

Leveraging NLP for Personality Prediction from Social Media Text

Megha

Dept. Computer Science and Engineering
Chandigarh University
Punjab, India
angelmegha1010@gmail.com

Narinder Yadav

Dept. Computer Science and Engineering
Chandigarh University
Punjab, India
narinder.e16474@cumail.in

Vandana

Dept. Computer Science and Engineering
Chandigarh University
Punjab, India
22bcs17096@cuchd.in

Sanjay Singla

Dept. Computer Science and Engineering
Chandigarh University
Punjab, India
Sanjay.e13538@cumail.in

Abstract—The ever-increasing nature of data generated by social media offers ample reasons for the extraction of meaningful information and proper categorization of the same information. One major challenge in natural language processing is text classification, as it allows the organizing of unstructured texts into categories of interest for value addition, such as sentiment. [1] Although deep learning models have brought so much breakthrough in this field, accuracy in classification is still a viable space for improvement. This study, therefore, applies the natural language processing techniques on a dataset from a WhatsApp group to determine the sentiment using five models of deep learning: Neural Network, Recurrent Neural Network, Long Short-Term Memory, Bidirectional LSTM, and Convolutional Neural Network. We have proposed a hybrid approach of CNN-BiLSTM that uses activations, dropouts, filters, kernel sizes, and numerous layers along with feature extraction to enhance the classification accuracy of sentiment. The proposed model's results are tested with comparisons to earlier studies. Among the individual models used, the highest result was 81 percent by LSTM and BiLSTM, and this hybrid model was much better as it achieved an accuracy of 88 percent. In fact, the proposed model can be seen to perform well in comparison with existing work concerning the task of sentiment classification.

Index Terms—Classification, WhatsApp group, Hybrid CNN-BiLSTM, NLP, Ocean Model

I. INTRODUCTION

classification of texts is an encoding process in which textual data is put into systematic categories using well-defined categories based on inherent properties of text. TC makes it possible to efficiently analyze huge volumes of raw, unstructured texts by automatically analyzing them. Thus, hidden insight may, after all, [2] be retrieved from there. Most of the typical classification of TC systems falls into hybrid, rule-based, and machine learning-based methods. As for the first type, rule-based systems utilize pre-defined hand-coded rules to be interpreted for classifying input strings. Machine learning systems, instead, rely on patterns observed in the training data and prior observations. The strategy to combine the benefits from both kinds of approaches usually yields better

efficiency and superior accuracy of classification in the hybrid system.

Accurate text classification is central to improved understanding and utilization of unstructured data. [3] Recent developments in Deep Learning (DL) have further propelled advances along this trajectory. Generally speaking, the common DL models applied to TC tasks are NNs, RNN, LSTM, and CNN. Very encouraging results have been found from these models regarding significant patterns they extract from text. Techniques like Word2Vec and GloVe improve the performance of classification at the sentence, paragraph, and document levels. [4] A fair amount of reports claimed that with these models combined, more accurate results are achieved. The remarkable success related to CNN models and LSTM models is especially noteworthy.

Many researchers have been working on improving DL frameworks to make NLP tasks more accurate. CNN, LSTM, and Bidirectional LSTM (BiLSTM) models have already been reported to achieve up to an accuracy of 77.4 percent. Hybrid models like BiLSTM-MLP have also shown improved performance by achieving up to an accuracy of 88.3 percent, depicting the success of merging multiple DL techniques to handle the complexity caused by textual data.

The proposed research aims at providing a new model of CNN-BiLSTM hybrid for sentiment analysis over a WhatsApp group dataset. The proposed model could benefit from both the extraction of spatial features of CNN and handling temporal dependency of BiLSTM for better classification accuracy. Unlike traditional LSTM, where only past data is being taken into consideration, BiLSTM models also consider past and future contexts leading to better overall predictive performance. [5]

To enhance the classification, the model uses multiple sizes of convolutional kernels to capture local temporal dependencies in the text. The proposed system is evaluated by using several performance metrics, namely accuracy, precision, recall, and F1-Score; thus, it is benchmarked against existing deep

learning models. The following are the major contributions of this study: development of a novel hybrid DL model for sentiment analysis, testing with a WhatsApp group dataset, and comparisons with the state-of-the-art DL model for text classification. [?]

II. LITERATURE REVIEW

Text classification and sentiment analysis play a fundamental role in many NLP applications. It enables the system to classify and automatically analyze the text. In this regard, the number of such developed models to increase the accuracy and effectiveness of such processes, especially DL models using various techniques for better performance, that have been introduced over the years.

A. Thai Language Sentiment Analysis Using DL Models

Among these studies, one that noted the application of sentiment analysis for the Thai language using DL classification models focuses on how word embedding, POS tagging, and sentic features improve the process of sentiment analysis. Sentic features give much insight to word polarity and subjectivity, thus greatly enhancing the understanding of the context in which sentiment will prevail. The three unique models were taken into consideration in this research, with LSTM generally reaching an F1-score of 0.726 featuring sentic characteristics and POS-tag embedding. In an interesting fashion, the CNN model performed the best with an F1-score of 81.7 percent, reflecting the idea that CNN uses the extraction of local features, making it profoundly effective in sentiment analysis. [6] Further, deeper learning features or hybrid models could be utilized for even more enhancement in the sense of sentiment classification in future research studies.

B. Improving Text Classification Accuracy with Bi-LSTM

Another study used some new advancement in the accuracy of text classification through the development of a Bi-LSTM model which incorporates Word2Vec, CNN, and an attention mechanism. The hybrid approach relies on Word2Vec for more improved word vector representation, CNN for feature extraction, as well as the application of an attention mechanism to focus where in the text is the most important section. This model yielded accuracy of 87.4 percent, achieving parameters: skipgram with size 300 embedding, Adam optimizer, and a batch size of 128. On 13k instances, it reached the peak value of F1-score, being 90.1 percent. Results indicated that when some combination of models such as CNN, Bi-LSTM, and attention mechanisms were applied, significant changes might be expected for a typical classification accuracy, increasing sufficiently when the dataset size was augmented. [7]

C. CNN and Bi-LSTM Hybrid Model for Sentiment Classification

The most recent research was presented by Salur and Aydin that introduced a highly innovative hybrid deep learning approach to the task of sentiment classification with a focus on

overcoming two major problems of the dataset: complexity and imbalance. They combined CNN for extracting local features from text with LSTM-based RNN for contextual long-term information. Beyond those, character embeddings and pre-trained embeddings such as Word2Vec, GloVe, and FastText were further used for improving representation features. The proposed hybrid CNN-BiLSTM model has outperformed a single BiLSTM model with an accuracy of 80.44 percent when tested on the dataset containing tweets related to a Turkish GSM operator; this stands at 82.14 percent. [8]

As described above, there are many applications of deep learning models. Among these applications, deep learning can be used in the domain of language to analyze and predict sentiment in text form. [9]

In this regard, this paper discusses Urdu text analysis and prediction of sentiment using deep learning models, including CNN, BiLSTM, and their hybrids. Naqvi et al. even discussed their attempt at Urdu text with sentiment analysis, wherein the author has applied DL techniques for achieving accuracy. Instead, they assume a technique by combining CNN for local features and LSTM-based RNN for a better understanding of long-term context in their methodology. Four different embedding models, Samar, CoNLL, pre-trained, and self-trained FastText, were found in the study, and the BiLSTM ATT model attained a maximum accuracy of 77.9 percent. On the other hand, the maximum precision was achieved by LSTM using Samar embedding with an accuracy of 85.16 percent. These results have proven the possibility of DL models in sentiment analysis on languages such as Urdu; and their adaptation is possible through some more techniques. [10]

D. Hybrid CNN-BiLSTM Model in the WAG Dataset Analysis

In this paper, an innovative hybrid model based on the research findings of pre-conducted studies that combines CNN along with Bi-LSTM techniques for WAG dataset sentiment analysis to improve predictive capabilities is proposed. It includes the ability of CNN to extract local features and the capability of the BiLSTM in managing long-term dependencies, which relies on both previous and future contexts for it to make better predictions of sentiments. Techniques of padding were used to deal with variable lengths of data; in addition, ReLU and Softmax activation functions were applied to enhance flexibility. The model is able to capture detailed patterns from the input data if set with an embedding size of 300, and applying differently sized convolutional kernels has been helpful in extracting local temporal dependencies. [?]

This is, in fact a very innovative approach wherein four LSTM layers are used and 64 units have been selected in the final layer. A hybrid model combining the CNN for extraction of local features and BiLSTM for sequential data helps improve the performance of such a hybrid model, which is found to be quite effective in handling complex textual datasets, including WAG, capturing dependencies with functions such as both local and global. Overall, the overall tendency of literature is oriented toward hybrid models that combine the strengths

of different DL architectures to develop text classification and sentiment analysis for various languages and datasets. [?]

III. MATERIALS AND METHODS

Within the scope of the TC domain, this research aims to fine-tune a hybrid model specifically designed for application to sentiment analysis. We critically tested this hybrid model through techniques borrowed from both CNN and BiLSTM architectures. Our approach addresses several salient points, such as the selection of appropriate data, correct labeling of datasets, strong construction of feature vectors, and the hybrid CNN-BiLSTM framework. All of these together enhance the precision as well as dependability in the context of sentiment analysis. We have applied this model for predicting the polarity of the textual data into respective categories based on the sentiments. [?]

A. Dataset

The data for this study was gathered from the "Forum DTC Riau" WhatsApp group, which consists of 134 Daihatsu Taruna owners located in the Riau region. It was selected mainly for fitting our research objectives and regional characteristics, where very particular user preferences, concerns, and opinions are reflected. We further tested the model using the ten thousand-row Amazon Product Summaries dataset from Kaggle. In order to get a comparable result, we prepared and tested the same methods both on the Amazon dataset and on the previous one.

The WhatsApp group was launched on 11/10/2018. Data for this paper were collected between 03/16/2023 and 03/15/2023. Data extraction activities were carried out using OPPO A15 smartphone, running Android Version 10, and the data gathering was done using 3 GB RAM, octa-core processor. The exported texts from the Whatsapp group conversations were sent through emails for further processing. [15]

The data from the WhatsApp group retrieved was messy; it included encrypted messages, call logs, and timestamps. In order to prepare this for sentiment analysis, we required making it machine readable. Simply put, we applied NLP techniques to interpret the content, classify it, and to glean opinions and sentiments expressed within the message. The NLP approach for this study involves two main stages: labeling and text preprocessing. [18]

B. Labeling

Raw data from the WhatsApp group is, for example, in timelines, phone numbers, member names, emojis, and multimedia tags. As the messages did not have any pre-established sentiment classification, emotion-based sentiment analysis was applied through SentiWordNet to set labels on the data. Labeling was done this way: 1. We preprocessed the data by using regular expressions to order our data concerning tokens date, time, author, and message content. [17] 2. The messages are translated to English using Google TransTranslator and tokenized. We used the nltk.sentiment.vader module to classify the messages as positive, negative, neutral, or compound based

on their sentiment scores. 3. For the Amazon dataset, the sentiment was labelled as: score ≥ 3 as +ve, 3 as neutral, and ≤ -3 as -ve.

C. Preprocessing

Preprocessing is one of the most critical processes in text classification, as data need to be prepared before further analysis is done on the data. The research presented here utilized a chain of preprocessing steps: 1. Removing HTML tags and URLs. 2. Replacement of negation words with their antonyms. 3. Elimination of neutral sentiments. 4. Removing punctuation and lemmatizing the text to abbreviate it into its root word.

A big challenge in NLP was how to design models that could capture the deep structure of sentences. Feature extraction helped to reduce the dimension of the data, making it efficient without loss of central information. Encoding techniques in this paper converted categorical data into numerical forms. Tokenization and padding facilitated consistency in dimensions for input and output. Data thus arranged in coherent clusters.

D. Proposed Model

We tested the dataset after labeling and preprocessing through several algorithms that included Neural Networks, RNN, LSTM, BiLSTM, and CNN. And after cleaning the data, tokenization, and encoding, we split our dataset into a training set and a testing set. We have proposed a hybrid CNN-BiLSTM model with an application of the embedding technique that presents text in low-dimensional vectors, while setting the size of padding to 300. It distinguishes itself from previous methods as it does not use embedding. The results were compared with other single-algorithm approaches. (see Table 1)

TABLE I
HYPERPARAMETERS OF THE PROPOSED HYBRID MODEL "SEQUENTIAL"

Layer (Type)	Output Shape	Param #
Embedding	(None, 100, 300)	7,500,000
Conv1D	(None, 98, 200)	180,200
Bidirectional	(None, 98, 128)	135,680
Dropout	(None, 98, 128)	0
Bidirectional	(None, 128)	98,816
Dense	(None, 50)	6,450
Dense	(None, 50)	2,550
Flatten	(None, 50)	0
Dense	(None, 100)	5,100
Dense	(None, 2)	202

This architecture was contrasted with former best performing studies as depicted in fig II. In this paper, the model architecture was designed independently for data-specific characteristics in each data set. For instance, Jang et al. Clothing and cameras product review — CNN-BiLSTM with Word2Vec Skip-Gram (Umarani et al. [55]); revise of Turkish GSM users tweets model embedding character-level and FastText-word — CNN-BiLSTM model (Salur Aydin [15]) Instead, Soumya and Pramod intended to perform sentiment analysis in Malayalam tweets using CNN-BiLSTM and CNN-LSTM. [19]

CNN layers are integrated for accurate extraction of text features with BiLSTM layers to seek contextual relationships between words, thus greatly improving overall understanding and interpretation of the text of the analysis. The CNN layers carry a feature filter and kernel, determine the important patterns of the text; while the BiLSTM layers process the word relationships in a bidirectional way. Activation functions such as ReLU introduce non-linearity while the dropout mechanisms keep away overfitting by randomly deactivating neurons during training.

To further prevent overfitting, we added a dropout layer immediately after the first BiLSTM layer. The feature drops out neurons during training and, if the dropout rate is 0, the active neurons ensure that the model does not overfit less than the data. Finally, including a dense layer with fewer units reduces the capacity of the model. The addition of CNN and BiLSTM in a single architecture will ensure that the data undergo more exhaustive processing. The hybrid CNN-BiLSTM model captures the spatial features whereas BiLSTM captures temporal and contextual information. The process, therefore, begins with labeling and preprocessing along with the feature extraction techniques discussed above in order to have enhanced analysis. [?]

Our model applies an embedding layer to feedforward networks in order to represent text as lowdimensional numerical vectors and uses padding to the size of 300. Embedding differs from methods used in, as demonstrated in table II and III.

TABLE II
THE FUNCTIONALITY OF THE CURRENT DL MODEL 1

Model	Embedding
CNN LSTM	POS Tagging
CNN Bi LSTM [9]	word2vec Sentiment Analysis
CNN BiLSTM [13]	Skip Gram
CNN BiLSTM [17]	Carakter + Fastext
CNN LSTM [26]	Tagging, Wordvector

TABLE III
DL MODEL 2'S EXISTING PERFORMANCE

F1-Score	Accuracy
81.7%	77.4%
88.0%	87.6%
89.0%	82.1%
75.0%	85.5%

IV. RESULT AND DISCUSSION

The VADER Sentiment library is used in order to classify the data into sentiments that are pleasant, neutral, or negative. In the context of this study, only positive and negative sentiments have been incorporated. The original dataset had 5,237 rows, but after filtering out the neutral sentiments, it was reduced to 1,089 rows. The data set was filtered, then stopwords were removed, and the data was labeled using an encoder to separate positive from negative sentiments. It was split into the training set and the testing set, with 80 percent assigned to training and 20 percent for testing. Table

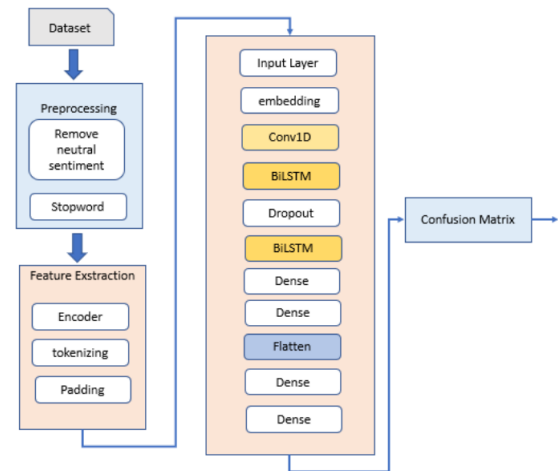


Fig. 1. The hybrid model's suggested architecture

IV displays the distribution of the training and testing sets with random state = 69 in detail.

TABLE IV
PARTITIONING THE TRAINING AND TESTING DATASETS

Feature	Training Dataset		Testing Dataset	
	DTC	PS	DTC	PS
Length	1,552	7,311	389	1,828
Sentences shape	1,520	7,311.0	389.0	1,828.0
Labels shape	1,522	7,311.2	389.2	1,828.2

Model with Softmax and ReLU activation functions; using a dropout of 0.5 to avoid overfitting and optimized with the Adam algorithm. Training occurs for 20 epochs. As shown in Table I, it integrates an embedding layer, a CNN layer, a bidirectional Long Short-Term Memory (BiLSTM) layer, and four dense layers. In contrast to prior works [13, 17, 18], the proposed method adopts different sizes of filters to capture various local dependencies, thus enriching the ability of feature extraction. Figure 1 proposes architecture of hybrid CNN-BiLSTM model. The performance of the single DL as well as the hybrid DL models is evaluated by their respective performances. [21]

The activation functions used are Softmax and ReLU. Overfitting has been minimized with a dropout of 0.5. The model trains for 20 epochs, and Adam is the optimizer. Table I: General Architecture of the model (embedding layer + CNN+Bi-LSTM + dense * 4). The model uses a methodology that has not been suggested in earlier studies [13, 17, 18]. Different filter sizes are used in order to capture varied local dependencies by the model, This will help in the feature extraction process. Figure 1 illustrates the architecture of the proposed CNN-BiLSTM hybrid model. Both single DL and hybrid DL models are tested for performance. [22] The confusion matrix presents the results while all metrics, such as accuracy, precision, recall, and F1-Score, are used within the model evaluation. Comparing these findings to earlier research using the same methodology, but improvement was

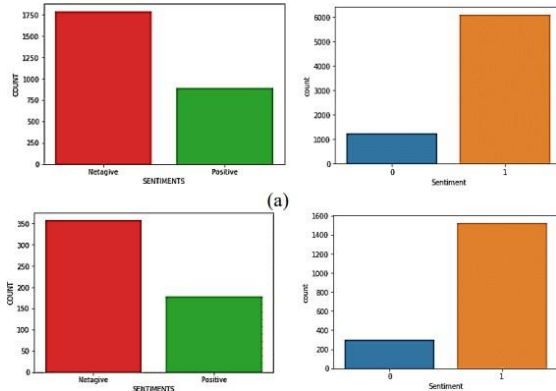


Fig. 2. DTC and PS dataset sentiment distribution: (a) training dataset (b) testing dataset

incorporated by slightly different approaches in this work. The proposed models in this research work are tested using the next performance metrics of performance to evaluate the efficiency:

Accuracy = $\frac{TP}{TP + TN + FP + FN}$ (1) where TP stands for True Positives, FN stands for False Negatives, TN stands for True Negatives and FP for False Positives. Precision = $\frac{TP}{TP + FP}$ (2) Precision is considered as the ratio of the samples correctly classified to all the samples that are predicted to be positive. It measures the fraction of true positive samples among all the actual positive samples. Recall = $\frac{TP}{TP + FN}$ (3) F1-Score is a harmonic mean of precision and recall as follows:

F1-Score = $\frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$ (4)
(4) The PS dataset was larger Compared to the DTC dataset, both during training and test set, providing more potent representation. On the other hand, the DTC dataset was found to be smaller. Therefore, the time taken toward the training process was much quicker, though unable to generalize complex patterns as shown in Table III, since the PS set is more positive and the DTC set is more negative, the split results of the set also reflect the difference in the sentiment distribution of the two sets. [?]

After sampling the test and train datasets, other preprocessing steps were used to prepare the data for analysis. This entailed tokenizing and padding with a specific feature with trunc type = post, max length = 100, embedding dim = 300, vocab size = 25000, and oov tok = OOV. It was used for efficient representation and standardization into models for the purpose of analyses. Each sequence in the dataset was equalized with regard to their length through tokenization and padding procedures. Although the PS dataset is very large with 7,311 train samples and 1,828 test samples padded to a length of 100, DTC has only 871 training samples and 218 testing samples. In this context, the large size of the PS dataset offers the possibility of better generalization to complex patterns. This shows the case of being limited for the much smaller DTC dataset. The extracted dataset was tested using multiple single deep learning classification

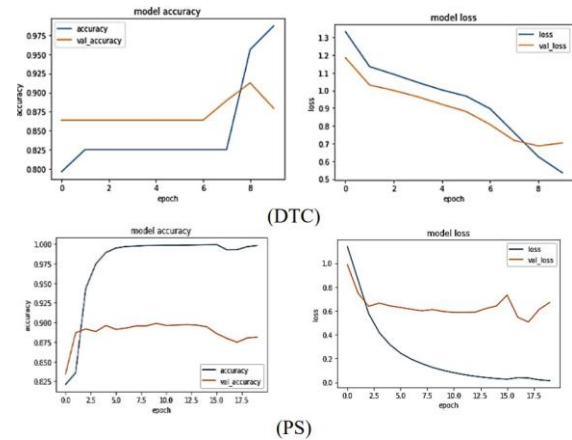


Fig. 3. Plots of accuracy versus epochs and loss versus epochs were derived from the suggested model.

models, like CNN, RNN, LSTM, Bidirectional LSTM, and Neural Network. The preprocessing stage as well as feature extraction stages of these datasets were modified to evaluate these models' efficacy. The hybrid model, with a score of 0.88 in both datasets, has the highest accuracy compared with the other models; however, precision was less than that obtained by LSTM and BiLSTM, that had a value of 0.76. Precision dropped down to 0.76 on the DTC dataset, while it performed surprisingly well on the PS dataset with F1-Scores of roughly 0.79 and 0.78 percent, respectively, indicating strong accuracy and recall. This implies that the model correctly categorized data correctly with precision between both datasets. Encoder labeling, data splitting, tokenization, and padding were all included in feature extraction. The deep learning models have split the data for testing in 80:20 ratio. As illustrated in table IV, the CNN-BiLSTM model obtained an accuracy score of 0.81 percent using 871 test samples, including 218 test data points from the WAG 'Forum DTC Riau' dataset with the potential to be improved further by this hybrid model.

V. CONCLUSION

This paper describes the dedicated hybrid model of CNN-BiLSTM, specifically for application to WAG data analysis. The CNN exploits its skill in feature extraction, whereas the BiLSTM captures the deep dependencies to represent both interdependencies sequential data, both short-term and long-term. Preprocessing, feature extraction, activation, and dropout are all incorporated into hybrid architecture with multiple layers using filters of kernel size 3. Experimental results have pointed out that this new CNN-BiLSTM attained an accuracy rate of 88 percent. Moreover, a 7-point improvement in processing WAG data has further validated the effectiveness of this hybrid model in comparison to the standalone Bidirectional LSTM model.

REFERENCES

- [1] M. P. Akhter, Z. Jiangbin, I. R. Naqvi, M. Abdelmajeed, A. Mehmood, and M. T. Sadiq, "Document-level text classification using single-layer

- multisize filters convolutional neural network,” IEEE Access, vol. 8, no. 1, pp. 42689–42707, 2020. doi: 10.1109/ACCESS.2020.2976744.
- [2] A. Wahdan, S. Hantoobi, S. A. Salloom, and K. Shaalan, “A systematic review of text classification research based on deep learning models in Arabic language,” Int. J. Electr. Comput. Eng., vol. 10, no. 6, pp. 6629–6643, 2020. doi: 10.11591/IJECE.V10I6.PP6629-6643.
 - [3] W. Fang, H. Luo, S. Xu, P. E. D. Love, Z. Lu, and C. Ye, “Automated text classification of near-misses from safety reports: An improved deep learning approach,” Adv. Eng. Informatics, vol. 44, no. March 2019, 101060, 2020. doi: 10.1016/j.aei.2020.101060.
 - [4] Z. Liu, C. Lu, H. Huang, S. Lyu, and Z. Tao, “Hierarchical multi-granularity attention-based hybrid neural network for text classification,” IEEE Access, vol. 8, pp. 149362–149371, 2020. doi: 10.1109/ACCESS.2020.3016727.
 - [5] H. Yang, L. Luo, L. P. Chueng, D. Ling, and F. Chin, “Deep learning and its applications to natural language processing,” in Deep Learning: Fundamentals, Theory and Applications, 2019, pp. 89–109.
 - [6] Q. Li et al., “A survey on text classification: From shallow to deep learning,” IEEE Trans. Neural Networks Learn. Syst., vol. 31, no. 11, pp. 1–21, 2020.
 - [7] F. Zaman, M. Shardlow, S. Hassan, and N. Radi, “HTSS: A novel hybrid text summarisation and simplification architecture,” Inf. Process. Manag., vol. 57, no. 6, 102351, 2020. doi: 10.1016/j.ipm.2020.102351.
 - [8] K. Pasupa, T. Seneewong, and N. Ayutthaya, “Thai sentiment analysis with deep learning techniques: A comparative study based on word embedding, POS-tag, and sentic features,” Sustain. Cities Soc., vol. 50, no. 7, 101615, 2019. doi: 10.1016/j.scs.2019.101615.
 - [9] K. Miok, D. Nguyen-Doan, B. S. Krlj, D. Zaharie, and M. Robnik-Sikonja, “Prediction uncertainty estimation for hate speech classification,” Statistical Language and Speech Processing, pp. 286–298, 2019.
 - [10] H. Faris, I. Aljarah, M. Habib, and P. A. Castillo, “Hate speech detection using word embedding and deep learning in the Arabic language context,” in Proc. 9th International Conference on Pattern Recognition Applications and Methods, 2020, pp. 453–460. doi: 10.5220/0008954004530460.
 - [11] A. Garain, “The titans at SemEval-2019 task 6: Offensive language identification, categorization and target identification,” in Proc. 13th International Workshop on Semantic Evaluation, 2019, pp. 759–762.
 - [12] B. Jang, M. Kim, G. Harerimana, S. Kang, and J. W. Kim, “Bi-LSTM model to increase accuracy in text classification: Combining Word2Vec CNN and attention mechanism,” Appl. Sci., vol. 10, no. 17, 5841, 2020.
 - [13] N. Jin, J. Wu, X. Ma, K. Yan, and Y. Mo, “Multi-task learning model based on multi-scale CNN and LSTM for sentiment classification,” IEEE Access, vol. 8, pp. 77060–77072, 2020. doi: 10.1109/ACCESS.2020.2989428.
 - [14] F. E. Ayo, O. Folorunso, F. T. Ibharalu, and I. A. Osinuga, “Machine learning techniques for hate speech classification of Twitter data: State-of-the-art, future challenges and research directions,” Comput. Sci. Rev., vol. 38, 100311, 2020. doi: 10.1016/j.cosrev.2020.100311.
 - [15] S. Kumar, C. Akhilesh, K. Vijay, and B. Semwal, “A multi-branch CNN-BiLSTM model for human activity recognition using wearable sensor data,” Vis. Comput., vol. 38, no. 12, pp. 4095–4109, 2021. doi: 10.1007/s00371-021-02283-3.
 - [16] M. U. Salur and I. Aydin, “A novel hybrid deep learning model for sentiment classification,” IEEE Access, vol. 8, pp. 58080–58093, 2020. doi: 10.1109/ACCESS.2020.2982538.
 - [17] U. Naqvi, A. Majid, and S. A. L. I. Abbas, “UTSA: Urdu text sentiment analysis using deep learning methods,” IEEE Access, vol. 9, pp. 114085–114094, 2021. doi: 10.1109/ACCESS.2021.3104308.
 - [18] J. Gaglani, Y. Gandhi, S. Gogate, and A. Halbe, “Unsupervised WhatsApp fake news detection using semantic search,” in Proc. International Conference on Intelligent Computing and Control Systems, 2020, pp. 285–289. doi: 10.1109/ICICCS48265.2020.9120902.
 - [19] H. T. Assaggaf, “A discursive and pragmatic analysis of WhatsApp text-based status notifications,” Arab World English J., vol. 10, no. 4, pp. 101–111, 2019. doi: 10.24093/awej/vol10no4.8.
 - [20] Y. Zhou, Q. Zhang, D. Wang, and X. Gu, “Text sentiment analysis based on a new hybrid network model,” Comput. Intell. Neurosci., vol. 2022, pp. 1–15, 2022.
 - [21] B. S. Rintyarna, R. Sarno, and C. Fatichah, “Evaluating the performance of sentence level features and domain sensitive features of product reviews on supervised sentiment analysis tasks,” J. Big Data, vol. 6, no. 1, 2019. doi: 10.1186/s40537-019-0246-8.
 - [22] H. Aljuaid, R. Iftikhar, S. Ahmad, M. Asif, and M. Tanvir Afzal, “Important citation identification using sentiment analysis of in-text citations,” Telemat. Informatics, vol. 56, 101492, 2021. doi: 10.1016/j.tele.2020.101492.
 - [23] N. Chintalapudi, G. Battineni, M. Di Canio, G. G. Sagaro, and F. Amenta, “Text mining with sentiment analysis on seafarers’ medical documents,” Int. J. Inf. Manag. Data Insights, vol. 1, no. 1, 100005, 2021. doi: 10.1016/j.jjime.2020.100005.
 - [24] A. U. Rehman, A.