

[15] P. Kaviani and S. Dhotre, "International Journal of Advance Engineering and Research Short Survey on Naive Bayes Algorithm," Int. J. Adv. Eng. Res. Dev., vol. 4, no. 11, pp. 607–611, 2017.

[16] T. M. Christian and M. Ayub, "Exploration of classification using NBTree for predicting students' performance," Proc. 2014 Int. Conf. Data Softw. Eng. ICODSE 2014, no. November 2017, 2014, doi: 10.1109/ICODSE.2014.7062654.

[17] S. Narkhede, "Understanding Confusion Matrix | by Sarang Narkhede | Towards Data Science," May 09, 2018. <https://towardsdatascience.com/understanding-confusion-matrix-a9ad42dcfd62> (accessed Feb. 04, 2021)