

Evaluating BERT and Neural Network Approaches for Tweet Sentiment Analysis

Amit Kumar Yadav¹, Mukta Bhatele², Akhilesh A. Wao³

^{1,2,3}Department of Computer Science and Engineering, AKS University, Satna, India
¹amitkyadav081@gmail.com, ²230.muktabhatele@gmail.com, ³3akhileshwao@gmail.com

Abstract: Deep Learning is being explored as a prominent tool in machine learning in the current era by researchers to a great extent. A lot of work being done in this area is opening avenues for consumers. Utilization and growth of Artificial Intelligence have leveraged the power of deep learning for its applicability. Since its inception, a lot of algorithms have been proposed and implemented in the area of deep learning. In this paper, BERT (Bidirectional Encoder Representations from Transformers) a model of Google, LSTM, DistilBERT and CNN models have been discussed. Deep learning models have a major challenge of low performance, which restricts their use in real-time applications. In this work, BERT, LSTM, DistilBERT and CNN have been implemented separately using the same dataset of Twitter to perform sentiment analysis and compared their performance and accuracies as a first step to make a new model that will not only provide good accuracy but will also have better accuracy in future.

Keywords: Deep Learning, BERT, LSTM, DistilBERT, CNN, Machine Learning, Sentiment Analysis

1. Introduction

Sentiment Analysis has become a necessity for the human community these days and the availability of data through social media platforms has made it more important. Users around the world are expressing their sentiments, thoughts and views through various media, viz. text, images, audio, video, animations, etc., which have been found to be very informative to understand the inclination of users to achieve various advantages. For analysis, the availability of technology in hardware, software and algorithmic ways have made it suitable for researchers to go ahead and achieve more and more accuracies in results with high performance and security. To process the huge amount of data, machine learning and a step ahead deep learning algorithms have been proven to be a boon.

Pre-processing is needed for raw data before applying any algorithms, which has imposed the requirements of data cleaning, stop words removal and formatting of data-related algorithms. This work is focused on comparing and using the various models available for data pre-processing and decision-making steps. Also, machine learning requires training and testing data for preparing the environments for conclusive systems.

This paper is organized to elaborate on Sentiment Analysis in the second section, and about Natural Language Processing in the third section. Fourth, fifth, sixth and seventh sections provide the details of the chosen models BERT, LSTM, DistilBERT and CNN. The eighth section provides the existing works separately literature by the various researchers in the recent past. A brief description of the proposed system, data set used, results obtained, conclusion, and future possibilities have been iterated in subsequent sections. Papers used for this work have been enlisted at the end.

2. Sentiment analysis

Sentiment Analysis is a way of predicting the inner tone of expressions through text messages, whether it is positive, negative or neutral. Sentiment Analysis has been proven to be a great help in different areas of human society,



e.g. health, marketing, behavioural, security and protection of the senders. There are different sources of the collection of message data in different media formats available in the current era. Media formats such as text, images, audio, video, animation etc. are available through emails, social media platforms, customer support chats, reviews and many other repositories. The messages are cut into chunks and analysed to get useful information from them in Sentiment Analysis. For analysis, there are several different algorithms available, and new ones are created that are not only more and more accurate but are efficient as well. Sentiment Analysis is becoming a most important tool for companies to keep their objectives intact, improve the quality of their products and services, and provide real-time solutions to their customers. By the proper use of Sentiment Analysis, companies further can improve their brand monitoring, perform market research, track campaign performances, and improve customer service.

3. Natural Language Processing

Computers can now understand, analyze, and alter human language thanks to a machine learning technique called natural language processing. Big amounts of text and speech data are available to organizations today from a variety of communication channels, including emails, text messages, social media newsfeeds, audio, video, and more. Natural language processing software is used to automatically assess this data, evaluate the message's sentiment or intent, and react to human conversation in real time.

Analysing text and speech data thoroughly and effectively requires natural language processing. It can resolve dialectal discrepancies, slang, and grammatical errors that are common in casual talks.

Among the many automated tasks that businesses utilize it for are: processing, analysing, and archiving huge documents; analysing call centre records or customer feedback; and implementing chatbots for automated customer support; responding to the who, what, when, and where inquiries; and identifying and extracting text.

Natural language processing can also be used in programs that interact with customers to improve customer service. A chatbot, for instance, can automatically answer frequently asked questions and route more complicated ones to customer service after analysing and classifying user inquiries. In addition to saving agents time on pointless inquiries, this automation lowers expenses and raises customer satisfaction.

4. Bert Model

BERT stands for Bi-directional Encoder Representation for Transformers, which was introduced by Google for performing algorithms related to natural language processing. It is mainly designed and developed for computer systems to reveal the context of the messages using the surrounding words and sentences. BERT Model works in bi-directional, meaning it considers forward and backwards words, unlike other models, which only work in one direction. The speciality of the BERT model is that it works on the transformer model of Google, which allows for bi-directionality. BERT, with its specialities, is applicable for large data processing, Sentiment Analysis, Text Classification, finding ambiguities arising due to multiple meanings of the words, labelling of roles based on semantics, natural language processing and many others. BERT leverages its power through pre-training with unsupervised learning methods from unlabelled text data.

Since BERT is open source, many other companies are contributing towards its applications in different areas of supervised learning, machine learning and artificial intelligence. With this, application-specific BERT variations have been developed e.g. DocBERT, BioBERT, PatentBERT, VideoBERT etc.

BERT has many aspects associated with it for its advantages:

A. Pre-training: Bert uses plain text which is unlabelled corpus to learn and applies feedback using unsupervised learning methods which are useful in mechanism such as search engine optimization.

B. Transformers: Transformer processes in bi-directions in such a way that it sets the relationship of each word with all other words in the text and hence it better understands the context of the sentence. It contrasted with all other traditional approaches by applying vector mapping of the word with other words in the sentences.

C. Masked Language Modeling: In BERT every word is defined by other words in its surrounding than the prefixed identity of itself unlike other languages and hence it provides better accuracies in contextual decisions.

D. Self-attention mechanism: BERT uses bi-directionality to include the augmentation effect of the other words in sentence to fix the meaning of any single word, which changes with the inclusion of more and more words. It never leads to a fixed meaning in forward only direction.

In this paper, Plain BERT is applied for Sentiment Analysis and is compared with the old model CNN. Further, BERT will be used in hybrid mode in future to test for better accuracy and performance.

5. LSTM Model

LSTM – Long Short-Term Memory, is a deep learning algorithm that works on sequential data. It also maintains the memory efficiently with the least usage in respect to other algorithms. LSTM has a variation named Bi-LSTM, which can work on bidirectional processing of sequential data to handle vital information in both directions. LSTM has various cells which work as memory elements and gates to handle the relevant information retained and used, or irrelevant information is removed. This way, LSTM reduces the memory requirements for large data. There are three different types of gates, viz. Input Gate, Forget Gate, and Output Gate. The input gate sends the useful information to the cells to remember it until it is used. Forget Gate applies a weighted mechanism to evaluate the non-useful information from the cells and removes them if they become worthless. Output Gate takes the information from the cell and forwards the same to output after removing it from the cell.

6. DistilBERT Model

DistilBERT applies a distillation process for reducing the size of the model to make it a cheap, faster, and smaller implementation of BERT. It is approximately 40% size of the original BERT and still have 97% of the understanding capability of BERT. It is approximately 60% faster than the original BERT model. It performs the knowledge distillation process during the pre-training stage from any already pre-trained larger model of BERT.

7. Cnn Model

CNN – Convolution Neural Network is an algorithm for processing visual data, mainly images. It is an artificial neural network, which is also a deep learning model made by mapping the way human eyes process images. Convolution Neural Network works using different layers, where each layer carries a different responsibility. Each intermediate layer till output layers takes its input from the previous layer whereas the first layer takes input data, which is a labelled and cleaned dataset. Therefore, Convolution Neural Network is termed as forward only. Main layers of the Convolution Neural Network are:

A. Convolution Layers: These layers take the responsibility of retrieving data (features) from the input data. They apply filtering for feature extraction. Extracted features are supplied to next set of layers for further processing.

B. Pooling Layers: Extracted features taken as input from convolution layers are too many and therefore they are reduced by applying the weights and weight adjustment. This is also known as reducing the spatial dimensions of the features.

C. Fully Connected Layers: These are mainly used for connecting filtered and reduced feature set with the output layers.

Due to weighing and weight adjustments Convolution Neural Network manages to show spatial invariance resulting in correct recognition of objects, this is not affected by the position or orientation of the objects. As Convolution Neural Network is forward only and therefore rely on hierarchical arrangements of the features from low level to high level, which makes it suitable for achieving the impressive results. For better results and accuracies

Convolution Neural Network required labelled and cleaned data set of vast dataset and hence needs very high power processing environment.

For processing of the text data, Convolution Neural Network applies a numerical value to the words and creates numerical matrix for the same. After creation of numerical matrix, different layers viz. convolution, pooling, and fully connected, are used to process the large set of data to generate the weighted outputs. Due to use to large input dataset Convolution Neural Network provides better accuracies then the other different classifiers such as SVM and Naïve Bayes.

8. Related Works

The text classification task, which is utilized in a variety of fields, including academia, social media, and medicine, is one of the subjects that receives the most research among studies on text mining. Sentiment analysis has been extensively studied as a subproblem of text classification to classify frequently opinion-based textual elements. Particularly, sentiment analysis efforts have relied on product or service user reviews and experiential feedback as fundamental data sources. Since the beginning of the 2000s, social media platforms like Twitter, Facebook, and Reddit have emerged as essential platforms for the exchange of opinions. In this way, in this work, we create a variety of machine-learning models to solve the problem of sentiment analysis on the Reddit comments dataset. Within intervals of 73–76 percent, the experimental models we constructed achieve F1 scores. As a result, we compare and contrast performance scores available from conventional machine learning models and deep learning models. [1]

Businesses and individuals incur annual losses in the millions of dollars caused by unsolicited emails such as phishing and spam. There have been a number of models and methods developed to detect spam emails automatically, but none of them have demonstrated 100% predictive accuracy. All others models that were proposed by many researchers are outperformed by the Machine learning and deep learning algorithms. The accuracy of the models was enhanced by natural language processing. Word embedding's efficacy in classification of spam emails is discussed in this work. The task of distinguishing non-spam emails and spam emails is carried out by fine-tuning the pre-trained transformer model Bidirectional Encoder Representations from Transformers. BERT puts the text's context into perspective by employing attention layers. A baseline DNN model with two stacked Dense layers and a BiLSTM (bidirectional Long Short Term Memory) layer is used to compare the findings to the baseline model. In addition, the results are compared to a set of traditional classifiers, including NB (Naive Bayes) and k-NN (k-nearest neighbors). The model is trained using one of two open-source data sets, and its persistence and robustness against unknown data are tested using the other. The proposed method had the highest F1 score and accuracy of 98.66%. [2]

Sentiment analysis using natural language processing is a difficult task, especially for social media texts that are too much informal, brief, and noisy. A comparative study of sentiment analysis of social media texts using deep learning models is presented in this paper. Based on deep neural networks, following three models have been developed: a CNN with layers of long short-term memory (CNN-LSTM), and a bidirectional LSTM with CNN layers (BiLSTM-CNN). As vector representations of words, we make use of the word embeddings GloVe and Word2vec. On two datasets, we assess how well the models work: IMDb Film Surveys and Twitter Opinion 140. As a baseline, we also compare the outcomes with a logistic regression classifier. On the IMDb dataset, the experimental results show that the BiLSTM-CNN model has the highest accuracy of 82.1%, while the CNN model has the highest accuracy of 90.1%. The proposed models are suitable for use in sentiment analysis of social media texts and are comparable to current models. [3]

Twitter is a vast archive and a treasure trove of human thoughts that convey a person's immediate emotions. A person's feelings can be better understood by looking back at tweets they posted during the COVID-19 pandemic. It is difficult to train a model to understand the exact feeling with a large amount of data. The field of deep learning has seen numerous advancements that point the way forward. Convolutional Neural Networks, attention-based Bidirectional LSTM, and Google BERT are all used to compare tweet sentiment classification. The embeddings are once more fine-tuned before the final models are trained on the Twitter dataset SemEval-2016. In contrast to machine learning methods, these models proved to be extremely effective and accurate in the study of emotions. [4]

Numerous natural language processing tasks rely on sentiment analysis of text content. Particularly in light of the growth of social media, sentiment analysis is essential for extracting meaningful information from Internet big data. We are interested in using deep learning models to handle the task of sentiment analysis because of the successes of deep learning. A framework we call word2vec + Convolutional Neural Network is what we propose in this paper. To begin, we will calculate vector representations of words using the Google-proposed word2vec, which will serve as the CNN's input. word2vec is used to get a vector representation of a word and to show how far apart words are. That will result in CNN initializing the parameters at a favorable point, which can effectively enhance the nets' performance in this issue. Second, we create a CNN architecture that is appropriate for the task of sentiment analysis. This architecture makes use of pooling layers and three pairs of convolutional layers. Word2vec and CNN are used for the first time to apply a 7-layer architecture model to sentiment analysis of sentences, according to our knowledge. In addition, we use Normalization, Dropout, and the Parametric Rectified Linear Unit (PReLU) to increase our model's generalizability and accuracy. In a public dataset, which is the corpus of movie review excerpts with five labels, we test our framework: neutral, negative, somewhat negative, positive, and negative in this dataset, our network outperforms other neural network models like the Recursive Neural Network and the Matrix-Vector Recursive Neural Network (MV-RNN) in terms of test accuracy. [5]

Spread of phony news on Twitter is a quickly developing issue, generally because of the rising number of bots. As a result, research into automatic bot detection is growing in importance. The BERT-based bot detection model and exploratory data analysis of tweets written by humans and bots are both presented in this work. We statistically demonstrate that contextualized embeddings and additional features improve model performance. In addition, we compare the difficulty of the two tasks and create a gender prediction model using derived features. Last but not least, we show that on both tasks, Logistic Regression outperforms Deep Neural Network. [6]

One of the most difficult issues in automated language understanding is emotion detection. Text without facial expression is considered difficult to comprehend human emotions. The machine learning community has been motivated recently by the development of a machine that can differentiate between emotions and comprehend the context of sentences. Using techniques from deep learning, we propose a method for identifying emotions. GloVe word embeddings, BERT Embeddings, and a collection of psycholinguistic features (such as those from the AffectiveTweets Weka-package) constitute the system's primary input. The proposed system, EmoDet2, combines a BiLSTM neural network with a neural network architecture with fully connected is used to achieve performance results that significantly outperform the baseline model provided by Semeval-2019 / Task-3 organizers (F1-score 0.58). [7]

Sentiment analysis using natural language processing is a method for recognizing and classifying viewpoints expressed in written or spoken language. A significant amount of emphasis has been placed on sentiment analysis from Twitter's Tweet data as a result of the widespread use of Twitter as a medium for the expression of opinions and feelings. A preexisting deep learning model known as BERT has demonstrated exceptional proficiency in tasks related to Natural Language Processing, such as sentiment analysis. The purpose of this investigation is to compare how well three distinct BERT models, namely RoBERTa, RoBERT, and fine-tuned BERT-base-multilingual-uncased, perform Twitter data sentiment analysis. Using a manually annotated dataset of 1,578,627 tweets, the models were graded based on their F1-score, recall, accuracy, and precision. With an accuracy score of 83.23%, the experimental results show that the RoBERTa model performs better than the other two models. The findings aid academics and professionals in selecting the best BERT model for their sentiment analysis task and provide valuable perspectives on the efficacy of BERT models when analyzing Twitter data based on how people feel. The study emphasizes that the evaluation metric of interest should dictate the choice of the BERT model for Twitter data sentiment analysis. This research sheds light on the suitability of BERT models for sentiment analysis of Twitter data for use in social media monitoring and analysis. [8]

In this study, an integrated sentiment analysis system specialized for movie reviews is developed. The proposed implementation utilizes LSTM-based sentiment analysis and simplifies model training using a pre-trained LSTM model, resulting in robust sentiment prediction capabilities. The architecture emphasizes abstraction,

separates front-end and back-end components, optimizes data flow with JSON format and pickle files, and provides a user-friendly interface. This study reveals the successful synergy of DL techniques, architectural design, and abstraction principles by highlighting the role of abstraction in improving user experience and system performance. The LSTM model achieved a test accuracy of 87.53%. [9]

To improve customer experience of a product, analyzing customer reviews is the best solution as it can identify the positive and negative opinions of customers. In our study, we developed two different supervised models for natural language tasks using (1) a deep neural long short-term memory model and (2) the advanced pre-trained Distil Bidirectional Encoder Representations from Transformers. In the first model, we analyzed how individual words in a sentence contribute to the sentiment of that sentence by training LSTM on the 50,000 Amazon review corpus publicly available on the Amazon website. In the second model, we explored how positional embeddings of words in a sentence and multi-head self-attention contribute to faster and better sentence representations for sentiment analysis by fine-tuning DistilBERT, which contains 66 million parameters and 6 encoder layers. The main contributions of our study include (1) a comparative analysis of the effectiveness of the above models, where the Transformer-based model outperforms the LSTM model in all evaluation metrics. (2) DistilBERT, a model for developing a faster and lighter sentiment analyzer. [10]

9. Proposed System

In this work, a comparative study of BERT, LSTM, DistilBERT and CNN is being done by applying them for sentiment analysis of the same dataset [11]. The work is intending to evaluate the accuracy and performance of the algorithms with the same environment. This will lead to apply the BERT, LSTM, DistilBERT and CNN[12] together to achieve better accuracy and performance in next step. Application of BERT, LSTM, DistilBERT and CNN will be evaluated using various parameters and adjustment of attributes such as training and testing split, block size, iteration count etc. of algorithms. The complete process is done in four steps:

1. BERT Processing
2. LSTM Processing
3. DistilBERT Processing
4. CNN Processing

From the proposed work and same input dataset, both the models shall be used to generate the results viz. accuracy and performance. These will be compared and shall be used to decide the weight of their application in next step.

10. Dataset

For implementation and testing of the proposed a dataset of twitter has been used taken from internet which contains 1,600,000 tweets extracted using the twitter api . The tweets have been annotated (0 = negative, 2 = neutral, 4 = positive) and they can be used to detect sentiment.

It contains target as polarity of the tweet, tweet id, date of tweet, a flag value, sender of the tweet and tweet message.

The dataset has been prepared at Stanford University by a team of researchers who iterated about the dataset: [11]

"Our approach was unique because our training data was automatically created, as opposed to having humans manually annotate tweets. In our approach, we assume that any tweet with positive emoticons, like :), was positive, and tweets with negative emotions, like :(, were negative. We used the Twitter Search API to collect these tweets by using keyword search" [3]

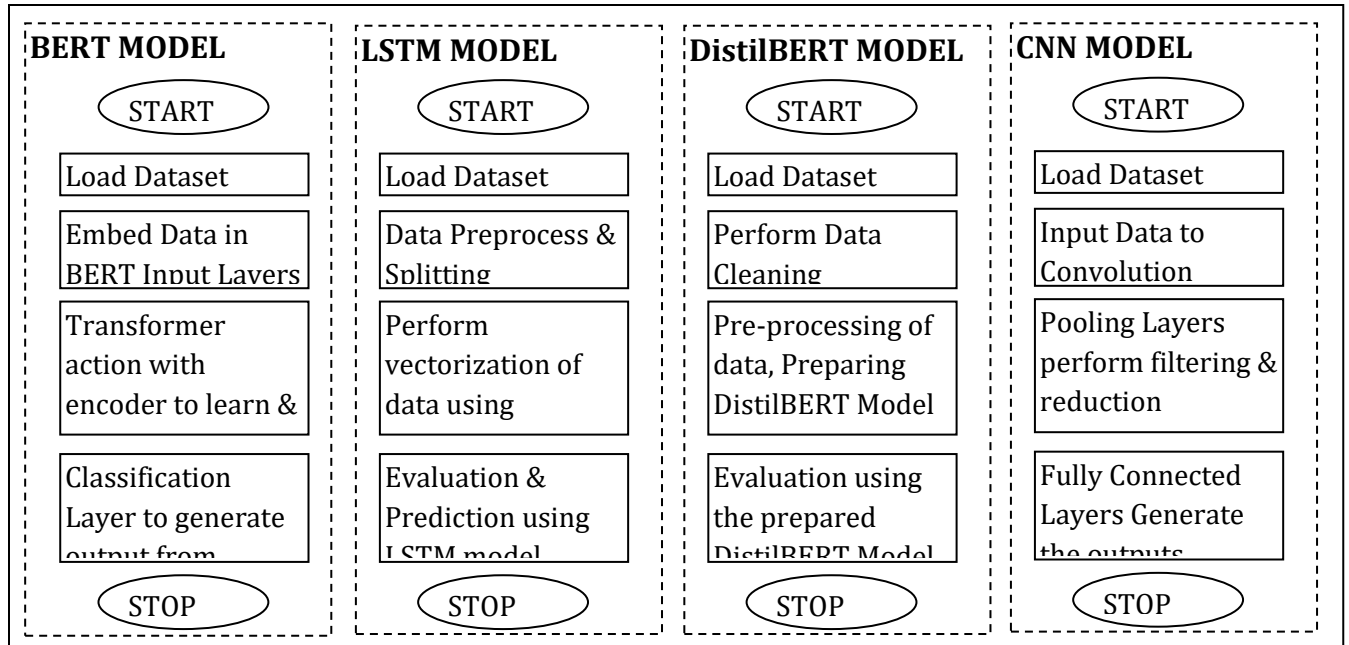


Figure 1: Flow Chart of the Proposed Models

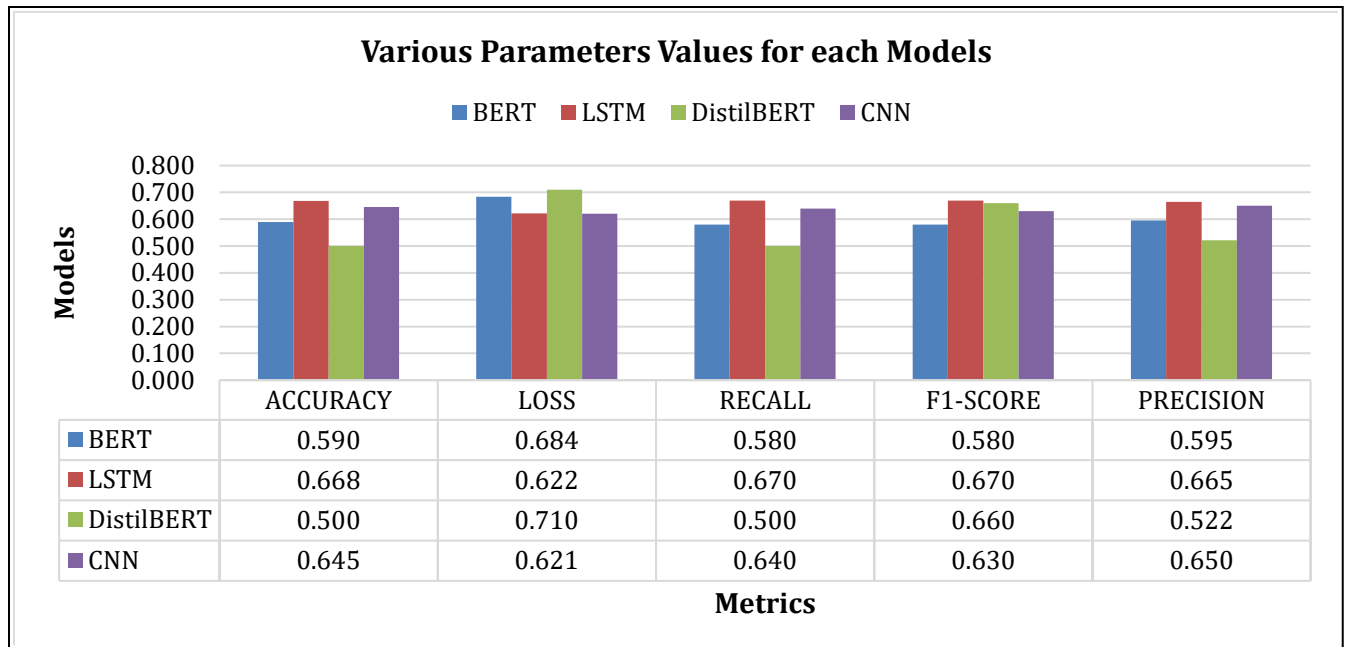
11. Results Obtained

Dataset as specified above [11] has been used to test the implementation of BERT, LSTM, DistilBERT, and CNN Models. The parameters have been set similarly for executing all of the algorithms and have been executed in the same environment so that accuracy, recall, precision, loss and F1-measure of algorithms can be evaluated. The reading obtained from the execution of the algorithms have been listed in Table 1 and shown in graphs in Figures 2 and 3 below.

Table 1: Readings of the various metrics obtained from execution of the BERT, LSTM, DistilBERT and CNN Deep Learning Models using the same dataset

MODEL	ACCURACY	LOSS	RECALL	F1-SCORE	PRECISION
BERT	0.590	0.684	0.580	0.580	0.595
LSTM	0.668	0.622	0.670	0.670	0.665
DistilBERT	0.500	1.420	0.500	0.366	0.522
CNN	0.645	0.621	0.640	0.630	0.650

Figure 2: Graph showing the values of the various metric results obtained using BERT, LSTM, DistilBERT and CNN Deep Learning Models



From the tables and graphs shown above, it is clear that in the same environment, LSTM outperforms all other models in respect of accuracy, F1-score, precision, recall, and has comparatively less validation loss values. CNN provides least validation loss and stands second in accuracy, recall and precision. BERT has least values in most of the metrics but has slightly better in respect of validation loss. CNN and DistilBERT perform moderately. Overall, all of the models are competitive in their performances and show better values for different metrics. From the results obtained, it can be concluded that all of the models have comparative values and can be applied together to overcome the deficiency of others. This conclusion leads us to our next step of combining them for the same dataset to generate quick and better results in the given situation.

Table 2: Graph showing the better accuracy of results obtained using BERT, LSTM, DistilBERT and CNN.

MODEL	ACCURACY
BERT	0.590
LSTM	0.668
DistilBERT	0.500
CNN	0.645

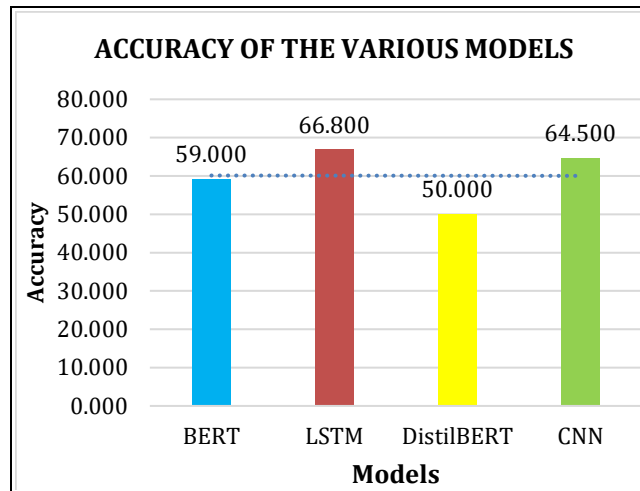


Figure 3: Graph showing the values of the various metric results obtained using BERT, LSTM, DistilBERT and CNN Deep Learning Models

The above figure depicts that accuracy of the LSTM is higher than the other models used in the case study. Although the DistilBERT is showing least accuracy the variation between the accuracies of all the models is not drastically varied. Hence all of the models can be used together to generate better results in our next research.

From the above two diagrams it can be concluded that when multiple models are applied in the similar situation on similar dataset the variation of metric values do not vary too much and as per their characteristics respectively. This leads and allows us to use these models in combination to generate the results which will be better and will provide high performance in future.

12. Conclusion

In this work, the focus is on evaluating and comparing the models for finding the possibilities of applying them together in future work. Out of the many different choices available, BERT LSTM, DistilBERT and CNN models have been chosen. Selection of BERT is based on its bi-directionality, LSMT due to its least memory utilization, DistillBERT due to its cheap and faster processing, whereas CNN has been chosen for its simplicity. The models are compared for their accuracies, recalls, precisions, F1-scores and validation losses on the same dataset of Twitter available online using the same hardware and software environments. The parameters of all have been kept the same (viz. training and testing ratio, block size, optimizer etc.) for making it more acceptable and technically comparable. The results obtained are explicit as all of them have moderate metric values and can work together to adjust deficiency of others. The performance of the models shall be leveraged in future work. Validation losses in all the models are almost similar, and hence they are acceptable in the given situations.

13. Future Work

This work is intended to find a set of matching models for sentiment analysis, which requires majorly to achieve high accuracy. The goal is to apply them together in future if they have comparable accuracies and performances in similar environments. The output of this work will be leveraged in future to use all the models together to achieve better accuracies with high performance. Other models can also be tested in future to include in collective functionalities.

REFERENCES

1. Erkantarci, B., Bakal, G. An empirical study of sentiment analysis utilizing machine learning and deep learning algorithms. J Comput Soc Sc 7, 241–257 (2024). <https://doi.org/10.1007/s42001-023-00236-5>

2. AbdulNabi, Isra'a & Yaseen, Qussai. (2021). Spam Email Detection Using Deep Learning Techniques. *Procedia Computer Science*. 184. 853-858. 10.1016/j.procs.2021.03.107.
3. Derbentsev, Vasily & Bezkorovainyi, Vitalii & Matviychuk, Andriy & Pomazun, Oksana & Hrabariev, Andrii & Hostryk, Alexey. (2023). A comparative study of deep learning models for sentiment analysis of social media texts. 3465. 168-188.
4. A. Roy and M. Ojha, "Twitter sentiment analysis using deep learning models," 2020 IEEE 17th India Council International Conference (INDICON), New Delhi, India, 2020, pp. 1-6, doi: 10.1109/INDICON49873.2020.9342279.
5. X. Ouyang, P. Zhou, C. H. Li and L. Liu, "Sentiment Analysis Using Convolutional Neural Network," 2015 IEEE International Conference on Computer and Information Technology; Ubiquitous Computing and Communications; Dependable, Autonomic and Secure Computing; Pervasive Intelligence and Computing, Liverpool, UK, 2015, pp. 2359-2364, doi: 10.1109/CIT/IUCC/DASC/PICOM.2015.349.
6. D. Dukić, D. Keča and D. Stipić, "Are You Human? Detecting Bots on Twitter Using BERT," 2020 IEEE 7th International Conference on Data Science and Advanced Analytics (DSAA), Sydney, NSW, Australia, 2020, pp. 631-636, doi: 10.1109/DSAA49011.2020.00089.
7. H. Al-Omari, M. A. Abdullah and S. Shaikh, "EmoDet2: Emotion Detection in English Textual Dialogue using BERT and BiLSTM Models," 2020 11th International Conference on Information and Communication Systems (ICICS), Irbid, Jordan, 2020, pp. 226-232, doi: 10.1109/ICICS49469.2020.239539.
8. A. Sahoo, R. Chanda, N. Das and B. Sadhukhan, "Comparative Analysis of BERT Models for Sentiment Analysis on Twitter Data," 2023 9th International Conference on Smart Computing and Communications (ICSCC), Kochi, Kerala, India, 2023, pp. 658-663, doi: 10.1109/ICSCC59169.2023.10335061.
9. Srinivas, Akana & Satyanarayana, Ch & Divakar, Ch & Sirisha, Katikireddy. (2021). Sentiment Analysis using Neural Network and LSTM. *IOP Conference Series: Materials Science and Engineering*. 1074. 012007. 10.1088/1757-899X/1074/1/012007.
10. Maisha Maliha, Vishal Pramanik, "Vehicle Detection and Identification System: Convolutional Neural Network with Self-Attention", *2023 14th International Conference on Computing Communication and Networking Technologies (ICCCNT)*, pp.1-7, 2023.
11. Go, A., Bhayani, R. and Huang, L., 2009. Twitter sentiment classification using distant supervision. *CS224N Project Report, Stanford, 1(2009)*, p.12.
12. Bhartiya, Bhatele, Wao. (2025). "Efficient traffic flow prediction with CNN-based deep learning techniques". *International Conference on Advances and Application in Artificial Intelligence*, 5(1), 1153-1173.
13. Al Amrani, Y., Lazaar, M., & El Kadiri, K. E. (2018). Random forest and support vector machine based hybrid approach to sentiment analysis. *Procedia Computer Science*, 127, 511-520.
14. Arias, M., Arratia, A., & Xuriguera, R. (2014). Forecasting with twitter data. *ACM Transactions on Intelligent Systems and Technology (TIST)*, 5(1), 1-24.
15. Bakal, G., Talari, P., Kakani, E. V., et al. (2018). Exploiting semantic patterns over biomedical knowledge graphs for predicting treatment and causative relations. *Journal of Biomedical Informatics*, 82, 189-199.
16. Y. Kim, "Convolutional neural networks for sentence classification," arXiv preprint arXiv:1408.5882, 2014. 858 Isra'a AbdulNabi / *Procedia Computer Science* 184 (2021) 853-858 Abdulnabiisraa et al. /*Procedia Computer Science* 00 (2019) 000-000
17. P. Zhou, W. Shi, J. Tian, Z. Qi, B. Li, H. Hao, and B. Xu, "Attention-based bidirectional long short-term memory networks for relation classification," in *Proceedings of the 54th Annual Meeting of the Association for Computational Linguistics (Volume 2: Short Papers)*, 2016, pp. 207-212.
18. V. Soloviev, A. Matviychuk, V. Kobets, L. Kibalnyk, H. Danylchuk, A. Kiv (Eds.), *Proceedings of 10th International Conference on Monitoring, Modeling & Management of Emergent Economy - M3E2, INSTICC, SciTePress*, 2023, pp. 163-175. doi:10.5220/0011932300003432.
19. M. Azlinah, B. W. Yap, J. M. Zain, M. W. Berry (Eds.), *Soft Computing in Data Science 6th International Conference, SCDS 2021*, volume 1489 of *SCDS: International Conference on Soft Computing in Data Science*, Springer, Singapore, 2021. doi:10.1007/978-981-16-7334-4.
20. M. Silberztein, F. Atigui, E. Kornysheva, E. Métais, F. Meziane (Eds.), *Natural Language Processing and Information Systems: Proceedings of 23rd International Conference on Applications of Natural Language to Information Systems*, volume 10859 of *Lecture Notes in Computer Science*, Springer, Cham, 2018. doi:10.1007/978-3-319-91947-8.
21. M. R. Hasan, M. Maliha and M. Arifuzzaman, "Sentiment Analysis with NLP on Twitter Data," 2019 International Conference on Computer, Communication, Chemical, Materials and Electronic Engineering (IC4ME2), 2019, pp. 1-4, doi: 10.1109/IC4ME247184.2019.9036670.
22. M. T. Varun Dogra Assvknj and, "Analyzing DistilBERT for sentiment classification of banking financial news," in *Intelligent Computing and Innovation on Data Science*, S.-L. Peng, S.-Y. Hsieh, S. Gopalakrishnan, and Balaganesh, Eds., p. 582, Springer Singapore, Singapore, 2021.