

Sentiment Analysis Using Transformer Based Model

Umesh Narayan¹, and Devendra Kumar²

¹ Computer Science and Engineering, B.N. College of Engineering & Technology (BNCET), Lucknow, India

² Assistant Professor, Computer Science and Engineering, B.N. College of Engineering & Technology (BNCET), Lucknow, India

Correspondence should be addressed to Devendra Kumar: devendrabncet2023@gmail.com

Copyright © 2024 Made Umesh Narayan et al. . This is an open-access article distributed under the Creative Commons Attribution License, which permits unrestricted use, distribution, and reproduction in any medium, provided the original work is properly cited

ABSTRACT- This paper delves into the realm of sentiment analysis, focusing on the state-of-the-art method of employing transformers for classification. Specifically, the study explores the efficient classification techniques utilizing transformers, with a primary focus on fine-tuning tasks. We have used Distil BERT for fine-tuned using a fine-grained emotion dataset and their performances are evaluated in terms of F1-score and processing time. Using data from the Hugging Face dataset, the study evaluates and benchmarks the performance of various transformer models. We have obtained a test loss of 0.2205, accuracy of 92.25%, and with F1-score of 92.26%.

KEYWORDS- Sentiment Analysis, Emotions, Distil BERT, ANN

I. INTRODUCTION

The Sentiment Analysis (SA) is a crucial aspect of Natural Language Processing (NLP) that focuses on examining individuals' sentiments, opinions, and emotions. This multifaceted process encompasses various stages, including data retrievals, extractions, preprocessing, and feature extraction. The exponential growth of social media comments in diverse industries has underscored the need to effectively navigate vast internet data and extract valuable insights automatically. SA models assume a crucial role in meeting this demand, influencing domains like politics, marketing, governmental policies, disaster management, and public health—all of which depend significantly on emotion detection [1].

The exploration of sentiment analysis began with Pang et al.'s [2] investigation into machine learning techniques using supervised classifiers like Naive Bayes(NB), Maximum Entropy, and Support Vector Machine (SVM) on a movie review dataset. However, these classifiers exhibited suboptimal performance, attributed to their reliance on traditional text categorization methodologies, particularly the bag-of-words (BOW) notion, which lacks grammatical structures and semantic relations crucial for sentiment analysis [3]. Past studies leaned on ml methods such as Bayesian Networks, NB SVM, DT, and Artificial Neural Networks [4-5], but their scalability was hindered by the requirement for labeled data, often costly or prohibitively expensive [6].

Addressing the constraints of conventional machine learning, deep learning (DL) surfaced as a potent substitute, providing benefits such as automated feature generation and scalability [7]. Investigators delved into various DL architectures, encompassing Convolutional Neural Networks (CNN), Long Short-Term Memory (LSTM), Gated Alternate Neural Network (GANN), and hybrid methodologies amalgamating multiple strategies [8-9]. The integration of ML and DL models has showcased enhanced accuracy, alleviating the shortcomings inherent in both approaches [10]. Nimmi et al. [11] introduced the AVELD Model, an ensemble deep learning approach incorporating transformer-based models like BERT, DistilBERT, and RoBERTa for emergency response call analysis. Adoma [12] fine-tuned pre-trained BERT, RoBERTa, DistilBERT, and XLNet to identify emotions in text, achieving notable accuracy. Uddagiri [13] integrated RoBERTa with Aspect-Based Sentiment Analysis (ABSA) and the LSTM, attaining a high accuracy of 94.7% on the Twitter Dataset for the Ukraine Conflict.

Researchers also explored attribute-based hybrid techniques [14], common-sense knowledge integration with LSTM [15], and feature extraction using SVM and PCA [16], each achieving commendable accuracies in sentiment analysis. Furthermore, the combination of Bangla-BERT with LSTM demonstrated 94.15% accuracy [17]. The BERT-DCNN model, stacking BERT with dilated CNN, achieved an 87.1% accuracy for Twitter airline sentiment data [18].

In the realm of airline sentiment analysis, Kian [19] utilized RoBERTa-LSTM, achieving 89.85% accuracy without data augmentation and 91.37% with augmentation. Another study proposed a CNN-LSTM architecture, achieving 91.3% accuracy on airline sentiment [20]. Barakat [21] introduced the ULMFit-SVM model with remarkable accuracy rates on diverse datasets, although it focused solely on document-level sentiment analysis, neglecting aspect-level sentiments.

Recognizing the unique challenges of sentiment analysis in tweets, a proposed hybrid model incorporating BERT, BiLSTM, and BiGRU aims to enhance accuracy by leveraging dynamic word vectors and extracting sequence characteristics. The study also investigates the impact of varying the number and location of (BiLSTM and BiGRU) layers to optimize performance. This comprehensive

overview underscores the diverse approaches and models employed in sentiment analysis, reflecting the evolving landscape of NLP and its applications across industries [22]. Despite advancements in contemporary living standards, the shifting dynamics of our living environments, workplaces, and interpersonal connections pose fresh challenges, contributing to mental exhaustion and the potential onset of depressive symptoms [23]. In response, our aim is to create a methodology enabling the early identification of depressive indicators based on individuals' conversational patterns.

Remaining of the paper describe as follows : section II involves the related work, section III involves the method and material, section IV contains result and discussion and section V contains conclusion.

II. RELATED WORK

The Numerous studies have explored the detection of depressive symptoms through user interviews and social media. Shen et al. [24] utilized Twitter to construct datasets distinguishing between depressed and non-depressed users, while Cohan et al. [25] categorized data from Reddit into nine groups aligned with DSM-5 criteria. Other researchers, such as Jain et al. [26] and Cha et al. [27], have focused on creating models for depression detection using machine learning and deep learning, analyzing content from platforms like 'r/SuicideWatch' and 'r/depression.' Some studies delve into multi-model datasets, as seen in Lin et al. [28], proposing an automated method for detection of depression by utilizing sound and language content with the help of the patient interviews.

Nevertheless, previous research works focused on identifying depressive emotions mainly center on binary categorization challenges, assessing whether individuals exhibit signs of depression or not. As a result, datasets are predominantly organized with this objective in mind, and there is a scarcity of studies addressing the prediction of depression intensity or the classification of intricate depressive emotions. To overcome this disparity, the investigation presents two innovative datasets: (BWS) [29] and DSM-5 dataset. These datasets seek to assess the magnitude of depressive emotions and intricate emotional states, respectively, aligning labels with the Diagnostic and Statistical Manual of Mental Disorders, Fifth Edition [30] (DSM-5) criteria. Furthermore, a tool for annotating Best-Worst Scaling [31] is crafted, leveraging MySQL and Flask, to improve the efficiency of the annotation process. Employing the DistilBERT [32] language model for both training and inference on these datasets, they evaluated its performance against various ML and DL algorithms. Their proposed AuD model, is meticulously designed to quickly evaluate the intensity of depression, intricate depressive emotions, and identify tokens with notable attention scores. By comparing its performance with other algorithms, including BERT [33] and ELECTRA [34], They establish the superiority of our model. Acknowledging the importance of both speed of swift inference and exceptional performance, especially in real-time service scenarios such as chatbots, They assert that their AuD model is well-suited for tasks associated with detection of depression.

In line with the Oxford English Dictionary, the term "emotion" is characterized as a potent sentiment arising from one's circumstances, mood, or interpersonal connections. The author of paper [35] view emotions as intrinsic

constructs integral to socialization, aiding in interactions. Emotions constitute a fundamental aspect of human existence, impacting choices, psychological well-being, and physical health. The task of classification of emotions seeks to identify potential emotions in text, reflecting the mental state of the users. Given their pervasive influence, models for emotion classification find applications across diverse fields, including marketing, medicine, and education.

Analyzing sentiment, a task within natural language processing, discerns the emotional tone present in a statement, determining if it is positive, negative, or neutral. Expanding beyond mere sentiment analysis, emotion classification strives for identifying the specific emotions conveyed in the statement, delving into underlying semantics and embracing feelings such as anger, sadness, and worry. In contrast to the conventional single-label emotion classification, where one emotion is linked to an instance, multi-label emotion classification allocates a subset of labels to each instance. This complex undertaking seeks to grasp the meanings of documents across diverse emotional dimensions. Prior research predominantly leaned on traditional machine learning methods, requiring extensive feature engineering. However, recent advancements in Artificial Neural Networks (ANN) for text classification, particularly CNN and LSTM, have outperformed traditional approaches. Notably, this progress extends beyond sentiment analysis for addressing the multi-label emotion detection challenges.

Inspired by Vaswani et al. [36], a distinctive model of multi-label emotion detection was introduced, incorporating multiple attention mechanisms, and its performance is compared against existing models. The state of art introduces various pre-trained models like , DistilBERT-MA, XLNet-MA and RoBERTa-MA, coupled with multi-attention architectures, showcasing superior outcomes for English datasets. Social networking platforms, notably Facebook and Twitter, have witnessed a surge in popularity, primarily attributed to their capacity for users to openly and freely express opinions, preferences, and dislikes on a global scale [37]. According to the data presented report by May 2022, about 867 million tweets are distributed each day [38]. The escalating abundance of online evaluations and perspectives has ignited heightened interest among service providers, manufacturers, and users. This, in turn, has elevated the significance of SA, a component of NLP. SA is dedicated to recognizing subjective content in expressions, encompassing opinions, evaluations, emotions, and attitudes aimed at a subject, individual, or entity. Since 2004, Sentiment Analysis has encountered substantial expansion, evolving into a dynamic and prolific realm of exploration [39].

Within literature, sentiment analysis (SA) is delineated by three fundamental approaches: Machine Learning [40] (inclusive of deep learning), Lexicon-Based [41], and Hybrid methodologies [42]. Machine Learning stands as the most extensively embraced and established technique. Nonetheless, effectively addressing negation [43] poses an ongoing and intricate challenge in sentiment analysis. Instances such as "I couldn't be prouder" and "I am not proud" might be categorized with negative sentiment, despite conveying opposing emotions. In existing

literature, negation words (not, never, can't, etc.) are frequently integrated into stop-word lists and removed during the pre-processing phase [39], a practice we deem indefensible in the SA context. Another nuanced aspect of sentiment analysis involves the identification of sarcasm, wherein language conveys disdain by indicating the opposite. For instance, the sentence "Of course, I am happy to spend all my money to buy a mobile!" encapsulates a sarcastic tone. The task of extracting accurate sentiment becomes more intricate when handling concise and noisy social media texts, replete with idiosyncrasies. Tweets, in particular, showcase unique features, such as the use of hashtags (e.g., "#sad") to express emotions, the "@" symbol to mention other users (e.g., "@user123"), and the inclusion of emoticons. Unfortunately, a considerable portion of current research overlooks these distinctive symbols within the SA context [44].

III. METHOD AND MATERIALS

• Dataset

The emotion dataset sourced from Hugging Face comprises three subsets: training, validation, and test, totaling 20,000 labeled instances. Each instance contains a 'text' field representing a tweet and a 'label' field classifying the tweet into one of six basic emotions: sadness, joy, love, anger, fear, or surprise. The sample of dataset is shown in Figure 1. The dataset is structured as a DatasetDict, offering convenient access to distinct subsets. The training set, consisting of 16,000 instances, is utilized for model training, while the validation set with 2,000 instances aids in model evaluation and hyperparameter tuning. The test set, also comprising 2,000 instances, serves as an independent dataset for final model performance assessment. The 'label' field employs a ClassLabel encoding with six classes corresponding to the aforementioned emotions. This dataset is well-suited for training and evaluating models aimed at classifying emotions in tweets. The word per Tweets and frequency of classes are shown in Figure 2 and Figure 3 respectively.

	text	label	label_name
0	i didnt feel humiliated	0	sadness
1	i can go from feeling so hopeless to so damned...	0	sadness
2	im grabbing a minute to post i feel greedy wrong	3	anger
3	i am ever feeling nostalgic about the fireplac...	2	love
4	i am feeling grouchy	3	anger

Figure 1: Samples of Dataset

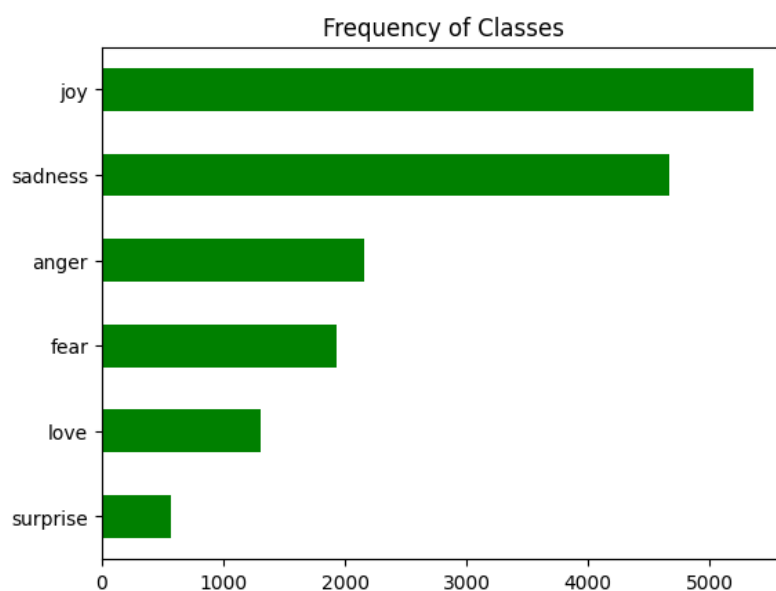


Figure 2: Frequency of Classes

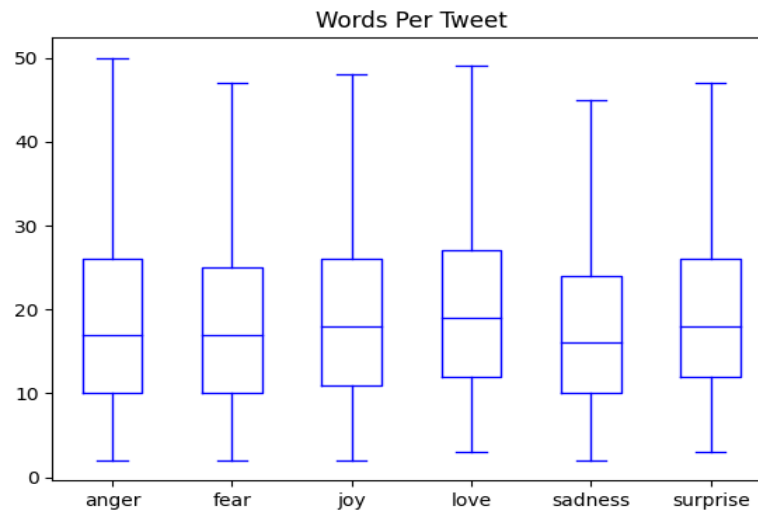


Figure 3: Words Per Tweet

• *Pre-processing Feature Extraction*

We have employed transformer models such as DistilBERT, which necessitate tokenized and numerically encoded input rather than raw strings. Tokenization, the process of breaking down text into atomic units, can follow various strategies. Character and word tokenization represent two extremes. Character-level tokenization involves feeding each character individually to the model, encoding each unique character with a numerical identifier. On the other hand, word tokenization involves splitting the text into words, mapping each word to an integer, and skipping the step of learning words from characters. The AutoTokenizer class, part of the "auto" classes, automatically retrieves the model's configuration, pretrained weights, or vocabulary from the checkpoint's name. For preprocessing, we utilized the map() method of our DatasetDict object, applying tokenization in batches with options for padding and truncation. This approach ensures consistency in input tensor shapes and attention masks across the entire dataset.

• *Fine-tuning Pre-trained Distil BERT Model*

We will leverage the AutoModel class from the Transformers library, specifically designed for convenient model loading, by using the from_pretrained() method to load the pre-trained DistilBERT model. The availability of a GPU is checked using PyTorch, and the model is configured to run on the GPU if present, else it defaults to CPU. The AutoModel class processes token encodings and passes them through the encoder stack, generating hidden states. To extract these states, the string is encoded, tokenized, and converted to PyTorch tensors. The resulting tensor has dimensions [batch_size, n_tokens], and after ensuring it is on the same device as the model, it is passed as input. The torch.no_grad() context manager is employed for inference to reduce memory usage. The hidden state tensor, shaped [batch_size, n_tokens, hidden_dim], is obtained. For classification, the practice is to use the hidden state associated with the [CLS] token as the input feature. This token is extracted by indexing into outputs.last_hidden_state. To process the entire dataset efficiently, the map() method of DatasetDict is used, enabling the extraction of all hidden states simultaneously.

The resulting preprocessed dataset contains the necessary information for training a classifier, with hidden states as input features and labels as targets. For the classifier, a pretrained DistilBERT model is employed, with AutoModelForSequenceClassification. This model includes a classification head for training on top of the pretrained outputs. The number of labels (six in this case) is specified to determine the number of outputs the classification head should predict. To monitor training metrics, a compute_metrics() function is defined for the Trainer, which receives EvalPrediction and returns a dictionary mapping metric names to values. For this application, the F1-score and accuracy are computed. The TrainingArguments class is utilized to define training parameters such as learning rate 2e-5, batch size =64 and epochs =2 which offering fine-grained control over the training and evaluation process. Key arguments include output_dir, which determines the storage location for training artifacts, and parameters like learning rate, batch size, and number of epochs. The training_args instance is configured to load the best model at the end of training. Finally, with all the components in place, the model is instantiated and fine-tuned using the Trainer.

IV. RESULT AND DISCUSSION

In analyzing the results of our fine-tuned DistilBERT model for emotion classification, a glimpse at the training logs reveals a remarkable improvement over the feature-based approach. The validation set showcases an impressive F1-score of approximately 92%, underscoring the model's enhanced performance. Diving deeper into the training metrics, a pivotal tool for evaluation is the confusion matrix. This allows us to visualize the model's predictions on the validation set and gain insights into its performance. Leveraging the predict() method of the Trainer class provides valuable objects for evaluation, such as the PredictionOutput containing predictions, label_ids, and pertinent metrics. The metrics obtained from the validation set include a test loss of 0.2205, a test accuracy of 92.25%, and an F1-score of 92.26%. These metrics signify the robustness and effectiveness of our

model in capturing the nuances of emotional content within tweets. Decoding the raw predictions using `np.argmax()` exposes the predicted labels, facilitating a granular examination of the model's performance. Intriguingly, the model exhibits some misclassifications, as evident in instances where predicted labels diverge from the true labels. For example, phrases expressing joy are occasionally mislabeled as sadness, suggesting challenges in distinguishing subtle emotional nuances. A closer look at specific examples from the dataset reinforces the model's capability but also highlights areas for improvement. Instances where the model mispredicted labels indicate complexities in disentangling emotions. For instance, a statement expressing admiration may be mistakenly labeled as sadness, underscoring the intricacies of emotional language. Moreover, there are instances where no clear class emerges, hinting at potential mislabeling or the need for a new class altogether. In particular, the emotion 'joy' appears to be mislabeled several times. This observation underscores the importance of dataset refinement, a crucial step in enhancing model performance. By addressing mislabeled instances and introducing new classes, the model's ability to discern subtle emotional nuances could be further refined.

It's essential to acknowledge the inherent challenges in emotion classification, given the subjectivity and variability in how individuals express and interpret emotions. While the model showcases a commendable level of accuracy, its misclassifications illuminate the nuances and intricacies inherent in understanding emotional context.

V. CONCLUSION AND FUTURE WORK

In conclusion, the fine-tuned DistilBERT model has proven to be a robust tool for emotion classification in tweets, exhibiting significant advancements over the feature-based approach. The achieved F1-score of approximately 92% on the validation set underscores the model's efficacy in discerning emotional nuances within diverse textual content. The insightful analysis of misclassifications sheds light on the complexities inherent in understanding and categorizing emotions, emphasizing the ongoing challenges in this domain. The results highlight the potential of leveraging advanced language models, such as DistilBERT, for emotion classification tasks, showcasing the capacity to capture subtle variations in emotional expression. The misclassifications observed provide valuable cues for further refinement of the dataset, emphasizing the iterative nature of model development. Addressing mislabeled instances and considering the introduction of new classes can enhance the model's ability to navigate the intricacies of emotional language. Future work in this domain should focus on several key aspects. Firstly, continuous refinement of the dataset is crucial to improve model generalization and ensure it adapts to the evolving landscape of online communication. Efforts to address mislabeling and incorporate new classes should be undertaken, drawing from ongoing advancements in emotion theory and linguistic understanding. Additionally, exploring ensemble models that combine the strengths of multiple architectures could further enhance the robustness of emotion classification systems. Integration with emerging transformer-based models and leveraging transfer learning from pre-trained models may contribute to more nuanced and context-aware emotional analysis.

Furthermore, investigating interpretability and explainability in the model's predictions is essential for building trust and understanding its decision-making processes. Developing methodologies to interpret the attention mechanisms and crucial input features can provide valuable insights into how the model perceives and classifies emotional expressions. Considering the dynamic nature of language and online communication, continuous monitoring and adaptation of the model to evolving linguistic trends and expressions are crucial. This involves regular updates to the model with fresh data and the integration of real-time learning mechanisms. In the broader context, the application of emotion classification models extends beyond sentiment analysis and can find utility in diverse fields such as mental health monitoring, customer feedback analysis, and social dynamics research. Exploring these applications and tailoring models for specific domains can unlock novel avenues for research and practical implementation. In essence, while the achieved results demonstrate promising strides in emotion classification, the journey towards a nuanced understanding of human emotion in textual data is an ongoing exploration. By combining advancements in natural language processing with domain-specific insights, we can continue to refine models that accurately capture the rich tapestry of human emotions in the ever-evolving landscape of digital communication.

CONFLICTS OF INTEREST

The authors declare that they have no conflicts of interest

REFERENCES

- [1] Mohammad S, et al. Semeval-2018 task 1: affect in tweets. In: Proceedings of the 12th international workshop on semantic evaluation. 2018
- [2] Pang B, Lee L, Vaithyanathan S, Thumbs up? Sentiment classification using machine learning techniques. arXiv preprint cs/0205070, 2002.
- [3] Xia R, Zong C, Li S. Ensemble of feature sets and classification algorithms for sentiment classification. Inf Sci. 2011;181(6):1138–52.
- [4] He Y. A Bayesian modeling approach to multi-dimensional sentiment distributions prediction. In: Proceedings of the First International Workshop on Issues of Sentiment Discovery and Opinion Mining. 2012
- [5] Sharma A, Dey S. A document-level sentiment analysis approach using artificial neural network and sentiment lexicons. ACM SIGAPP Appl Comput Rev. 2012;12(4):67–75.
- [6] Alessia D, et al. Approaches, tools and applications for sentiment analysis implementation. IJCA. 2015;125(3):26–33.
- [7] Biswas S. Advantages of deep learning, plus use cases and examples. <https://www.width.ai/post/advantages-of-deep-learning>. Accessed 10 Nov 2021.
- [8] Kalchbrenner N, Grefenstette E, Blunsom P. A convolutional neural network for modelling sentences. arXiv preprint [arXiv:1404.2188](https://arxiv.org/abs/1404.2188). 2014.
- [9] Hong J, Fang M. Sentiment analysis with deeply learned distributed representations of variable length texts. Stanford: Stanford University Report; 2015. p. 1–9.
- [10] Bhattacharya A. Deep hybrid learning—a fusion of conventional ML with state of the art DL. <https://towardsdatascience.com/deep-hybrid-learning-a-fusion-of-conventional-ml-with-state-of-the-art-dl-cb43887fe14>. Accessed 26 Jul 2020.

- [11] Nimmi K, et al. Pre-trained ensemble model for identification of emotion during COVID-19 based on emergency response support system dataset. *Appl Soft Comput.* 2022;122: 108842.
- [12] Adoma AF, Comparative analyses of bert, roberta, distilbert, and xlnet for text-based emotion recognition. 17th International Computer Conference on Wavelet Active Media Technology and Information Processing (ICCWAMTIP), 2020.
- [13] Sirisha U, Boleem SC. Aspect based sentiment & emotion analysis with ROBERTa, LSTM. *IJACSA.* 2022. <https://doi.org/10.14569/IJACSA.2022.0131189>.
- [14] Bansal B, Srivastava S. Hybrid attribute based sentiment classification of online reviews for consumer intelligence. *Appl Intell.* 2019;49(1):137–49.
- [15] Ma Y, et al. Sentic LSTM: a hybrid network for targeted aspect-based sentiment analysis. *Cogn Comput.* 2018;10(4):639–50.
- [16] Zainuddin N, Selamat A, Ibrahim R. Hybrid sentiment classification on twitter aspect-based sentiment analysis. *Appl Intell.* 2018;48(5):1218–32.
- [17] Prottasha NJ, Sami AA, Kowsher M, Murad SA, Bairagi AK, Masud M, Baz M. Transfer learning for sentiment analysis using BERT based supervised fine-tuning. *Sensors.* 2022;22:4157.
- [18] Jain PK, et al. Employing BERT-DCNN with sentic knowledge base for social media sentiment analysis. *J Ambient Intell Human Comput.* 2022. <https://doi.org/10.1007/s12652-022-03698-z>.
- [19] Tan KL, et al. RoBERTa-LSTM: a hybrid model for sentiment analysis with transformer and recurrent neural network. *IEEE Access.* 2022;10:21517–25.
- [20] Jain PK, Saravanan V, Pamula R. A hybrid CNN-LSTM: a deep learning approach for consumer sentiment analysis using qualitative user-generated contents. *ACM Trans Asian Low-Resour Lang Inf Process.* 2021;20(5):84.
- [21] AlBadani B, Shi R, Dong J. A novel machine learning approach for sentiment analysis on twitter incorporating the universal language model fine-tuning and SVM. *Applied System Innovation.* 2022;5(1):13.
- [22] Fredrick H. Why is twitter important?. <https://yourbusiness.azcentral.com/twitter-important-5023.html> Accessed Jan 2024.
- [23] A Red Light for Modern Mental Health and Stress Management Are Essential. Available online: <http://www.medical-tribune.co.kr/news/articleView.html?idxno=100431> (accessed on 28 December 2023).
- [24] Guangyao, S.; Jiang, J.; Liqiang, N.; Fuli, F.; Cunjun, Z.; Tianrui, H.; Tat-Seng, C.; Wenwu, Z. Depression detection via harvesting social media: A multimodal dictionary learning solution. In *Proceedings of the 26th International Joint Conference on Artificial Intelligence (IJCAI'17)*, Sidney, Australia, 19–25 August 2017; AAAI Press: Washington, DC, USA; 2017; pp. 3838–3844.
- [25] Cohan, A.; Desmet, B.; Yates, A.; Soldaini, L.; MacAvaney, S.; Goharian, N. SMHD: A Large-Scale Resource for Exploring Online Language Usage for Multiple Mental Health Conditions. In *Proceedings of the 27th International Conference on Computational Linguistics*, Association for Computational Linguistics, Santa Fe, NM, USA, 20–26 August 2018; pp. 1485–1497. [Google Scholar]
- [26] Pratyaksh, J.; Srinivas, K.R.; Vichare, A. Depression and suicide analysis using machine learning and NLP. *J. Phys. Conf. Series* 2022, 2161, 1. [Google Scholar]
- [27] Cha, J.; Kim, S.; Park, E. A lexicon-based approach to examine depression detection in social media: The case of Twitter and university community. *Humanit. Soc. Sci. Commun.* 2022, 9, 325.
- [28] Lin, L.; Chen, X.; Shen, Y.; Zhang, L. Towards Automatic Depression Detection: A BiLSTM/1D CNN-Based Model. *Appl. Sci.* 2020, 10, 8701.
- [29] Louviere, J.J.; Flynn, T.N.; Marley, A.A.J. *Best-Worst Scaling: Theory, Methods and Applications*; Cambridge University Press: Cambridge, MA, USA, 2015. [Google Scholar]
- [30] Kumar, D., & Ahamad, F. (2023, December). Opinion extraction from big social data using machine learning techniques: A survey. In *AIP Conference Proceedings* (Vol. 2916, No. 1). AIP Publishing.
- [31] Sanh, V.; Debut, L.; Chaumond, J.; Wolf, T. DistilBERT, a distilled version of BERT: Smaller, faster, cheaper and lighter. *arXiv* 2019, arXiv:1910.01108.
- [32] Devlin, J.; Chang, M.-W.; Lee, K.; Toutanova, K. BERT: Pre-training of Deep Bidirectional Transformers for Language Understanding. In *Proceedings of the 2019 Conference of the North American Chapter of the Association for Computational Linguistics: Human Language Technologies*, Volume 1 (Long and Short Papers), Minneapolis, MI, USA, 2–7 June 2019; Association for Computational Linguistics: Toronto, Canada, 2019; pp. 4171–4186.
- [33] Clark, K.; Luong, M.-T.; Le Quoc, V.; Christopher, D. Manning: Electra: Pre-training text encoders as discriminators rather than generators. *arXiv* 2020, arXiv:2003.10555.
- [34] Hwang, H. C., & Matsumoto, D. (2016). Facial expressions.
- [35] Vaswani, A., Shazeer, N., Parmar, N., Uszkoreit, J., Jones, L., Gomez, A. N., ... & Polosukhin, I. (2017). Attention is all you need. *Advances in neural information processing systems*, 30.
- [36] Kumar, D., & Ahamad, F. (2024, February). Application of Machine Learning Algorithm for Optimal Model Design for Opinion Extraction. In *TELEMATIQUE* (Volume 23 Issue 1).
- [37] M. Yaqub. (2022). How many tweets per day 2022 (new data). [Online]. Available: <https://www.renolon.com/number-of-tweetsper-day/>
- [38] M. Birjali, M. Kasri, and A. Beni-Hssane, “A comprehensive survey on sentiment analysis: Approaches, challenges and trends,” *Knowledge-Based Systems*, vol. 226, 107134, 2021.
- [39] Rashid and C.-Y. Huang, “Sentiment analysis on consumer reviews of amazon products,” *International Journal of Computer Theory and Engineering*, vol. 13, no. 2, p. 7, 2021.
- [40] H. A. Alhammi and K. Haddar, “Building a Libyan dialect lexiconbased sentiment analysis system using semantic orientation of adjective-adverb combinations,” *Int. J. Comput. Theory Eng.*, vol. 12, no. 6, pp. 145–150, 2020.
- [41] Z. Madhoushi, A. R. Hamdan, and S. Zainudin, “Sentiment analysis techniques in recent works,” in *Proc. 2015 Science and Information Conference (SAI)*, 2015, pp. 288–291.
- [42] Sangeeta and N. G. Singh, “Review of factors affecting efficiency of twitter data sentiment analysis,” *International Journal of Computer Theory and Engineering*, vol. 12, no. 1, pp. 53–58, 2020.
- [43] P. Dandannavar, S. Mangalwede, and S. Deshpande, “Emoticons and their effects on sentiment analysis of twitter data,” in *Proc. EAI International Conference on Big Data Innovation for Sustainable Cognitive Computing*, 2020, pp. 191–201.