

Depression Detection Using Machine Learning Techniques on X/Twitter Data

Mohammed Faizan Ahmed, Mohammed Arsalan Ansari, Abdul Qadir, Dr. Mohammed Jameel Hashmi

^{1,2,3} B.E Students, Department Of CSE, ISL Engineering College, India.

⁴ HOD & Associate Professor Department Of CSE, ISL Engineering College, India

faizanahmedd246@gmail.com

ABSTRACT

Depression has become a serious problem in this current generation and the number of people affected by depression is increasing day by day. However, some of them manage to acknowledge that they are facing depression while some of them do not know it. On the other hand, the vast progress of social media is becoming their “diary” to share their state of mind. Several kinds of research had been conducted to detect depression through the user post on social media using machine learning algorithms. Through the data available on social media, the researcher can able to know whether the users are facing depression or not. Machine learning algorithm enables to classify the data into correct groups and identify the depressive and non-depressive data. The proposed research work aims to detect the depression of the user by their data, which is shared on social media. The Twitter data is then fed into two different types of classifiers, which are Naïve Bayes and a hybrid model, NBTree. The results will be compared based on the highest accuracy value to determine the best algorithm to detect depression. The primary objective of this proposed research work is to detect depression through users’ posts shared on social media platforms, with a specific focus on Twitter. Twitter was chosen due to its open-access data policies and the concise nature of user posts, making it a rich source of emotional and behavioral indicators. In this study, the collected Twitter/X data is pre-processed and then fed into two different classification models: the Naïve Bayes classifier and a hybrid model known as NBTree, which combines the Naïve Bayes algorithm with decision tree methodologies to enhance prediction capabilities.

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

1-INTRODUCTION

Depression has become a serious problem affecting mental health globally. It is recognized as a major

illness impacting more than 264 million people worldwide . There are multiple factors that can lead to depression, including sudden changes in surroundings, imbalances in brain neurotransmitters due to emotional distress, and genetic influences . Depression can be treated through therapy sessions or medication. However, despite the availability of treatments, a significant number of people remain undiagnosed due to a lack of awareness about mental health. Undiagnosed depression can lead to severe consequences such as self-isolation, risky behavior, suicidal thoughts, and dependency on anti-depressant medications. It can also negatively impact daily activities, reducing focus and interest, and over time, can cause long-term damage to both the body and brain if left untreated With the advancement of technology, the number of social media users has grown significantly, reaching approximately 3.8 billion people globally . Social media platforms encourage users to share their moods, emotions, and opinions instantly, often serving as an outlet for feelings they find difficult to express otherwise. These platforms provide a valuable space for users to document their emotional states in written form, which in turn, offers researchers valuable insights into user mental health. The text shared on social media can be analyzed and utilized for various intelligent applications, such as in healthcare, entertainment, political communication, and tourism blood. As social media evolved, users became more comfortable sharing their mental health struggles publicly. This trend provides an opportunity for researchers to analyze available social media data to detect users' mental health status and potentially offer early interventions to aid in recovery. Twitter, a popular micro-blogging platform, allows users to post short messages, or "tweets," limited to 140 characters. Its open nature — with most tweets being publicly accessible — makes it an ideal source for collecting user data using the Twitter API. This API supports complex queries, such as pulling tweets related to specific topics , and offers multilingual support in over 35 languages. Given its global reach and ease of data accessibility, Twitter has become a powerful platform for analyzing trends and events on both

local and global scales. By extracting and analyzing Twitter data, researchers can detect signs of depression among users based on the sentiments expressed in their tweets. Sentiment analysis is the

method used to determine the tone of a text — whether it is positive, negative, or sentiment score, helping to classify the emotional state of the user.

2-LITERATURE SURVEY

Title: Depression Detection Using Machine Learning Techniques on Twitter Data

AUTHORS: Kuhaneswaran A/L Govindasamy, Naveen Palanichamy.

Year: 2021

Depression has become a serious problem in this current generation and the number of people affected by depression is increasing day by day. However, some of them manage to acknowledge that they are facing depression while some of them do not know it. On the other hand, the vast progress of social media is becoming their “diary” to share their state of mind. Several kinds of research had been conducted to detect depression through the user post on social media using machine learning algorithms. Through the data available on social media, the researcher can able to know whether the users are facing depression or not. Machine learning algorithm enables to classify the data into correct groups and identify the depressive and non-depressive data.

Title: Depression Detection Using Machine Learning

AUTHORS: Punam Vishnu Chavan, Aishwarya Masne, Sanjana Nadgouda.

Year: 2023

According to the World Health Organization, depression is expected to be the second leading cause of disability by 2030. Depression is a state of mental illness. It is characterized by long-lasting feelings of sadness and despair. Most people with depression do not report that they are depressed. If a person remains grieving for a very long time, the person may be called depressed. Such a person needs the help of a psychiatrist to make the correct diagnosis. You can check the emotions of people by their facial expressions. Facial expressions are very useful for examining a person's emotional state. So this project will help such people check for depression in themselves. To see if this person is sad most of the time, we can assume that he is a depressed person. Once this is confirmed, a correct diagnosis can be made. Depression can be recognized by facial expressions and through texts. Deep learning algorithms can help us understand a person's emotions better by analyzing their facial expressions. In this article, we proposed a CNN model for analyzing human emotions.

Title: A textual-based featuring approach for depression detection using machine learning classifiers and social media texts

AUTHORS: Raymond Chiong, Gregorius Satia Budhi, Sandeep Dhakal, Fabian Chiong.

Year: 2021

Depression is one of the leading causes of suicide worldwide. However, a large percentage of cases of depression go undiagnosed and, thus, untreated. Previous studies have found that messages posted by individuals with major depressive disorder on [social media platforms](#) can be analysed to predict if they are suffering, or likely to suffer, from depression. This study aims to determine whether [machine learning](#) could be effectively used to detect signs of depression in [social media users](#) by analysing their social media posts—especially when those messages do not explicitly contain specific keywords such as ‘depression’ or ‘diagnosis’. To this end, we investigate several text preprocessing and textual-based featuring methods along with [machine learning](#) classifiers, including single and ensemble models, to propose a generalised approach for depression detection using social media texts. We first use two public, labelled Twitter datasets to train and test the machine learning models, and then another three non-Twitter depression-class-only datasets (sourced from Facebook, Reddit, and an electronic diary) to test the performance of our trained models against other social media sources. Experimental results indicate that the proposed approach is able to effectively detect depression via social media texts even when the training datasets do not contain specific keywords (such as ‘depression’ and ‘diagnose’), as well as when unrelated datasets are used for testing.

Title: Machine Learning for Depression Detection on Web and Social Media:

AUTHORS: Lin Gan, Yingqui Guo

Year: 2024

Depression, a significant psychiatric disorder, affects individuals' physical well-being and daily functioning. This focused analysis provides a comprehensive exploration of contemporary research conducted between 2012 and 2023 that delves into the utilization of sophisticated machine learning methodologies aimed at identifying correlates of depression within social media content. Our study meticulously dissects various data sources and performs a comprehensive examination of different machine learning algorithms cited in the researched articles and literature, aiming to pinpoint an approach that can enhance detection accuracy. Furthermore, we have scrutinized the use of varied data from social media platforms and pinpointed emerging trends, notably spotlighting novel applications of artificial neural networks for image processing and classification, along with advanced gait image

models. Our results offer essential direction for future research focused on enhancing detection precision, acting as a valuable reference for academic and industry scholars in this field.

Title: Machine Learning-based Approach for Depression Detection in Twitter Using Content and Activity Features

AUTHORS: Hatoon AlSagri, Mourad Ykhlef

Year: 2021

Social media channels, such as Facebook, Twitter, and Instagram, have altered our world forever. People are now increasingly connected than ever and reveal a sort of digital persona. Although social media certainly has several remarkable features, the demerits are undeniable as well. Recent studies have indicated a correlation between high usage of social media sites and increased depression. The present study aims to exploit machine learning techniques for detecting a probable depressed Twitter user based on both, his/her network behavior and tweets. For this purpose, we trained and tested classifiers to distinguish whether a user is depressed or not using features extracted from his/her activities in the network and tweets. The results showed that the more features are used, the higher are the accuracy and F-measure scores in detecting depressed users. This method is a data-driven, predictive approach for early detection of depression or other mental illnesses. This study's main contribution is the exploration part of the features and its impact on detecting the depression level.

Title: Prediction of Mental Disorder for employees in IT Industry

AUTHORS: Sandhya P, Mahek Kantesaria

Year: 2019

Mental health is nowadays a topic which is most frequently discussed when it comes to research but least frequently discussed when it comes to the personal life. The wellbeing of a person is the measure of mental health. The increasing use of technology will lead to a lifestyle of less physical work. Also, the constant pressure on an employee in any industry will make more vulnerable to mental disorder. These vulnerabilities consist of peer pressure, anxiety attack, depression, and many more. Here we have taken the dataset of the questionnaires which were asked to an IT industry employee. Based on their answers the result is derived. Here output will be that the person needs an attention or not. Different machine learning techniques are used to get the results. This prediction also tells us that it is very important for an IT employee to get the regular mental health check up to tract their health. The employers should have a medical service provided in their company and they

should also give benefits for the affected employees. There are many suggestions that employers and employees could keep in mind. Employers need to keep track of number of their employees having mental disorder. Employers should allow flexible work environment with flexible work scheduling and break timings. They should allow employees to work from home or have flexible place of work.

Title: Depression Detection System Using Machine Learning

AUTHORS: Samruddhi Patil, Pratiksha Nikam, Swejal Kobal, Prof. Shital Jade.

Year: 2023

Depression is a common mental disorder, which can happen to anyone, at any age. The problem lies with the fact that it is still not identified and treated in many cases, leading to severely affecting the mental and physical health of a person. It is generally characterized by a loss of interest, feelings of guilt or low self-worth, disturbed sleep or appetite, feelings of tiredness, poor concentration, social withdrawal and slowed speech. It adversely affects the physical health of a patient, such as increased aches and pain, insomnia or oversleeping and weight problems. According to the Harvard Mental Health Letter, Heart disease is linked to Depression. The recurrence of cardiovascular problems is linked more closely to depression than to high blood pressure, smoking, diabetes, or high cholesterol. If untreated, depression raises the risk of dying after a heart attack. This project used data from different social media networks to explore new methods of early detection of MDDs based on machine learning. We performed a thorough analysis of the datasets to characterize the subjects' behaviour based on different aspects of their PHQ9 depression detection system uses a text-based depression system in a smart application in which a user provides a query in text format and the server process the text to apply a feature extraction and deep learning.

Title: Depression detection using ML

AUTHORS: Nikhil Chauhan, Swati, Divya Rani

Year: 2023

Nowadays most people suffer from depression because of so many reasons, like were taking the example of students they mainly take stress regarding of their results, extracurricular activities, humiliated by others, pressure and so many other things. The purpose of this research is we are going to ask some questions regarding to human stress. Nowadays so many people do not have time to take an appointment of a doctor in their daily life schedule and also help that people those who are not showing their stress but they are stressed. So, we can make a program there is a Questionnaire to ask questions one by one of a

person. On the basis of answers of that person those who want to check the stress level. We are using MachineLearning for measuring the human depression using python, HTML, CSS.

3-METHODOLOGY

The framework begins with data collection by using Twitter scrapper tools and is stored in a .csv file. The raw data will be cleaned and start to do data pre-processing. The data will be tokenized, stemming, and lemmatization as a part of normalizing the data. Next, the data will be analysed 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 - Naive Bayes - Logistic Regression	Naive Bayes gives high accuracy	- Usage of the old dataset - Focus on user confession
[14]	- SVM - Naive Bayes	Naive 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 - Naive Bayes	SVM gives high accuracy	- Usage of the different kernel in SVM - Cannot avoid overfit data

Table 1 Summary of Previous Work

The data is split into the train and test set. The training data is used for model development to make sure the classifier learns. The test data will feed into the model once the model learned about the data for evaluation. The accuracy results of both classifiers will be compared to determine which algorithm outperforms well. Fig. 1 shows the overall implementation framework is visualized as per the flowchart below.

A. Data Collection

The data is collected in two different volumes which are 1000 data and 3000 data via Tweepy. Tweepy is a Python library that can be used to access using the Twitter API. The data is extracted by using Tweepy needs API requirements from the Twitter Developer. The Twitter Developer needs some information before they give the access key and token API to

extract data. The data is then extracted and stored in the format of CSV as compatible for analysis. The dataset contains one column named “text” which contains the tweets which were extracted from Twitter.

B. Data Preprocessing

The collected dataset is loaded on Jupyter Notebook which is an open-source code editor which can open through Anaconda prompt. The dataset was cleaned by removing the URLs, retweets, mention, and stop words.

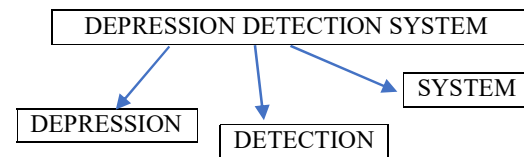


Fig 1. Tokenization

Then, the dataset is tokenized by splitting the sentence to each token or word in an array format. The tokenized sentences are then processed to stemming and lemmatization. Stemming is used to eliminate the affixes from a word so that it can return the original word. For example, the word “Walking” will be normalizing to “Walk”. Lemmatization able to return a word to its base word. It finds the lemma of a word depending on its meaning. The difference between stemming and lemmatization is stemming removes the last few characters like “s”, or “-ing” which can lead to incorrect meanings and spelling errors. However, lemmatization understands the word and returns the word to its meaningful base form.

C. Sentiment Analysis

The cleaned dataset is then converted into the string format. The cleaned data is then sentiment analyzed using the TextBlob. Textblob can be imported as a library in Python to count the sentiment score with subject and polarity. Each sentence in the dataset is scored by using subjectivity and polarity. Polarity score is then labelled in new attributes with the criteria of -1,0 and 1. The labelled data determined by the polarity score in which if the score is below 0, then it is labelled as negative. If the score is exactly 0, then it is labelled as neutral. If the score is above 0, then it is labelled as positive.

The polarity score defined that if the value is more than 0, the sentence score is considered as not depressed. However, if the sentence score is below 0, then it is declared as a depressive sentence.

D. Classification

The scored dataset is fed into the WEKA. WEKA is a GUI open-source software that has a collection of machine learning algorithms for data mining tasks and analysis. The train-test split ratio for the experiment is 70-30. The scored dataset is split according to the ratio. The classifiers are trained and tested using the data with labels. The NBTree and Naïve Bayes classifier is used on Weka. The process is repeated for the 3000 data dataset as well to identify any significant changes.

a) Naïve Bayes: It is a statistical classification origin from Bayes Theorem and it treats the feature in the class as independent of other features. Naïve Bayes classifier assumes that the effect of one feature in a class is independent of another feature. Naïve Bayes practices a probabilistic approach which to predict the class of the test data and excel on predicting multiple class. However, the flaw of Naïve Bayes is it assumes the values are independent among the features which could not apply in the real-life problem. Equation (1) shows the formula of the Bayes and Fig. 2 shows the Naïve Bayes model. Let note C is a class label and x_n are the feature inputs.

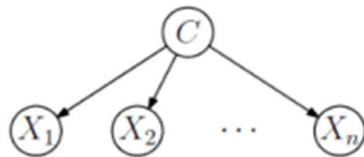


Fig. 2. Naïve Bayes Model [15]

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

- $P(c|x)$ = 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

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

Step 1 - From the given class labels, calculate the prior probability

Step 2 - For each class, calculate the likelihood probability in each attribute in the dataset

Step 3 - Calculate posterior probability by adding value in Bayes formula

Step 4 - The class which has a higher probability based on given input belongs to the higher probability class.

b) NBTree :

It is a hybrid model made up of Naïve Bayes and Decision Tree. The algorithm is introduced by Ron Khavi and applied in an open-source software named WEKA. NBTree algorithm looks like a kind of decision tree that has Naïve Bayes features at the leaf node. The parent node is an attribute that divides the dataset value into two groups. The rule is generated for each path starting from the root to the leaf node. From each splitting condition as a given path form antecedent rule, and class assigned by the leaf at the consequent rule. Fig.3 shows the NBTree model. The oval shape is the attribute that split the data into a narrower data group and the rectangular shape is the leaf node that is classified by the Naïve Bayes classifier. The “v” indicates the value of the attribute of the node.

Fig. 4 shows the insight of the Naïve Bayes Model. The first column shows the attributes and the nominal value of each attribute. The following column records the class value and the frequency count or value of normal distribution for the numeric attributes. The final stage is to identify the algorithm which provides higher accuracy in detecting depression. Accuracy can be calculated by the fraction of correctly predicted class per total number of predictions. The calculation is applied to both classification performance measures. The flow ends once, the best algorithm will be identified to detect the depression based on their accuracy value.

Fig. 5 shows the confusion matrix that is used as the evaluation metrics. Confusion Matrix is a technique to summarize the performance of the selected algorithm which will decide the best classifier for a problem set.

The confusion matrix enables us to pitch the idea to find the best accuracy of an algorithm.

The rows in the confusion matrix show the actual class while the columns show the predicted class.

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

- True Positive: indicated the output where a model is correctly predicted to the positive class.

- False Positive: indicates the output of a model incorrectly predicts the positive class.

- True Negative: indicated the output where a model is correctly predicted to negative class.

- **False Negative:** indicates the output of a model incorrectly predicts the negative class. The four parts are very important to calculate the accuracy, precision, recall of the implemented algorithms. The confusion matrix can be used for both binary and multi-class classification.

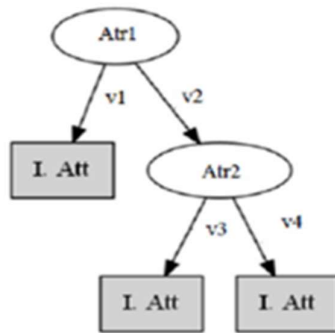


Fig. 3. NBTtree model [16]

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. Naive Bayes leaf model [16]

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

Fig. 5. Confusion Matrix[17]

Table II shows the formula table for the confusion matrix calculation. From the confusion matrix, the *Accuracy*, *Recall*, and *Precision* can be calculated. Accuracy is a common performance measure ratio of correctly predicted observations to the total observations.

Recall is a ratio of correctly predicted positive observations to all observations in actual class which are class of yes/positive/1.

Precision is a ratio of correctly predicted positive observations to the total predicted positive observations. The formula of the calculation is as per the table below.

4-IMPLEMENTATION

Algorithms Used:

- **Enhanced NBTree:** An improved hybrid model combining Naïve Bayes' probabilistic strengths with decision tree structure.
- **Deep Learning Models:** Integration of LSTM (Long Short-Term Memory networks) or Transformer models like BERT for deep contextual understanding and superior sentiment classification.

Flowchart:

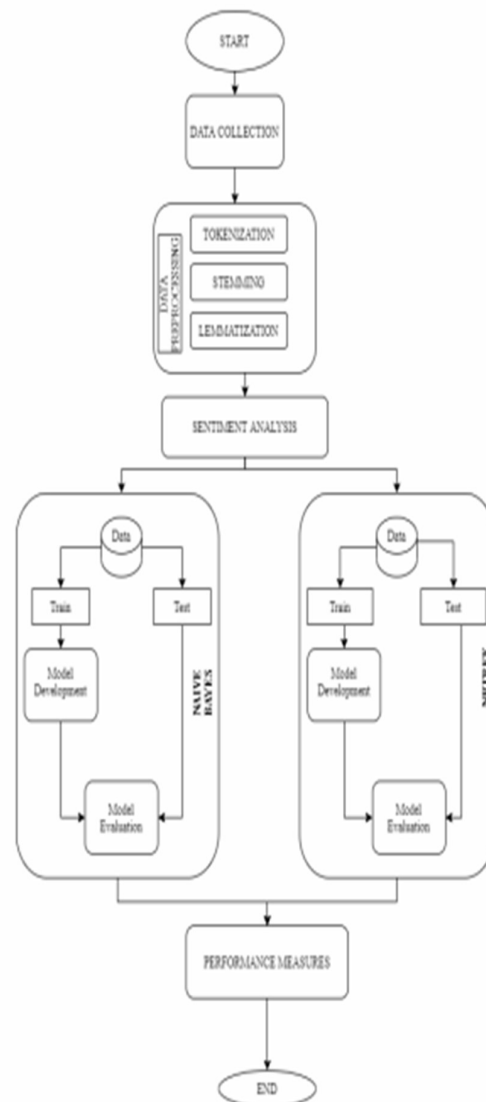


Fig 6. General Flow of Data

5-RESULTS

The section discussed the results obtained from the input of cleaned twitter datasets into the machine algorithm to determine the accuracy score.

ML Results:



Fig 7. ML results

Prediction:



Fig 8. Prediction

7-CONCLUSION

In conclusion, depression has become a serious problem and affecting mental health. Besides that, the fast growth of social media enables an abundance of data in one's social media. Twitter is an application which is focused on 140 characters per tweets enables the user to share their opinion and thought shortly and directly. Each tweet enables the researchers to extract and analyse the information shared in the tweet. The sentiment analysis technique is applied to each tweet to identify the sentiment score and labelled them as positive, negative, or neutral. The labelled tweets are fed into a machine learning algorithm that enables the classification of the tweets into correct groups. Naive Bayes and NBTree which are the selected machine learning algorithm have been implemented on two different sizes of tweet datasets to find the accuracy of the algorithm on classifying the depressive and non-depressive tweets. The NBTree algorithm gives an accuracy of 97.31% to classify the depressive and non-depressive tweets on the 3000 tweets dataset and 92.34% on the 1000 tweets dataset. On the other hand, Naïve Bayes

show 97.31% on the 3000 tweets dataset and 92.34 % on 1000 tweets datasets.

Naïve Bayes and NBTree are equally efficient by giving the same accuracy value in the experiment. However, the project is limited to the text only. The work can be enhanced in future work by targeting a selective user and their tweets at certain time and determine either the status of depression.

8-FUTURE SCOPE

The proposed system for depression detection through Twitter using Naïve Bayes and NBTree algorithms lays a solid foundation for further advancements in mental health analysis using social media data. While the current system demonstrates promising results, there is significant potential for future enhancement, scalability, and interdisciplinary applications. One of the key areas for future development is the integration of deep learning models such as LSTM (Long Short-Term Memory) and transformer-based architectures like BERT (Bidirectional Encoder Representations from Transformers). These models are capable of capturing more complex linguistic patterns, understanding context with greater accuracy, and handling long-term dependencies in user behaviour. Their inclusion could significantly improve the precision and recall rates of depression detection. Another major advancement lies in multi-modal data analysis. The current system relies primarily on textual data; however, the inclusion of images, videos, and voice notes shared on social media could provide richer insights into user emotions. Leveraging computer vision and audio processing technologies in conjunction with textual analysis can open up new dimensions for mental health diagnosis. Multilingual support is also an essential aspect for future consideration. Presently, the system focuses on English-language tweets, but expanding it to support regional and international languages would increase its global applicability. This would require the development of language-specific models and sentiment lexicons, ensuring accurate detection across different cultures and demographics.

9-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-matrixa9ad42dcfd62> (accessed Feb. 04, 2021).
- [18] Cohn, J F, T S Kruez, I Matthews, Y Yang, M H Nguyen, M T Padilla, F Zhou, and F De la Torre. 2009. —Detecting Depression from Facial Actions and Vocal Prosody. In 2009 3rd International Conference on Affective Computing and Intelligent Interaction and Workshops, 1–7. doi:10.1109/ACII.2009.5349358.
- [19] Cruz, Joseph A, and David S Wishart. 2006. —Applications of Machine Learning in Cancer Prediction and Prognosis. In Cancer Informatics 2. SAGE Publications Sage UK: London, England: 117693510600200030.
- [20] Dhall, Abhinav, Roland Goecke, Jyoti Joshi, ichael Wagner, and Tom Gedeon. 2013. —Emotion Recognition in the Wild Challenge 2013. In Proceedings of the 15th ACM on International Conference on Multimodal Interaction, 509–16. ACM.
- [21] Mr. Pathan Ahmed Khan, Dr. M.A Bari,: Impact Of Emergence With Robotics At Educational Institution And Emerging Challenges", *International Journal of Multidisciplinary Engineering in Current Research(IJMEC)*, ISSN: 2456-4265, Volume 6, Issue 12, December 2021,Page 43-46
- [22] Shahanawaj Ahamad, Mohammed Abdul Bari, Big Data Processing Model for Smart City Design: A Systematic Review ", *VOL 2021: ISSUE 08 IS SN : 0011-9342 ;Design Engineering (Toronto) Elsevier SCI Oct : 021;Q4 Journal*
- [23] M.A.Bari & Shahanawaj Ahamad, "Object Identification for Renovation of Legacy Code", in *International Journal of Research and Reviews in Computer Science (IJRRCS)*,ISSN:2079-2557,Vol:2,No:3,pp:769-773,Hertfordshire,U.K., June 2011