

A DEEP LEARNING-BASED HYBRID MODEL FOR AUTOMATICALLY DETECTING DEPRESSION IN SOCIAL MEDIA POSTS

T. SRAJAN KUMAR¹, DR. M. NARAYANAN², DR. HARIKRISHNA KAMATHAM³

¹Research Scholar, Department of CSE, School of Engineering, Malla Reddy University, Hyderabad, India.

ORCID_ID: 0009-0009-5364-9616

²Supervisor, Professor, Department of CSE, School of Engineering, Malla Reddy University, Hyderabad, India.

³Co-Supervisor, Professor, Department of CSE, School of Engineering, Malla Reddy University, Hyderabad, India. Email-ID: 2232CS010020@mallareddyuniversity.ac.in, narayananmecse@gmail.com, kamathamhk@gmail.com

Corresponding author's email address: 2232CS010020@mallareddyuniversity.ac.in

ABSTRACT

In recent times, mental health issues have been on the rise, influenced by various factors, including lifestyle changes. With the widespread use of social media, individuals from different backgrounds can openly share their thoughts and emotions, providing valuable data for research. This has opened the opportunity to analyse social media discussions to evaluate the potential presence of depression by examining the sentiments conveyed in the text. Although various heuristic methods for depression detection (DD) are available, the rise of AI has enabled the development of more efficient learning-based methods. However, since a single approach may not be universally applicable, there is a need to refine DL models to improve their performance in detecting depression. In this paper, we introduce a DL context that integrates CNN and BiLSTM networks, enabling the model to capture both features from the data and temporal dependencies. We present an algorithm called Learning Based Depression Detection (LBDD), which analyses Twitter posts to classify them based on the probability of depression. After evaluating the approach on a standard dataset, the projected model outclasses several prevailing methods, achieving an accuracy of 96.32%.

Keywords: *Deep Learning (DL), Depression Detection, Convolutional Neural Networks (CNN), Bi-Directional Long Short-Term Memory (Bilstm)*

1. INTRODUCTION

Accounting for a significant portion of worldwide illnesses, depression is a prevalent mental health disorder that impacts a substantial percentage of the global population. More than 350 million individuals experience depression. Sadly, two-thirds of those affected do not seek treatment. The key challenge is that depression often disrupts an individual's personal and social life, and if left unchecked, it can lead to more severe consequences such as other mental health issues and even suicide. On average, nearly one person in their 40s dies by suicide every 40 seconds, resulting in approximately 800,000 suicide-related deaths globally each year. Adolescents are especially vulnerable to depression, with suicide being a foremost source of death amid young people. In

India, where suicide rates are alarmingly high, research on recognizing and addressing depression is critical. A 2012 study published in The Lancet highlighted that in India, a student dies by suicide every hour due to depression. Many of these cases went unreported, with a total of 39,775 students dying by suicide over the five years preceding 2015. The WHO data reveals that about 8,934 students in India took their own lives in 2015 alone. This alarming trend underscores the urgent need for increased attention and action to address this issue. Addressing mental health issues across all stages of life like childhood, adolescence, and adulthood is essential. People affected by depression often experience a temporary or long-lasting sense of sadness that hampers their enthusiasm and ability to engage in everyday activities [1]. Chronic or recurring periods of stress and prolonged low mood

can develop into serious health problems [2]. Individuals struggling with depression often experience symptoms such as insomnia, social isolation, appetite loss, difficulty concentrating, and, in some cases, an increased risk of suicide [3]. Long-term depression can stem from various factors, with a problematic childhood, substance abuse, sexual abuse, physical health situations, work-related stress, and the enduring effects of racism, colonialism, and caste discrimination [4]. If left untreated, depression and anxiety can worsen, leading to further complications like heart problems, memory issues, sleep disturbances, and other health conditions. In response, several countries, and organizations, including the WHO, have initiated programs to address depression. However, many individuals affected by depression, particularly those from lower- and middle-income backgrounds, face barriers to accessing these treatments due to financial constraints [5][6]. Additionally, limited resources and funding in developing nations prevent the establishment of effective depression treatment programs.

Identifying individuals with depression can be challenging, as there is no reliable method to distinguish them from those who are not depressed. Furthermore, there is a shortage of resources and qualified professionals to effectively treat depression. The lack of accurate diagnostic procedures leads to many current prediction models being unreliable. However, given the extensive user engagement on communal media platforms such as Instagram, Twitter, Snapchat, and Facebook, these platforms can provide valuable data to help predict depression. For example, Twitter generates around 6,000 tweets per second, amounting to approximately 200 billion tweets each year, and this data is publicly accessible [7]. Research indicates that DL models can be effective for detecting depression from text. The paper provides a DL framework based on BiLSTM networks to capture temporal patterns in the data. The contributions of this paper are outlined as follows.

1. To capture features from data and temporal relationships, a DL framework based on a hybrid model that seamlessly combines CNN and LSTM models.
2. The proposed algorithm LBDD which takes Twitter tweets as input and classify the tweets reflecting the probability of depression.
3. The proposed hybrid DL model is trained and evaluated by various evaluation metrics.

The paper is organized as follows: A review of existing research on DD is provided in Section 2. Section 3 introduces the proposed DL-based framework and the algorithm designed for the automatic DD from social media content. We compare the performance of our hybrid DL model with various methods in Section 4, where we present the experimental results. Section 5 outlines potential directions for future research and concludes the paper.

2. RELATED WORK

Prior works on the automatic detection of depression using social media conversations are reviewed in this section.

The increased prevalence of depression during the COVID-19 pandemic has led to advancements in detection techniques, as highlighted by Ghosh *et al.* [8], who noted that the use of social media data for depression detection and treatment has been significantly boosted. Malviya *et al.* [9] highlighted that technological and social media advancements enable broader expression, particularly during pandemics, enhancing the identification of depression through DL. Lin *et al.* [10] utilized Twitter data and multimodal learning, where the Sense Mood system proved effective in diagnosing depression and improving care. Social media trends offer valuable insights into depression, which affects millions globally, often going undetected. Sentiment analysis on communal webs can help detect mood disorders associated to depression, as noted by Giuntini *et al.* [11], who focused on text analysis from platforms like Facebook and Twitter, highlighting the challenges of temporal analysis. Renjith *et al.* [12] acknowledged the challenges of DD on social media but recognized it as a promising area for DL and natural language processing (NLP) applications.

New frameworks for DD can advantage from the data available on social media, as noted by Yang *et al.* [13], who emphasized the importance of confess for the analysis of depression. In a subsequent study, The KC-Net model leverages mental state information to improve DD and achieve promising results, as highlighted by Yang *et al.* [14], who emphasized the significance of identifying stress and sadness on communal media. Ghosh *et al.* [15] offered valuable insights into mental health, demonstrating that BiLSTM-CNN outperforms preceding models in detecting depression in Bangla texts. Rissola *et al.* [16] examined how linguistic analysis techniques can support initial disclosure by

detecting indicators of depression conditions. Ziwei *et al.* [17] pointed out that limited access to therapy often drives the use of social media for identifying depression, with symptoms such as sadness, disinterest, and physical ailments being key indicators of the condition.

Tariq *et al.* [18] noted that social media analysis plays a significant role in decision-making, with mental health being one key area of application. They also highlighted the potential of training with classifiers like NB, SVM, and RF. Peng *et al.* [19] emphasized that emotions on textual analysis, which employs DL techniques, is useful for extracting emotions from text, aiding sentiment analysis and development. Skaik and Inkpen [20] applied NLP and ML techniques to evaluate social media data, identifying mental health concerns, while pointing out the challenges of sampling and feature selection. Adikari *et al.* [21] introduced an AI outline that utilizes advanced NLP techniques to analyse complex emotions in SMP. By joining semantic and sentiment data for more effective emotion recognition, SS-BED, a DL approach introduced by Chatterjee *et al.* [22], outperforms traditional models.

Yang *et al.* [23] applied a neural model to two tasks while analyzing Chinese microblog posts to investigate depression, a significant concern, and achieved results that outperformed baseline performance in predicting depression severity and causes. Suicidality, a significant symptom of depression, is often observed alongside other conditions like insomnia and anxiety, which Yao *et al.* [24] identified while exploring an Online Depression Community using a coding method. Cao *et al.* [25] revealed that social media conversations could be accurately identified using a personal suicide-oriented knowledge graph combined with an attention mechanism. A DL approach reliably mined medical-related data from SMP across different datasets, as discovered by Scepanovic *et al.* [26]. Pran *et al.* [27] analysed Bangladeshi sentiment regarding COVID-19, where CNN models achieved high accuracy, with most sentiments being analytical in nature.

The study emphasized the impact of social influence, as Blanco *et al.* [28] evaluated emotional shifts and model performance in their analysis of pessimism and optimism in COVID-19-related Twitter conversations using a DL approach. The goal of supporting crisis intervention organizations was addressed by Subramani *et al.* [29], who introduced a DL technique to distinguish domestic abuse signs on SMP. Farruque *et al.* [30] developed a procedure to classify clinical depression at the

user level from temporal SMP. Showcasing advanced performance with a BERT-Bi-LSTM pipeline, Kumar *et al.* [31] investigated the use of SMP, focusing on Arabic content, to sense sadness. Ahmed and Lin [32] investigated phrase analysis for DD on social media, recommending text classification through Graph Attention Networks (GATs). Using Reddit Depression data, their model achieved a 0.91 ROC score.

Yuki *et al.* [33] discovered that LSTM models achieved the maximum accuracy in detection, with emphasizing that conversations play a key role in the early identification of depression, which often remains undetected until physical symptoms emerge. They described the dataset and collection methods used in their DD study, with Rissola *et al.* [34] highlighting that the absence of available datasets incurs novelty in mental health research. Giuntini *et al.* [35] introduced a method to evaluate the emotional behaviour and engagement of individuals with depression on social networks, combining network analysis with the extraction of emotional features from text to assess mood stages and communication patterns. Mendu *et al.* [36] proposed a hierarchical structure that links private message characteristics to mental health, offering insights into individual behaviour. Govindasamy and Palanichamy [37] examined depression as a major yet often overlooked issue, suggesting that machine learning algorithms could be used to detect sadness in SMP. The literature shows that DL models are highly efficient in identifying depression from textual data. To capture temporal relationships, we propose a Bi-LSTM based DL framework in this paper.

3. PROPOSED METHODOLOGY

This section describes the proposed state of the work, covering the DL framework, algorithm, dataset information, and evaluation strategy.

3.1 Problem Definition

This study tackles the challenge of creating a DL framework to automatically assess the probability of depression based SMP texts.

3.2 Proposed Framework

We propose a DL-based framework for the automatic detection of depression from SMP. The system architecture is shown in Figure 1. Data is gathered from Twitter posts and is processed to improve its excellence. The system then executes attribute mining and produces word embeddings, followed by training the proposed hybrid DL model

for DD. After training the model, it categorizes the test data samples. The proposed system consists of several modules: extracting data generated by internet users, cleaning raw data, and performing feature extraction to convert the data into a machine-readable format.

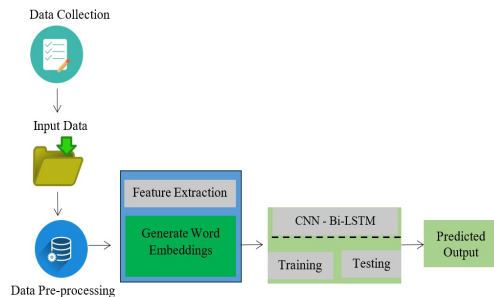


Figure 1: Summary Of The Proposed System Architecture

Depression categorization involves differentiating between tweets that suggest depression and those that do not. Several parameters are evaluated to comparability the performance of the proposed model with present classification techniques. In this study, a hybrid CNN-BiLSTM approach is utilized to forecast depression using Twitter datasets. This approach improves precision and prediction accuracy while minimizing classification errors.

Figure 2 illustrates, the DL based methodology follows a systematic process with several essential steps. Initially, a list of Twitter data, containing both normal and depressive symptoms, is imported. Preprocessing techniques are then applied to remove noise from the data. Proper data processing has an important positive effect on the quality of feature mining. The text data undergoes various preprocessing steps, including stop word removal, tokenization, data normalization, and punctuation removal. Feature extraction is performed on the cleaned data using an algorithm designed to identify relevant and meaningful features, after preprocessing.

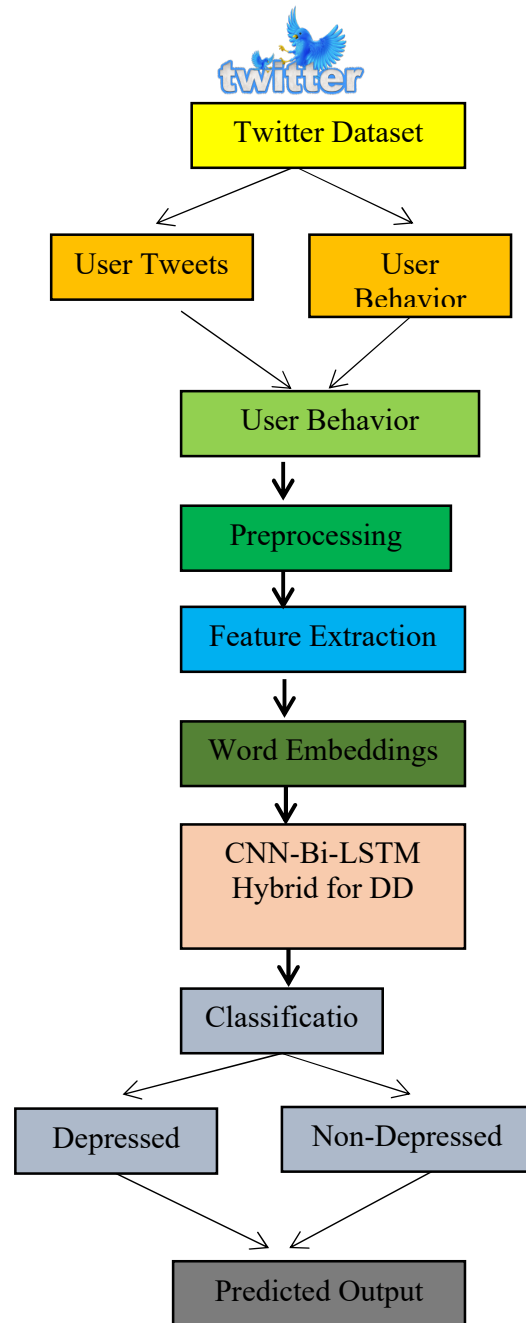


Figure 2: Proposed Functional Flow Of The DL Framework

These extracted features highlight the important data dimensions, enhancing the performance of categorization algorithms. A hybrid CNN-BiLSTM classification technique is employed to achieve improved accuracy. During both the training and validation phases, classifiers use the refined attributes derived from the feature extraction process. The next pace comprises assessing the

system's performance by computing metrics using the proposed DD analysis framework.

3.3 Pre-processing of Data

Practical data is typically composed through various methods and may not be domain-specific, leading to incomplete, erroneous, or unstructured data, so data preprocessing is a crucial step. Directly analysing such data often results in inaccurate predictions. Our framework incorporates several techniques during the preprocessing phase. The first method removes user-defined text patterns. It also removes blank strings, rows with NaN values, and replica entries. This step ensures that URLs are eliminated from all tweets, as they are irrelevant to the prediction and increase processing complexity. Following this, time, date, numbers, and hashtags are removed, as they do not contribute to depression prediction. While hashtags can be useful in some cases, they have been found to significantly reduce accuracy in this context, so they are excluded. Lastly, emojis and unnecessary whitespace are removed from the text to ensure cleaner data for analysis.

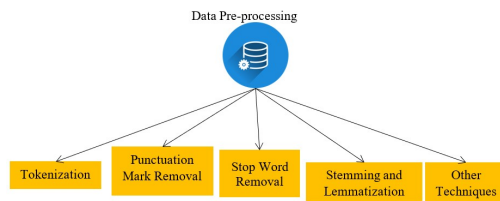


Figure 3: Data Pre-Processing

The following step involves removing stop words and performing stemming. Stop words, such as "are," "was," "at," "if," and others, do not contribute significant meaning to a statement, so they are eliminated. We use a set of stop words from the NLTK package for this purpose. Stemming is a method that reduces words to their root form, extracting prefixes or suffixes like "-ize," "-ed," "-s," or "-de." After cleaning the text, the next step is tokenization. Tokenization is an essential part of pre-processing in NLP, where a tokenizer breaks the text into smaller segments, such as words or phrases, using regular expressions. Figure 3 depicts the different pre-processing procedures employed.

The tokenization process begins by allocating the tokenizer with cleaned datasets of negative, and positive tweets by using the tokenization functions, through NLTK package. The `fit_on_texts()` method is then applied to create a vocabulary index based on word frequency. It appraises the core vocabulary by processing a list of texts, assigning the lowest

index value to the most frequently occurring word. This method generates an index for each word, with a maximum word count of 10,275. Following this, the `texts_to_sequences()` method is used to convert the words into numeric sequences. This method maps each word in a tweet to its equivalent numerical value from the word index, transforming the tweets into categorizations of numerical of varying lengths. Finally, any tweets shorter than the predefined maximum tweet length of 25 are padded with zeros to maintain uniformity.

3.4 Word Embeddings

Embedding techniques of NLP can help ML applications manage large datasets by representing words in low-dimensional vectors. These embeddings are dense, low-dimensional representations of words that would otherwise exist in high-dimensional, sparse vectors. Recent methods that use word vectors for learning, based on a given text corpus, often result in high-dimensional solutions, typically matching the size of the entire corpus. Words with similar meanings tend to be placed close to each other in the embedding space. For example, words like "happy" and "sad" are positioned far apart due to their contrasting meanings, making their semantic differences more distinguishable. By using an embedding layer, these relationships are captured, transforming the tokenized vectors to reflect semantic connections. Initially, tokenized vectors lack these relationships, but the embedding process helps reveal them based on the proximity of words in the space. As models like CNNs or RNNs are trained, they can better identify and separate features, enhancing their predictive capabilities. Embeddings are a powerful tool for encoding text, such as sentences or paragraphs, and are considered one of the major breakthroughs in deep learning for solving complex NLP challenges.

For each pre-processed data point, we generated a numerical vector using the "Word Embeddings" technique. Using the Keras text tokenizer, we initially transformed each to word indexes. We confirmed that the vocabulary length was properly adjusted and that no word received a zero index. Each word in the dataset was then allocated a exclusive index, which was used to create integral vectors sample for each text. To construct text sequences, we first determined the total length of all tweets. As illustrated by the histogram in Figure 6, utmost tweets in the training set contain fewer than 25 words, with the number of tweets rising as the word count increases. Consequently, text categorizations are converted into integer

categorizations, and zero padding is applied. Given that utmost tweets are shorter than 25 words, this length was set as the maximum allowed for the dataset. Since tweets exceeding this length are rare and only introduce zeros into the vector categorization, which can sluggishly model training and reduce performance, we removed any additional five words. In this study, we constructed an embedding matrix with dimensions $\text{Max_unique_words} * \text{Embedding_dim}$. The embedding_dim is 300, which represents the length of the vector. Initially, the matrix is filled with zeros. We then populate this matrix with vectors corresponding to 10,275 exclusive words, each vector having 300 features. Using an embedding layer of length 50 in our DL network, we produced output embedding vectors of size 25×50 for each tokenized vector. The DL network model aims to reconstruct the linguistic context of words by leveraging a large text corpus (the EMBEDDING_FILE) as input. This process results in a unique vector for each word in the corpus, creating a vector space that typically has hundreds of dimensions. Finally, the data was divided into negative and positive sets for training and validation, with 30% of the data used for testing and 70% for training.

3.5 Hybrid Deep Learning Model

We proposed a method that combines CNN with Bi-LSTM, to achieve improved classification performance for predicting depression in Twitter users (as shown in Figure 4). After conducting several experiments, we discovered that CNNs are effective when contextual information from previous sequences is not necessary, and are efficient in mining spatial features. In contrast, RNNs excel in situations where classification depends on the context provided by surrounding elements. The process begins by feeding multidimensional data directly into the CNN as low-level input. Each layer extracts relevant features during the convolution and pooling operations. Unlike traditional CNNs, both the output and hidden layers are fully connected. By employing multiple convolutional layers (CLs), pooling layers, and modified convolutional kernels (CNs), depression-related tweets are identified with greater detail, resulting in enhanced accuracy. However, the complexity of the network can lead to the risk of overfitting.

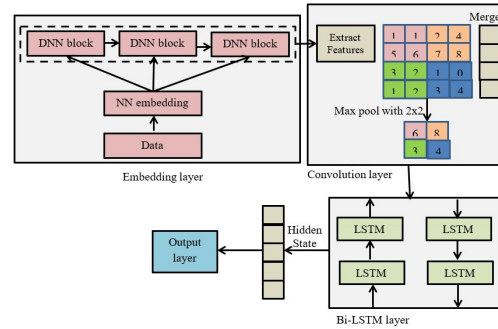


Figure 4: Proposed Hybrid DL Model

To address the sequence problem, the CNN model integrates the LSTM network. This allows important information to be retained in the state cell for extended periods and extracted alongside related or redundant data using CKs. The combination of CNN and Bi-LSTM contributes to most of the results. The Bi-LSTM architecture is illustrated in Figure 5.

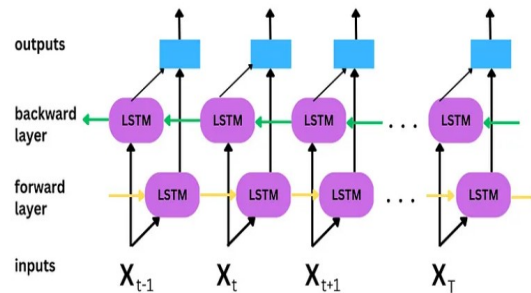


Figure 5: Architectural Overview Of Bi-LSTM Model

After the CL reduces the dimensions, Bi-LSTM is employed to facilitate the extraction of low-dimensional semantic features from the text. Additionally, the text is processed by Bi-LSTM as a sequence of inputs. This approach improves performance on the input vectors by combining several 1D CKs. As defined in Eq. 1, sequential input data is represented by the average of the embedding vector for each word. By means of the various CK sizes, 1D CNN is applied the $X_{1:T}$ features for Unigram, Bigram, and Trigram. The input is made up of the features generated in the t^{th} convolution, which is the process of taking a window of d words extending from t_0 to $t_0 + d$. The convolution procedure given in Eq. (2) derives features for the window.

$$X_{1:T} = [x_1, x_2, x_3, x_4, \dots, x_T] \quad (1)$$

$$h_{d,t} = \text{tan}^{-1} h(W_d x_{(t:t+d-1)+b_d}) \quad (2)$$

h_d represents the embedding vector for each unique word within the context window defined by $x_{t-d+1:t}$, where the parameters include a learnable weight matrix and b_d is the bias term. Individual filter applies convolution to diverse segments of the text, and ensuing feature map from the filter, as described in Eq. 3, corresponds to a convolution of size d .

$$h_d = [h_{d1}, h_{d2}, h_{d3}, h_{d4}, \dots, x_{(T-d+1)}] \quad (3)$$

CNNs can capture the hidden relationships amid adjacent words by using multiple CKs of different sizes. Convolutional filters are effective for extracting features from text because they reduce the number of trainable parameters during the attribute learning process. To enhance this, a max-pooling layer is added after the CLs. The input is first processed through multiple convolutional channels, each containing its own set of values. The max-pooling operation then selects the highest value from each CL to form a new set of features. Max-pooling is applied to the attribute maps of each CK with a size of d , as shown in Eq. 4. Concatenating p_d for each of the filter sizes $d = 1, 2, \text{and } 3$ results in the final features of each window are recovered. The concealed features of the bigram, unigram, and trigram are then mined as shown by Eq. 5.

$$p_d = \text{Max}_t(h_{d1}, h_{d2}, h_{d3}, h_{d4}, \dots, x_{(T-d+1)}) \quad (4)$$

$$h_d = [p_1, p_2, p_3] \quad (5)$$

The key advantage of using a CNN-based feature extraction method compared to traditional LSTM is the substantial reduction in the overall number of features. Once the features are extracted, they are further utilized by the depression prediction model. To address the "vanishing gradient" issue often encountered with sequential data, LSTM architectures incorporate gate structures such as forget, output, and input gates along with cell states. These serve as a shared long-term memory for the LSTM unit and provide additive connections between the different states. Eq. 6-11 determine the output state of an LSTM cell at a specific time t , based on the input x_t and intermediate state h_t .

$$f_t = \sigma(W_f x_t + U_f h_{(t-1)} + b_f) \quad (6)$$

$$i_t = \sigma(W_i x_t + U_i h_{(t-1)} + b_i) \quad (7)$$

$$o_t = \sigma(W_o x_t + U_o h_{(t-1)} + b_o) \quad (8)$$

$$g_t = \tanh(W_g x_t + U_g h_{(t-1)} + b_g) \quad (9)$$

$$c_t = f_t \odot c_{(t-1)} + i_t \odot g_t \quad (10)$$

$$h_t = o_t \odot \tanh(c_t) \quad (11)$$

In this model, the parameters that can be learned are denoted as W, U , and b , with the convolution operation ($\odot$) and the sigmoid activation function (σ) being clearly defined. The LSTM gates output, input, and forget are signified as o_t, i_t and f_t individually. The memory or cell state is

represented by c_t . The ability of LSTMs to handle longer sequences is largely due to the cell state, which captures long-term dependencies in the input data. The CNN component of the network consists of three CLs with varying filter counts. The first two layers each have 128 filters with a 3×3 kernel size and use sigmoid and ReLU activation functions. The third CL includes 64 filters with a sigmoid function as activation and a 3×3 size kernel. A 4×4 kernel-sized Max-pooling layer follows this. The model also includes a BiLSTM with slightly different hidden computations. To prevent over fitting on the training set, a dropout layer with a 0.1 is used. The model is optimized using the RMSprop algorithm and employs the binary cross-entropy loss function. The ReLU activation function is applied in the output layer.

3.6 Proposed Algorithm

We introduced an algorithm called LBDD, which takes tweets of Twitter as input data and classifies them based on the probability of indicating depression.

Algorithm 1: LBDD

Input: Twitter dataset D

Output: DD results R , performance statistics P

Begin

1. $D \leftarrow \text{Preprocess}(D)$
2. $\text{features} \leftarrow \text{FeatureExtraction}(D)$
3. $\text{embeddings} \leftarrow \text{WordEmbeddings}(\text{features})$
4. $(T1, T2) \leftarrow \text{DataSplit}(D, \text{embeddings})$
5. Configure CNN – BiLSTM hybrid model m (as in Figure 4)
6. Compile m
7. Train m using $T1$
8. Save m for future reuse
9. Load m
10. $R \leftarrow \text{Test}(m, T2)$
11. $P \leftarrow \text{Evaluate}(R, \text{ground truth})$
12. Display R , and P

End

As presented in Algorithm 1, it takes Twitter dataset as input and performs various mechanisms to detect depression probabilities based on SMP. The data undergoes pre-processing to enhance its quality for supervised learning. The hybrid DL model is trained with the word embeddings and the model is saved for reuse. The test data, consisting

of tweets, is processed through DD by the trained hybrid DL model. This leads to the classification of the validation data. Performance is assessed by comparing the predicted labels with the ground truth.

3.7 Evaluation Methodology

Since we used to learn based approach (supervised learning), metrics derived from confusion matrix, shown in Figure 6 are used for evaluation our methodology.

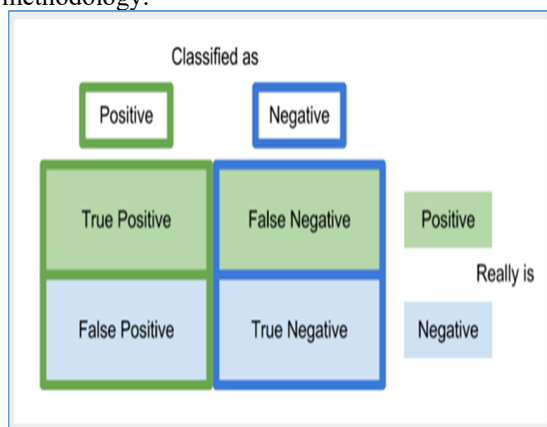


Figure 6: Confusion Matrix

Statistical performance is derived by associating the predicted labels with the truth labels using the confusion matrix. Equations 12-15 define the various metrics utilized in the performance assessment.

$$\text{Accuracy} = (\text{TP} + \text{TN}) / (\text{TP} + \text{TN} + \text{FP} + \text{FN}) \quad (12)$$

$$\text{Precision (p)} = \text{TP} / (\text{TP} + \text{FP}) \quad (13)$$

$$\text{F1-score} = 2 * ((p * r) / ((p + r))) \quad (14)$$

$$\text{Recall (r)} = \text{TP} / (\text{TP} + \text{FN}) \quad (15)$$

4. EXPERIMENTAL RESULTS



Figure 7: Random Tweets Word Cloud

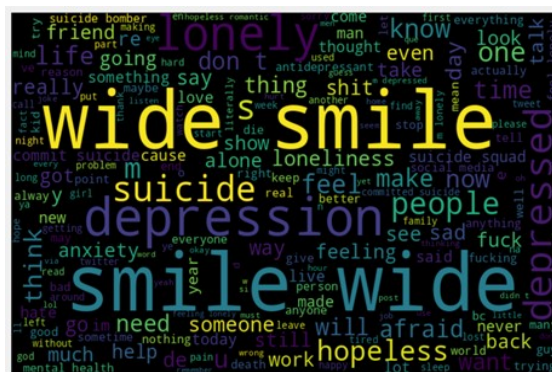
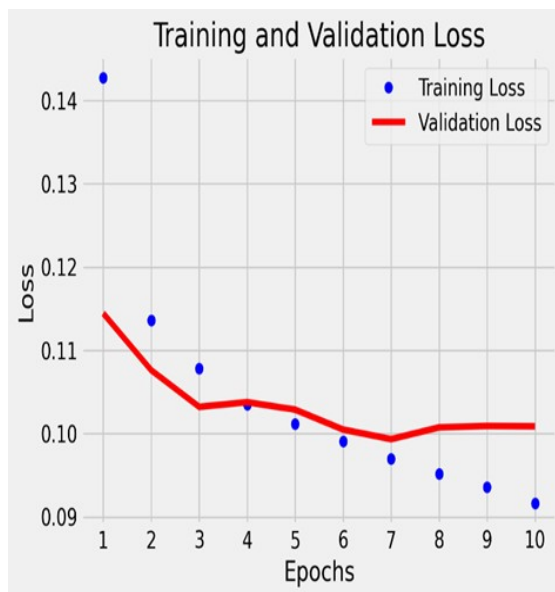
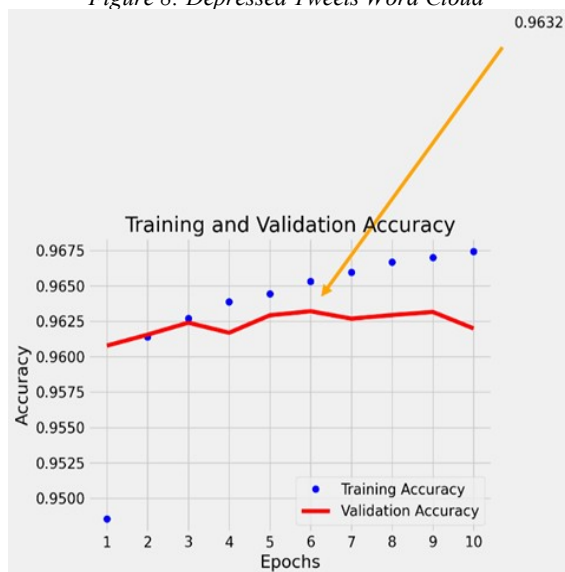


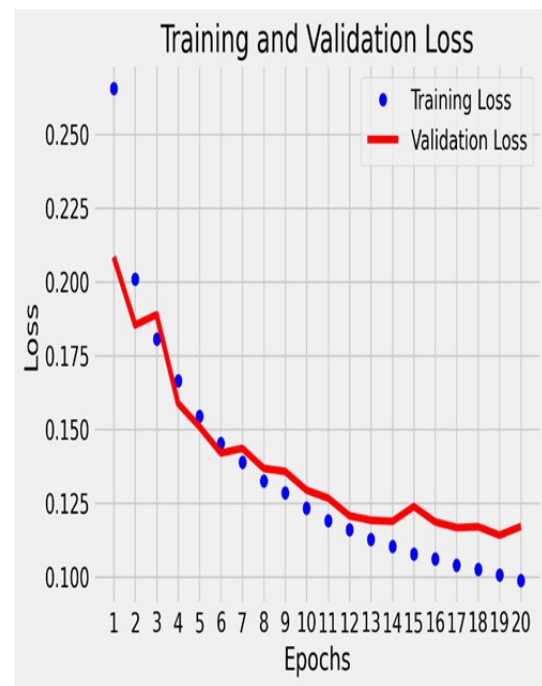
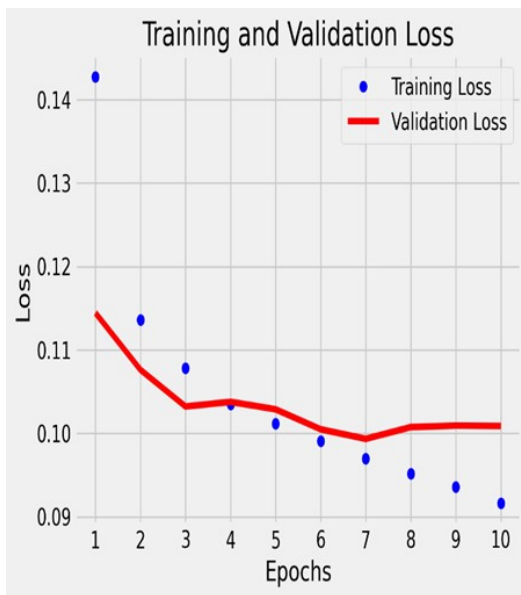
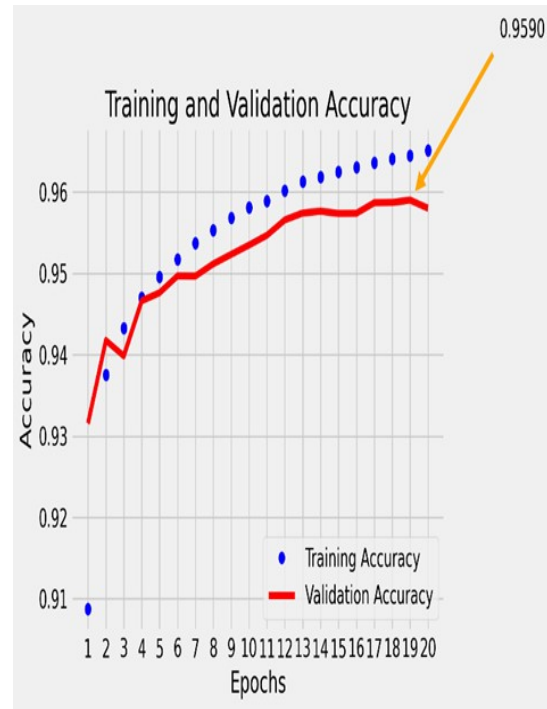
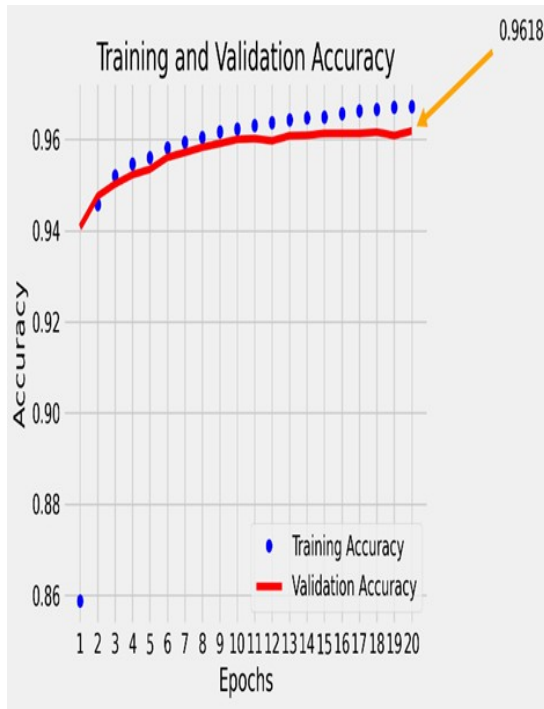
Figure 8: Depressed Tweets Word Cloud



(A) Accuracy During Training And Validation

(B) Loss During Training And Validation

Figure 9: LSTM Model Summary



(A) Accuracy During Training And Validation

(B) Loss During Training And Validation

Figure 10: Simple RNN

(A) Accuracy During Training And Validation

(B) Loss During Training And Validation

Figure 11: CNN-LSTM

Table 2: Model Comparison

Model	Accuracy (%)	Precision	Recall	F1-Score
Simple RNN	93.04	0.9023	0.895	0.899
LSTM	93.98	0.925	0.91	0.9023
CNN-LSTM	96.32	0.96	0.96	0.96

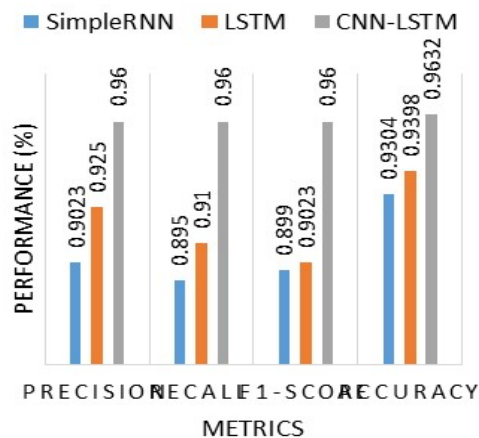


Figure 12: Model Comparison

5. DISCUSSION

This paper presents a DL based framework for the automatic DD in SMP. A review of existing literature revealed that DL models like CNN are effective for extracting contextual information from tweets, while LSTM models are more suited for analysing data related to DD. Building on this insight, we proposed a novel hybrid model of DL that integrates CNN and Bi-LSTM to enhance detection performance. The framework leverages preprocessing techniques and word embeddings to generate data suitable for supervised learning. The dataset, collected from Twitter, is annotated for supervised learning purposes. The proposed hybrid DL model surpasses existing models in performance, as indicated by the experimental results. However, certain limitations are discussed in Section 5.1.

5.1 Limitations

The proposed DL framework for the automatic DD in SMP shows promising performance. However, it has certain limitations. The generalizability of the results may be hindered by the relatively small size of the dataset used in the experiments. A concern is

the absence of data multiplicity, as the data was composed from only a solitary source. A significant drawback is the lack of a distributed programming framework and the absence of big data concepts. Considering the vast amount of social media data, integrating big data techniques and distributed programming models would enhance the scalability and efficiency of the framework.

5.2 Light of Findings

The previous work investigated phrase analysis for DD on social media, recommending text classification through Graph Attention Networks (GATs). Using Reddit Depression data, their model achieved a 0.91 ROC score and DL framework created on a hybrid model that seamlessly integrates CNN and BiLSTM to capture both feature representations and temporal relationships in the data. Using algorithm, LBDD, which processes Twitter tweets achieving an accurateness of 96.32%.

6. CONCLUSION AND FUTURE WORK

The paper presents a DL framework created on a hybrid model that seamlessly integrates CNN and BiLSTM to capture both feature representations and temporal relationships in the data. Our approach combines CNN with Bi-LSTM to enhance classification performance for predicting DD in Twitter users. Over extensive investigation, we detected that CNN excels when contextual information from previous sequences is not essential and is effective in extracting spatial features. In contrast, RNNs are better suited for tasks that require understanding the context of surrounding data for classification. The proposed framework enables automatic DD from SMP. The algorithm, LBDD, which processes Twitter tweets and classifies them based on the likelihood of depression. After evaluating our method using a promising dataset, we evaluated that our DL model outperformed several prevailing models, achieving an accurateness of 96.32%. The strengths of my work are LBDD algorithm & dataset. The weakness of my work is the absence of big data concepts. The future work is to enhance our framework to analyze human expressions in images or videos and audio content for automatic depression detection

REFERENCES

- [1] Paffenbarger RS Jr, Lee IM, Leung R (1994) Physical activity and personal characteristics associated with depression and suicide in American college men. *Acta Psychiatrica Scandinavica* 89:16–22. <https://doi.org/10.1111/j.1600-0447.1994.tb05796.x>.
- [2] Zafar A, Chitnis S (2020) Survey of depression detection using social networking sites via datamining. In: 2020 10th International Conference on Cloud Computing, Data Science & Engineering (Confluence). IEEE, pp 88–93. <https://doi.org/10.1109/Confluence47617.2020.9058189>.
- [3] Kohrt BA, Speckman RA, Kunz RD, Baldwin JL, Upadhyaya N, Acharya NR, ... Worthman CM (2009) Culture in psychiatric epidemiology: using ethnography and multiple mediator models to assess the relationship of caste with depression and anxiety in Nepal. *Annals of Human Biology* 36(3):261–280. <https://doi.org/10.1080/03014460902839194>.
- [4] Oquendo MA, Ellis SP, Greenwald S, Malone KM, Weissman MM, Mann JJ (2001) Ethnic and sex differences in suicidal ideation and major depression in the United States. *American Journal of Psychiatry* 158(10):1652–1658. <https://doi.org/10.1176/appi.ajp.158.10.1652>.
- [5] Beard C, Millner AJ, Forgeard MJ, Fried EI, Hsu KJ, Treadway MT, ... Björgvinsson T (2016) Network analysis of depression and anxiety symptom relationships in a psychiatric sample. *Psychological Medicine* 46(16):3359–3369. <https://doi.org/10.1017/S0033291716002300>.
- [6] Biradar A, Totad SG (2018) Detecting Depression in Social Media Posts Using Machine Learning. In: International Conference on Recent Trends in Image Processing and Pattern Recognition. Springer, Singapore, pp 716–725. https://doi.org/10.1007/978-981-13-9187-3_64.
- [7] Arora P, Arora P (2019) Mining twitter data for depression detection. In: 2019 International Conference on Signal Processing and Communication (ICSC). IEEE, pp 186–189. <https://doi.org/10.1109/ICSC45622.2019.8938353>.
- [8] Shreya Ghosh and Tarique Anwar; (2021). *Depression Intensity Estimation via Social Media: A Deep Learning Approach*. *IEEE Transactions on Computational Social Systems*. <http://doi:10.1109/TCSS.2021.3084154>
- [9] Keshu Malviya; Bholanath Roy and SK Saritha; (2021). *A Transformers Approach to Detect Depression in Social Media*. 2021 International Conference on Artificial Intelligence and Smart Systems (ICAIS). <http://doi:10.1109/icaais50930.2021.9395943>
- [10] Chenhao Lin; Pengwei Hu; Hui Su; Shaochun Li; Jing Mei; Jie Zhou and Henry Leung; (2020). *SenseMood: Depression Detection on Social Media*. *Proceedings of the 2020 International Conference on Multimedia Retrieval*. <http://doi:10.1145/3372278.3391932>
- [11] Giuntini, Felipe T.; Cazzolato, Mirela T.; dos Reis, Maria de Jesus Dutra; Campbell, Andrew T.; Traina, Agma J. M. and Ueyama, Jô (2020). *A review on recognizing depression in social networks: challenges and opportunities*. *Journal of Ambient Intelligence and Humanized Computing*. <http://doi:10.1007/s12652-020-01726-4>
- [12] Shini Renjith, Annie Abraham, Surya B. Jyothi, Lekshmi Chandran and Jincy Thomson. (2022). An ensemble deep learning technique for detecting suicidal ideation from posts in social media platforms. *Elsevier*. 34(10), pp.9564-9575. <https://doi.org/10.1016/j.jksuci.2021.11.010>
- [13] Yang, Xingwei; McEwen, Rhonda; Ong, Liza Robee and Zihayat, Morteza (2020). *A big data analytics framework for detecting user-level depression from social networks*. *International Journal of Information Management*, 54, 102141–. <http://doi:10.1016/j.ijinfomgt.2020.102141>
- [14] Kailai Yang, Tianlin Zhang and Sophia Ananiadou. (2022). A mental state Knowledge-aware and Contrastive Network for early stress and depression detection on social media. *Elsevier*. 59(4), pp.1-16. <https://doi.org/10.1016/j.ipm.2022.102961>
- [15] Tapotosh Ghosh, Md. Hasan Al Banna, Md. Jaber Al Nahian, Mohammed Nasir Uddin, M. Shamim Kaiser and Mufti Mahmud. (2023). An attention-based hybrid architecture with explainability for depressive social media text detection in Bangla. *Elsevier*. 213(C), pp.1-17. <https://doi.org/10.1016/j.eswa.2022.119007>
- [16] Esteban A. Rissola; David E. Losada and Fabio Crestani; (2021). *A Survey of Computational Methods for Online Mental State Assessment on Social Media*. *ACM*

- Transactions on Computing for Healthcare.* <http://doi:10.1145/3437259>
- [17] Ziwei, Bernice Yeow and Chua, Hui Na (2019). *Proceedings of the 2nd International Conference on Computing and Big Data - ICCBD 2019 - An Application for Classifying Depression in Tweets.* 37–41. <http://doi:10.1145/3366650.3366653>
- [18] Tariq, Subhan; Akhtar, Nadeem; Afzal, Humaira; Khalid, Shahzad; Mufti, Muhammad Rafiq; Hussain, Shahid; Habib, Asad and Ahmad, Ghufuran (2019). A novel Co-training based approach for the classification of mental illnesses using social media posts. *IEEE Access*, 1–1. <http://doi:10.1109/ACCESS.2019.2953087>
- [19] Sancheng Peng, Lihong Cao, Yongmei Zhou, Zhouhao Ouyang, Aimin Yang, Xinguang Li, Weijia Jia and Shui Yu. (2022). A survey on deep learning for textual emotion analysis in social networks. *Elsevier*. 8(5), pp.745-762. <https://doi.org/10.1016/j.dcan.2021.10.003>
- [20] Ruba Skaik and Diana Inkpen; (2021). Using Social Media for Mental Health Surveillance. *ACM Computing Surveys*. <http://doi:10.1145/3422824>.
- [21] Adikari, Achini; Gamage, Gihan; de Silva, Daswin; Mills, Nishan; Wong, Sze-Meng Jojo and Alahakoon, Duminda (2020). A self structuring artificial intelligence framework for deep emotions modelling and analysis on the social web. *Future Generation Computer Systems*, S0167739X20330053–. <http://doi:10.1016/j.future.2020.10.028>
- [22] Chatterjee, Ankush; Gupta, Umang; Chinnakotla, Manoj Kumar; Srikanth, Radhakrishnan; Galley, Michel and Agrawal, Puneet (2018). Understanding emotions in text using deep learning and big data. *Computers in Human Behavior*, S0747563218306150–. <http://doi:10.1016/j.chb.2018.12.029>
- [23] Yang, T., Li, F., Ji, D., Liang, X., Xie, T., Tian, S., and Liang, P. (2021). Fine-grained depression analysis based on Chinese micro-blog reviews. *Information Processing & Management*, 58(6), 102681. <http://doi:10.1016/j.ipm.2021.10.2681>
- [24] Xiaoxu Yao; Guang Yu; Jingyun Tang and Jialing Zhang; (2021). Extracting depressive symptoms and their associations from an online depression community. *Computers in Human Behavior*. <http://doi:10.1016/j.chb.2021.106734>
- [25] Lei Cao; Huijun Zhang and Ling Feng; (2022). Building and Using Personal Knowledge Graph to Improve Suicidal Ideation Detection on Social Media. *IEEE Transactions on Multimedia*. <http://doi:10.1109/tmm.2020.3046867>
- [26] Sanja Scepanovic; Enrique Martin-Lopez; Daniele Quercia and Khan Baykaner; (2020). Extracting medical entities from social media. *Proceedings of the ACM Conference on Health, Inference, and Learning*. <http://doi:10.1145/3368555.3384467>
- [27] Pran, Md. Sabbir Alam; Bhuiyan, Md. Rafiuzzaman; Hossain, Syed Akhter and Abujar, Sheikh (2020). 11th International Conference on Computing, Communication and Networking Technologies (ICCCNT) - Analysis of Bangladeshi People's Emotion During Covid-19 In social media Using Deep Learning. 1–6. <http://doi:10.1109/ICCCNT49239.2020.9225500>
- [28] Guillermo Blanco and Analia ' Lourenço. (2022). Optimism and pessimism analysis using deep learning on COVID-19 related twitter conversations. *Elsevier*. 59(3), pp.1-15. <https://doi.org/10.1016/j.ipm.2022.102918>
- [29] Subramani, Sudha; Michalska, Sandra; Wang, Hua; Du, Jiahua; Zhang, Yanchun and Shakeel, Haroon (2019). *Deep Learning for Multi-Class Identification from Domestic Violence Online Posts*. *IEEE Access*, 7, 46210–46224. <http://doi:10.1109/ACCESS.2019.2908827>
- [30] Nawshad Farruque, Randy Goebel, Sudhakar Sivapalan and Osmar Zaiane. (2024). Deep temporal modelling of clinical depression through social media text. *Elsevier*. 6, pp.1-14. <https://doi.org/10.1016/j.nlp.2023.100052>
- [31] ABHINAV KUMAR, JYOTI KUMARI and JIESTH PRADHAN. (2023). Explainable Deep Learning for Mental Health Detection from English and Arabic Social Media Posts. *ACM*, pp.1-18. <https://doi.org/10.1145/3632949>
- [32] USMAN AHMED, JERRY CHUN-WEI LIN and GAUTAM SRIVASTAVA. (2023). Graph Attention Network for Text Classification and Detection of Mental Disorder. *ACM*. 17(3), pp.1-31. <https://doi.org/10.1145/3572406>
- [33] Jesia Quader Yuki; Md. Mahfil Quader Sakib; Zaisha Zamal; Sabiha Haque Feel and Mohammad Ashrafuzzaman Khan; (2020). *Detecting Depression from Human Conversations*. *Proceedings of the 8th*

International Conference on Computer and Communications

Management. <http://doi:10.1145/3411174.3411187>

- [34] Esteban A. Rissola; Seyed Ali Bahrainian and Fabio Crestani; (2020). A Dataset for Research on Depression in Social Media . Proceedings of the 28th ACM Conference on User Modeling, Adaptation and Personalization. <http://doi:10.1145/3340631.3394879>
- [35] Felipe Taliar Giuntini; Kaue L. de Moraes; Mirela T. Cazzolato; Luziane de Fatima Kirchner; Maria de Jesus D. Dos Reis; Agma J. M. Traina; Andrew T. Campbell and Jo Ueyama; (2021). Modeling and Assessing the Temporal Behavior of Emotional and Depressive User Interactions on Social Networks. IEEE Access. <http://doi:10.1109/access.2021.3091801>
- [36] Sanjana Mendu; Anna Baglione; Sonia Bae; Congyu Wu; Brandon Ng; Adi Shaked; Gerald Clore; Mehdi Boukhechba and Laura Barnes; (2020). A Framework for Understanding the Relationship between Social Media Discourse and Mental Health . Proceedings of the ACM on Human-Computer Interaction. <http://doi:10.1145/3415215>
- [37] Kuhaneswaran AL Govindasamy and Naveen Palanichamy; (2021). Depression Detection Using Machine Learning Techniques on Twitter Data. 2021 5th International Conference on Intelligent Computing and Control Systems (ICICCS). <http://doi:10.1109/ICICCS51141.2021.9432203>