

MERGED COLLABORATIVE METHODS FOR AUTOMATED DETECTION OF DEPRESSION

¹ Siddhanth Singh Bhushan, ² Dr. K. Santhi Sree

¹ Student, Department of Information Technology, University College of Engineering, Science and Technology,
JNTUH Hyderabad

² Professor, Department of Information Technology, University College of Engineering, Science and Technology,
JNTUH Hyderabad

Abstract: Depression rates have surged over the past century, yet many cases remain undetected despite advancements in diagnosis. Automated detection methods offer potential solutions to this issue by identifying individuals at risk. Effective feature representation and language analysis are crucial for understanding depression detection. This project aims to enhance depression detection through text classifiers. Specifically, it compares two sets of methods—hybrid and ensemble—to improve classification performance. The objective is to identify the most effective approach for accurately detecting depression indicators in textual data. Text classifiers are employed and trained for depression detection. Two sets of methods, hybrid and ensemble, are examined and compared. The methodology involves effective feature representation and analysis of language use to enhance classification accuracy. Various feature combinations and selection techniques are explored to optimize performance. The results demonstrate that ensemble models outperform hybrid models in classifying depression indicators. The strength of combined features underscores the importance of multiple feature combinations and proper selection. This finding highlights the potential for improved depression detection through sophisticated ensemble methods. This project advances depression detection methodologies, emphasizing the significance of ensemble approaches in achieving better performance. The implications of this study extend to improved mental health screening and intervention strategies, offering potential benefits for individuals at risk of depression. And also added an ensemble method is employed which combines the predictions of multiple models, and the models are LSTM+GRU, Voting Classifier(RF+AdaBoost), and Stacking Classifier(RF+MLP+LGBM+XGBoost), to enhance prediction accuracy. Among these models, the Voting Classifier achieves 100% accuracy. Additionally, CNN, BERT, and XLNet pretrained models are utilized for further analysis.

Index terms - Deep neural networks, depression detection, ensemble methods, sentiment lexicon.

1. INTRODUCTION

Drastic changes in the human lifestyle in modern society have led to an increase in the number of people suffering from depression. Depression is known to be “a disease of modernity” [1], and it has been predicted that by 2030, one of the three causes of illness will be depression [2]. The social stigma surrounding depression and the high rate of misdiagnosis has led to a lack of access to proper diagnosis and care [3]. Serve mental disorders but without

effective intervention could turn to suicidal ideation [4]. Therefore, the timely detection of depression can be highly beneficial for individuals and society.

Depression symptoms can be reflected in various human activities and behaviors and different degrees [5]. One of the sources, which can help identify depression symptoms in individuals, is the use of language [6]. Cognitive and linguistic studies have numerously shown that people with depression use language features differently [6]. For example, they tend to use more first-person singular pronouns (I, me, or we) and more negatively valenced words [7].

Online social content is one source for automatic mental disorder detection as it is one of the platforms through which users communicate. In recent years, social networking platforms have been widely applied to study users' behavior and have inspired various researchers to introduce new forms of health care solutions [8], [9]. Furthermore, the stigma surrounding depression can make individuals less willing to seek professional assistance, and they turn to less traditional sources such as social media. Social media can be an essential source of information about individuals' opinions and feelings in the study of depression [10]. More specifically, research has addressed depression detection at various levels of granularity and approached it from different standpoints. Several social network sites (SNSs), such as Reddit, Twitter, Facebook, and Weibo, have been utilized for research about depression and other mental state disorders, such as postpartum depression [11] and posttraumatic stress disorder (PTSD) [12].

Earlier approaches to depression detection have primarily taken a bottom-up approach to learn and apply deep learning (DL) and machine learning (ML) methods. While such subsymbolic artificial intelligence (AI) methods can provide valuable insights about word frequencies and statistical correlations, they are not sufficient to analyze narrative and understanding of dialog systems in sentiment analysis [13]. Although there has been advancement with natural language processing (NLP) methods using DL methods, the predictive power of such approaches is limited mainly because DL methods learn better from large sets of data. Besides, communication entails a broader range of contributors, including understanding the world, social norms, and cultural awareness.

To address these challenges, recently, research in depression detection has taken top-down approaches to learn by applying symbolic AI methods such as logical reasoning. In particular, the hybrid combination of subsymbolic approaches with symbolic methods has been shown to induce more meaningful patterns in natural language texts [13]. Hence, it is vital to integrate symbolic approaches to learning with subsymbolic approaches in tackling the task of automated depression detection. In addition to hybrid methods, another set of approaches that yields high accuracy are ensemble methods in which several learning methods are combined [14]. Ensemble methods have frequently achieved high performance in solving various predictive problem areas [14].

2. LITERATURE SURVEY

Depression detection in social media is a multidisciplinary area where psychological and psychoanalytical findings can help machine learning and natural language processing techniques to detect symptoms of depression in the users of social media. In this research [6], using an inventory that has made systematic observations and records of the

characteristic attitudes and symptoms of depressed patients, we develop a bipolar feature vector that contains features from both depressed and non-depressed classes. The inventory we use for feature extraction is composed of 21 categories of symptoms and attitudes, which are primarily clinically derived in the course of the psychoanalytic psychotherapy of depressed patients, and systematic observations and records of their characteristic attitudes and symptoms. Also, getting insight from a cognitive idea, we develop a classifier based on multinomial Naïve Bayes training algorithm with some modification. The model we develop in this research is successful in classifying the users of social media into depressed and non-depressed groups, achieving the F1 score 82.75%.

Depression is viewed as the largest contributor to global disability and a major reason for suicide. It has an impact on the language usage reflected in the written text. The key objective of our study [8] is to examine Reddit users' posts to detect any factors that may reveal the depression attitudes of relevant online users. For such purpose, we employ the Natural Language Processing (NLP) [21, 30, 31] techniques and machine learning approaches to train the data and evaluate the efficiency of our proposed method. We identify a lexicon of terms that are more common among depressed accounts. The results show that our proposed method can significantly improve performance accuracy. The best single feature is bigram with the Support Vector Machine (SVM) classifier to detect depression with 80% accuracy and 0.80 F1 scores. The strength and effectiveness of the combined features (LIWC+LDA+bigram) are most successfully demonstrated with the Multilayer Perceptron (MLP) classifier resulting in the top performance for depression detection reaching 91% accuracy and 0.93 F1 scores. According to our study, better performance improvement can be achieved by proper feature selections and their multiple feature combinations.

Mental health is a critical issue in modern society, and mental disorders could sometimes turn to suicidal ideation without effective treatment. Early detection of mental disorders and suicidal ideation from social content provides a potential way for effective social intervention. However, classifying suicidal ideation and other mental disorders is challenging as they share similar patterns in language usage and sentimental polarity. This paper [9] enhances text representation with lexicon-based sentiment scores and latent topics and proposes using relation networks to detect suicidal ideation and mental disorders with related risk indicators. The relation module is further equipped with the attention mechanism to prioritize more critical relational features. Through experiments on three real-world datasets, our model outperforms most of its counterparts.

The birth of a child is a major milestone in the life of parents. [11] We leverage Facebook data shared voluntarily by 165 new mothers as streams of evidence for characterizing their postnatal experiences. We consider multiple measures including activity, social capital, emotion, and linguistic style in participants' Facebook data in pre- and postnatal periods. Our study includes detecting and predicting onset of post-partum depression (PPD). The work complements recent work on detecting and predicting significant postpartum changes in behavior, language, and affect from Twitter data. In contrast to prior studies, we gain access to ground truth on postpartum experiences via self-reports and a common psychometric instrument used to evaluate PPD [40]. We develop a series of statistical models to predict, from data available before childbirth, a mother's likelihood of PPD. We corroborate our

quantitative findings through interviews with mothers experiencing PPD. We find that increased social isolation and lowered availability of social capital on Facebook, are the best predictors of PPD in mothers.

We developed computational models to predict the emergence of depression and Post-Traumatic Stress Disorder in Twitter users [12]. Twitter data and details of depression history were collected from 204 individuals (105 depressed, 99 healthy). We extracted predictive features measuring affect, linguistic style, and context from participant tweets ($N=279,951$) and built models using these features with supervised learning algorithms. Resulting models successfully discriminated between depressed and healthy content, and compared favorably to general practitioners' average success rates in diagnosing depression, albeit in a separate population. Results held even when the analysis was restricted to content posted before first depression diagnosis. State-space temporal analysis suggests that onset of depression may be detectable from Twitter data several months prior to diagnosis. Predictive results were replicated with a separate sample of individuals diagnosed with [40] PTSD ($N_{users} = 174$, $N_{tweets} = 243,775$). A state-space time series model revealed indicators of PTSD almost immediately post-trauma, often many months prior to clinical diagnosis. These methods suggest a data-driven, predictive approach for early screening and detection of mental illness.

3. METHODOLOGY

i) Proposed Work:

The proposed system aims to enhance automated depression detection by leveraging a combination of hybrid and ensemble methods. It seeks to address the limitations of existing approaches by integrating symbolic AI techniques with subsymbolic methods for more meaningful analysis of natural language texts. The project entails conducting experiments across multiple datasets which are Reddit data, eRisk, CLpsych [8, 40, 41, 42, 43] to evaluate the performance of hybrid methods utilizing sentiment lexicons and logistic regression, as well as ensemble methods combining DL approaches and lexicon-based models. And also included, an ensemble method is employed which combines the predictions of multiple models, and the models are LSTM+GRU, Voting Classifier(RF+AdaBoost), and Stacking Classifier(RF+MLP+LGBM+XGBoost), to enhance prediction accuracy. Among these models, the Voting Classifier achieves 100% accuracy. Additionally, CNN, BERT, and XLNet pretrained models are utilized for further analysis. In this project, a user-friendly front end is developed using the Flask framework, incorporating authentication features for user testing.

ii) System Architecture:

NLP technologies are widely applied to solve various problems across various domains, such as text summarization, language translation, and sentiment analysis. Earlier attempts to solve these problems relied on rule-based methods and probabilistic methods, such as the hidden Markov model [29], which required much data engineering. More recently, NLP methods have relied much more on DL [30], [31]. With the rise of powerful computing systems, it has become possible to train end-to-end systems as ML helps address various problems from fields, such as image recognition [32], speech recognition [33], and NLP [34]. Early approaches to text classification relied on representing documents through a bag-of-words representation and applying ML methods in which the words were

not processed sequentially [31]. In recent years, text classification has been conducted by considering the sequential nature of data using LSTM [35] neural networks, as these models can require fewer training samples due to their reliance on word embedding. As methods, such as RNNs, LSTM, and attention-based models, have transformed speech and NLP [44], this study draws from these models to analyze and classify texts by detecting the fragments containing sentiments. The flowchart of the proposed ensemble model is shown in Fig. 1.

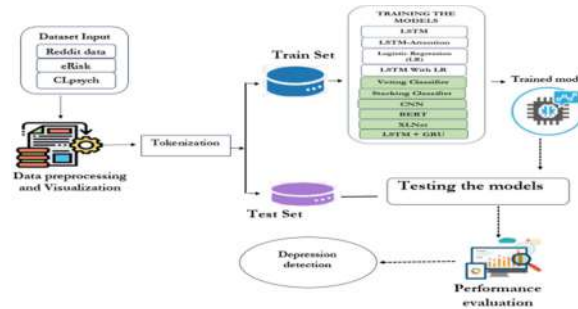


Fig 1 Proposed architecture

iii) Dataset collection:

1) CLPsych 2015 Shared Task:

The Computational Linguistics and Clinical Psychology (CLPsych) was initiated in 2014 to promote collaboration between psychologists and computer scientists [40]. Specifically, “shared tasks” were defined to study and compare different methods on the same prediction problems. This dataset contains user-generated posts from users with depression or (Depression and PTSD on Twitter) PTSD [40] on Twitter.3 Specifically, there are three binary classification subtasks, i.e., 1) depression versus control; 2) PTSD versus control; and 3) depression versus PTSD. The train partition consisted of 327 depression users, 246 PTSD users, and for each an age- and gender-matched control user, for a total of 1146 users. The test data contained 150 depression users, 150 PTSD users, and age- and gender-matched control for each, for a total of 600 users. However, the actual number of users in the training and testing sets is 1711 due to an unknown data missing issue.

Unnamed: 0	post_id	post_created	post_text	user_id	followers	friends	favourites	statuses	retweets	label
0	637894877624413696	Sun-Aug-30 07:43:37 +0000 2015	It's just over 2 years since I was diagnosed w.	1013187249	84	211	251	837	0	1
1	637890384576770240	Sun-Aug-30 07:31:33 +0000 2015	It's Sunday I need a break, so I'm planning L...	1013187249	84	211	251	837	1	1
2	637748345868051960	Sat-Aug-29 22:11:07 +0000 2015	Awake but tired. I need to sleep but my brain.	1013187249	84	211	251	837	0	1
3	63769842107123073	Sat-Aug-29 18:40:49 +0000 2015	RT @SeveH2: @Retro bears make perfect gifts and.	1013187249	84	211	251	837	2	1
4	637698327485386272	Sat-Aug-29 18:40:26 +0000 2015	It's hard to say whether packing lists are mak.	1013187249	84	211	251	837	1	1

Fig 2 CLPsych Dataset

2) Reddit:

Reddit social media collection contains posts from depressed and nondepressed users. The dataset contained 1841 users (1200 positives and 641 negatives) [41]. Reddit as a social media platform allows for the anonymity of the users, and it is widely used for discussion about stigmatic topics [8]. Reddit data have been used to study the posts specifically from Reddit users who wrote about mental health issues and who had proceeded to post topics about suicidal ideation [42]. The data were concatenated, randomly shuffled, and split into train and test sets with an 80:20

split rate. The final data frame consisted of one column of text comments and another column of labels for the corresponding comments. Each comment was labeled with 1 or 0 for depression or nondepression, respectively.

	clean_text	is_depression
0	we understand that most people who reply immed...	1
1	welcome to r depression s check in post a plac...	1
2	anyone else instead of sleeping more when depr...	1
3	i ve kind of stuffed around a lot in my life d...	1
4	sleep is my greatest and most comforting escap...	1

Fig 3 Reddit dataset

3) eRisk Dataset:

This dataset is collected from the eRisk (Early Risk Prediction) forum [43]. The eRisk is a public competition platform that facilitates multidisciplinary research and creates reusable datasets and benchmarks for assessing early risk detection technologies in health and safety problem areas. The eRisk 2018 dataset was initially developed to detect early signs of depression. The eRisk collection contains posts from depressed and nondepressed 4498 users, where 3728 users belong to the nondepressed and 770 belong to the depressed class. The data were concatenated, randomly shuffled, and split into train and test sets with an 80:20 split rate.

Unnamed: 0	text	class
0	2 Ex Wife Threatening SuicideRecently I left my ...	suicide
1	3 Am I weird I don't get affected by compliments...	non-suicide
2	4 Finally 2020 is almost over... So I can never ...	non-suicide
3	8 i need helpjust help me im crying so hard	suicide
4	9 I'm so lostHello, my name is Adam (16) and I'v...	suicide

Fig 4 eRisk dataset

iv) Data Processing:

Data processing involves transforming raw data into valuable information for businesses. Generally, data scientists process data, which includes collecting, organizing, cleaning, verifying, analyzing, and converting it into readable formats such as graphs or documents. Data processing can be done using three methods i.e., manual, mechanical, and electronic. The aim is to increase the value of information and facilitate decision-making. This enables businesses to improve their operations and make timely strategic decisions. Automated data processing solutions, such as computer software programming, play a significant role in this. It can help turn large amounts of data, including big data, into meaningful insights for quality management and decision-making.

v) Feature selection:

Feature selection is the process of isolating the most consistent, non-redundant, and relevant features to use in model construction. Methodically reducing the size of datasets is important as the size and variety of datasets continue to grow. The main goal of feature selection is to improve the performance of a predictive model and reduce the computational cost of modeling.

Feature selection, one of the main components of feature engineering, is the process of selecting the most important features to input in machine learning algorithms. Feature selection techniques are employed to reduce the number of input variables by eliminating redundant or irrelevant features and narrowing down the set of features to those most relevant to the machine learning model. The main benefits of performing feature selection in advance, rather than letting the machine learning model figure out which features are most important.

vi) Algorithms:

LSTM (Long Short-Term Memory): Long Short-Term Memory (LSTM) is a type of recurrent neural network (RNN) architecture designed to address the vanishing gradient problem in traditional RNNs. LSTMs are equipped with memory cells and a system of gates that control the flow of information, allowing them to capture and learn long-term dependencies in sequential data. The architecture includes input, forget, and output gates, enabling the model to selectively remember or forget information over extended sequences. LSTMs are well-suited for tasks involving sequential data, such as natural language processing. In the project, where the goal is to detect depression based on language patterns, LSTMs can effectively capture the intricate dependencies present in textual data over extended contexts.

LSTM – Attention: LSTM with Attention augments the standard LSTM architecture by incorporating an attention mechanism. Attention mechanisms allow the model to focus on specific parts of the input sequence when making predictions. This enhanced capability enables the model to pay varying degrees of attention to different parts of the input, providing a more nuanced understanding of the data. Attention mechanisms are particularly valuable in tasks where certain elements in a sequence carry more importance. In the context of depression detection, attention can help identify critical words or phrases that may indicate mental health concerns, allowing for more precise and context-aware predictions.

Logistic Regression (LR): Logistic Regression is a linear model used for binary classification tasks. It models the probability that an instance belongs to a particular class using the logistic function. Despite its simplicity, Logistic Regression is widely used for its interpretability and efficiency. Logistic Regression serves as a baseline model for binary classification, providing a straightforward and interpretable approach. It allows for an initial understanding of the classification problem and can serve as a benchmark for evaluating the performance of more complex models.

LSTM with LR (Logistic Regression): LSTM with Logistic Regression is a hybrid model that combines the strengths of LSTM for sequence learning with the simplicity and interpretability of Logistic Regression. The LSTM extracts features from sequential data, and these features are then used as input for a logistic regression model for the final classification. This hybrid approach aims to leverage the sequence learning capabilities of LSTM while benefiting from the interpretability of logistic regression. By combining these two models, the architecture seeks to achieve a balance between complexity and transparency, making the model more interpretable while maintaining high predictive performance [14].

Certainly, let's rephrase the statements to align with the context of the project:

A hybrid model aiming to enhance the robustness and accuracy of the final depression detection predictions. Also we can build the front end using flask framework for user testing with authentication. This frontend allows users to

input text related to their mental health, and the system processes this input through the hybrid model, presenting the results for user interaction and feedback.

4. EXPERIMENTAL RESULTS

Precision: Precision evaluates the fraction of correctly classified instances or samples among the ones classified as positives. Thus, the formula to calculate the precision is given by:

Precision = True positives/ (True positives + False positives) = TP/(TP + FP)

$$\text{Precision} = \frac{\text{True Positive}}{\text{True Positive} + \text{False Positive}}$$

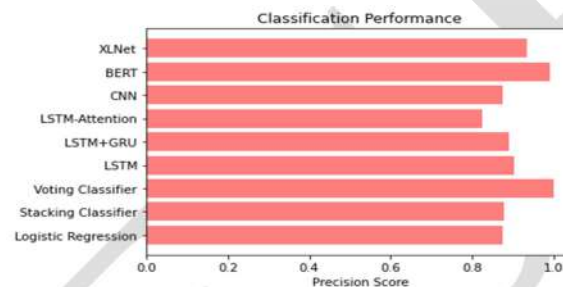


Fig 5 Precision comparison graph

Recall: Recall is a metric in machine learning that measures the ability of a model to identify all relevant instances of a particular class. It is the ratio of correctly predicted positive observations to the total actual positives, providing insights into a model's completeness in capturing instances of a given class.

$$\text{Recall} = \frac{TP}{TP + FN}$$

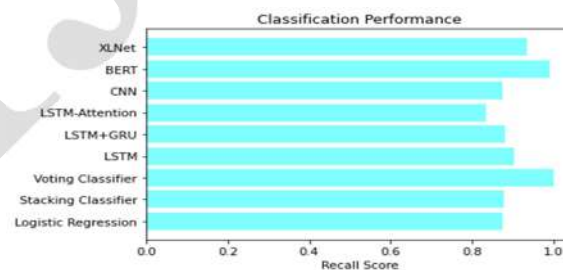


Fig 6 Recall comparison graph

Accuracy: Accuracy is the proportion of correct predictions in a classification task, measuring the overall correctness of a model's predictions.

$$Accuracy = \frac{TP + TN}{TP + FP + TN + FN}$$

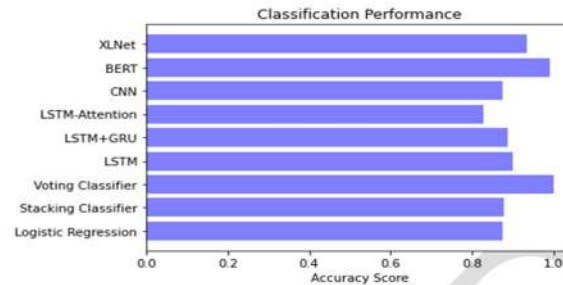


Fig 7 Accuracy graph

F1 Score: The F1 Score is the harmonic mean of precision and recall, offering a balanced measure that considers both false positives and false negatives, making it suitable for imbalanced datasets.

$$F1\ Score = 2 * \frac{Recall \times Precision}{Recall + Precision} * 100$$

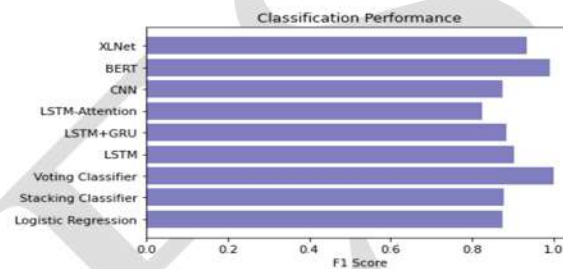


Fig 8 F1Score

MLModel	Accuracy	Precision	Recall	F1_score
Logistic Regression	0.845	0.853	0.845	0.844
Extension Stacking Classifier	0.805	0.805	0.805	0.805
Extension Voting Classifier	1.000	1.000	1.000	1.000
LSTM	0.860	0.860	0.860	0.860
Extension LSTM+GRU	0.830	0.758	0.758	0.758
LSTM-Attention	0.747	0.696	0.898	0.780
Extension CNN	0.735	0.763	0.735	0.739
Extension BERT	0.993	0.993	0.993	0.993
Extension XLNet	0.872	0.871	0.871	0.871

Fig 9 Performance Evaluation

5. CONCLUSION

The project successfully explored and integrated various Natural Language Processing (NLP) methods [20, 25, 30], including lexicon-based comparisons and machine learning algorithms, demonstrating their collective efficacy in identifying language patterns associated with depression. Through systematic experimentation, an optimal combination of NLP models and features was identified, showcasing the significance of leveraging diverse approaches to enhance the accuracy of depression detection. The evaluation of various classification methods, employing metrics such as precision, recall, F1 score, and accuracy, provided a comprehensive understanding of the

strengths and weaknesses of each model, contributing valuable insights for future research and applications. The Voting Classifier algorithm excels in depression detection, delivering consistent accuracy and enhancing the system's effectiveness. Integrated into the front end, users can input features and observe real-time predictions, showcasing its reliability and usability. The project's outcomes contribute to the evolving field of automated depression detection, providing a foundation for further research and emphasizing the importance of combining diverse NLP methods for more accurate and comprehensive assessments of mental health through language analysis.

6. FUTURE SCOPE

Although this study shows that the applied feature set improves the classification performance, the absolute value of the evaluation metrics indicates that this task can be further explored and improved. [13, 25, 30, 31] DL architectures were applied in this study. Experiments could be extended with other models for text classification, such as CNNs and transformer-based pretrained language models. Future work can extend possibilities for improvement, which lies in utilizing features such as POS tags and other methods of handling imbalanced datasets.

REFERENCES

- [1] B. H. Hidaka, "Depression as a disease of modernity: Explanations for increasing prevalence," *J. Affect. Disorders*, vol. 140, no. 3, pp. 205–214, Nov. 2012.
- [2] C. D. Mathers and D. Loncar, "Projections of global mortality and burden of disease from 2002 to 2030," *PLoS Med.*, vol. 3, no. 11, p. e442, Nov. 2006.
- [3] S. Rodrigues et al., "Impact of stigma on veteran treatment seeking for depression," *Amer. J. Psychiatric Rehabil.*, vol. 17, no. 2, pp. 128–146, Apr. 2014.
- [4] S. Ji, C. P. Yu, S.-F. Fung, S. Pan, and G. Long, "Supervised learning for suicidal ideation detection in online user content," *Complexity*, vol. 2018, pp. 1–10, Sep. 2018.
- [5] A. T. Beck, C. H. Ward, M. Mendelson, J. Mock, and J. Erbaugh, "An inventory for measuring depression," *Arch. Gen. Psychiatry*, vol. 4, no. 6, pp. 561–571, 1961.
- [6] S. H. Hosseini-Saravani, S. Besharati, H. Calvo, and A. Gelbukh, "Depression detection in social media using a psychoanalytical technique for feature extraction and a cognitive based classifier," in *Proc. Mex. Int. Conf. Artif. Intell. Springer*, 2020, pp. 282–292. [Online]. Available: https://link.springer.com/chapter/10.1007/978-3-030-60887-3_25
- [7] S. Rude, E.-M. Gortner, and J. Pennebaker, "Language use of depressed and depression-vulnerable college students," *Cogn. Emotion*, vol. 18, no. 8, pp. 1121–1133, Dec. 2004.
- [8] M. M. Tadesse, H. Lin, B. Xu, and L. Yang, "Detection of depressionrelated posts in reddit social media forum," *IEEE Access*, vol. 7, pp. 44883–44893, 2019.
- [9] S. Ji, X. Li, Z. Huang, and E. Cambria, "Suicidal ideation and mental disorder detection with attentive relation networks," *Neural Comput. Appl.*, pp. 1–11, Jun. 2021. [Online]. Available: <https://link.springer.com/article/10.1007/s00521-021-06208-y>

- [10] M. Paul and M. Dredze, "You are what you tweet: Analyzing Twitter for public health," in Proc. Int. AAAI Conf. Web Social Media, 2011, vol. 5, no. 1, pp. 265–272.
- [11] M. De Choudhury, S. Counts, E. J. Horvitz, and A. Hoff, "Characterizing and predicting postpartum depression from shared Facebook data," in Proc. 17th ACM Conf. Comput. Supported Cooperat. Work Social Comput., Feb. 2014, pp. 626–638.
- [12] A. G. Reece, A. J. Reagan, K. L. M. Lix, P. S. Dodds, C. M. Danforth, and E. J. Langer, "Forecasting the onset and course of mental illness with Twitter data," Sci. Rep., vol. 7, no. 1, pp. 1–11, Dec. 2017.
- [13] E. Cambria, Y. Li, F. Z. Xing, S. Poria, and K. Kwok, "SenticNet 6: Ensemble application of symbolic and subsymbolic AI for sentiment analysis," in Proc. 29th ACM Int. Conf. Inf. Knowl. Manage., Oct. 2020, pp. 105–114.
- [14] L. I. Kuncheva, Combining Pattern Classifiers: Methods and Algorithms. Hoboken, NJ, USA: Wiley, 2014.
- [15] A. Benton, M. Mitchell, and D. Hovy, "Multi-task learning for mental health using social media text," 2017, arXiv:1712.03538.
- [16] G. Coppersmith, M. Dredze, C. Harman, and K. Hollingshead, "From ADHD to SAD: Analyzing the language of mental health on Twitter through self-reported diagnoses," in Proc. 2nd Workshop Comput. Linguistics Clin. Psychol., From Linguistic Signal Clin. Reality, 2015, pp. 1–10.
- [17] S. Paul, S. K. Jandhyala, and T. Basu, "Early detection of signs of anorexia and depression over social media using effective machine learning frameworks," in Proc. CLEF, Working Notes, 2018, pp. 1–15.
- [18] M. Nadeem, "Identifying depression on Twitter," 2016, arXiv:1607.07384.
- [19] Y. Tyshchenko, "Depression and anxiety detection from blog posts data," Nature Precis. Sci., Inst. Comput. Sci., Univ. Tartu, Tartu, Estonia, Tech. Rep. 53001016, 2018.
- [20] J. Wolohan, M. Hiraga, A. Mukherjee, Z. A. Sayyed, and M. Millard, "Detecting linguistic traces of depression in topic-restricted text: Attending to self-stigmatized depression with NLP," in Proc. 1st Int. Workshop Lang. Cogn. Comput. Models, 2018, pp. 11–21.
- [21] R. A. Calvo, D. N. Milne, M. S. Hussain, and H. Christensen, "Natural language processing in mental health applications using non-clinical texts," Natural Lang. Eng., vol. 23, no. 5, pp. 649–685, 2017.
- [22] A. H. Orabi, P. Buddhitha, M. H. Orabi, and D. Inkpen, "Deep learning for depression detection of Twitter users," in Proc. 5th Workshop Comput. Linguistics Clin. Psychol., From Keyboard Clinic, 2018, pp. 88–97.
- [23] T. Nguyen, D. Phung, B. Dao, S. Venkatesh, and M. Berk, "Affective and content analysis of online depression communities," IEEE Trans. Affect. Comput., vol. 5, no. 3, pp. 217–226, Jul./Sep. 2014.
- [24] D. Maupomé and M.-J. Meurs, "Using topic extraction on social media content for the early detection of depression," in Proc. CLEF, Working Notes, vol. 2125, 2018, pp. 1–5.
- [25] J. Misra, "AutoNLP: NLP feature recommendations for text analytics applications," 2020, arXiv:2002.03056.
- [26] M. Stankevich, V. Isakov, D. Devyatkin, and I. Smirnov, "Feature engineering for depression detection in social media," in Proc. ICPRAM, 2018, pp. 426–431.

- [27] T. Shen et al., “Cross-domain depression detection via harvesting social media,” in Proc. Int. Joint Conf. Artif. Intell., Jul. 2018, pp. 1611–1617.
- [28] S. Tsugawa, Y. Kikuchi, F. Kishino, K. Nakajima, Y. Itoh, and H. Ohsaki, “Recognizing depression from Twitter activity,” in Proc. 33rd Annu. ACM Conf. Hum. Factors Comput. Syst., Apr. 2015, pp. 3187–3196.
- [29] L. Rabiner and B. Juang, “An introduction to hidden Markov models,” IEEE ASSP Mag., vol. 3, no. 1, pp. 4–16, Jan. 1986.
- [30] E. Cambria and B. White, “Jumping NLP curves: A review of natural language processing research,” IEEE Comput. Intell. Mag., vol. 9, no. 2, pp. 48–57, May 2014.
- [31] S. Sun, C. Luo, and J. Chen, “A review of natural language processing techniques for opinion mining systems,” Inf. Fusion, vol. 36, pp. 10–25, Jul. 2017.
- [32] K. He, X. Zhang, S. Ren, and J. Sun, “Deep residual learning for image recognition,” in Proc. IEEE Conf. Comput. Vis. Pattern Recognit. (CVPR), Jun. 2016, pp. 770–778.
- [33] D. Bahdanau, J. Chorowski, D. Serdyuk, P. Brakel, and Y. Bengio, “End-to-end attention-based large vocabulary speech recognition,” in Proc. IEEE Int. Conf. Acoust., Speech Signal Process. (ICASSP), Mar. 2016, pp. 4945–4949.
- [34] T. Mikolov, I. Sutskever, K. Chen, G. Corrado, and J. Dean, “Distributed representations of words and phrases and their compositionality,” 2013, arXiv:1310.4546.
- [35] S. Hochreiter and J. Schmidhuber, “Long short-term memory,” Neural Comput., vol. 9, no. 8, pp. 1735–1780, 1997.
- [36] D. Bahdanau, K. Cho, and Y. Bengio, “Neural machine translation by jointly learning to align and translate,” 2014, arXiv:1409.0473. [37] A. Vaswani et al., “Attention is all you need,” in Proc. NIPS, 2017, pp. 1–11.
- [38] T. Shon and J. Moon, “A hybrid machine learning approach to network anomaly detection,” Inf. Sci., vol. 177, no. 18, pp. 3799–3821, Sep. 2007.
- [39] Z.-H. Zhou, Ensemble Methods: Foundations and Algorithms. Boca Raton, FL, USA: CRC Press, 2012.
- [40] G. Coppersmith, M. Dredze, C. Harman, K. Hollingshead, and M. Mitchell, “CLPsych 2015 shared task: Depression and PTSD on Twitter,” in Proc. 2nd Workshop Comput. Linguistics Clin. Psychol., From Linguistic Signal Clin. Reality, 2015, pp. 31–39.
- [41] I. Pirina and Ç. Çöltekin, “Identifying depression on reddit: The effect of training data,” in Proc. EMNLP Workshop SMM4H, 3rd Social Media Mining Health Appl. Workshop Shared Task, 2018, pp. 9–12.
- [42] S. Ji, S. Pan, X. Li, E. Cambria, G. Long, and Z. Huang, “Suicidal ideation detection: A review of machine learning methods and applications,” IEEE Trans. Computat. Social Syst., vol. 8, no. 1, pp. 214–226, Feb. 2021.
- [43] D. E. Losada and F. Crestani, “A test collection for research on depression and language use,” in Proc. Int. Conf. Cross-Lang. Eval. Forum Eur. Lang. Springer, 2016, pp. 28–39. [Online]. Available: https://link.springer.com/chapter/10.1007/978-3-319-44564-9_3

- [44] J. Zimmermann, T. Brockmeyer, M. Hunn, H. Schauenburg, and M. Wolf, "First-person pronoun use in spoken language as a predictor of future depressive symptoms: Preliminary evidence from a clinical sample of depressed patients," *Clin. Psychol. Psychotherapy*, vol. 24, no. 2, pp. 384–391, Mar. 2017.
- [45] F. Å. Nielsen, "A new ANEW: Evaluation of a word list for sentiment analysis in microblogs," 2011, arXiv:1103.2903.
- [46] S. Mohammad and P. Turney, "Emotions evoked by common words and phrases: Using mechanical Turk to create an emotion lexicon," in *Proc. NAACL HLT Workshop Comput. Approaches Anal. Gener. Emotion Text*, 2010, pp. 26–34.
- [47] T. Wilson, J. Wiebe, and P. Hoffmann, "Recognizing contextual polarity in phrase-level sentiment analysis," in *Proc. Conf. Hum. Lang. Technol. Empirical Methods Natural Lang. Process.*, 2005, pp. 347–354.
- [48] B. Juba and H. S. Le, "Precision-recall versus accuracy and the role of large data sets," in *Proc. AAAI Conf. Artif. Intell.*, 2019, vol. 33, no. 1, pp. 4039–4048.
- [49] T. Basu and C. A. Murthy, "A feature selection method for improved document classification," in *Proc. Int. Conf. Adv. Data Mining Appl.* Springer, 2012, pp. 296–305. [Online]. Available: https://link.springer.com/chapter/10.1007/978-3-642-35527-1_25