

Depression Detection from Social Media Interactions using Machine Learning Approaches

Swathi Y¹ and Girish M²

^{1,2}Gopalan College of Engineering and Management/CSE, Bengaluru, India
Email: gcemcs@gopalancolleges.com, girishgcm@gmail.com

Abstract— Depression, a leading cause of global disability and a significant risk factor for suicide, remains untreated in many due to various barriers. Recognizing the impact of depression on language patterns, this paper explores the innovative use of machine learning to analyze social media communications for early signs of depression. By harnessing advanced models to sift through digital interactions, this paper aims to identify linguistic indicators of depressive states, offering a novel approach to early detection. This study not only underscores the potential of social platforms as valuable data sources for mental health monitoring but also contributes to the development of proactive support systems, leveraging technology to bridge gaps in mental health care.

Index Terms— Depression, early detection, social platforms

I. INTRODUCTION

The World Health Organization (WHO) reports that over 300 million people worldwide—roughly 4.4% of the world's population—struggle with depression. Although the frequency of various forms of depression varies globally and is more frequent in women (5.1%) than in men (3.6%), it can affect anyone of any age and is not limited to any particular circumstance in life [1]. Due to a discrepancy between a depressed person's perception of themselves and the real world, depression is commonly associated with paradoxes. According to the most recent findings from the 2016 National Survey on Drug Use and Health in the United States, 6.7% of adults and 12.8% of teenagers between the ages of 12 and 17 reported having experienced a major depressive episode (MDE) in 2016.

The word "depressed" is now frequently used in casual speech. Depression is generally thought to be associated with mood swings and should even be accompanied by negative self-perception, feelings of wanting to run away or hide, vegetative changes, and decreased activity levels. The symptoms that depressed people know will have a significant impact on their capacity to disrupt any situation in their level of living, which makes them very different from the typical mood swings that everyone goes through. Although depression and other mental illnesses can cause social disengagement and isolation, it has been observed that social media platforms are being used by affected individuals so much that these findings are supported [2]. Peer-to-peer communities on social media may be able to combat stigma, raise the likelihood that people with mental illnesses will find qualified assistance, and directly offer online assistance to those in need. An equivalent study conducted in the United States found that internet users with stigmatized illnesses such as depression or incontinence are more likely than those with other chronic illnesses to use online resources for health-related information and to communicate about their condition. All of this highlights how critical it is to analyze strategies for providing depression support to users of social

media and the internet more broadly.

This paper will use social media text mining to try and predict early indicators of depression. To construct the reasoning behind our victimization neural systems. information set extraction from social media to feed into neural Networks. The goal of this proposed system is to create an application that can forecast a person's time in the past. So that they can receive treatment to recover. This study aims to forecast an individual's depression based on the text messages they send.

This document is divided into four pieces for the remaining content. Section II provides a brief overview of current methodologies, Section III outlines our proposed methodology, Section IV offers the analysis results that bolster our proposed approach, and Section V concludes with a conclusion.

II. LITERATURE SURVEY

1. Primary care physicians' early diagnosis of depression Graham Worrall, MB, BS, CCFP, MRCP; John W. Fleischner, MD, MSc, FCFP. This essay offers specifics regarding early indicators of depression and its management. Compared to other disorders observed in the primary care physician's office, the early identification of depression is far more complicated and difficult. The part that symptoms play in depression is a second issue. Presymptomatic disorders are to be identified as the aim of early detection. As a result, creating a reliable and precise instrument is very challenging.

2. Using Linguistic Data and Neural Networks to Early Identify Depression Indications in Text Sequences Sven Koitka, Marcel Trotzek, and IEEE Member Christoph M. Friedrich. The fine print regarding emotions and sentiment is provided in this publication. Given that sentiment analysis is primarily concerned with obtaining viewpoints, emotions, and sentiments from written texts, it makes sense that data from this domain would be quite beneficial for locating emotional statements in the context of mental state text categorization. The emotion and sentiment options were excluded from the final set of data characteristics used in this work's tests since they were useless in this particular situation.

3. Map-supported Decision Tree algorithmic program optimization analysisXing Associate Business & Technical College, Wulanhaote 137400, China; CAS Key Laboratory of ordering Sciences and data Peiping Institute of Genomics, Chinese Academy of Sciences, Beijing 100101, China; Reduce faculty of laptop and data engineering, province traditional university, Hohhot 010010, China. This paper describes the implementations of decision trees. One of the inductive learning algorithms supporting a collection of coaching samples is the call tree algorithmic program. It is derived from a set of no rules, which is essentially derived from the order of sample set derived classification rules of decision tree representation.

4. Investigation of Depression Analysis Utilizing Machine Learning Methods Priyanka Zar, Vaibhav Sharma, and Devakunchari Ramalingam, International Journal of Innovative Technology and Exploring Engineering (IJITEE), Volume 8, Issue 7C2, ISSN 2278-3075, may 2019. The fine print on acoustic depression detection is provided in this study.

III. PROPOSED SYSTEM

What might the outcome signify? Positive indicates that there is little chance of anxiety or despair in that person. The medium level, known as neutral, is when a person may or may not be depressed but may also be more prone to depression. At that point, the person might exhibit symptoms similar to depression. Finally, the user's message reveals despair and anxiety symptoms at the lowest level, which is called Negative. The message conveys more unpleasant feeling the more negative terms the user inserts. The goal of this study is to use social media text mining to anticipate early indicators of depression. Neural networks are what we are employing to develop the reasoning for it. Social media data extraction for neural network feeds [3].

The following is a summary of the "Early Depression Detection" goals:

1. Create a model that can categorize social media posts into groups related to depression or not.
2. To do the same, create a Convolutional Neural Network (CNN) model that is capable of carrying out the classification task and comprehending the underlying linguistic data [4].

A. Prediction Models

The goal of the paper is to use social media text mining to identify early indicators of depression [5]. The steps to identifying a person's early indicators of depression are listed below in fig.1.

B. About The Dataset

The dataset, comprising 1,048,575 tweets, designed for sentiment analysis or mood expression studies on Twitter. With six columns, it details tweet IDs, timestamps, a static field likely indicating query absence, user handles, and the tweet content itself. The first column, potentially a sentiment label, lacks descriptive headers [6], making precise interpretation challenging. Its structure suggests utility in understanding social media behaviour, particularly regarding sentiments, though the exact nature or sentiment scale represented remains unclear without further context or a data dictionary.

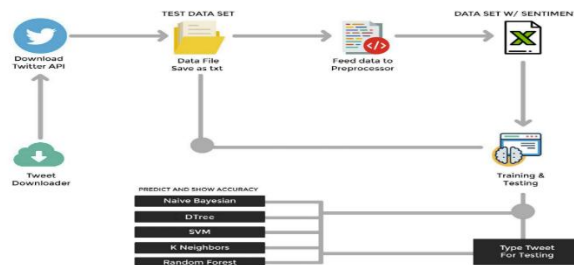


Fig.1.Process to identify depression

Distribution of Sentiment Values: The sentiment distribution plot indicates a predominant focus on tweets with a specific sentiment value, likely indicating a binary classification of sentiments (e.g., positive vs. negative). However, without specific labels or a legend in the plot, the precise nature of these sentiment categories remains unclear. This distribution suggests that the dataset may have been curated for sentiment analysis [7], focusing on distinguishing between different emotional states or opinions expressed in the tweets as in fig.2.

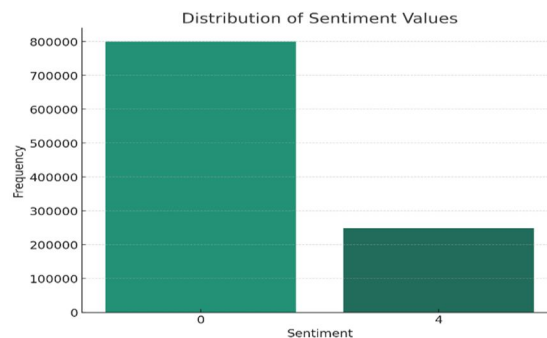


Fig.2.Distribution of sentiment values

Common Themes and Topics: The analysis of the most common words in the processed tweets highlights frequent terms that might reflect general discussions or themes present in the dataset [8]. The top words could suggest topics of interest or concern among the users, though the exact words are not specified in this summary. Identifying these common terms is crucial for understanding the dataset's overarching themes and can inform further text analysis or feature engineering for machine learning models as in fig.3.

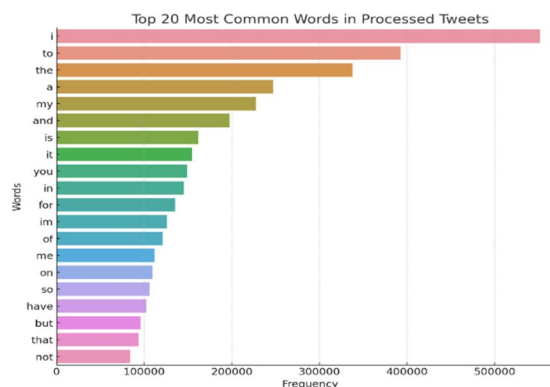


Fig.3.Top 20 most common words in processed Tweets

The word cloud visualization serves as a graphical representation of the frequency of words related to mental health discussions, illustrating a range of topics from challenges like "depression" and "anxiety" to positive aspects such as "hope" and "wellness." The size of each word in the cloud indicates its prevalence in the conversation, emphasizing the community's focus on both the struggles and the supportive, healing aspects of mental health [9]. This visualization highlights the comprehensive nature of social media discourse on mental well-being, encompassing not only the difficulties individuals face but also the collective efforts toward recovery, support, and positive mental health practices as in fig.4.

Fig.4. Word Cloud Visualization

The summary of the steps for conducting sentiment analysis using Twitter data:

- *Create a Twitter Developer Account:* Obtain the necessary credentials (consumer_key, consumer_secret, access_token, access_secret) by creating a Twitter Developer account.
- *Data Approval and Collection:* After gaining approval, collect Twitter data that will be used for text mining.
- *Data Processing and Sentiment Calculation:* Analyze the tweets against a dictionary that associates words with their sentiment polarity. Tokenize the tweets, assign polarity to each word, and calculate the overall sentiment of each tweet by summing the polarities and normalizing by the tweet's word count as in fig.5.

Preprocessing and Output File: Post-processing, the data will be available in an Excel file (processed_data/output.xlsx) with two columns: one for the tweet ID and the other for its sentiment, categorized into "Positive," "Neutral," or "Negative" based on depression-related keywords as in fig.6.

weaksubj		abandoned	adj	n	negative
weaksubj	1	abandonment	noun	n	negative
weaksubj	1	abandon	verb	y	negative
strongsubj	1	abase	verb	y	negative
strongsubj	1	abasement	anypos	y	negative
strongsubj	1	abase	verb	y	negative
weaksubj	1	abate	verb	y	negative
weaksubj	1	abdicade	verb	y	negative
strongsubj	1	aberration	adj	n	negative
strongsubj	1	aberration	noun	n	negative
strongsubj	1	abhor	verb	y	negative
strongsubj	1	abhor	adj	y	negative
strongsubj	1	abhorred	adj	n	negative
strongsubj	1	abhorrence	noun	n	negative
strongsubj	1	abhorrent	adj	n	negative
strongsubj	1	abhorrently	anypos	n	negative
strongsubj	1	abhors	adj	n	negative
strongsubj	1	abhors	noun	n	negative
strongsubj	1	abidance	adj	n	positive
strongsubj	1	abidance	noun	n	positive
strongsubj	1	abide	anypos	n	positive
strongsubj	1	abject	adj	n	negative
strongsubj	1	abjectly	adverb	n	negative
weaksubj	1	abjure	verb	y	negative
weaksubj	1	abilities	noun	n	positive
weaksubj	1	ability	noun	n	positive
weaksubj	1	able	adj	n	positive

Fig.5: Dictionary

#	id	sentiment
1	670430762255953920	-
2	670430770141253632	-
3	670430771592413187	-
4	670430772880479233	-
5	670430776432644096	-
6	670430779439841280	-
7	670430780425723904	-
8	670430781927288833	-
9	670430784678703104	-
10	670430796485726208	-
11	670430797723633603	-
12	6704307995905004224	-
13	670430802575793216	-
14	670430804446289921	-
15	670430804896313920	-
16	670430807592140801	-
17	670430807344779264	-
18	670430814431344832	-
19	670430815473332224	-
20	670430819134734336	-
21	670430820791726080	-
22	670430822129582080	-
23	670430829331316736	-
24	670430838915293184	-
25	670430839343124481	-
26	670430840412566900	-
27	670430843713540098	-
28	670430845147914240	-
29	670430848338713384	-

Fig.6: Sentimental

Training and Prediction: With all files properly placed, execute the script `depression_sentiment_analysis.py`. This script processes the output file [10], matches tweets with their IDs, and feeds the data into various classifiers. The console will display the results of the analysis in real-time as in fig.7.

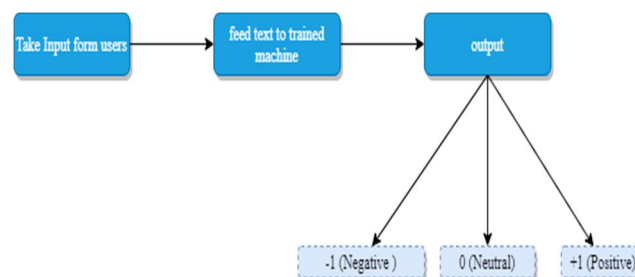


Fig.7. Analysis of the result

C. The Confusion Matrix And Disagreement Score

A confusion matrix of size $n \times n$ associated with a classifier shows the predicted and actual classification, where n is the number of different classes.

The pseudocode provided outlines a method for selecting attributes based on the confusion matrix in a binary classification context. The method aims to create groups of attributes that both possess strong individual classification abilities and complement each other by differing in their misclassification patterns [11]. The process can result in multiple candidate attribute subsets, from which the subset that demonstrates the highest classification accuracy should be chosen for further examination. The algorithm for Confusion Matrix-based Attribute Selection is as given below.

```
Require: 2-class data of  $n$  attributes
Require: classification technique
Require:  $k$  - number of member subset
Ensure: Output  $k$ -attribute subset as tuple  $S_k = (A_1, A_2, \dots, A_k)$ 
Compute classifier  $C_i$  based on feature  $A_i, i = 1..n$ 
Obtain:  $Accuracy(C_i)$  and  $ConfMatrix(C_i)$ 
Rank  $A_i$  according to  $Accuracy(C_i) \Rightarrow R_A$ 
for  $i = 1 \dots n$  do
    Compute disagreement based on  $ConfMatrix(C_i)$ 
    as:  $D_i = \frac{|b-c|}{\max\{b,c\}}$ 
end for
Rank  $A_i$  according to  $D_i \Rightarrow R_D$ 
Select top  $k$  (according to  $R_A$ ) attributes having large  $D$  (according to  $R_D$ ) but in different classes:  $\Rightarrow S_k = (A_1, A_2, \dots, A_k)$ 
```

IV. RESULTS

The image shows an IPython console with user interaction and a confusion matrix for a Naive Bayes (NB) classifier. In the console, a user has input a tweet for sentiment analysis, which the system has classified as "Neutral." Below the text interaction, there is a confusion matrix visualization that represents the performance of the classifier. The matrix shows four quadrants with different shades of blue indicating the number of predictions. The darker the shade, the higher the number of instances for the true label vs. predicted label categories. The quadrants represent true positives, false positives, true negatives, and false negatives, essential for understanding the classifier's performance in terms of precision, recall, and accuracy as in fig.8.

The tabular results display the performance metrics of five different machine learning classifiers used for a classification task:

- **Naive Bayes** - This classifier achieved a high accuracy of 96.004%. It also had the fastest completion speed among all classifiers, taking only about 0.47179 seconds to make predictions, indicating it is both accurate and efficient for this particular dataset.
- **Decision Tree** - The Decision Tree classifier outperformed all others in terms of accuracy, reaching nearly perfect accuracy at 99.9868%. It was also relatively fast, with a completion speed of 1.81819 seconds, which suggests that it is highly effective for the dataset while still being reasonably efficient.
- In this paper, we examine the parallel implementation of decision tree based on genetic algorithm (GA) optimization approach using MapReduce framework on the basis of opensource cloud computing platform. Inspired by the principles of genetics and evolution, the genetic algorithm is an optimization technique for swarm intelligence searches. The process of finding globally optimal rule set as follows:

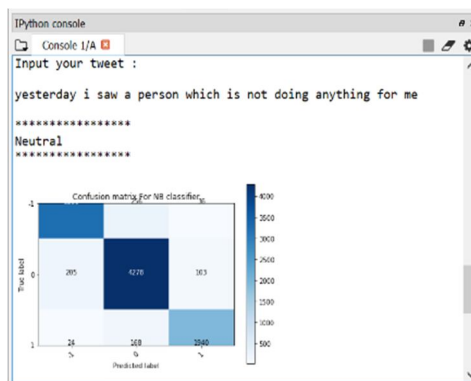


Fig.8.Confusion Matrix for NB

- (1) Convert the tree to a rule set, then encoding rules. Corresponding to the entire coding rule set as a chromosome.
- (2) Get decision tree initial population, randomly select several subsets of sample in the rule set as N chromosomes, using ID3 building decision tree.
- (3) Use the fitness function to evaluate the fitness of the individual, judge whether or not to meet the end condition, if it is not satisfied then performed as follow:
 - a. selection operation
 - b. crossover operation
 - c. mutation operation
- (4) Determine the reasonableness of the decision tree. Selective operations are as follows:
 - a. increase rules
 - b. consolidation rules
 - c. delete rules
- (5) Go to step (3).

- **Support Vector Machine (SVM)** - The SVM had the lowest accuracy at exactly 50.00%, which is equivalent to random guessing in a binary classification task. Its completion speed was significantly slower than the others, taking 22.17733 seconds. This indicates that SVM may not be suitable for this dataset or might require parameter tuning for better performance.
- **K-neighbours** - With an accuracy of 85.6475%, the K-neighbours classifier performed moderately well, suggesting it is somewhat effective for the dataset. Its completion speed was 3.26495 seconds, which is slower than Naive Bayes and Decision Tree but faster than SVM.
- **Random Forest** - The Random Forest classifier had an accuracy slightly above random chance at 51.6480%. Its completion speed was quite fast, almost on par with Naive Bayes, at 0.47549 seconds. Despite its speed, the accuracy indicates that the model may not have been well-tuned or that Random Forest is not the best fit for this data as in fig.9.

Classifier	Accuracy (%)	Completion Speed (seconds)
Naive Bayes	96.0040	0.47179
Decision Tree	99.9868	1.81819
Support Vector Machine	50.0000	22.17733
K-neighbors	85.6475	3.26495
Random Forest	51.6480	0.47549

Fig.9. Various algorithms performance

In summary, while Decision Tree provided the highest accuracy, it was not the fastest. Naive Bayes offered a good balance between accuracy and speed, making it potentially the best choice for scenarios where both factors are important. SVM and Random Forest showed poor performance on this dataset in terms of accuracy, with SVM also lagging in speed. K-neighbours offered a middle ground in terms of both accuracy and speed. These results could vary based on the nature of the dataset, the features used, and how each model was configured and tuned.

The sentiment analysis system processed two different tweets. The first tweet contained a mixture of uncertainty and a positive recommendation, leading the system to classify it as "Positive" due to its optimistic conclusion. The second tweet began with a clear negative statement [12] but was followed by an ambiguous phrase, causing the system to classify it as "Neutral" likely because the sentiment could not be conclusively determined as in fig.10.

Input your tweet : ""

I dont know this is right for me but if think this is right then follow me for better life

Positive

Input your tweet :

you ruined my life but i am

Neutral

Fig.10. Sentiment analysis based on Tweets

VI. CONCLUSION AND FUTURE SCOPE

What the result could mean? Positive, this mean that person is unlikely to have depression or anxiety. Neutral, this is the middle level wherein the user may or may not have depression but may also be more prone to being depress. At that stage the user may display some depression like symptoms. lastly, Negative is the lowest level where depression and anxiety symptoms are being detected through the user's tweets. The more negative words the user uses mean the more negative emotion the tweet has. The paper emphasizes the potential of leveraging social media data for identifying signs of depression through advanced machine learning techniques. It suggests that refining the analysis methods, such as enhancing the language processing capabilities and expanding the sentiment analysis lexicon, could significantly improve the accuracy of depression detection. The findings advocate for the continued exploration of digital platforms as valuable resources for mental health monitoring and highlight the importance

of technological advancements in facilitating early detection and intervention strategies for depressive symptoms. To improve the accuracy, we can add n-gram techniques for word mining, reduce the stop word eliminating frequency.

REFERENCES

- [1] Ghosh S, Anwar T. Depression intensity estimate via social media: a deep learning technique. 2021;8(6):1465–1474, IEEE Trans Comput Social Syst. doi:10.1109/TCSS.2021.3084154
- [2] Aragon ME, Gonzalez-Gurrola LCG, Lopez-Monroy AP, et al. Emotional patterns in social media can be used to identify mental health issues, such as depression and anorexia. 2021 IEEE Transactions.
- [3] Aragón ME, González FA, Lara JS, et al. Deep bag of sub-emotions for social media sadness identification. Proceedings of the 24th International Conference on Text, Speech, and Dialogue (TSD 2021).
- [4] Salas-Zárate MDP, Salas-Zárate R, Alor-Hernández G, et al. A comprehensive study of the literature to identify indicators of depression on social media. 10(2):291 in Healthcare (2022);
- [5] Rustagi A, Manchanda C, Sharma N, et al. Combinational deep neural network-based approach to depression anatomy. Published by Springer Singapore in Volume 2 of the proceedings of the International Conference on Innovative Computing and Communications (ICICC 2020), pages 19–33.
- [6] Ye J, Yu Y, Wang Q, et al. Emotional audio and assessment text-based multi-modal depression detection. 2021;295:904–913 in J Affect Disord. 1016/j.jad.2021.08.090
- [7] Angskun J, Tipprasert S, Angskun T. Real-time depression identification using big data analytics on social networks. 2022;9(1):69, doi:10.1186/s40537-022-00622-2, J Big Data.
- [8] Govindasamy KA, Palanichamy N. Machine learning algorithms for depression identification using Twitter data. The Fifth International Conference on Intelligent Computing and Control Systems (ICICCS) is scheduled for 2021. IEEE, 2021, pp. 956–967.
- [9] William D. and Suhartono D. A comprehensive evaluation of the research on text-based depression detection on social media posts. 2021;179:582–589. Procedia Comput Sci. 1016/j.procs.2021.01.043
- [10] Smys S, Raj JS. A comparative analysis of deep learning methods for early social media network depression diagnosis. 2021;3(01):24–39. J Trends Comput Sci Smart Technol (TCSST). The doi is 10.36548/jtcsst.2021.1.003.
- [11] Giallanza T., Reddi V., and Haque A. Deep learning for depression and suicide detection with unsupervised label correction. In: Proceedings of the 30th International Conference on Artificial Neural Networks, held in Bratislava, Slovakia, September 14–17, 2021, Part V 30. Springer International Publishing, 2021, pp. 436–447. Artificial Neural Networks and Machine Learning–ICANN 2021.