

International Journal of System of Systems Engineering

ISSN online: 1748-068X - ISSN print: 1748-0671

<https://www.inderscience.com/ijsse>

Bidirectional LSTM and self-attention mechanisms based multi-Label sentiment analysis

A.R. Arunarani

DOI: [10.1504/IJSSE.2025.10059046](https://doi.org/10.1504/IJSSE.2025.10059046)

Article History:

Received:	31 May 2023
Last revised:	01 June 2023
Accepted:	07 June 2023
Published online:	21 February 2025

Bidirectional LSTM and self-attention mechanisms based multi-Label sentiment analysis

A.R. Arunarani

Department of Computational Intelligence,
School of Computing,
College of Engineering and Technology,
SRM Institute of Science and Technology,
Kattankulathur, Chennai, 603203, India
Email: arunaraghav83@gmail.com

Abstract: This study proposes the implementation of a novel optimisation-depend on deep learning algorithm for multi-label sentiment analysis (MLSA). The goal of the algorithm is to improve the accuracy of sentiment analysis, particularly in the context of e-commerce-related applications. This technique effectively categorise the text data into multiple sentiment classes, such as positive, negative, neutral, or other emotions, and to determine the overall sentiment expressed in a given text document. The challenge of MLSA on e-commerce data lies in the informal and often cryptic nature of the text, which can make sentiment analysis difficult. To address this, a novel optimisation-empowered bidirectional long-short term memory (Bi-LSTM) system with self-attention mechanisms is proposed in this research work. It uses the Bi-LSTM network to capture the sequential relationships between words in the manuscript and the self-attention mechanism to dynamically weigh the importance of different words in determining the overall sentiment expressed in the text.

Keywords: self-attention; bi-directional long short-term memory; multi-label sentimental analysis; deep learning; sentimental analysis; NLP; natural language processing.

Reference to this paper should be made as follows: Arunarani, A.R. (2025) 'Bidirectional LSTM and self-attention mechanisms based multi-Label sentiment analysis', *Int. J. System of Systems Engineering*, Vol. 15, No. 1, pp.67–78.

Biographical notes: A.R. Arunarani is the Assistant Professor, Department of Computational Intelligence, School of Computing, College of Engineering and Technology, SRM Institute of Science and Technology, Kattankulathur Campus, Chennai. Former Teaching Fellow at the Faculty of Department of Computer Science and Engineering, CEG Campus, Anna University, Chennai, India. She obtained PhD degree in Information and Communication Engineering from Anna University in 2018. She has nearly 10 years of teaching experience. Her research interests spans over cloud computing, social network analysis, information retrieval, data and text mining, machine learning, artificial intelligence and big data.

1 Introduction

Identifying the emotional state of a text is done by sentiment analysis, commonly referred to as opinion mining. It is a subfield of natural language processing (NLP) (Zhang et al., 2018) and containing a text analysis technique with algorithms to extract and categorise emotions, opinions, and a particular data through the given text. Evaluate if a text's emotion is good, negative, neutral, or a combination of these feelings by using sentiment evaluation (Jagdale et al., 2019). It plays an important role in understanding and analysing public opinions and emotions and can provide valuable insights for organisations and individuals. Sentiment analysis is used in various applications such as e-commerce applications, social media monitoring, customer feedback analysis, and product review analysis. The goal of sentiment analysis is typically done by examining the words, phrases, and other features of the text, and using algorithms and machine learning models to make predictions about the sentiment expressed in the text (Deng et al., 2021).

The process contains some stages, including text pre-processing, feature extraction, and classification. Sentiments expressed in the text can be ambiguous and subjective, making it difficult for algorithms to accurately categorise them as positive, negative, or neutral. To address these limitations, multi-label sentiment analysis (MLSA) has been developed, which allows for the categorisation of text into multiple sentiment labels (Priyadarshini and Cotton, 2021). It can deliver an additional nuanced and comprehensive understanding of the sentiments articulated in manuscript and can be especially useful in applications where the expression of multiple and complex emotions is relevant. Multi-label sentiment analysis can be used to categorise text data into a wide range of sentiment labels, such as happy, sad, angry, excited, scared, and more (Liu and Guo, 2019).

The labels used in MLSA can vary depending on the detailed request with the category of information being analysed. Multi-label sentiment analysis typically involves using machine learning algorithms, such as decision trees, random forests, and neural networks, to learn from annotated training data and make predictions about the sentiment categories for new, unseen data (Rehman et al., 2019). Overall, MLSA is a useful tool for organisations and individuals to gain a more nuanced and comprehensive understanding of the attitudes and opinions expressed in text data. Hence in this research work, a combination of novel deep learning and optimisation algorithm is introduced to attain a better performance of multi-label sentimental analysis (Ruz et al., 2020). The contribution of the proposed model is illustrated as follows,

- The contribution of SentiWordNet in feature extraction for sentiment analysis is significant as it provides a rich source of information about the sentiment of words.
- The sentiment scores assigned by SentiWordNet can be used as features in deep learning algorithms to train models for sentiment analysis. This can lead to improved performance compared to using only the raw text as features.
- Pre-processing is a crucial step in sentiment analysis as it prepares the raw text data for further analysis and improves the performance of the sentiment analysis algorithms.

- A combination of Bi-Directional LSTMs and self-attention can be used for MLSA by combining the ability of LSTMs to capture contextual information with the ability of self-attention to assign different weights to different parts of the sequence. This can lead to improved performance in MLSA as the model can capture the context and importance of each word in the sequence.

The manuscript is organised as follows; Section 2 discusses the relevant literature. Section 3 details the proposed methodology and its implementation. Section 4 presents the results and discusses the implications. Finally, Section 5 offers the conclusion and future scope.

2 Literature review

Sentiment analysis has numerous applications in various domains, including marketing, social media analysis, customer service, and political science. The research community has recently shown a growing interest in creating novel methods and strategies for sentiment analysis. Some of the recent works are taken for analysis in our study.

2.1 Sentiment analysis on product reviews

Sentiment analysis can be used to analyse product reviews on e-commerce platforms, such as Amazon or eBay, to determine the overall sentiment of a product. This can help businesses to identify customer satisfaction levels, identify product issues, and make improvements to their products. Sivakumar and Uyyala (2021) suggested an intelligent system employing long-term short-memory with fuzzy logic. The proposed approach was tested using benchmark datasets from Amazon's evaluations of cell phones, video games, and consumer products.

2.2 Customer service-based sentimental analysis

Cross-lingual sentiment analysis was studied by Yilmaz et al. (2021) who also offer a creative framework for this task in a multi-label environment. But the model often struggles to correctly identify sentiments when the context of the text is not taken into consideration. Similarly, Jain et al. (2022) researches are on a strategy for making predictive suggestions for travel and tourism, notably for airlines. The authors used numerous fundamental classification models, including K-nearest neighbours, support vector machine, multi-layer perceptron, logistic regression, random forest, and ensemble learning, to construct their method as a multi-label classification system.

2.3 Sentimental analysis on social media

Sentiment analysis can be used to analyse social media posts related to e-commerce products or brands. This can help businesses to monitor and track brand reputation, identify customer needs and preferences, and improve their marketing strategies. A residual attention-based deep learning network (RA-DLNet) for visual sentiment analysis is suggested in the research by Yadav and Vishwakarma (2020). The HEMOS (Humor-EMOji-Slang-based) method for fine-grained sentiment categorisation in the

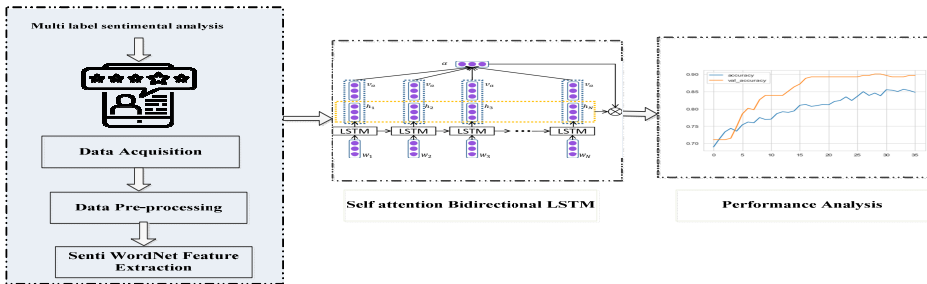
Chinese language was presented in the work by Li et al. (2020). The goal of Alsayat et al study's (Alsayat, 2022) was to improve sentiment classification performance using a tailored deep learning model that combined an advanced word embedding technique with LSTM network. Ombabi et al. (2020) developed a deep-learning method for Arabic sentiment analysis. The authors suggest a model that incorporates a convolutional neural network (CNN) for extracting local features, an LSTM network for keeping track of long-term dependencies, and an SVM classifier for producing the final classification.

From the overall analysis, it is observed that MLSA is challenging because it requires the model to be able to handle multiple labels and identify their relationships with each other. The accuracy of the model is crucial in this problem statement, as incorrect labelling of sentiments can have serious implications in various applications. Thus the existing research studies often struggle to accurately identify complex emotions because some words can express both positive and negative sentiments simultaneously. There is a need for research that addresses this issue and develops models that can handle complex emotions more effectively Overall, there is a need for further research in MLSA to address these research gaps and develop more effective and robust models that can handle the complexity of sentiments expressed in natural language. Current research work in MLSA aims to address the drawbacks of existing models by developing new, more effective models that can better handle the complexity of sentiments expressed in natural language.

3 Proposed methodology

Multi-label sentiment analysis is a rapidly growing field that holds great promise for improving sentiment analysis and understanding the complex relationships between different sentiments expressed in text data. The main challenge of MLSA is to accurately identify and classify the multiple sentiments expressed in a single text document, as opposed to traditional sentiment analysis, which typically focuses on a single sentiment classification. MLSA algorithms use a deep algorithm to analyse and understand the relationships between the different sentiments expressed in a text document. The layout of the proposed methodology is displayed in Figure 1.

Figure 1 Layout of the proposed methodology (see online version for colours)



Data pre-processing is an important step in sentiment analysis to prepare the data for analysis and improve the accuracy of the model. This may involve various techniques such as removing irrelevant information, handling missing data, correcting spelling and grammatical errors, tokenisation, stop word removal, stemming or lemmatisation, and

normalisation of the text data. By pre-processing the data, we can reduce noise and irrelevant information, simplify the input for analysis, and ensure that the model focuses on the most important features that determine the sentiment of the text. This can result in better accuracy and more meaningful insights from the analysis. SentiWordNet is a lexical resource that can be used for feature extraction in sentiment analysis. It is a lexicon that assigns sentiment scores to words based on their semantic orientation, allowing the sentiment of the text to be determined based on the sentiment scores of the words it contains. By using SentiWordNet to extract features from the text data, we can identify the sentiment of the words and use this information to determine the overall sentiment of the text. This can improve the accuracy and reliability of the sentiment analysis model by providing a more nuanced understanding of the sentiment expressed in the text. However, it's worth noting that SentiWordNet may not always provide accurate sentiment scores for all words, especially in cases where the context of the text is complex or ambiguous. Therefore, it is important to consider further the self-attention model in this research work. Bidirectional long short-term memory (LSTM) can be used in MLSA to capture contextual dependencies in both forward and backward directions, allowing the model to better understand the meaning and sentiment of the input text. This can improve the model's ability to recognise and classify multiple labels or categories of sentiment in the text. Along with this, the self-attention mechanism works by computing the attention scores for each word in the text based on its relationship with the other words and then using these scores to dynamically adjust the weights of the words in the final sentiment representation. This allows the network to focus on the most important words in the text and make a more accurate prediction of the overall sentiment.

3.1 Data acquisition

Data acquisition is a crucial aspect of sentimental analysis. The accuracy and effectiveness of a sentiment analysis model depend largely on the quality and quantity of data used for training. In this research work, e-commerce data is taken for analysis to understand people's perspectives about a wide range of products. E-commerce data is a popular source of data for sentimental analysis due to the large volume of user-generated content and the wide range of emotions and sentiments expressed in reviews on online platforms.

3.2 Data pre-processing using hybrid approaches

Online forums data is often unstructured, noisy, and diverse. Therefore, it is essential to understand the data and its limitations before analysing it. Therefore, in our proposed research several hybrid approaches are carried out in the pre-processing phase which makes the text in a structured format.

- *HTML/XML Tag Removal*: E-commerce product data often contains HTML/XML tags that can interfere with the analysis. Therefore, it is necessary to remove these tags before further processing the data.
- *Tokenisation*: Tokenisation involves splitting the text into individual tokens or words. This step is important because sentiment analysis models typically work with individual words rather than the whole text. In this process, the larger chunk of text is separated into documents, paragraphs, or sentences, into smaller units called

tokens. The resulting tokens can then be further processed using techniques such as stop word removal and stemming.

- *Stop word removal*: Stop words are common words such as ‘and’, ‘the’, and ‘a’, that does not carry much meaning in sentiment analysis.
- *Stemming*: Stemming involves reducing words to their base or root form. This process can help to reduce the number of features and account for variations in the text, such as plurals or verb tenses.

Using a combination of these pre-processing approaches can help to create a structured and normalised dataset for sentiment analysis.

3.3 *Senti-WordNet-based feature extraction*

The objective of feature extraction in NLP is to create a set of feature vectors from the pre-processed text data that can be fed into a deep learning algorithm. By reducing the size of the feature set, we can improve the efficiency and accuracy of the classification model. In the proposed methodology, the feature extraction method is based on Senti-WordNet, which is a lexical resource that assigns sentiment scores to words in the WordNet repository. Senti-WordNet provides a dictionary of opinion words that can be used to calculate the overall sentiment of a piece of text. The proposed feature extraction model uses Senti-WordNet to extract features from e-commerce data, which can be utilised to classify the sentiment of the reviews. The feature extraction process involves mapping each word in the pre-processed text to its corresponding sentiment score in Senti-WordNet. The sentiment scores are then aggregated to produce a feature vector for the entire text.

3.4 *Self-attention based bi-directional LSTM*

The recommended ILW-LSTM classification is fed with the characteristics that were chosen throughout the feature selection phase in order to categorise sentiment. To solve the vanishing/exploding issue that RNNs experience, a special type of RNN called LSTM was developed. In the same way as other RNN varieties, LSTMs generate their result depending on the information with the present step in time and the results through the preceding time step, and they then send their current result to the following time step. Each LSTM unit consists of a memory cell that may maintain its state for any duration and three non-linear gates: an input gate (in_t), a forget gate (fg_t), and an output gate (out_t). These gates control how information enters and exits the memory cell. Consider applying a sigmoid function $\sigma(\cdot)$, $\odot$ a dot product, and $\tanh(\cdot)$ a hyperbolic tangent function to the elements. Describe in_t and hid_t the input and concealed state vectors at time t , respectively. Bias vectors are represented by bias, whereas gate weight matrices are shown by X and Y . The forget gate decides what data has to be forgotten by generating a number between $[0,1]$.

$$fg_t = \sigma(Y_{fg} hid_{t-1} + X_{fg} in_t + bias_{fg}) \quad (1)$$

The input gate computes in_t and $\hat{c}_t$, syndicates them, and determines what extra data requires to be preserved using the following equations.

$$in_t = \sigma(Y_{in} hid_{t-1} + X_{in} in_t + bias_{in}) \quad (2)$$

$$\hat{c}_t = \tanh(Y_c hid_{t-1} + X_c in_t + bias_c) \quad (3)$$

$$c_t = fg_t \otimes c_{t-1} + in_t \otimes \hat{c}_t \quad (4)$$

Which elements of the cell state should be output by the output gate is determined by the following equations.

$$out_t = \sigma(Y_{out} hid_{t-1} + X_{out} in_t + bias_{out}) \quad (5)$$

$$hid_t = out_t \otimes \tanh(c_t) \quad (6)$$

The self-attention mechanism adapts the weighted attention feature vector A self-attention mechanism can be used in combination with bidirectional LSTM models for sentiment analysis tasks. In this approach, the bidirectional LSTM model is used to encode the input sequence and capture its contextual information. Then, a self-attention mechanism is applied to the encoded sequence to emphasise important information and suppress irrelevant information for sentiment analysis. The step-by-step process is illustrated as follows,

- *Input sequence*: The input sequence of words or tokens is fed into the bidirectional LSTM model.
- *BiLSTM encoding*: The bidirectional LSTM model encodes the input sequence and captures the contextual information of each word in the sequence.
- *Self-attention mechanism*: The encoded sequence is then passed through a self-attention mechanism that calculates a weight for each word in the sequence based on its relevance to the sentiment analysis task.
- *Weighted sequence representation*: The weights are used to compute a weighted sum of the encoded sequence, which generates a final representation of the sequence that emphasises the important parts and suppresses irrelevant parts for sentiment analysis.
- *Classification layer*: The final sequence representation is then passed through a classification layer to expect the sentiment of the input sequence.

This method has proven successful for sentiment assessment problems because it enables the system to concentrate on the input sequence's most important elements despite disregarding noise and unimportant data.

3.5 Parameter tuning using LFMO

The weighting parameters present in the deep learning algorithm are tuned using the proposed levy flight-based may fly optimisation algorithm. The behaviour of mayflies, particularly during mating, served as the model for the mayfly algorithm. The d-dimensional vector represents each mayfly as a randomly placed potential solution in the

issue space $Q_{Gx} = (Q_{G1}, Q_{G2}, \dots, Q_{Gd})$. The performance is then assessed in accordance with the defined goal function $F(C_T(Q_{Gx}))$. During a shift in location, a mayfly's velocity $vel = (vel_1, \dots, vel_d)$ is initialised. Every mayfly has a tendency to alter its course to follow its current, ideal location (P_{best}). It also varies based on the highest location that any other mayfly in the swarm has obtained thus far (G_{best}). The number of male mayflies as $Q_{Gmx} (x=1, 2, \dots, IG)$ is known as having initialised speeds of vel_{mx} . The swarms of male mayflies show that each mayfly's location shifts based on its own experiences and those of its neighbours. Q_{Gx}^T is assumed to be mayfly x 's current location in search space at time step T . As velocity is introduced vel_x^{T+1} , the position of mayfly x changes. This notation is presented here precisely as it is written.

$$Q_{Gmx}^{T+1} = Q_{Gx}^T + vel_x^{T+1} \quad (7)$$

few metres above the water, with $Q_{Gxm}^0 U(Q_{Gm \min}, Q_{Gm \max})$, male mayflies are thought to be present and engaged in a nuptial dance.

$$vel_{xy}^{T+1} = g * vel_{xy}^T + m_1 e^{-\beta r_p^2} (P_{best_{xy}} - Q_{Gmxy}^T) + m_2 e^{-\beta r_g^2} (G_{best_y} - Q_{Gmxy}^T) \quad (8)$$

Here, the effects of the social and mental elements are scaled up using positive attraction constants is $m_1 \cdot vel_{xy}^T$ correlates to the speed in dimension $y = 1, \dots, i$ at time step T of the mayfly x 's velocity. Q_{Gmxy}^T specifies the mayfly's x^{th} position in dimension j at time step T, m_1 . Similarly, r_g denotes the Cartesian separation between Q_{Gx} and G_{best} while r_p denotes the separation between Q_{Gx} and P_{best_x} . In addition, P_{best_x} indicate the mayfly's best location to date. This step involves calculating the velocity of a mayfly candidate solution utilising the Levy flight methodology. The Levy flight technique is used in this stage to move the position of the global finest constituent. The system's weight factor was constructed and used to evaluate the deep learning framework after a set number of iterations. This has the ability to increase accuracy by lowering the error function of wrongly categorised sentiment.

4 Results and discussion

In the proposed research work, the bidirectional long-short term memory (Bi-LSTM) network is used to capture the sequential relationships between words in a text document, while the self-attention mechanism is used to dynamically weigh the importance of different words in determining the overall sentiment expressed in the text. The experimental assessment described in the scenario involves the collection of text data about sentiments, which is then processed using Python. Several related phrases are calculated and examined employing the provided approach for data classification, and the general feeling of different evaluations is expressed employing the sentiment analysis technique. The evaluation process is done with 20% of the datasets used for testing and 80% used for training, and 0.2% of the testing data being used as a validation subset. The experimental procedure described in the scenario is approved out consuming the Python

tool on a machine with 12 GB of RAM and an Intel TM core (7M) i3-6100 CPU running at 3.70 GHz. Various Python programs, such as OpenCV, NumPy, TensorFlow, Keras, and Matplotlib, are used to handle data processing and analysis. The performance of the proposed approach is evaluated based on accuracy, precision, recall, and F-measure.

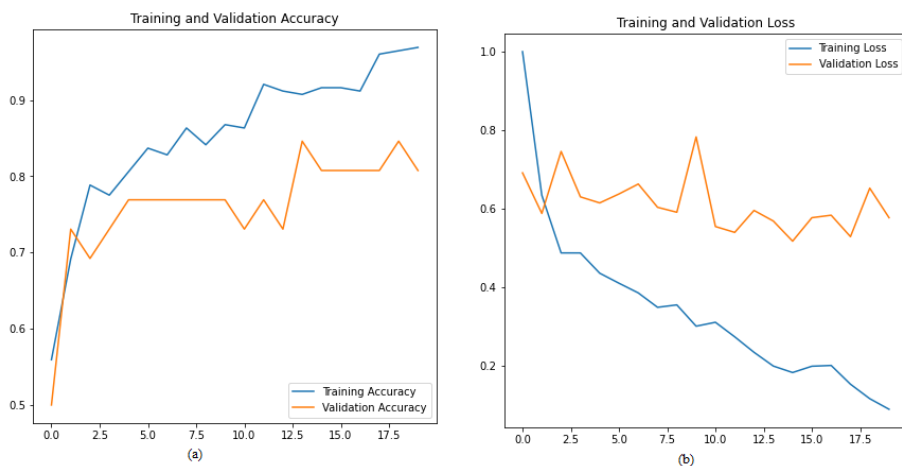
4.1 Dataset description

In this case, 0.2% of the testing data is validation subset, 80% of the datasets are used for training, and 20% for testing. To validate the offered method, simulations are done on the Python working platform using the program's multi-class opinions. 20% of the Kaggle dataset originated from the testing stage, whereas 80% originated from the training process. While the suggested emotional analysis assignment took 9.71 s to process, the training method took 103.86432 s to finish. Evaluating the suggested strategy's evaluation measures to those of the existing methodologies further demonstrates the predictive model's efficacy. The confusion matrix produced from the testing findings is used to construct metrics for evaluation. Accuracy, Precision, F1-Score, Specificity, Sensitivity, Processing Time, and Error are employed as evaluation criteria for this analysis. The dataset is accessed via (<https://www.kaggle.com/datasets/roopalik/amazon-baby-dataset>) which encompass amazon product review with its product source.

4.2 Performance analysis of the proposed model

The recommended process begins with the collecting of data from e-commerce webpages, which includes reviews of products for babies on Amazon. Online reviews are gathered from various groups of the 17168 customers who bought the relevant goods. Typically, input data is frequently composed of tags, conjunctions, papers, adverbs, blank spaces, prepositions, superfluous and overused words, colloquial language, and unidentified terms, which is unsuitable for further learning. The next sections calculate the assessment metrics sensitivity, specificity, accuracy, precision, processing time for the F1-score, and loss functions to assess the proposed model by contrasting it with the current approaches. The next sections provide a summary of the findings as well as comparison graphs between the suggested and actual categorisation schemes.

Figure 2(a) displays the outcomes of the accuracy metric's training and validation. In the supplementary figure, training and validation accuracy are represented by the blue and orange lines, respectively. It appears that altering the iteration count leads the developed accuracy metric to gradually increase and then remain to increase, demonstrating superior calculation consequences. For instance, the training and validation accuracy is generally above 95% when the iteration reaches the 15th number of iterations. The subsequent classification results for training and validation loss is clearly seen in Figure 2(b). The results of the proposed deep learning technique's loss function (prediction error), which yields the lowest error function, are shown in Figure 2(b). As an illustration, going from 0 to 17.5 iterations lowers the prediction error of the recommended solution from 1.0 to 0.2. As a result, its accuracy and error function are improved. The proposed methodology is examined and compared to current methods in the parts that follows.

Figure 2 Simulation outcome (a) accuracy measure and (b) loss (see online version for colours)

4.3 Comparative analysis of proposed and existing system

The results from the proposed study are contrasted with those from several different current approaches to show the performance efficacy of the suggested model. These results demonstrated the effectiveness of the ILW-DNN method and sparked several research initiatives to develop a deep learning-based system for sentiment categorisation. As a result, in this section, the outcomes of the proposed model have been contrasted with those of existing techniques.

Table 1 provides a clear comparison of the different approaches utilising the suggested technique. The comparison analysis uses accuracy, precision, recall, and F1 score as its determining variables. When compared to CNN, LSTM, Bi-LSTM, CNN-LSTM, ConvBi-LSTM, and ILW-DNN, the suggested methodology's values for accuracy, precision, recall, and F1-score are 97.61, 97.24, 99.61, and 98.27%, respectively. Using the suggested strategy, it appears that the loss function obtained is small. Overall research showed that the recommended deep learning approach yields a superior emotional model analysis result. Using the given framework, the sentimental analysis technique was discussed in this study paper. The author also demonstrates how this technique may be used to categorise various feelings for the input that has been gathered.

Table 1 Overall performance analysis of proposed and existing methods

<i>Techniques</i>	<i>Accuracy</i>	<i>Precision</i>	<i>Recall</i>	<i>F1-Score</i>
CNN (Jianqiang et al., 2018)	91.89	89.19	90.92	89.26
LSTM (Huang et al., 2021)	90.94	89.48	88.06	87.86
Bi-LSTM (Liu and Guo, 2019)	90.57	84.42	93.99	88.1
CNN-LSTM (Rehman et al., 2019)	90.81	89.96	92.02	90.13
ConvBi-LSTM (Ruz et al., 2020)	93.76	90.5	94.42	91.81
Proposed model	97.61	97.24	99.35	98.27

5 Conclusions and future scope

The major aim of the proposed system is to build a multi-label sentimental analysis system using the proposed deep learning scheme. For developing such a system, pre-processing and feature extraction techniques are done before the classification. Then the proposed deep learning model is used to classify the sentimental analysis task. Experimental results have shown that this approach outperforms traditional MLSA methods, such as support vector machines and decision trees, in terms of accuracy and F1 score, especially when dealing with complex texts that express multiple sentiments. In conclusion, the Bi-LSTM with Self Attention Mechanism approach is a promising technique for MLSA and has shown great potential for improving the accuracy and interpretability of sentiment analysis models. Additionally, the use of the self-attention mechanism has been shown to improve the interpretability of the predictions by highlighting the most important words in the text that contribute to the overall sentiment. This research has the potential to make a significant contribution to the field of sentiment analysis and NLP, as well as provide valuable insights into public opinions and emotions.

References

- Alsayat, A. (2022) 'Improving sentiment analysis for social media applications using an ensemble deep learning language model', *Arabian Journal for Science and Engineering*, Vol. 47, No. 2, pp.2499–2511.
- Deng, J., Cheng, L. and Wang, Z. (2021) 'Attention-based BiLSTM fused CNN with gating mechanism model for Chinese long text classification', *Computer Speech and Language*, Vol. 68, p.101182.
- Huang, F., Li, X., Yuan, C., Zhang, S., Zhang, J. and Qiao, S. (2021) 'Attention-emotion-enhanced convolutional LSTM for sentiment analysis', *IEEE Transactions on Neural Networks and Learning Systems*, Vol. 33, No. 9, pp.4332–4345.
- Jagdale, R.S., Shirsat, V.S. and Deshmukh, S.N. (2019) 'Sentiment analysis on product reviews using machine learning techniques', *Cognitive Informatics and Soft Computing: Proceeding of CISC. (2017) Springer Singapore*, pp.639–647.
- Jain, P.K., Pamula, R. and Yekun, E.A. (2022) 'A multi-label ensemble predicting model to service recommendation from social media contents', *The Journal of Supercomputing*, Vol. 78, pp.1–18.
- Jianqiang, Z., Xiaolin, G. and Xuejun, Z. (2018) 'Deep convolution neural networks for twitter sentiment analysis', *IEEE Access*, Vol. 6, pp.23253–23260.
- Li, D., Rzepka, R., Ptaszynski, M. and Araki, K. (2020) 'HEMOS: a novel deep learning-based fine-grained humor detecting method for sentiment analysis of social media', *Information Processing and Management*, Vol. 57, No. 6, p.102290.
- Liu, G. and Guo, J. (2019) 'Bidirectional LSTM with attention mechanism and convolutional layer for text classification', *Neurocomputing*, Vol. 337, pp.325–338.
- Ombabi, A.H., Ouarda, W. and Alimi, A.M. (2020) 'Deep learning CNN–LSTM framework for Arabic sentiment analysis using textual information shared in social networks', *Social Network Analysis and Mining*, Vol. 10, pp.1–13.
- Priyadarshini, I. and Cotton, C. (2021) 'A novel LSTM–CNN–grid search-based deep neural network for sentiment analysis', *The Journal of Supercomputing*, Vol. 77, No. 12, pp.13911–13932.
- Rehman, A.U., Malik, A.K., Raza, B. and Ali, W. (2019) 'A hybrid CNN-LSTM model for improving accuracy of movie reviews sentiment analysis', *Multimedia Tools and Applications*, Vol. 78, pp.26597–26613.

- Ruz, G.A., Henríquez, P.A. and Mascareño, A. (2020) 'Sentiment analysis of twitter data during critical events through Bayesian networks classifiers', *Future Generation Computer Systems*, Vol. 106, pp.92–104.
- Sivakumar, M. and Uyyala, S.R. (2021) 'Aspect-based sentiment analysis of mobile phone reviews using LSTM and fuzzy logic', *International Journal of Data Science and Analytics*, Vol. 12, No. 4, pp.355–367.
- Yadav, A. and Vishwakarma, D.K. (2020) 'A deep learning architecture of RA-DLNet for visual sentiment analysis', *Multimedia Systems*, Vol. 26, No. 4, pp.431–451.
- Yilmaz, S.F., Kaynak, E.B., Koç, A., Dibeklioglu, H. and Kozat, S.S. (2021) 'Multi-label sentiment analysis on 100 languages with dynamic weighting for label imbalance', *IEEE Transactions on Neural Networks and Learning Systems*, Vol. 34, No. 1, pp.331–343.
- Zhang, L., Wang, S. and Liu, B. (2018) 'Deep learning for sentiment analysis: a survey', *Wiley Interdisciplinary Reviews: Data Mining and Knowledge Discovery*, Vol. 8, No. 4, pe1253.

Website

<https://www.kaggle.com/datasets/roopalik/amazon-baby-dataset>