

# Sentiment Analysis of IMDB Movie Reviews Based on LSTM

Jiahao Lu\*, Hongji Fan, Yali Zhang

Software School, HenanNormalUniversity ,XinXiang, 453000, China

\*Corresponding author: Jiahao Lu , 3329362823@qq.com

**Copyright:** 2025 Author(s). This is an open-access article distributed under the terms of the Creative Commons Attribution License (CC BY-NC 4.0), permitting distribution and reproduction in any medium, provided the original author and source are credited, and explicitly prohibiting its use for commercial purposes.

**Abstract:** This study focuses on sentiment analysis of IMDB movie review data. Firstly, by leveraging the CBOW model and dynamically adjusting the context window, the distributed semantic features of words in movie reviews are deeply explored to construct precise word vector representations. To comprehensively evaluate the effectiveness of the model, a series of comparative experiments are designed. These experiments include not only traditional machine learning algorithms such as Naive Bayes, Decision Tree, and Random Forest but also a basic RNN model. The experimental results show that the LSTM combined with the Word2vec model achieves an accuracy of 0.9265, a recall of 0.9471, and an F1-score of 0.9353, demonstrating high precision, strong generalization ability, and good applicability.

**Keywords:** IMDB Dataset; Sentiment Analysis; Word2vec Model; CBOW Algorithm; LSTM Model

**Published:** Jun 6, 2025

**DOI:** <https://doi.org/10.62177/jaet.v2i2.429>

## 1.Introduction

The global film industry is continuously growing under the joint influence of policy promotion and market demand. As a flagship platform for movie information exchange, IMDB has become a crucial place for movie enthusiasts to obtain information and exchange opinions. This platform aggregates over fifty thousand user reviews daily, with diverse forms ranging from brief star ratings to detailed long reviews. These reviews not only reflect the audience's immediate feelings during movie-watching but also construct a rich semantic system of movie evaluations through distinctive language expressions. This system is an indispensable valuable resource for the field of natural language processing, especially for sentiment analysis tasks. Currently, with the rapid development of streaming services and online movie review communities, the in-depth mining of the emotional features of comments on IMDB is of great value to multiple parties. Movie creators can optimize their work directions based on this, distributors can accurately formulate marketing strategies, and audiences can make more informed movie-watching choices.

Text sentiment analysis is a key method in the study of movie reviews, relying heavily on the deep application of natural language processing technology. This method is dedicated to the meticulous feature analysis of text information with emotional connotations, extracting the core semantics of the text through complex algorithm models, and thus achieving in-depth mining of information. Wang Renwu<sup>[10]</sup> improved the accuracy of sentiment analysis by constructing a sentiment dictionary using Word2vec compared to traditional methods.

In the process of sentiment analysis, NLP technology not only parses the literal meaning of text but also perceives the underlying emotional tendencies, providing a comprehensive and in-depth understanding and analysis of review texts. As an important branch of natural language processing, sentiment analysis technology can be broadly divided into two technical

schools: traditional machine learning methods and deep learning methods. These two approaches exhibit distinctly different strategies in feature engineering and model construction. Traditional machine learning methods in sentiment analysis focus more on the meticulous construction of feature engineering.

Based on statistical sentiment analysis methods, the two main technical paradigms can be systematically divided into traditional machine learning methods and deep learning methods, which form a sharp contrast in feature engineering and model construction. Traditional methods emphasize the optimization of feature engineering. For example, Wang Mingyang et al.<sup>[3]</sup> found that expanding the sentiment dictionary and introducing the Word2vec model both help improve the effect of the k-nearest neighbor (K-Nearest Neighbor, KNN) method on multi-emotion classification of Weibo text; Wang Haiyan<sup>[12]</sup> et al. proposed a text classification algorithm based on a rough set-based bag-of-words model combined with support vector machines.

Deep learning<sup>[13]</sup> methods have gradually become mainstream due to their end-to-end learning advantages, effectively avoiding manual intervention by automatically generating text representation features. Zhang Qian<sup>[11]</sup> et al. applied the Word2vec model to transform text data into real-number vectors; Reference<sup>[5]</sup> used the Word2Vec tool to train a corpus, obtained text word vector representations, and then used an LSTM model with an attention mechanism for text feature extraction, combined with cross-entropy training, and applied the model to the classification of tourism issue texts; Reference<sup>[7]</sup> proposed a new short text classification method by combining context-dependent features with a multi-stage attention model based on time convolutional networks and CNNs; The main features of convolutional neural networks<sup>[1]</sup> lie in weight sharing and local connections; Reference<sup>[7]</sup> proposed a new short text classification method by combining context-dependent features with a multi-stage attention model based on time convolutional networks and CNNs; Reference<sup>[2]</sup> proposed a dual graph convolutional network model for aspect-based sentiment analysis, which considers both the complementarity of syntactic structure and semantic relevance; Liu Hui<sup>[4]</sup> et al. proposed an aspect-level sentiment analysis model that integrates matching long short-term memory networks and grammatical distance; Liu Jun<sup>[8]</sup> et al. created a long text analysis model using CNN and bidirectional long short-term memory networks (Bi-directional Long Short-Term Memory, Bi-LSTM); Ding Feng et al.<sup>[6]</sup> proposed a BiLSTM-CRF model and applied it to aspect-based sentiment analysis tasks. The bidirectional long short-term memory network can capture long-distance bidirectional semantic dependencies and learn text semantic information, predicting the global optimal label sequence at the sentence level.

Through scientific sentiment analysis methods, in-depth mining of off-site data in the movie industry has value in three dimensions: First, it can provide precise assistance for audience group movie-watching decisions and, through visual analysis results, perceive public feedback on movies, promoting multi-dimensional emotional resonance among movie enthusiasts. Secondly, film producers can gain data support for precise market positioning through insights into emotional data, accurately grasping the aesthetic preferences and emotional needs of the audience. More importantly, this analysis mechanism not only provides an in-depth analysis of contemporary public aesthetic value orientations but also has important strategic reference value for the market potential mining of documentary and other sub-genres, as well as for innovating film content production and optimizing the industrial ecosystem.

This study aims to develop an LSTM-based sentiment classification model using the IMDB movie review dataset. Under this framework, the CBOW algorithm is used to transform movie review text into high-dimensional distributed word vectors through pre-training, thereby constructing a semantic information-rich word representation space. Subsequently, the LSTM network is introduced, leveraging its unique temporal memory ability to deeply extract emotional features from the text, continuously optimizing network parameters to enhance model performance. Finally, a horizontal comparison strategy is adopted, selecting classic models such as SVM and CNN as references, and using precision, recall, and F1-score as three key indicators to comprehensively evaluate the effectiveness of the LSTM model. In addition, detailed visual analysis of the experimental data is conducted to reveal the model's performance and characteristics from multiple dimensions.

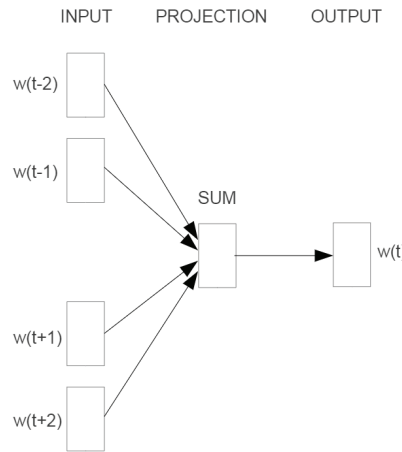
## 2. Algorithm Introduction

### 2.1 Word2vec Model

In the field of natural language processing, word vector representation has always been a challenge. However, the Word2vec

framework has successfully broken through the limitations of traditional One-Hot encoding through its innovative distributed representation method. Word2vec includes two classic algorithms—CBOW and Skip-gram models—which adopt different contextual learning strategies. The CBOW algorithm uses a predefined context window to aggregate semantic information from surrounding words to predict the features of the target word; in contrast, the Skip-gram model starts from the central word and infers the semantic representations of its surrounding words in reverse. Although these two training methods complement each other in achieving their goals, both can construct dense vector spaces rich in semantic associations. These distributed word vectors have demonstrated extremely high engineering application value in practical applications such as text classification and semantic similarity calculation.

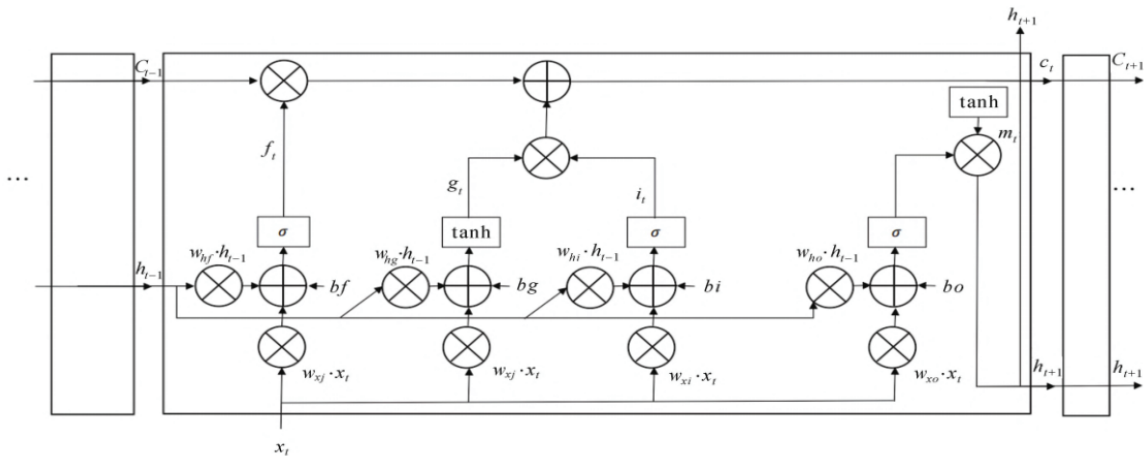
Figure 1 The structure of CBOW



When discussing the architecture implementation of the CBOW model, we note that its parameter passing mechanism follows a carefully designed hierarchical computation strategy. The model starts from the input layer, mapping each central word's contextual information into a specific embedding space. During this process, each context word is converted into a vector representation. Then, the hidden layer aggregates these input vectors, generating an intermediate semantic representation through weighted summation, which fuses contextual information. The output layer adopts an optimized binary tree topology, where leaf nodes correspond one-to-one with words in the corpus, and non-leaf nodes form feature combination paths. This design helps improve computational efficiency. In model input, one-hot encoded vectors are used as a starting point, transformed through a weight matrix  $Y$  to generate  $N$ -dimensional hidden layer vectors. Subsequently, this vector is reconstructed through another weight matrix  $W$ , returning to  $M$ -dimensional output space, and finally normalized through Softmax to obtain the predicted distribution of the target word. By setting an appropriate window size  $C$ , the CBOW model can effectively control the context range, achieving semantic conversion from discrete symbols to continuous vector spaces. This process is clear and efficient.

## 2.2 LSTM Model

Figure 2 LSTM specific flow chart



To achieve long-term dependency goals, we designed a selective memory cell structure and introduced a forget gate mechanism to filter and retain key information from the input sequence. The computation formulas precisely define this process.

$$f_t = \sigma(f^{(t)}) = \sigma(W_{xf}x_t + W_{hf}h_{t-1} + b_f), (1)$$

The input gate is responsible for integrating information from the previous and current moments, with the specific computation formula provided, precisely controlling the inflow of information.

$$g_t = \tanh(g^{(t)}) = \tanh(W_{xg}x_t + W_{hg}h_{t-1} + b_g), (2)$$

Subsequently,  $i_t$  is computed to prepare for the final cell state  $C_t$ , with the formula as follows:

$$i_t = \sigma(i^{(t)}) = W_{xi}x_t + W_{hi}h_{t-1} + b_i, (3)$$

The output gate  $O_t$  will further filter the results of the input gate, retaining useful information. Specifically,  $x_t$  and  $h_{t-1}$  are multiplied by weight matrices in the output gate, summed, and a bias  $b_o$  is added. The computation formula is as follows:

$$O_t = \sigma(O^{(t)}) = W_{xo}x_t + W_{ho}h_{t-1} + b_o, (4)$$

At time  $t$ , the system receives three inputs: the current value  $x_t$ , the previous state  $h_{t-1}$ , and the memory cell state  $C_{t-1}$ . These three inputs jointly affect the forget gate, input gate, and output gate, participating in their computations involving  $x_t$  and  $h_{t-1}$ , thereby influencing the overall state output of the system.

In the forget gate, these two vectors are multiplied by a weight matrix, summed, and a bias term  $b_f$  is added. Finally, a sigmoid activation function is applied to obtain the result of the forget gate  $f_t$ . In the update gate,  $x_t$  and  $h_{t-1}$  are again multiplied by preset weight matrices, summed, and a bias term  $b_g$  is added, followed by processing through a tanh activation function layer. Subsequently, this result is added to the output of the forget gate through matrix addition to obtain the intermediate state  $g_t$  of the update gate.

Next, the input gate multiplies the input and the previous state  $x_t$  and  $h_{t-1}$  by preset weight matrices, sums them, and applies a sigmoid function activation to obtain the result  $i_t$ . Finally, the update gate and the forget gate outputs are combined to form the new cell state  $C_t$ , with the formula as follows:

$$C_t = C_{t-1} \odot f_t + g_t \odot i_t, (5)$$

where  $\odot$  denotes element-wise multiplication.

To further filter out non-critical information, the data  $C_t$  is processed through a tanh layer to obtain a refined matrix  $m_t$ :

$$m_t = \tanh(C_t), (6)$$

Then, the output gate result  $O_t$  is element-wise multiplied with  $m_t$  to obtain the final state vector  $h_t$ :

$$h_t = O_t \cdot m_t, (7)$$

Finally, the state vector  $h_t$  is multiplied by a weight matrix  $W_{yh}$  and a bias term  $b_y$  is added to obtain the prediction result  $y_t$  at time  $t$ :

$$y_t = W_{yh}h_t + b_y, (8)$$

Thus, the forward propagation process within the memory cell is completed.

Figure2: Comparison of Michael Jordan and LeBron James' shooting percentages in the regular season

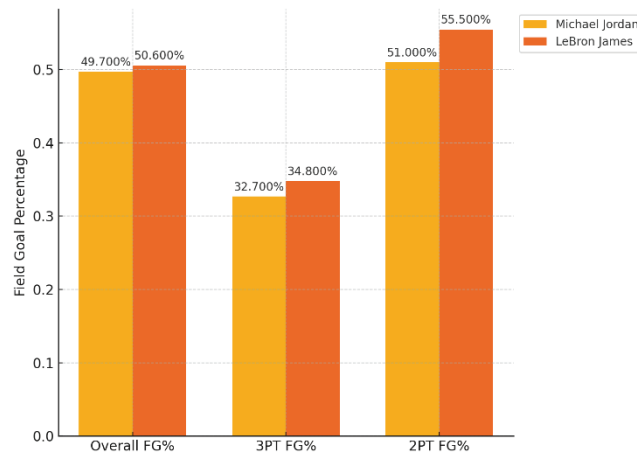


Figure2 is a comparison of the shooting percentages of Michael Jordan and LeBron James in the regular season (up to the 2023 season), including overall shooting percentage, three-point shooting percentage, and two-point shooting percentage.

### 3.Data Acquisition and Preprocessing

#### 3.1 Experimental Data Acquisition

This study uses the internationally recognized IMDB benchmark dataset as the experimental object, obtained from the official website of Stanford University and following standard data collection protocols. It contains two independent modules, stored in a separate Excel file, namely the training set and the test set. The training set file contains 25,000 labeled samples (12,500 positive and negative reviews each), and the test set documents construct an equally proportioned balanced data distribution, forming a total of 50,000 high-quality labeled materials.

The dataset adopts a two-pole classification storage mechanism, with each Excel file internally divided into "POS" and "NEG" logical units, storing reviews with positive emotional labels (Positive) and negative emotional identifiers (Negative), respectively. Each review contains structured fields such as a unique identifier, text number, star rating, and original review text. The text content is cleaned through standardized processing to effectively eliminate HTML tag interference using special characters.

#### 3.2 Data Preprocessing

##### 3.2.1 Discrete Variable Labeling

For the binary emotional labeling system targeting the IMDB dataset, a numerical encoding strategy is adopted, decentralizing the category. In the original data, the "POS" folder storing positive evaluations (Positive) is uniformly represented by the numerical label "1"; the "NEG" folder storing negative evaluations (Negative) corresponds to the numerical label "0".

*Table 1 Discrete Variable Labeling*

|     | Description                    | Numerical Label |
|-----|--------------------------------|-----------------|
| POS | Positive Evaluation (Positive) | 1               |
| NEG | Negative Evaluation (Negative) | 0               |

##### 3.2.2 Data Tokenization

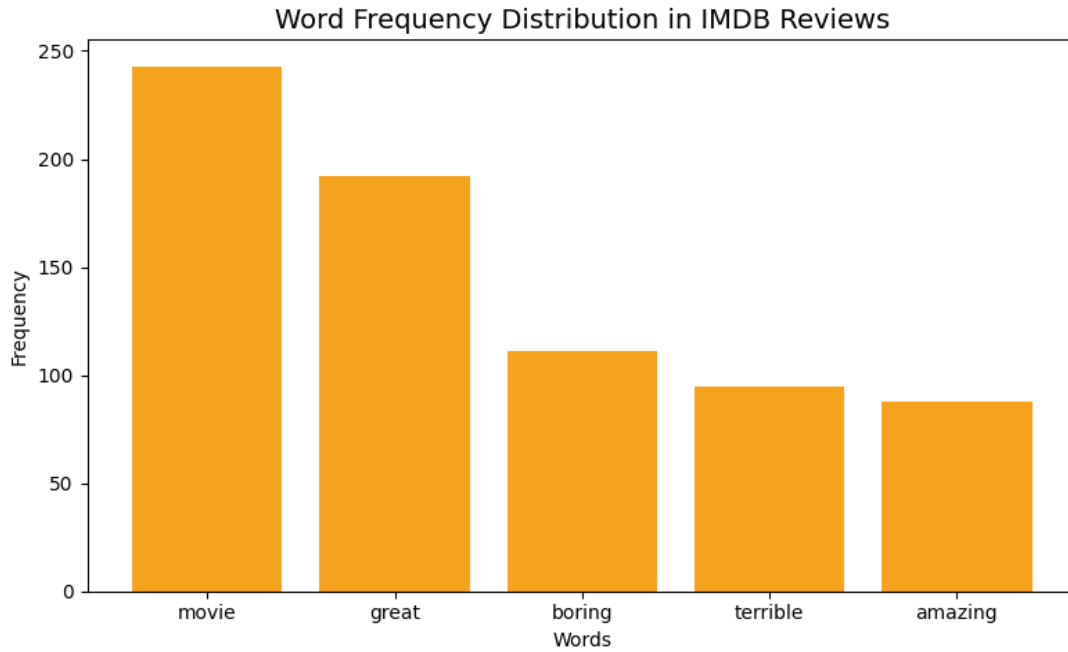
This study identifies high-frequency semantic units by constructing a vocabulary distribution histogram, adopting a feature engineering optimization strategy based on word frequency statistics. By calculating the probability distribution of term occurrences, key word indicators are screened during the corpus parsing stage, resulting in a distinct differentiation. This statistic-based feature screening mechanism not only generates a high-density information feature vector space but also effectively enhances the text representation's ability to capture contextual semantics, thereby establishing a vocabulary association model with strong explanatory power and providing an input matrix for subsequent supervised learning models after denoising.

*Table 2 Word Frequency Distribution*

| Word     | Frequency |
|----------|-----------|
| movie    | 243       |
| great    | 192       |
| boring   | 111       |
| terrible | 95        |
| amazing  | 88        |

Based on the quantitative data in Table 2, a visualization chart is constructed using word frequency statistical analysis to systematically present the distribution patterns of high-frequency vocabulary.

Figure 3 Word frequency statistics



#### 4. Word2vec Model for Constructing Word Vector Space

This study aims to utilize the Continuous Bag-of-Words (CBOW) model to construct a distributed semantic space for text. During the construction process, the dimension parameters of the word vector space are first defined, determined by the vocabulary size  $a$  and the preset semantic embedding dimension  $K$ . In this space, each vocabulary is mapped to a row vector, and the geometric layout of these vectors reflects the semantic relevance between words. For the review text corpus  $M$ , the study converts it into a sequence form through context window sliding. The total number of words in each review is denoted by  $T$ , and the position of words in the vocabulary is represented by an index value  $\in [1, a]$ . To accurately capture word relationships, this conversion introduces dynamic window sampling techniques, gradually adjusting the geometric structure of the word vector space based on the relevance between central words and context words. This modeling process includes multiple steps, aiming to deeply explore text semantic information through the CBOW model, thereby enhancing the accuracy and efficiency of text processing and analysis.

##### 4.1 Data Transformation Expression

A global corpus encoding strategy is adopted for the study. First, the preprocessed review text is integrated into a comprehensive text representation matrix through the CBOW algorithm, enhancing data generalization ability. Subsequently, a stratified random subdivision strategy is employed, dividing the integrated feature matrix into training and test subsets in a 1:1 ratio, ensuring consistency in dimensions such as emotional label distribution and text length in statistics. This method not only retains the integrity of contextual semantics but also avoids model overfitting risks due to data leakage, eliminating numerical tensor generation caused by feature deviation between datasets.

##### 4.2 Model Definition

In neural network processing, this study conducts experiments on the IMDB movie review dataset. To convert text vocabulary into a numerical form recognizable by computers, the efficient Word2Vec word embedding technique is utilized to achieve effective conversion of vocabulary into numbers for subsequent data processing and analysis.

Specifically, the Gensim library for natural language processing in Python is used, directly calling the Word2Vec model within it to train the corpus and generate word vectors. This approach maintains the relationships between words while effectively transforming each word into a low-dimensional continuous space.

The optimal configuration in the parameter space is determined using the grid search method during the model training process. By adjusting the word vector dimension parameters, the model complexity is reduced to accommodate scenarios with smaller training data scales.

Table 3 Parameters in the Word2vec Model

| Parameter Name | Description               | Setting |
|----------------|---------------------------|---------|
| vector_size    | Word Vector Dimension     | 300     |
| window         | Context Window Size       | 5       |
| min_count      | Minimum Occurrence Count  | 2       |
| epochs         | Number of Training Epochs | 10      |

## 5.Sentiment Analysis Based on LSTM Model

### 5.1 Data Format Conversion

#### 5.1.1 Calling Word Vector Space Results

Here, the vector matrix dictionary obtained from the Word2vec model is directly used.

Table 4 Index List and Corresponding Review Text Content

| Index | Word     |
|-------|----------|
| 0     | movie    |
| 1     | great    |
| 2     | boring   |
| 3     | terrible |
| 4     | amazing  |

#### 5.1.2 Creating Tuples

##### 1) Data Structure Initialization

To create a semantic mapping interface, a system based on the pre-trained word embedding model is constructed, loading the processed review text and corresponding emotional labels. Simultaneously, the vocabulary features generated by Word2vec are integrated, calculating the total number of samples to set appropriate batch training parameters, providing an accurate data scale reference for loss evaluation, ensuring efficient and accurate training processes.

##### 2) Semantic Mapping and Missing Value Handling

A word-by-word search is conducted on the review text. If a word exists in the embedding matrix, its corresponding vector is extracted.

To handle out-of-vocabulary (OOV) words, a zero initialization strategy is adopted, generating zero vectors matching the embedding dimensions. This ensures that OOV words are correctly processed during model training, enhancing the model's generalization ability.

Using serialized word vectors, they are concatenated into a sequential feature matrix, i.e., sample size multiplied by sequence length multiplied by embedding dimension. This method better represents the temporal information in the text, providing more possibilities for subsequent processing.

##### 3) Dimensional Standardization Processing

The maximum sequence length of the training set is calculated as the padding baseline. Short texts are padded with zeros at the end, while long texts are truncated from the tail. Emotion labels are converted into 64-bit integer tensors to adapt to the PyTorch framework. The final generated tuple data structure (feature matrix, label tensor) format meets the dimensionality requirements of deep learning models for input data. Meanwhile, through length standardization processing, the training efficiency loss caused by differences in text length is eliminated.

## 5.2 Defining the LSTM Model Structure

### 5.2.1 Defining Model and Functional Layers

The core architecture of this study focuses on constructing an emotion classification model, divided into two key stages:

parameter initialization and hierarchical feature mapping. First, based on task requirements, the network's topological parameters are configured, including the setting of input, hidden layer, and output dimensions, as well as the selection of activation functions. The tanh activation function is adopted to model temporal dependencies, and gating mechanisms are introduced to dynamically regulate information flow. The final layer is a fully connected classification layer, which maps the hidden layer state to the probability space using the sigmoid function to obtain the final emotional polarity prediction values. The entire design is implemented using the PyTorch framework, achieving automatic gradient computation and thereby improving training efficiency. Through this structure, an efficient emotion classification model is successfully constructed, providing strong support for research and application in the field of sentiment analysis.

### 5.2.2 Defining the Forward Propagation Function

In the process of temporal feature modeling, its implementation relies on a structured information flow mechanism. At the initial state setting of the model, both the memory cell state vector and the hidden state tensor need to undergo a strict zero-value initialization. This step ensures consistency and a clear baseline condition at the starting point of temporal processing. Subsequently, the forward computation process begins. The input sequence undergoes a series of gating operations in the Long Short-Term Memory (LSTM) model, including the input gate, forget gate, and output gate, which perform multi-scale and hierarchical feature fusion. Through iterative updates and continuous corrections of the memory cell, the hidden contextual dependencies in the text sequence are deeply captured and represented. At the final moment of temporal feature extraction, after dimensionality reduction and nonlinear activation by the fully connected layer, the probability distribution of emotion classification is formed. The entire computation process is designed with a dynamic computational graph, achieving end-to-end gradient propagation from word sequence input to emotion judgment output. Additionally, gradient normalization strategies are employed to ensure the robustness and reliability of the training process, thereby establishing a complete mapping framework from input to judgment.

### 5.2.3 Setting Random Seeds

In this exploration, we address the issue of convergence uncertainty caused by the inherent randomness in neural network weight initialization. A hardware-independent random control mechanism is carefully designed and implemented. By leveraging the PyTorch deep learning framework, we precisely set its deterministic computation mode to ensure that, even on high-performance computing devices like GPUs, the weight matrices can be reproducibly generated during parameter initialization. This approach not only significantly influences and restricts the search path in the parameter space but also effectively prevents the training process from prematurely falling into local minima. It further ensures consistency in the gradient descent paths across different experimental batches, laying a solid mathematical foundation for the steady improvement of model performance in the extensive exploration of the global space.

### 5.2.4 Building Datasets and Iterators

This study constructs a data supply pipeline oriented towards deep learning frameworks. It first loads preprocessed word embeddings and sentiment-annotated vectors, instantiating the training set's feature matrix and data loader. Efficient memory management and data throughput optimization are achieved by configuring the number of samples selected in each training session, whether to shuffle the data, the number of processes for parallel loading, and selecting sample quantities that are multiples of the remaining samples. The data loader employs an on-demand generation mechanism, dynamically performing tensor conversion and batch combination during the training process. This ensures stable transmission of large text data within GPU memory constraints, providing a standardized data input interface for model iterations. The parameters of the training set iterator are set as shown in Table 5.

Table 5 Iterator Parameters

| Parameter Name | Description                        | Setting Value |
|----------------|------------------------------------|---------------|
| batch_size     | Batch Size                         | 32            |
| shuffle        | Whether to Shuffle Data            | True          |
| drop_last      | Whether to Drop Incomplete Batches | True          |

### 5.2.5 Constructing the LSTM Model

Based on the training and inference of the LSTM network in a GPU-accelerated environment, hyperparameters such as hidden layer dimensionality and learning rate are jointly optimized through hyperparameter space traversal. A predefined network topology configuration scheme is adopted for the LSTM model, with parameters set as shown in Table 6.

Table 6 LSTM Model Parameters

| Parameter Name | Description                    | Setting Value |
|----------------|--------------------------------|---------------|
| input_size     | Input Dimension                | 300           |
| hidden_size    | Number of Hidden Layer Neurons | 128           |
| output_size    | Number of Output Classes       | 2             |

### 5.2.6 Establishing Optimizer and Loss Function

This study employs the Adaptive Moment Estimation (Adam) optimizer for parameter updates, training the optimizer's parameter set. The initial learning rate is set to  $2 \times 10^{-6}$  to control parameter updates, tailored to the characteristics of the emotion classification task. Additionally, the cross-entropy loss function is constructed as the supervisory signal. Its mathematical properties effectively measure the discrepancy between the true labels and the probability distribution. The training scheme incorporates a dynamic learning rate decay coefficient of 0.

### 5.3 Model Prediction

During the model training process, the number of iterations is empirically set to 32. As training progresses, model parameters are continuously adjusted and optimized, leading to a steady decline in training loss values and a gradual increase in the training set's accuracy, ultimately stabilizing. Detailed loss values and accuracy data are shown in Table 8.

### 5.4 Model Comparison

Table 7 Training Loss and Accuracy During Training

| Iteration | Loss   | Training Accuracy |
|-----------|--------|-------------------|
| 1         | 0.5732 | 0.8794            |
| 5         | 0.5243 | 0.8952            |
| 10        | 0.4881 | 0.9108            |
| 15        | 0.472  | 0.9181            |
| 20        | 0.4635 | 0.9215            |
| 1         | 0.5732 | 0.8794            |

The computation processes for accuracy, recall, and F1-score are as follows:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN}, \text{Recall} = \frac{TP}{TP + FN}, \text{F1} = \frac{2PR}{P + R}$$

TP (True Positive) is one of the important indicators for evaluating model performance. In machine learning algorithms, TP represents the model's ability to correctly predict positive instances. A high TP value indicates a high recognition accuracy of positive instances by the model, reflecting excellent model performance. Therefore, when evaluating a model's performance, attention should be paid to the TP indicator to ensure the accuracy and reliability of the model's predictions.

FP (False Positive) refers to the situation where the model incorrectly predicts negative instances as positive. Such errors may lead to misleading conclusions and decisions. To avoid FP, it is necessary to improve the model's accuracy and precision, reducing false predictions. During model construction and validation, special attention should be paid to the FP rate to ensure the model's stability and reliability.

FN (False Negative) is a common and impactful error type in data analysis, referring to the situation where the model

incorrectly predicts positive instances as negative. This may lead to severe consequences; therefore, during model establishment and optimization, it is essential to focus on and identify the causes of FN to improve the model's accuracy and reliability.

TN (True Negative) is one of the crucial indicators for evaluating classification model performance, representing the model's ability to correctly predict negative instances. The TN value reflects the model's accuracy and reliability in identifying negative instances. When evaluating classification models, the TN value is critical for assessing the model's performance and is vital for enhancing the model's effectiveness.

*Table 8 Model Result Comparison*

| Model         | Accuracy | Recall | F1-Score |
|---------------|----------|--------|----------|
| Naive Bayes   | 0.8712   | 0.8931 | 0.8819   |
| Decision Tree | 0.8834   | 0.9075 | 0.8922   |
| Random Forest | 0.8991   | 0.9162 | 0.9027   |
| RNN           | 0.9033   | 0.925  | 0.9135   |
| LSTM          | 0.9265   | 0.9471 | 0.9353   |

## Conclusion

Through data preprocessing, the IMDB dataset employs Word2vec to construct a word vector space, followed by training an LSTM model. Comparative experiments with other models show that the LSTM combined with the Word2vec model achieves an accuracy of 0.9265, a recall of 0.9471, and an F1-score of 0.9353. This represents an improvement of at least 0.0232 in accuracy, 0.054 in recall, and 0.0218 in F1-score compared to other models. The model demonstrates high precision, strong generalization ability, and good applicability.

## Funding

no

## Conflict of Interests

The author(s) declare(s) that there is no conflict of interest regarding the publication of this paper.

## References

- [1] Liu, H. D., Sun, X. H., Li, Y. B., et al. (n.d.). A Review of Deep Learning Models for Image Classification Based on Convolutional Neural Networks. *Computer Engineering and Applications*, 1–29. Retrieved April 25, 2025, from <http://kns.cnki.net/kcms/detail/11.2127.TP.20250213.1223.013.html>
- [2] Ting, Z., Ying, S., Kang, C., et al. (2023). Hierarchical dual graph convolutional network for aspect-based sentiment analysis. *Knowledge-Based Systems*, 276.
- [3] Chai, Y. (2022). Research on Sentiment Analysis of Book Review Texts Based on LSTM and Word2vec. *Information Technology*, (07), 59–64+69.
- [4] Liu, H., Ma, X., Zhang, L. Y., et al. (2023). An aspect-level sentiment analysis model integrating matching long short-term memory network and syntactic distance. *Computer Applications*, 43(01), 45–50.
- [5] Wanli, L., & Lei, Z. (2022). Question Text Classification Method of Tourism Based on Deep Learning Model. *Wireless Communications and Mobile Computing*, 2022.
- [6] Ding, F., & Sun, X. (2022). Negative emotion opinion target extraction based on attention mechanism and BiLSTM-CRF. *Computer Science*, 49(02), 223–230.
- [7] Yingying, L., Peipei, L., & Xuegang, H. (2022). Combining context-relevant features with multi-stage attention network

- for short text classification. *Computer Speech & Language*, 71.
- [8] Liu, J., Cao, J. X., Ding, W. N., et al. (2022). Research on reservoir porosity prediction method based on bidirectional long short-term memory neural network. *Progress in Geophysics*, 37(05), 1993–2000.
- [9] (2018). Text Categorization with Improved Deep Learning Methods. *Journal of Information and Communication Convergence Engineering*, 16(2), 106–113.
- [10] Wang, R. W., Song, J. Y., & Chen, C. B. (2017). Research on the application of sentiment analysis based on Word2vec in brand cognition. *Library and Information Service*, 61(22), 6–12.
- [11] Zhang, Q., Gao, Z. M., & Liu, J. Y. (2017). Research on microblog short text classification based on Word2vec. *Information Network Security*, (01), 57–62.
- [12] Wang, H. Y., Li, J. H., & Yang, F. L. (2014). A review of support vector machine theory and algorithms. *Journal of Computer Applications*, 31(05), 1281–1286.
- [13] Hinton, G. E., Osindero, S., & Teh, Y. W. (2006). A fast learning algorithm for deep belief nets. *Neural Computation*, 18(7), 1527–1554.
- [14] Li, J. H., Ahmed, A. S., Chai, Q., et al. (2021). Emotionally charged text classification with deep learning and sentiment semantics. *Neural Computing and Applications*, 34(3), 2341–2351.