

Depression Classification on Social Media Posts using Deep Learning

Thimmaiah K B¹ and Mrs. Rashmi M. R²

¹Department of Information Technology, National Institute of Engineering, Mysuru, Karnataka, India
Email: ratanthimmaiah@gmail.com

²Department of CSE, National Institute of Engineering, Mysuru, Karnataka, India
Email: Rashmi.mr7@nie.ac.in

Abstract—To use a hybrid methodology of Convolutional Neural Networks (CNN) and Long Short-Term Memory (LSTM), Natural Language Processing (NLP) techniques like Word encoding, tokenization, and factorization to recognise words that convey emotional responses of depressive symptoms and connected sentiments, we propose a novel approach for recognising posts on social network that are still indicative of anxiety. By analysing the nearby environment of phrases, deep learning is used to accurately predict sentiment in words and eliminate false - positive. This research used information that was harvested from well social networking site Kaggle. Labels are used to gather posts around same subject collectively when the information is evaluated. The negative class provides postings from those other, randoms, whereas the positive class has comments across different clusters addressing hopelessness and self-harm. As a result of the variety in the negative class, it is possible to claim that the dataset is representative of an actual situation. The proposed method can be used in a variety of contexts, such as the identification of possibly harmful behaviors on provided by social media services that are popular among young people.

Index Terms— NLP, CNN, LSTM.

I. INTRODUCTION

Contributions Social networking have become the go-to tool for communication, information, and eventually political participation; therefore, content filtering has now become crucial for social media organizations. Contents moderation—the process of finding and removing or exposing potentially harmful contents on social networking platforms—is widely seen as a crucial responsibility of social networking companies. Many government and judiciary bodies are started investigating the negative effects of social networking on false information, anti-social conduct, harassment, as well as other problems.

Such a domain wherein information moderating is essential is the identification and monitoring of social networking comments those are suggestive of psychological disease and self. Research show that persons commonly freely communicate their struggles on social networks in conditions of depressed and self, making rapid and precise identifying of this data vital for protection and a great help. The existences of a sad, empty, or irritable mood along with physical and cognitive problems that significantly limit the person's ability to operate are characteristics of depressed diseases. Social networking has emerged as a crucial forum for thought-sharing via posting thanks to specialised comments threads involving the like people and for providing a place where in

one need not worry about getting judged. With Twitter being with us preferred social media platform, this paper has inspired us to examine online tweets as a first step in identifying behavioural patterns associated with anxiety.

II. RELATED WORK

Including an efficiency rate of 80% on overall trained examples, Xinyu Wang et al. proposed a sentiment analysis-based method for identifying depressed individuals on social networking platforms utilising content from Chinese Micro-Blogs. They discovered that all multiple categorisation methods Bayes Network, a Tree Classification, and a Rule-Based Decision Table—behaved similarly [1].

On Tweets, Vishal A. Kharde examined various sentiment analytical techniques and compared Naive Bayes (NB) and Support Vector Machines (SVM) to Ngrams. They achieved correctness with Individual words and Bigrams utilizing NB of 74.56 % & 76.68 %, both, as contrasted to precision utilizing SVM of 76.44 % & 77.73% [2].

Yogesh Garg focused on comparing results from various data sources while utilizing Highest Entropy to analyse the sentiment of twitter posts. It teaches us how to use and analyze N-grams. By integrating Unigrams, Bigrams, Trigrams, and a Negation cue using a NB Classifier, they were able to get a 75.3 % accuracy [3].

XGBoost (XCB) surpassed the competition having an accuracy level of almost 70%, according to Aziguli Wulamu's evaluation of the outcomes of XGB, SVM, and RF Classifiers in combination using CountVectorizer's N-gram modeling for sentiment classification [4].

The outcomes of Decision Tree (DT) F-measure (F-m) 71%, K-nearest neighbours F-m 60%, SVM F-m 71%, and Ensemble classification F-m 64% methods for identifying depressive comments on Fb were compared by Md. Rafqul Islam et al., who discovered that DT would have the higher precision [5].

By having students respond to a survey, Swati Jain et al. were able to get current information and evaluate both data groups. Researcher utilised ML techniques like Logistic Regression (LR), RF Classifier, XGB Classifier, and SVM to categorize the information. On dataset 1, XGB had the higher precision (83.87%), although on datasets 2, the highest precision was achieved by LR. Additionally, researchers found the LR started to fail whenever there were a lot of aspects, lost information sets, and too numerous category variables [6].

Ahmed Mohammed compared CNN, Multi-Channel CNN, Multi-Channel-Pooling CNN, and BiLSTM for the purpose of detecting anxiety in People on twitter. Their top classifier, a Multi-Channel CNN, had an accuracy rating of 83.1% [7].

The application of an LSTM neural network for sentiment classification on internet information was developed by researchers including Baziotis with encouraging outcomes [8].

With a categorization efficiency of 70%, Munmun De Choudhury used SVM to identify sadness using social networks. To predict when anxiety may attack, researchers examined at the biological disposition to the condition and the linguistic characteristics of sad individuals [9].

III. ALGORITHMS

A. CNN-LSTM hybrid model

Due to CNN's adaptive & self-learning extracting features, our method could get around the reliance of traditional recognizing techniques on features collection and information rebuilding related to personal knowledge as personal perception. Proposed method kernels are utilised to search the entire molten pool image for duplication attributes in all potential members. Now at product hybrid phase, a one-dimensional information vectors is extended out of a three-dimensional information tensor that CNN's last level generated. In addition to any blank, similar, and unnecessary values, this rate moves all featured data that the convolutional kernels were able to collect. Each feature is next converted to two-dimensional area before being fed into LSTM.

Since CNN's features extractor is flexible and self-learning, this could remove the dependence of classical recognized methods on features harvest and information rebuilding related to personal knowledge and personal awareness. Multiple convolutional kernels are utilised to search the entire molten pool image for duplication attributes in all potential members. At the features hybridization phase, a one information set is extended out of the three-dimensional features tensor that CNN's last layer produces. In addition to a certain blanks, repetitive, and unnecessary input, this vector includes all features extracted that the convolution kernels were able to collect. The extracted features are next formed by combining data before being fed in through LSTM.

Every line to in features space corresponds to different features. The treatment of every data set rows as just a basic unit to also be hybridized. Every sequence divides features matrices into multiple time scales & examines a

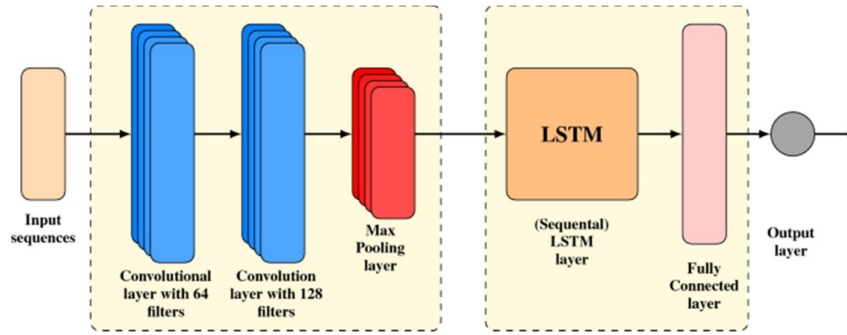


Figure 1. CNN-LSTM Hybrid Architecture

line of features information. The molten pool's unique visual is transform to "sequencing" info to use this method. The LSTM network is used to find relationships amongst every row of such information matrices in order to reduce and hybridize the information produced by the CNN network.

Figure 1 shows the CNN system as an illustration of invention. In the timespan of a Convolutionary neural identifying molten pool image, the intake of a LSTM at period t include the outputs h_{t-1} and unit condition $ct-1$ at time $t-1$ and even the show's intake x_t of moment. These useful selected properties from of the CNN could been preserved in cells mode for as longer period whereas a incorrect features are ignored thanks to multiple correctly constructed thresholds structure that can filtrate and hybridize the featured tensor with cell state.

IV. PROPOSED SYSTEM

The dataset from Kaggle was utilized in this study. NLP approaches are employed in the proposed study for data cleansing and vectorization. The following are the steps involved:

A. Data pre-processing:

NLP comprises a number of text categorization pre-processing methods. The following four major pre-processing tasks are carried out in this work:

1. Clean up your data: The data cleansing in this example is accomplished by changing all uppercase letters in the mail to lowercase ones.
2. Removal of special characters: One of its most important text processing operations is the removal of special characters; there seem to be a total of 32 fundamental punctuation marks that must be handled. The string module has 32 punctuation options, which are listed below:
`!"#$%&'()*+,-./:;<=>?@[\\]^_`{|}~'`
3. Tokenization: The process of breaking up the string of data using items could be word, letters, sentences, or any more expressing elements—is known as tagging. Tickets are separated by "white space letters," including which "spaces," "ending of a route "or" grammar and spelling character types."
4. Removal of Stopwords: Stopwords are the most commonly occurring words in a paragraph that provide no useful information. Stop - words include words like then, that, when, these, and many others. For deleting stopwords, the Spacy library is a valuable NLP tool.
5. Stemming: The process of separating involves distilling a word down to its fundamental stems source. For example, the words run, running, runs, and runned all come from those similar Latino root, "run." Usually stem removes prefixes and suffixes from phrases like s, ing, es, and other similar terms. The NLTK package is used to stem the phrases.

B. Vectorization

The term "vectorization" describes a conventional technique for converting text-based source information into vector of real numbers, this is the structure supported by ML models. A step in the ML feature extraction process is vectorization. The intention is to identify certain distinctive properties in the material is for system to trained on by tokenizing phrases using Tokenizer method.

C. Data splitting

In information these stage, we shall separate the information into two sets: like training dataset as well as the testing dataset.

D. Model education

The Machine Learning algorithm is taught by providing datasets to it at this phase. This is the period during which the student learns. Regular training can increase the Machine Learning model's prediction rate greatly.

E. Data Visualization

Data visualization is the visual representation of material and data in a graphically or physical way (Example: charts, graphs, and maps). To adjust the hyperparameter or examine underfitting or overfitting, visualization is employed. For purpose of detecting depressions, a pie graph and a horizontally bars graphs were made using the Matplotlib visualization library.

Pie chart: A pie graph displays the share of a total number across ranges of a category dataset as a circular divided into radially segments. Every categories number corresponds to a unique circle slices, and the size & arcs lengths of every piece indicate how much of overall total every classification value takes up. Figure 2 shows a pie graph showing the volume disparity between recordings with and without depressed. More than 50 percent of a items are found in first grey slice, that stands for Non-Anxiety. Anxiety (green), which makes up almost of a whole, is in 2nd spot.

Size Comparision of Depression and Non-Depression records

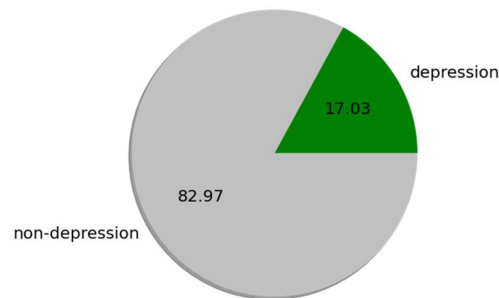


Figure 2. Pie chart

Bar chart: The title of the horizontal bar graph describes the data that the graph represents. The data categories are represented on the vertical axis. The data categories in Figure 3 are Depression and Non-Depression. The values corresponding to every data value are represented on the horizontal axis. The number of records is the data value here.

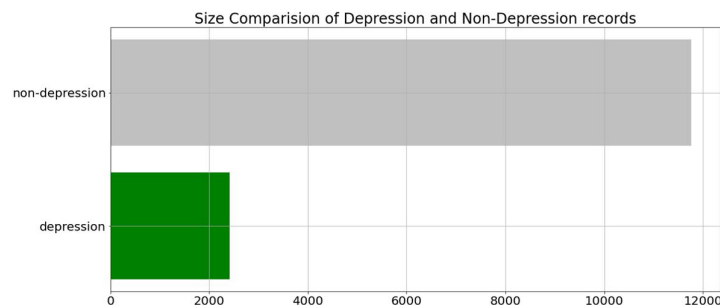


Figure 3. Horizontal bar graph

Word cloud: A text, often referred to simply a tagging cloud, is a phrase visualization which lists words in a phrase according to how often they are utilized, form tiny to big. Figure 4 displays the words clouds for the information.

V. RESULT AND ANALYSIS

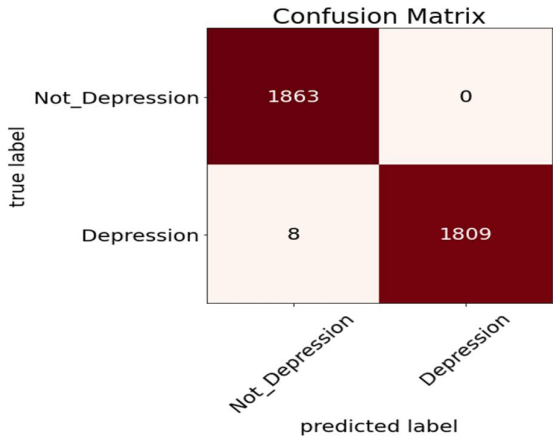
The machines models would be examined using the verification information. These helps determine whether model is valid. Finding the performance metrics depending on where the system is intended to perform is essential for establishing connection.



The Categorization Reports and Confusion Matrix for hybrid CNN+LSTM are created depending on a number of success indicators, including F1-score, Recall, Precision, and Support. Table I displays the hybrid CNN+LSTM classification for the anxiety class.

A classifier's performance on a testing set using actual values is shown in a graphs is called the confusion matrix.

Metrics	Precision	Recall	F1-score	Support
0	1.00	1.00	1.00	1863
1	1.00	1.00	1.00	1817
Accuracy			1.00	
Macro avg	1.00	1.00	1.00	3680
Weighted avg	1.00	1.00	1.00	3680



The training form validity and modeling lost graphs for the Classification algorithm are displayed in Fig 6 & Fig 7. Train accuracy and validating accuracy are both included in the predictive performance. Additionally, train loss vs. validating loss is included in the modeling loss.

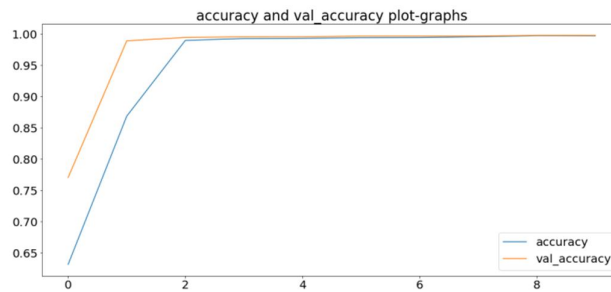


Figure 6. Model Accuracy for CNN+LSTM

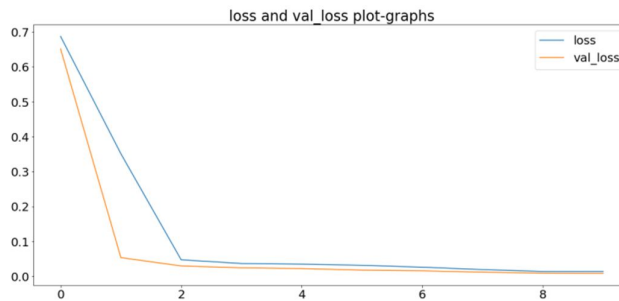


Fig. 7. Model Loss for CNN+LSTM

VI. CONCLUSION

The combination CNN-LSTM approach employed in this research to extract sentiment through texts could be applied to an adaptable, highly functional system having long-term reliance. A specific text's context - specific information in the concept is necessary. The classifier has reached outstanding accuracy of 99.78% in a variety of test dataset. Given the significance of the escalating mental health problems, this is a crucial step in identifying individuals on social networking who are having problems in their early stages. This content could also be utilized to draw attention to postings that are suicidal in nature, as individuals having been found to share pro-suicide articles on social networks. The recommended approach may be expanded to include multi-class categorisation of numerous other psychological disorders.

It must be improved to incorporate a pattern-finding study of the posts as the many additional variables that affect it. A doctor physically validating the information might make it even better by guaranteeing that every client receives assistance.

In the approach, this machine could be integrated with just about any social networking platform, making it easier for administrators to authenticate users who are depressed and take appropriate action. The development of such framework for web phones is a good idea.

REFERENCES

- [1] Xinyu Wang, Chunhong Zhang, Yang Ji, Li Sun, Leijia Wu, and Zhana Bao. A Depression Detection Model Based on Sentiment Analysis in Micro-blog Social Network. In Revised Selected Papers of PAKDD 2013 International Workshops on Trends and Applications in Knowledge Discovery and Data Mining - Volume 7867. Springer- Verlag, Berlin, Heidelberg, 201–213. (2013) https://doi.org/10.1007/978-3-642-40319-4_18
- [2] Kharde, Vishal & Sonawane, Sheetal. Sentiment Analysis of Twitter Data: A Survey of Techniques. International Journal of Computer Applications. (2016). doi:139. 5-15 10.5120/ijca.2016908625
- [3] Garg, Yogesh & Chatterjee, Niladri. Sentiment Analysis of Twitter Feeds. 33-52. (2014). doi: 10.1007/978-3-319-13820-6_3
- [4] Gaye, Babacar, and Aziguli Wulamu. "Sentimental Analysis for OnlineReviews using Machine learning Algorithms." International Research Journal of Engineering and Technology. (2019).
- [5] Islam, Md Rafiqul & Kabir, Ashad & Ahmed, Ashir & Kamal, Abu & Wang, Hua & Ulhaq, Anwaar. (2018). Depression detection from social network data using machine learning techniques. Health Information Science and Systems. doi:10.1007/s13755-018-0046-0

- [6] S. Jain, S. P. Narayan, R. K. Dewang, U. Bhartiya, N. Meena and V.Kumar, "A Machine Learning based Depression Analysis and Suicidal Ideation Detection System using Questionnaires and Twitter," 2019 IEEE Students Conference on Engineering and Systems (SCES), Allahabad, India, pp. 1-6, (2019) doi:10.1109/SCES46477.2019.8977211
- [7] Hussein Orabi, Ahmed & Buddhitha, Prasadith & Hussein Orabi, Mahmoud & Inkpen, Diana. Deep Learning for Depression Detection of Twitter Users. 88-97. (2018). doi:10.18653/v1/W18-0609
- [8] Baziotis, Christos & Pelekis, Nikos & Doukeridis, Christos. DataStories at SemEval-2017 Task 4: Deep LSTM with Attention for Message-level and Topic-based Sentiment Analysis. 747-754. (2017). doi:10.18653/v1/S17-2126
- [9] Gamon, Michael & Choudhury, Munmun & Counts, Scott & Horvitz, Eric. Predicting Depression via Social Media. Association for the Advancement of Artificial Intelligence. (2013).
- [10] Murthy, Dr & Allu, Shanmukha & Andhavarapu, Bhargavi & Bagadi, Mounika. Text based Sentiment Analysis using LSTM. International Journal of Engineering Research and. V9. (2020). doi:10.17577/IJERTV9IS050290
- [11] Dyson MP, Hartling L, Shulhan J, Chisholm A, Milne A, Sundar P, et al. A Systematic Review of Social Media Use to Discuss and View Deliberate Self-Harm Acts. PLoS ONE 11(5): e0155813. doi:10.1371/journal.pone.0155813 Editor: Soraya Seedat, University of Stellenbosch, SOUTH AFRICA Received: December 1, 2016. Clerk Maxwell, A Treatise on Electricity and Magnetism, 3rd ed., vol. 2. Oxford: Clarendon, 1892, pp.68–73. (2016)
- [12] American Psychiatric Association. (2013). Anxiety Disorders. In Diagnostic and statistical manual of mental disorders (5th ed.). <https://doi.org/10.1176/appi.books.9780890425596.dsm05>
- [13] "Depression", Who.int, [Online]. Available: <https://www.who.int/news-room/fact-sheets/detail/depression>. [Accessed: 25- Aug- 2020]. (2020)
- [14] Judy Hanwen Shen and Frank Rudzicz. 2017. Detecting anxiety through Reddit. In Proceedings of the Fourth Workshop on Computational Linguistics and Clinical Psychology – From Linguistic Signal to Clinical Reality. Association for Computational Linguistics, pages 58–65. doi:http://aclweb.org/anthology/W17-3107 Gamon, Michael & Choudhury, Munmun & Counts, Scott & Horvitz, Eric. Predicting Depression via Social Media. Association for the Advancement of Artificial Intelligence (2013).