

Automated Classification of Depressive Tweets Detection Using Machine Learning and Deep Learning Techniques

Abdul Mohhammed Jasmine^a, Sandeep Yelisetti^b

^{a,b}Department of Information Technology, V R Siddhartha Engineering College, Vijayawada, India

Abstract

The most occurring emotional mental disorder now a days is Anxiety and depression. Where anxiety and depression are related and directly proportional to each other. Anxiety and Depression are two very dangerous and harmful diseases or disorders which occurs in children, undergraduates, teenagers, employees and old people. More anxiety and depression have suicidal tendency. Depression occurs due to many parameters like tensions, lack of sleep, lack of confidence and etc. The person who suffers from depression have the higher probability to suffer with anxiety. The main problem with these disorders is that detecting anxiety and depression is a very difficult task, therefore one of the way to detect the people who are suffering from the anxiety and depression is with the help of social support by using the twitter data and extracting the data from the tweets which they post and analysing the data and detecting whether the tweets posted by the people on twitter are depressive or not by using best machine learning techniques like naïve Bayes, logistic regression and also the deep learning techniques like CNN and LSTM. The naive Bayes classifier is primarily used in sentiment classification, and it is based on the Bayesian network. The Naive Bayes method predicts text based on the frequency of words in the text. In contrast, logistic regression analysis is a statistical probabilistic classifier that is used in deep learning to estimate the probability of a binary response based on one or more predictors. With the aid of social media (twitter) data, the final result of this research is to classify whether a given tweet from the dataset is depressive or not, and applying feature extraction methods to this classifier helps to improve the accuracy results of the classifier. The TF-IDF feature extraction method was used in this study (Term frequency and inverse document frequency). And whereas in deep learning, LSTM and CNN are used to predict the depressive tweets with best accuracy.

Keywords: Depression, Machine Learning, Deep Learning

1. Introduction

Depression is a serious mental health problem that affects people of all ages, genders, and communities all over the world. People feel more at ease sharing their opinions on social networking sites almost every day in this age of modern communication and technology. Undiagnosed and untreated depression-induced conditions may lead a patient to chronic diseases, including suicide. About 300 million people worldwide are believed to suffer from depression, 4.4% of the world's population is publicly known. [1]. Suicide is a possibility when depression is serious. Suicide claims the lives of about 800,000 people each year. Between the ages of 15 and 29, suicide is the second leading cause of death [2]. Depression is also blamed for numerous other diseases that are also physical. Depression can cause an increase in appetite, which can lead to unwanted weight loss or gain. They can also experience unexplained aches and pains, such as joint or muscle pain, sore breasts, and headaches [3]. While depression has a high rate of treatment, nearly two out of every three people who suffer from it do not actively pursue or receive proper treatment [4]. Depression has long been described as a single disorder with a set of medical criteria as a general mental health condition. It is also linked to anxiety and other neurological and physical disorders, as well as having an emotional effect. And the affected people's behaviour [5]. According to the WHO report, 322 million people worldwide suffer from depression, accounting for 4.4 percent of the global population. Depression is a mood condition that affects people. It will compete with the day-to-day job, contributing to wasted time and reduced productivity.

This paper aims to detect the depressive tweets using twitter data as twitter data is one of best form of data to be taken from social media source where there are tweets posted by the users. The main problem in today's generation is stress, anxiety and depression among the people who are working/employees, undergraduate or teenagers. To get a happy and healthy life one should prevent from this mental disorder (stress, anxiety and depression). Bigram with the Support Vector Model (SVM) classifier was found to be the most effective single function in detecting depression with an accuracy of 80% in previous reports, but it may be less accurate in identifying suicidal tweets. Hence applying a classifier with TF-IDF or any other feature extraction which increase the accuracy rate helps to detect depressive tweets with high accuracy and also by applying deep learning techniques such as CNN with LSTM in this paper we can get more accurate results.

2. Related Work

Traditionally, anxiety and depression have been diagnosed using such predictive criteria such as heartbeats, skin conductance, and pupil diameter. Another method is questionnaires, which may help recognise a person who is vulnerable to stress.

The main goal of the developer MICHAEL M. TADESSE [6] research is to look at Reddit user posts to see if there are any factors that might reveal local online users' depression attitudes. To this end, the authors use Natural Language Processing (NLP) techniques and machine learning approaches to train the data and test the efficiency of the proposed framework. Bigram with the Support Vector Machine (SVM) classifier is the best single feature for diagnosing depression with 80% accuracy and 0.80 F1 score.

Mike Thelwall, the author, used WSD paper [7] technology as a pre-processing stage and a stress-based lexicon or therapeutic technique that improves precision TensiStrength. Thus, the data set of 1000 tweets with the word "Fine" reflects uncertainties in statement description in different cases. This paper further eliminates obsolete terms from the retrieved tweets such as prepositions, interjections, and conjunctions until working on the remaining words and categorising them in -5 and +5 based on dictionary indications.

The paper [7] examines the interaction between psychological stress states and social experiences using a coherent hybrid paradigm that combines a factor graph model with convolutionary neural and cross auto encoders to create user-level interface attributes from a tweet level of attribute and partially labelled factor graph to integrate psychological experiences at the consumer stage to identify the tension Using the Twitter and SinaWeibo datasets to compare outcomes for improved consistency also operates on certain comparison approaches such as logistic regression, SVM, gradient boosted decision tree, and recurrent neural network.

The paper [7] uses social media to recognise and diagnose depression in adults. Social networking may reveal real or false attributes of an individual. Depression is measured and detected in this tweet. First, they collected crowdsourcing data to identify participants to determine whether or not they have been adequately treated for depression. SVM classifier was used to predict a certain user's depression, indicating that signs of depression tend to be busier at night and in the evening.

SentiStrength's newly discovered algorithm, which measures emotion and severity of feeling from an informal English text, is illustrated in paper [8]. Pronunciation is impossible to spot in social media, and spelling is ignored because of abbreviations, thoughts, and truncated words. As a result, rather than fixing the pronunciation, it is necessary to correct the spelling.

The authors of this paper [9] suggest a mood analysis and suicidal ideation warning method for predicting suicidal behaviour based on depression levels. They collected data in real time from students and parents by having them fill out PHQ-9-style questionnaires (Parent Health Questionnaire) that included questions like "How old are you?" Or do you go to school / college on a daily basis?, etc. Then, using machine classification algorithms, it is trained and rated in five levels of depression based on severity: minimal, mild, severe, and intense. In this dataset, the maximum precision was achieved by using XGBoost, which yielded an 83.87 percent classifier. Depression research has been around a lot longer than the Internet, and it has become a much bigger subject. The detection of depression from records, in particular, has grown in importance as a research area, with intriguing methods and studies reported on Facebook, Twitter, and a variety of other forums [8].

This paper [10] discusses depression research, which began much earlier and has grown to be a larger topic than the Internet. The diagnosis of depression from records, in particular, has grown in importance as a research area, with intriguing approaches and results reported on Facebook, Twitter, and a variety of other forums. Several measures for predicting suicidal tendencies in people from questionnaire findings have been presented in the literature [11, 12]. Zung's Self-rated Depression Scale (SDS) and Hamilton's Depression Rating Scale (HDRS) are two widely used tools for assessing suicidal tendencies. The patient's answer to 20 questions yields Zung's SDS.

Automated Classification of Depressive Tweets Detection Using Machine Learning and Deep Learning Techniques

In this paper [13], they have proposed a new methodology for analysing users with depressive moods through long-term analysis of their daily tweets. Then, to further comprehend their messages, we use both media forms of tweets, powerful features and emoticons, as well as messages, and create a multimodal approach to analyse them. Three single-modal analyses are conducted in the proposed system first to derive the moods of the hidden users from text, emoticon and images: a Learning based text processing, an emoticon processing based on words and an image classifier based on SVM. This is also to observe the transformation of the mental state of the individual and to recognise the psychological state as stable or unstable, which can be used to detect behavioural changes in people with depression [14], predict postpartum changes in the emotion and behaviour of the individual [13] and detect stress [15].

The goal of this paper [16] is to suggest a model focused on data analytics to detect any psychological suffering. This proposed model was built using data from respondents' posts on two popular social media platforms: twitter and facebook. The severity of a user's depression was assessed using his social media posts. Linguistics (NLP) was used in collaboration with the Support Vector Machine (SVM) and the Nave Bayes algorithm to logically diagnose pain in a more convenient and accurate manner. As of 2018, Twitter had over 290 million daily users worldwide [19], and Facebook had over 1.9 billion users worldwide [20]. These people want to post their social media lives, feelings, sorrow,pleasure.

Choudhury et al. [21] Crowdsourcing to compile opinions of several thousand People who have now been diagnosed with chronic depression on Twitter. They also shown a contrast of behaviour between happy user and depressed user, which shows several variations. They have also suggested creating an MDD classifier to determine whether or not an individual is prone to depression The model had a 70% accuracy rate.

Another study by Hassan et al.[22] introduced the techniques of machine learning to evaluate sentiment and compared on the subject, classification models were used: SVM, Navie Bayes (NB), and Maximum Entropy (ME). SVM performed higher than NB and ME, according to the researchers.

Lots of relevant studies seen above depression screening methods but they did have certain drawbacks. Such as figuring out the feasibility of the suggested versions, there was no actual world evaluation. In this proposed model, posts are created from twitter and a machine learning algorithm is used to diagnose a person's susceptibility to depression by determining whether a tweet shared by that person is suicidal or not, with the highest accuracy and precision of all models.

In this paper [23] it has been investigated the psychological significance of syntactic constructs in the estimation of depressive symptom change. He believed that a written text differed only slightly in word use, but that it differed significantly in syntactic structure, especially in the creation of relationships between occurrences. Examining the roles of causation and insight terms, he noticed that complex usage of complex words was diminished in the text written by suicidal persons. Compared to non-depressed ones, grammar.

Twitter is among the most popular social media platforms [24], with nearly 326 million users worldwide and 1 billion posts publicly transmitted to a broad audience. Many scholars have used Twitter data effectively as a basis of insights into attitudes, stress and stress epidemiology.

The people who tweet have other psychiatric illnesses. Linguistic characteristics were used by De Choudhury et al. [25] to prepare an algorithm to look for depression-related posts on Social media.

Relaxation gives you a health while stress, anxiety and depression affects your health fitness. Currently, depression is the fastest situation Growing-up. For this reason people are never happy, despite Well-being. Positive mental wellbeing should be enabling people to achieve their capacity. It'll even help them deal with both the pressures at home and at work. It'll improve productivity and economic growth. Now you should get support from each other to maintain this good mental health, communicate with others, help others, have routine and improve coping skills. With increasing population the depression is also rapidly increasing in the people. The main problem in today's generation is stress, anxiety and depression among the people who are working/employees, undergraduate or teenagers. To get a happy and healthy life one should prevent from this mental disorder (stress, anxiety and depression). There are both traditional and there are scientific methods for identifying individuals who are anxious. A) Survey: Psychiatrists use a broad questionnaire to determine whether or not anyone is depressed based on the answers. This approach has its own limitations and weaknesses, as the responses are not accurate at times. Occasionally any of the questions in the questionnaire don't make sense. B) The other approach is the system used to test sensors. This method's drawback is, it's time consuming, and a little costly. Another and recent way of measuring tension is by social media [1]. As in today's generation all people are aware of social media (twitter), it is best way to collect data from the social media (twitter) in order to classify people who are suffering from anxiety and depression. Sentiment 140 is a dataset present in Kaggle in which the positive and negative messages or posts of users are scraped by using Twint and labelled

them as 0 and 1. The results show that tweet posted by the user is depressive tweet or positive tweet. If it is depressive tweet it will display true or if it is positive tweet it will display the output as false. The summary of the literature survey is that using social media data gives the better and accurate data set by which one can get better accurate results for predicting the depressive tweet.

3. Methodology

Machine Learning

A. DATASET COLLECTION:

The dataset chosen in this paper is sentiment 140 which is available in kaggle. But for the binary classification Twint (scraping tool) is used to collect the negative weighted tweets and positive weighed tweets are scraped from the sentiment 140. There are three columns in this dataset, one is user id (unique id), second is message or the tweet and third is label. Label 0 indicates the positive weighted sentence or message and label 1 indicates the negative weighted message.

B. WORD CLOUD ANALYSIS:

The size of each word represents its occurrence or value in a Word Cloud, which is a computer visualisation tool used to display text data. We're highlighting important textual data points with a term cloud. Word clouds are often used to analyse data from social media websites. By using word cloud analysis here depressive words and positive words are visualized.

C. DATA PRE-PROCESSING:

In this data pre-processing stage all the unwanted data is removes with help of following:

1) **Lower case conversion:** the first step involved in data pre-processing is lower case conversion which means all the messages in the dataset are converted to lowercase.

2) **Punctuation removal:** Improper punctuation is also used as a symptom of lack of knowledge and a decrease in quality. All the punctuations from the dataset are removed in this step.

3) **Stop words removal:** In this data pre-processing level, stop words are words that are filtered. Stop words will cause trouble when looking for sentences.

4) **Removing URL's:** This Dataset have several URL's. Therefore removing URL's are done in this stage.

5) **Removing HTML tags:** Another common method of pre-processing is to delete the HTML tags. Typically displayed HTML tags while scraping results.

6) **Tokenization:** The method of substituting a sensitive data variable with a non-sensitive alternative, referred to as a token. Tokenization plays a vital role in data pre-processing stage.

7) **Stemming and lemmatization:** Stemming and lemmatization is commonly used in tagging schemes, indexing, keywords, Site search results, and retrieving information. And here lemmatization is done using POS tagging.

D. FEATURE EXTRACTION:

Operating with real datasets millions of attributes is now becoming increasingly popular. If the selection of attributes in a dataset is comparable (or perhaps even higher!) than the sum of samples in the dataset, a Machine Learning algorithm would most likely be overfit. Feature extraction starts from a collection of calculated data and creates extracted values (features) that are meant to be descriptive and non-redundant, allowing for faster learning and gross exaggeration and, in some cases, improved human understanding. The dimensionality is reduced when the function is extracted.

TF-IDF which stands for Inverse Text Frequency – Term Frequency. Representing how important a particular word or expression is to a source context is one of the most common methods used for information retrieval. The TF-IDF value rises in terms of the number of times a word appears in the text, although this is often offset by the corpus's word frequency, which seems to refer to the fact that such words are used more often. TF-IDF employs two mathematical methods: Word Frequency first, and Inverse Text Frequency the second. The reciprocal word frequency calculation of the sum of information given by the word corresponds to the total number of times the expression t appears in the document doc against (per) the total number of terms in the document. It calculates the weight of each word throughout the document. IDF shows that a given word is common or rare across all documents.

E. CLASSIFIERS

NAÏVE BAYES

With very large data sets the Naive Bayes approach is adequate. Naive Bayes is considered to surpass even extremely advanced methods of classification, along with flexibility. Naive Bayes is a probability-driven classification algorithm commonly used by the Bayes Theorem based scientific group. The premise behind the Naive Bayes definition is that the existence of one attribute in a class is unrelated to the presence of some other attribute. Naive Bayes holds strong assumptions of independence between both the features thus, it also makes use of sentiment analysis as SVM efficient process.

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

ALGORITHM FOR DETECTING DEPRESSION USING NAÏVE BAYES CLASSIFIER:

Input: tweets

Output: Detecting whether the person is having depression or not

1: Gather information

2: Pre-processing the information

3: Compute Word Cloud Analysis

4: For a term 't' in Document 'd', the weight 'wt', 'd' of term t in document is given by TF-IDF

Compute: $w_t, d = TF_t, d \log [N/DF_t]$

Where

- The number of times t appears in document 'd' is represented by TF_t, d .
- The number of documents uses the word 't' is represented by DF_t .
- The cumulative number of documents in the collection is N.

5: compute Naïve Bayes classifier:

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

Posterior $\propto$ Prior $\times$ Likelihood

Where

- As we know that stress can cause depression 50% of time
- A patient's chance of developing depression is 1 in 50,000.
- The possibility of any person developing stress is 1/20.

LOGISTIC REGRESSION

The regression classification system of regression models (LR) is used to approximate the probability of a conditional response based on one or more indicators [20], [21] and features. The two meanings are denoted by the numbers 0 and 1. The logistic model's cumulative chances (logarithm of the odds) also for number labelled '1' are a weighted average of one or more response variable ('predictors,' one of which may be a categorical variable (two classes coded by an indicator variable) or a continuous variable (any actual value).

$$l = \log b \frac{P}{1-P} = \beta_0 + \beta_1 x_1 + \beta_2 x_2$$

DEEP LEARNING

F. DATASET

The dataset chosen in this paper is sentiment 140 which is available in kaggle. But for the binary classification Twint (scraping tool) is used to collect the negative weighted tweets and positive weighed tweets are scraped from the sentiment 140. There are three columns in this dataset, one is user id (unique id), second is

message or the tweet and third is label. Label 0 indicates the positive weighted sentence or message and label 1 indicates the negative weighted message.

G. WORD2VEC MODEL

Word2vec is a natural language processing technique. Using a neural network architecture, the word2vec algorithm learns word connexions from a wide corpus of text. Once learned, for a partial sentence, such a model can detect synonymous terms or recommend additional words. As the name implies, Word2vec describes that different phrase with a common series of numbers called a vector. Word2vec represents each individual word using a dynamic list of numbers known as a vector. The vectors are deliberately chosen such that the degree of semantic similarity between the terms represented by those vectors is implied by a basic mathematical feature (the cosine similarity between the vectors). In this work the word2vec produces word embedding which converts word into vector and we can do multiple operations on this vectors like addition, subtraction, etc.

H. DATA PRE-PROCESSING

This stage involves eliminating all unnecessary data from the dataset, such as removing duplicates, links and images, removing hashtags, removing @ mentions, removing emojis, removing stop words, removing punctuation.

I. CLASSIFIERS

LSTM and CNN:

CNNs are regularised copies of multilayer perceptrons. Typically, multilayer perceptron refers to networks that are totally linked, meaning that each neuron in one layer is connected to all neurons in the next layer. This networks are susceptible to data overfitting due to their "completely interconnectivity."

The algorithm takes inputs and returns a particular statistic indicating the likelihood that the tweet implies depression. Each input expression is replaced by its integrating, and then a convolutional neural layer is applied to the new clustering element. CNNs excel at learning global characteristics from data; the convolutional layers network takes full advantage of this and learns certain complexity from data sets before passing through a standard LSTM layer

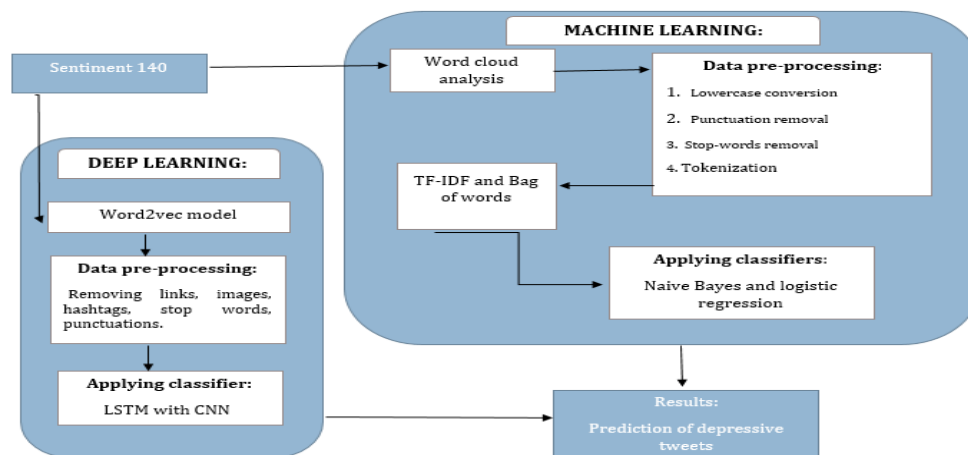


FIGURE1: Architecture for Predicting Depressive Tweets Using Tweeter Data

4. Results

Table1: performance Results of machine learning and deep learning classification model.

MACHINE LEARNING			DEEP LEARNING
	Logistic Regression	Naïve Bayes	LSTM+CNN

Automated Classification of Depressive Tweets Detection Using Machine Learning and Deep Learning Techniques

Accuracy	52.1	88.0	98.0
Precision	54	100	97.0
F-score	50.3	59.6	98.0
Recall	63.5	42.5	99.0

Accuracy:

- Accuracy is the closeness of the measurements to a given value when measuring a set, the accuracy rate using logistic regression is 52.1% which is least accuracy rate compared to other algorithms or classifiers.
- Accuracy rate using naïve Bayes classifier to predict the depressive tweets is 88.0%
- Accuracy rate using LSTM and CNN classifier to predict the depressive tweets is 98.0%

Precision:

- While precision is the closeness of the measurements to each other. Precision rate using logistic regression is 54% which is lowest accuracy arte compared to other classifiers.
- The Precision rate using naïve Bayes is 100%, which is highest compared to other classifiers.
- The precision rate using LSTM and CNN is 97%

F-score:

- The f-score using logistic regression is 50.3%, which is lowest then compared to all classifiers.
- The f-score using naïve Bayes classifier is 59.6%
- The f-score using LSTM and CNN is 98.00% which is highest among all classifiers

Recall:

- The recall rate using logistic regression is 63.5%
- The recall rate using naïve Bayes is 42.5% which is lowest then compared to all classifiers.
- The recall rate using LSTM and CNN is 99.0% which is highest among all classifiers.

5. Result Analysis:

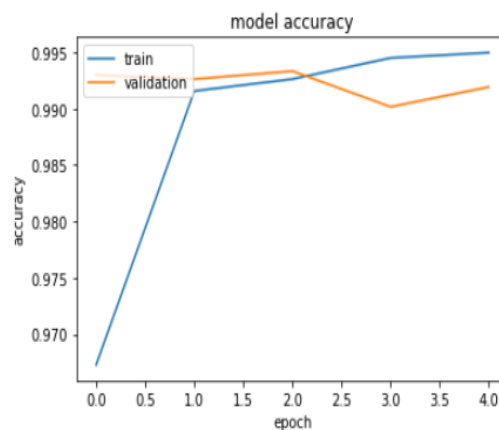


FIGURE 2: MODEL ACCURACY

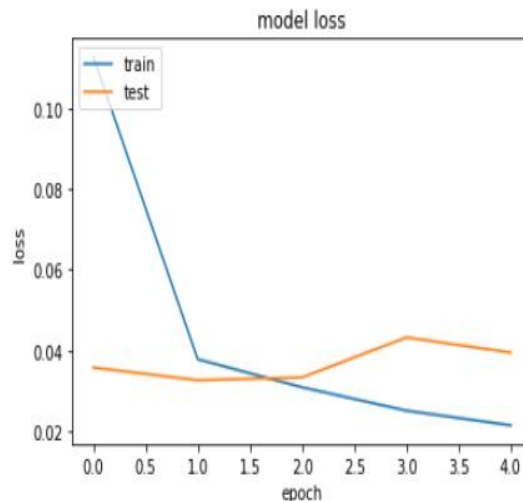


FIGURE 3: MODEL LOSS

6. Conclusion

In this work, Prediction of depressive tweets are done using machine learning and deep learning where TF-IDF feature extraction is used along with the machine learning classifiers (logistic regression and naïve Bayes) and LSTM and CNN (convolution neural networks) are used in deep learning and gives highest accuracy then compared to the machine learning classifiers. Whereas logistic regression gives the lowest accuracy to predict the depressive tweets with 59% of accuracy rate. Naïve Bayes using TF-IDF gives the 88.0% of accuracy rate.

References

- [1] MICHAEL M. TADESSE¹, HONGFEI LIN, BO XU , AND LIANG YANG,," Detection of Depression-Related Posts in Reddit Social Media Forum", DOI 10.1109/ACCESS.2019.2909180, IEEE Access, 2019.
- [2] AkshiKumara ,AditiSharmab , Anshika Arora", Anxious Depression Prediction in Real-time Social Data", f International Conference on Advanced Engineering, Science, Management and Technology – 2019.
- [3] "Number of people with depression increases," SouthEast Asia Regional Office., 1 january 2019. [Online]. Available: <http://www.searo.who.int/bangladesh/enbandepressionglobal/en/>. [Accessed 27 August 2019].
- [4] Timothy J. Legg, "The effects of depression on the body and physical health. [online]," Medical News Today , 1 january 2019. [Online]. Available: <https://www.medicalnewstoday.com/articles/322395.php>. [Accessed 21 November 2019].
- [5] "Media - DBSA Facts about Depression -," Secure2.convio.net., 1 january 2019. [Online]. Available: https://secure2.convio.net/dabsa/site/SPageServer/?jsessionid=00000000.app20101a?NONCE_TOKEN=B2DB3A070406934661E7561D. [Accessed 21 November 2019].
- [6] W. H. Organization, "Depression and other common mental disorders: Global health estimates. geneva: World health organization; 2017. licence: Cc by-nc-sa 3.0 igo." <http://www.who.int/en/news-room/factsheets/detail/depression>, 2017.
- [7] MICHAEL M. TADESSE¹, HONGFEI LIN , BO XU , AND LIANG YANG,," Detection of Depression-Related Posts in Reddit Social Media Forum", DOI 10.1109/ACCESS.2019.2909180, IEEE Access, 2019.
- [7] ReshmiGopalakrishna Pillai, Mike Thelwall, ConstantinOrasan, "Detection of Stress and Relaxation Magnitudes for Tweets", International World Wide Web Conference Committee ACM, 2018.
- [8] Huijie Lin, JiaJia, JiezhonQiu, " Detecting stress based on social interactions in social networks", IEEE Transactions on Knowledge and Data Engineering, 2017.

Automated Classification of Depressive Tweets Detection Using Machine Learning and Deep Learning Techniques

- [9] Munmun De Choudhury, Michael Gamon, Scott Counts, Eric Horvitz, "Predicting Depression via Social Media", Seventh International AAAI Conference on Weblogs and Social Media, 2013.
- [10] Thelwall M., Buckley, Paltoglou, Cai, Kappas, "Sentiment strength detection in short informal text". Journal of the American Society for Information Science and Technology, 2010.
- [11] Swati Jain, Suraj Prakash Narayan, Rupesh Kumar Dewang, UtkarshBhartiya, NaliniMeena and Varun Kumar," A Machine Learning based Depression Analysis and Suicidal Ideation Detection System using Questionnaires and Twitter", SCES 2019 5th Students Conference on Engineering and Systems,2019.
- [12] P. Tzirakis, G. Trigeorgis, M. A. Nicolaou, B. W. Schuller, and S. Zafeiriou, "End-to-end multimodal emotion recognition using deep neural networks," IEEE Journal of Selected Topics in Signal Processing, vol. 11, no. 8, pp. 1301–1309, 2017.
- [13] ShoTsugawa , Yukiko Mogi , Yusuke Kikuchi , Fumio Kishino , Kazuyuki Fujita , Yuichi Itoh\$, and Hiroyuki Ohsaki," On Estimating Depressive Tendencies of Twitter Users Utilizing their Tweet Data", IEEE Virtual Reality 2013 16 - 20 March, Orlando, FL, USA 978-1-4673-4796-9/13/\$31.00 ©2013 IEEE, 2013.
- [14] W. W. K. Zhung, "A self-rating depression scale," Archives of General Psychiatry, vol. 12, no. 1, pp. 63-70, 1965.
- [15] M. Hamilton, "Development of a rating scale for primary depressive illness," British Journal of Social & Clinical Psychology, vol. 6, no. 4, pp. 278-296, 1967.
- [16] Keumhee Kang, Chanhee Yoon, Eun Yi Kim," Identifying Depressive Users in Twitter Using Multimodal Analysis", 978-1-4673-8796-5/16/\$31.00 2016 IEEE, BigComp 2016.
- [17] Doryab, Afsaneh, et al. "Detection of behavior change in people with depression." AAAI, 2014.
- [18] De Choudhury, Munmun, Scott Counts, and Eric Horvitz. "Predicting postpartum changes in emotion and behavior via social media." Proceedings of the SIGCHI Conference on Human Factors in Computing Systems. ACM, 2013.
- [19] Lin, Huijie, et al. "User-level psychological stress detection from social media using deep neural network." Proceedings of the ACM International Conference on Multimedia. ACM, 2014.
- [20] Nafiz Al Asad, Md. Appel Mahmud Pranto, SadiaAfreem, Md. Maynul Islam," Depression Detection by Analyzing Social Media Posts of User", 2019 IEEE International Conference on Signal Processing, Information, Communication & Systems(SPICSCON) 28-30 November,2019, Dhaka, Bangladesh.
- [21] "[online]," Usatoday.com., 1 January 2019. [Online]. Available: <https://www.usatoday.com/story/tech/news/2017/10/26/twitterovercounted-active-users-since-2014-sharessurge/801968001/>. [Accessed 21 November 2019].
- [22] "Facebook Reports Fourth Quarter and Full Year 2017 Results,," [Online]. Available: <https://investor.fb.com/investor-news/pressreleasedetails/2018/Facebook-Reports-Fourth-Quarterand-Full-Year-2017-Results/default.aspx>. [Accessed 21 November 2019].
- [23] J. Zinken, K. Zinken, J. C. Wilson, L. Butler, and T. Skinner, "Analysis of syntax and word use to predict successful participation in guided selfhelp for anxiety and depression," Psychiatry research, vol. 179, no. 2, pp. 181–186, 2010.
- [24] T. S. PortalStatistics and Studies, "Social media usage worldwide," <https://www.statista.com/statistics/272014/global-social-networks-ranked-by-number-of-users/>, 2019.
- [25] M. De Choudhury, S. Counts, and E. Horvitz, "Predicting postpartum changes in emotion and behavior via social media," in Proceedings of the SIGCHI conference on human factors in computing systems. ACM, 2013, pp. 3267–3276.