

Article

A Study on the Emotional Tendency of Aquatic Product Quality and Safety Texts Based on Emotional Dictionaries and Deep Learning

Xingxing Tong, Ming Chen *  and Guofu Feng

Key Laboratory of Fisheries Information, Ministry of Agriculture and Rural Affairs, Shanghai Ocean University, Hucheng Ring Road 999, Shanghai 201306, China; m210911534@st.shou.edu.cn (X.T.); gffeng@shou.edu.cn (G.F.)

* Correspondence: mchen@shou.edu.cn

Abstract: The issue of aquatic product quality and safety has gradually become a focal point of societal concern. Analyzing textual comments from people about aquatic products aids in promptly understanding the current sentiment landscape regarding the quality and safety of aquatic products. To address the challenge of the polysemy of modern network buzzwords in word vector representation, we construct a custom sentiment lexicon and employ the Roberta-wwm-ext model to extract semantic feature representations from comment texts. Subsequently, the obtained semantic features of words are put into a bidirectional LSTM model for sentiment classification. This paper validates the effectiveness of the proposed model in the sentiment analysis of aquatic product quality and safety texts by constructing two datasets, one for salmon and one for shrimp, sourced from comments on JD.com. Multiple comparative experiments were conducted to assess the performance of the model on these datasets. The experimental results demonstrate significant achievements using the proposed model, achieving a classification accuracy of 95.49%. This represents a notable improvement of 6.42 percentage points compared to using Word2Vec and a 2.06 percentage point improvement compared to using BERT as the word embedding model. Furthermore, it outperforms LSTM by 2.22 percentage points and textCNN by 2.86 percentage points in terms of semantic extraction models. The outstanding effectiveness of the proposed method is strongly validated by these results. It provides more accurate technical support for calculating the concentration of negative emotions using a risk assessment system in public opinion related to quality and safety.

Keywords: aquatic product safety; text classification; sentiment analysis; deep learning; BERT; bidirectional LSTM



Citation: Tong, X.; Chen, M.; Feng, G. A Study on the Emotional Tendency of Aquatic Product Quality and Safety Texts Based on Emotional Dictionaries and Deep Learning. *Appl. Sci.* **2024**, *14*, 2119. <https://doi.org/10.3390/app14052119>

Academic Editor:
Douglas O'Shaughnessy

Received: 1 February 2024
Revised: 22 February 2024
Accepted: 28 February 2024
Published: 4 March 2024



Copyright: © 2024 by the authors. Licensee MDPI, Basel, Switzerland. This article is an open access article distributed under the terms and conditions of the Creative Commons Attribution (CC BY) license (<https://creativecommons.org/licenses/by/4.0/>).

1. Introduction

In recent years, with increasing attention on food safety issues, the safety of aquatic product quality has gradually become a focal point of societal concern. Aquatic products are a crucial source of protein. Quality and safety issues with these products have a significant impact on people's health. The reputation of the food industry and related businesses is also affected by them. Therefore, research and regulation concerning the quality and safety of aquatic products have become particularly important [1].

Due to the rapid development of the internet and social media, an abundance of internet slang emerges with the rise of trending online events. Various types of information and comments about aquatic products are shared online by many users. A significant amount of internet slang is often used to express consumer sentiment toward aquatic products in these pieces of information. These emotions have a significant impact on the direction of online public opinion. Building upon the current status of research on public opinion risk assessment in food quality and safety, a comprehensive assessment system for public opinion on risks to agricultural product quality and safety was proposed by Deng Yu et al. [2] They focused on four aspects, including the concentration of negative

emotions, in establishing this system. Technical support for calculating the concentration of negative sentiments can be provided according to accurate analysis and prediction of these emotional tendencies. Consequently, this serves as a basis for the assessment of public opinion on risks related to product quality and safety.

In the field of natural language processing, sentiment analysis is a crucial research area. However, traditional sentiment analysis methods often fail to capture the semantic information and context relevance of a text adequately. Reduced accuracy and stability of sentiment analysis results can result from this limitation [3]. To overcome these challenges, researchers have increasingly turned their attention to sentiment analysis models based on deep learning.

In the past few years, research on text sentiment analysis has garnered widespread attention and significant progress in the field of natural language processing. An overview of the current status and key achievements of some relevant research is provided in the following.

Research based on traditional methods: Dictionaries and machine learning methods are primarily utilized in traditional studies on text sentiment analysis. Sentiment dictionaries, in particular, are commonly used tools that classify sentiment by associating words with emotional polarities. For instance, Liu et al. [4] proposed a sentiment lexicon called “Opinion Lexicon” for sentiment analysis of text. However, substantial human resources are required for the establishment and maintenance of sentiment lexicons, and they lack strong scalability. Moreover, conventional machine learning algorithms like Support Vector Machines (SVMs) and Naive Bayes have been extensively used in sentiment classification tasks [5]. For example, Ahmad et al. [6] utilized the SVM method for sentiment analysis on texts from Twitter. These methods excel in powerful multi-feature modeling but require manual feature engineering, making model generalization challenging.

Research based on deep learning methods: In recent years, deep learning methods have achieved significant progress in the study of text sentiment analysis. Models such as recurrent neural networks (RNNs) and long short-term memory (LSTM) are widely used to model sequential information in text. For instance, Kim et al. [7] proposed a model based on Convolutional Neural Networks (CNNs) for sentiment classification tasks. Additionally, the LSTM model was introduced by Hochreiter et al. [8] to address the challenge of capturing long-distance dependencies. These methods exhibit an excellent contextual understanding and possess the ability to automatically learn features from a text. Complex semantic structures are more easily comprehended using methods of this kind. However, these methods rely on robust word embedding models. The semantic and contextual relationships in vocabulary may be effectively captured using traditional word embedding models, such as Word2Vec, which are mostly unidirectional.

Research based on pre-trained language models: The emergence of pre-trained language models has brought new opportunities to the study of text sentiment analysis. Devlin et al. [9] introduced the BERT (Bidirectional Encoder Representations from Transformers) model, achieving excellent performance on multiple sentiment classification datasets. BERT is a pre-trained language model based on the Transformer architecture. A Transformer is an architecture that employs an attention mechanism. The central concept of a Transformer is to capture dependencies in sequential data using self-attention and positional encoding. This results in sequential data being processed more effectively compared to using the traditional recurrent neural networks (RNNs) and long short-term memory networks (LSTMs). The Transformer architecture has been widely applied in various fields. For example, in the field of computer vision (CV), Transformers are used for tasks such as image classification, object detection, and image generation [10]. In the field of speech processing, Transformers are used for speech recognition and speech synthesis [11]. Additionally, Transformers are also applied in recommendation systems, time-series analysis, and bioinformatics. BERT represents a significant application of the Transformer architecture in the field of natural language processing (NLP). The encoder part of the Transformer is utilized by BERT to learn bidirectional language representations. The innovation of BERT lies in its pre-training tasks,

which mainly include the Masked Language Model (MLM) and Next Sentence Prediction (NSP). BERT pre-trains on a large number of text data. As a result, rich language features can be learned. These features can be transferred to downstream natural language processing (NLP) tasks. Examples of such tasks include text classification, named entity recognition, and question-answering. The performance of these tasks is significantly improved by this transfer [12]. However, as a model for text sentiment analysis, BERT has high data and computational resource requirements, making it less suitable for small-scale applications.

Roberta (A Robustly Optimized BERT Pretraining Approach) was introduced by Facebook AI in 2019 as a significantly improved pre-training language model based on BERT [13]. By conducting thorough optimizations and improvements on the BERT architecture, Roberta successfully achieved superior performance. This improvement in Roberta is mainly attributed to a series of intelligent training strategies and carefully tuned hyperparameters. The model's representational and generalization capabilities are further boosted by these enhancements, leading to outstanding performance in various natural language processing tasks. On the other hand, a bidirectional long short-term memory (BiLSTM) model is a classic recurrent neural network that effectively captures the temporal dependencies in text sequences [14].

Currently, there is limited research on sentiment analysis of specific aquatic product review data, and the methods mostly used are the aforementioned general sentiment analysis methods. These methods include traditional approaches, deep learning techniques, and methods based on pre-trained language models. However, in sentiment analysis of aquatic product review datasets, these methods may encounter some disadvantages. Firstly, aquatic product reviews contain specific industry terms, domain knowledge, and the expression of sentiments. General sentiment analysis models may struggle to accurately capture the semantics and sentiments within them. Additionally, people may use various expressions to convey emotions when reviewing aquatic products, including a plethora of internet slang, making it challenging for sentiment analysis models to accurately capture all emotions. In addressing the advantages and disadvantages of the aforementioned methods, the current study further explores a combination of sentiment lexicons, pre-trained language models, and deep learning networks. A text sentiment analysis model is proposed based on the integration of a sentiment lexicon and a deep learning network. This study aims to utilize textual data related to aquatic products to build a sentiment classification model for training and evaluation.

In conclusion, sentiment analysis research encompasses traditional methods, deep learning approaches, and methods based on pre-trained language models. The modeling capabilities for sentiment classification tasks have been significantly enhanced by the advent of pre-trained language models such as BERT. Based on the above, we chose a variant of the pre-trained BERT model, Roberta-wwm-ext, as the foundational word embedding model. Roberta-wwm-ext is based on the Transformer architecture and has been pre-trained according to extensive unsupervised learning, acquiring rich language representations. It has demonstrated excellent performance across various Chinese natural language processing tasks. The capability to better comprehend contextual information in text is possessed by Roberta-wwm-ext. By considering the position of words in a sentence or text and their relationships with surrounding words, it allocates more expressive vector representations to each word. This enables the model to capture semantic information in Chinese text more effectively, leading to improved performance across various natural language processing tasks. Roberta-wwm-ext is specifically trained for the Chinese context, taking into account the characteristics and structures of the Chinese language. Therefore, superior performance in understanding semantic and syntactic structures is exhibited by it when processing Chinese text. To address the challenge of the internet slang in online platform data, which can be difficult for models to understand, we incorporated a sentiment lexicon. Combining the Roberta-wwm-ext model with a sentiment lexicon allows for more accurate extraction of semantic information from text. The current study further explores the combination of sentiment lexicons, pre-trained language models,

and deep learning networks. A text sentiment analysis model is proposed based on the integration of a sentiment lexicon and a deep learning network. The aim of this study is to build a sentiment classification model for training and evaluation using textual data related to aquatic products. The emotional analysis performance of the sentiment lexicons and deep learning models is assessed according to analysis of the experimental results. A comparison with traditional methods is conducted to validate the effectiveness and superiority of the proposed model in sentiment analysis of textual content related to the safety of aquatic products. The model provides technical support for calculating the concentration of negative emotions, thereby serving as a basis for assessing the public sentiment on risks related to the safety of aquatic products. The main contributions of the work in this paper are as follows:

1. Firstly, a dataset of consumer reviews regarding the safety of aquatic product quality was established. This dataset is a valuable resource for sentiment analysis. It contains consumer opinions concerning the safety of aquatic products. An effective basis is provided by it for assessing public sentiment and the risks associated with aquatic product quality and safety.
2. Based on Dalian University of Technology's sentiment lexicon ontology, recently emerged internet slang is incorporated into the sentiment lexicon ontology. This approach addresses issues in text sentiment recognition. These issues include rapid updates to internet slang and the presence of multiple meanings for individual words. As a result, errors in sentiment identification can be caused by these issues. The aim of this addition is to enhance the accuracy of text sentiment recognition by improving the sentiment feature extraction.
3. Considering the characteristics of textual data, a custom sentiment lexicon is combined with the Roberta-wwm-ext network for word embedding. Additionally, bidirectional LSTM networks are employed for semantic feature extraction, contributing to a comprehensive improvement in the accuracy and stability of the sentiment analysis.

2. Materials and Methods

2.1. Data Acquisition

The dataset used in this study is sourced from JD.com. JD.com is one of the largest e-commerce platforms in China with a diverse user base, representing a comprehensive spectrum of consumers. Review data from e-commerce platforms offer genuine evaluations from consumers after their purchases. These data are considered to provide the most authentic feedback on products. Analyzing reviews for a specific type of product on the entire e-commerce platform helps us understand the general public sentiment toward such products. Distinct characteristics are exhibited by this dataset compared to comment datasets on other products. Firstly, aquatic product comment datasets often involve knowledge from specialized fields such as water quality, water treatment technology, and aquaculture. Therefore, more professional vocabulary is contained in them compared to datasets for other products. Secondly, aquatic products are related to human health and safety. Therefore, comments on these products tend to focus more on aspects such as quality, safety, and freshness. These aspects may not be as prominent in comments on other products. Additionally, compared to other products, the production of aquatic products may be more susceptible to environmental factors such as water pollution and climate change. Hence, the dataset may involve more discussions on the impact of environmental factors on product quality and sustainability. Furthermore, as the dataset is sourced from an online platform, it is likely to contain a large amount of internet slang while possessing these characteristics. Based on the above, this dataset is used for the experiments in this study.

Web scraping techniques were utilized to extract reviews of aquatic products from various stores on JD.com. Specifically, we collected positive reviews, labeled as 1, and negative reviews, labeled as 0, from different stores. The main focus of this paper is two types of aquatic products: consumer reviews of fresh shrimp from different stores and

consumer reviews of salmon. The combined dataset comprises 127,349 reviews for the two products, including 86,325 positive reviews and 41,024 negative reviews.

The user reviews on JD.com may reflect the opinions and preferences of specific user groups, rather than the entire audience. Some users may prefer shopping on JD.com, while others may prefer other e-commerce platforms. Less comprehensive evaluations of products or services may result from such biases. Additionally, the quality of reviews from different users may vary, with some reviews being overly subjective or lacking detailed information. Some reviews may be fake or manipulated by competitors, affecting the credibility of the data. In the obtained review dataset, there were default comments, which are automatically generated positive reviews given by the system when users did not provide timely feedback. Additionally, duplicate comments were present in the dataset. This is likely due to certain merchants engaging in fake positive review practices, which led to repetitive standardized positive comments. These default and duplicate comments lack substantive content and do not represent genuine consumer evaluations. As a result, they potentially influence the experiments. Therefore, we removed these comments. Furthermore, some short-length comments, most of which consisted of only one or two words, were filtered out. After undergoing the aforementioned steps, we finally selected 35,000 positive comments and 35,000 negative comments, totaling 70,000 comments, as the final experimental dataset. Some examples of comments in the dataset in this article are shown in Table 1.

Table 1. Examples from experimental dataset.

Label	Comment
1	The first impression upon opening the box is that these shrimp are really large, with heads even bigger than the Qingdao large shrimp I used to buy. They also appear quite fresh, with no unusual odor. I will make a call to this store.
1	Selected from the finest Qingdao large shrimp, freshly caught, and quickly locked-in freshness within 8 min. These shrimp have a large size, thin shells, firm flesh, do not shrink when cooked, and have a sweet and savory taste. They are tender, smooth, suitable for all ages, and high in protein.
0	They are really not good, it's definitely not frozen. After eating them a few times, the head and body were separated for most of the shrimp, and for some, the meat was even scattered.
0	Don't buy this shrimp because its quality is extremely poor, some of it smells bad, and the non-odorous taste is also very unpleasant. It tastes like chewing wax and has no shrimp-like taste.

Due to the presence of numerous meaningless words and expressions in the comment dataset, after selecting the final dataset, the text data are initially processed by removing special characters and punctuation. Non-alphanumeric characters, punctuation, and other special characters are eliminated from the text. This helps reduce the data dimensions, making them easier for the model to handle. Since the text in this dataset is in Chinese, tokenization is necessary. The Jieba segmentation tool was utilized to tokenize the text, breaking down long sentences into individual words. This process provides the basic units for subsequent processing. After tokenizing the text, not all the obtained words contributed substantially to our research. For example, words like “is” and “this” have no apparent impact on textual expression but occur frequently in text. They are classified as stop words. By excluding such words, we effectively reduce the size of the vocabulary dataset, thereby achieving compression of the dimensions of the word vectors. Lowering the computational complexity of the model training and enhancing the accuracy of the text analysis are achieved through this operation.

2.2. Experimental Model

The model structure of this paper is shown in Figure 1. Firstly, we use the WordPiece tokenizer to tokenize the preprocessed input text into subwords. Then, these tokenized

subwords are fed into the Roberta-wwm-ext model, converting each subword into a dense vector representation. Roberta-wwm-ext is capable of capturing contextual information by considering the surrounding words in the text sequence, thereby generating embeddings that can represent the contextual meaning of the entire sentence. Utilizing the Roberta-wwm-ext embeddings, we obtain word vectors with context-aware capabilities. This means that each word vector not only encodes the meaning of an individual word but also considers the surrounding words and their relationships in the sentence. This context-aware ability is crucial for capturing subtle semantic and syntactic information in natural language text. As a result, more accurate downstream tasks, such as sentiment analysis, are enabled. Secondly, a self-constructed sentiment dictionary is utilized to extract the sentiment features from the text. This dictionary includes words or phrases, along with their corresponding sentiment scores or labels. For each word in the input text, we search for and calculate its sentiment score in the sentiment dictionary. The emotional polarity of each word in the text can be quantified using this method. After obtaining the output from Roberta-wwm-ext and the sentiment features extracted from the sentiment dictionary, we concatenate them. This concatenation combines the rich semantic information captured using Roberta-wwm-ext with the emotional polarity information extracted from the sentiment dictionary. The resulting feature vector contains both contextual and emotional information, providing a comprehensive representation of the input text for sentiment analysis. Then, the concatenated feature vector is put into a bidirectional LSTM (BiLSTM) network. The BiLSTM architecture allows the model to capture both the forward and backward contextual dependencies in the input sequence. LSTM units are capable of maintaining information over long sequences, making them suitable for modeling sequential data such as text. The bidirectional LSTM processes the concatenated feature vector across multiple time steps, updating the hidden states at each time step. The LSTM is enabled to capture sequential patterns and semantic information in the input text according to this process. By analyzing the hidden states at each time step, semantic features that can encode the semantics of the text are extracted. After processing the input sequence using the bidirectional LSTM, we extract the hidden state corresponding to the last time step. This final hidden state contains summary information on the entire input sequence and captures its semantic representation. The final hidden state is further processed and transformed using a fully connected layer, also known as a dense layer. Finally, the output of the fully connected layer is passed through a softmax function. The softmax function normalizes the output scores for multiple classes and generates a probability distribution for these classes. The class with the highest probability is selected as the predicted sentiment label for the input text. Using these steps, sentiment analysis can be effectively performed on the input text. This is carried out by utilizing the contextual embeddings from Roberta-wwm-ext and the sentiment features extracted from the self-constructed sentiment dictionary. Ultimately, accurate sentiment classification results are produced.

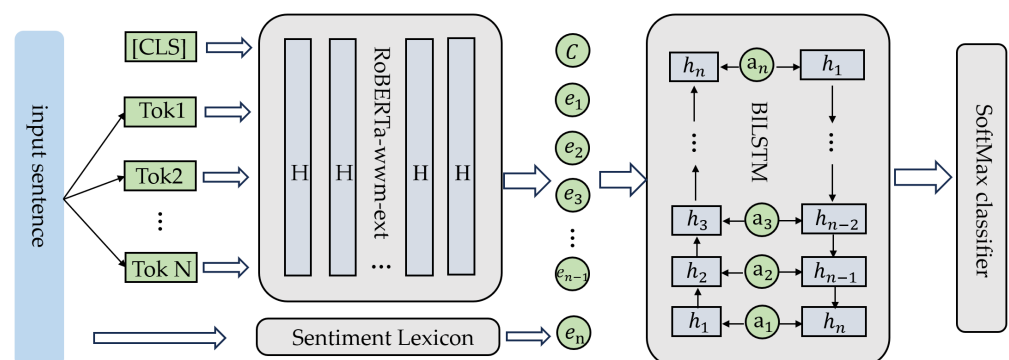


Figure 1. Based on emotional dictionary and deep learning network structure.

2.2.1. Word Embedding Layer

The word embedding layer is utilized to map words or characters from the text into a real-number vector space (embedding vector). Various models exist for the word embedding layer, such as Word2Vec, GloVe, FastText, etc. [15–17]. The word embedding layer utilized in this study is a variant of the BERT model known as chinese-roberta-wwm-ext (<https://huggingface.co/hfl/chinese-roberta-wwm-ext/tree/main>, accessed on 1 May 2023). BERT (Bidirectional Encoder Representations from Transformers) is a deep learning model based on the Transformer architecture, proposed by Google in 2018. As illustrated in Figure 2, the structure of the BERT model primarily consists of the encoder part of the Transformer. The encoder employs a self-attention mechanism to interact with each element in the input sequence, calculating relevance scores between each element and others [18]. Allowing the encoder to focus on essential contextual information from each element enhances our overall understanding of the semantic and contextual relationships within the input sequence. Repeated stacking of the encoder layers contributes to the extraction of deep-level feature representations from the input sequence. In BERT, multiple such encoder layers are stacked, forming a deep network structure. In the original BERT paper, the authors utilized 12-layer and 24-layer Transformer encoders to build models, with parameter totals of 110 million and 340 million, respectively. Compared to traditional word embedding models, BERT exhibits contextual awareness, allowing us to better capture the semantic and contextual relationships of vocabulary. This is attributed to BERT's bidirectional encoding approach, which considers both the left and right context of words simultaneously, while models like Word2Vec and GloVe are predominantly unidirectional. Thus, BERT is enabled to comprehensively understand the context of the vocabulary. Moreover, rich language knowledge is acquired by BERT using pretraining tasks like masked language modeling and next sentence prediction, enhancing its generalization capabilities for downstream tasks. In contrast, Word2Vec and GloVe primarily focus on local word co-occurrence patterns.

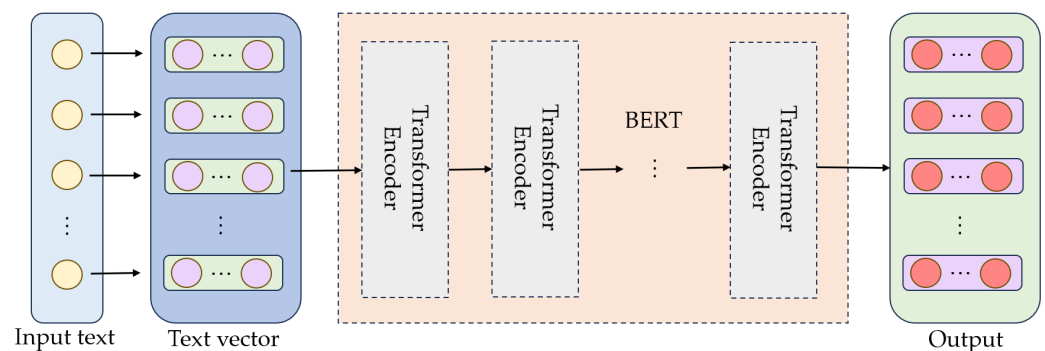


Figure 2. BERT structure diagram.

The input design of the BERT model is cleverly crafted. As shown in Figure 3, the BERT input is constructed into a linear sequence, dividing two sentences using a separator. Two special tokens are added at the beginning and end of the sequence, ensuring accurate capture of the overall context. In BERT, there are three key embeddings at the word level. The first key embedding is positional information embedding, which is used to encode the position of words in a sequence. This is important because word order is a critical feature in natural language processing. Word embedding is the second key embedding. The semantic information of the words can be captured in this step. The final embedding is sentence embedding. When the training data consist of two or more sentences, the overall context for each word is provided by this embedding. These three embeddings are stacked together, forming the input structure of BERT.

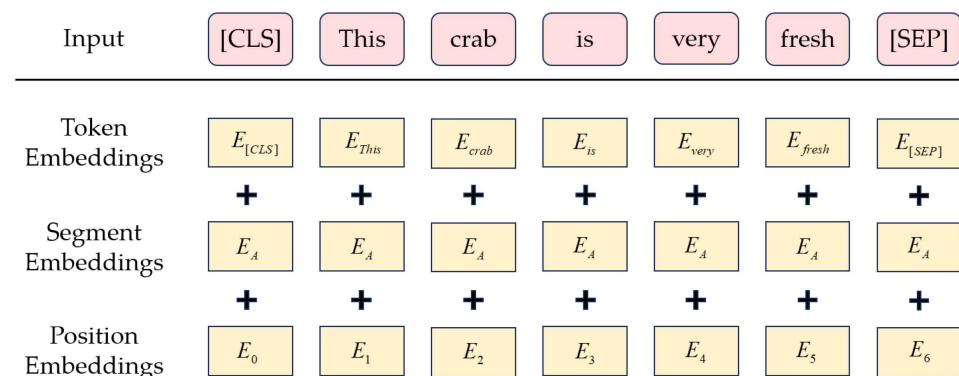


Figure 3. Input of BERT.

In terms of the model input, BERT converts each character in the text into a one-dimensional vector by querying a character vector table, serving as the fundamental input for the model. Simultaneously, the model's output consists of vector representations for each character, incorporating fused semantic information from the entire text. The input encompasses not only character vectors but also text vectors and positional vectors. The text vectors are learned during the model training process, capturing the global semantic information from the text and merging it with the semantic information on individual characters. Positional vectors, on the other hand, are utilized to distinguish the semantic information differences carried by characters at different positions. When handling English vocabulary, BERT employs the WordPiece method, which involves a finer granularity of tokenization to better capture the semantic units. For Chinese, the current BERT models do not tokenize the input text but treat individual characters directly as the basic units forming the text. These design choices and strategies collectively form the foundation of BERT's deep understanding of natural language.

Roberta adopts a Transformer architecture, similar to BERT, but with adjustments to the details. One of the most notable changes is the elimination of sentence-level segment embeddings, opting for a simpler approach without gaps. This modification enables Roberta to better capture the long-distance dependency relationships in the text, enhancing the model's expressive capacity. In terms of the training data, a larger dataset is utilized by Roberta compared to BERT, providing the model with more extensive and richer language knowledge. By learning from a greater variety of language phenomena, Roberta demonstrates superior performance in downstream tasks. Regarding sequence length handling, Roberta allows for the processing of longer sequences, which is crucial for tasks that involve more contextual information in a single input. The ability to handle longer sequences contributes to enhancing the model's understanding of complex contexts. Furthermore, Roberta introduces a dynamic masking approach, enhancing the model's generalization by dynamically selecting masked words in each iteration. This flexible masking strategy aids the model in better adapting to various textual inputs. In summary, these advantages enable Roberta to achieve superior performance across various natural language processing tasks, providing robust support for the model research in this paper.

In the original BERT and Roberta models, as they were trained on English language corpora, they adopted the "WordPiece" tokenization method for tokenization. This method involves a finer level of tokenization than words. For instance, the word "emotion" is tokenized into three tokens: "e", "##mo", "##tion". This tokenization method effectively reduces the size of the pretraining vocabulary (represented as $|V|$). It mitigates the issue of out-of-vocabulary words by combining word roots, allowing for a more effective representation of a broader vocabulary. However, for Chinese corpora, the concept of word roots, as seen in English, does not exist in the same manner. To address this issue, the model introduces a technique called "Whole Word Masking" (wwm) [19]. Taking the word "emotion" as an example, using the whole word masking method, it is tokenized into "e", "##mo", and "##tion". Then, all three tokens are masked. The handling of Chinese

tokenization issues is improved by this approach, enhancing the model's performance on Chinese tasks.

2.2.2. Emotional Feature Construction

This study employs an emotion lexicon to assign corresponding emotional features Sw_i to each word w_i in the text. Based on the emotional lexicon ontology from Dalian University of Technology [20], we not only expanded upon its existing content but also incorporated 300 recent slang terms that have gained popularity among young people on the internet. This expansion aims to enhance the lexicon's suitability for text-based comments on the safety of aquatic products. The emotion lexicon ontology in this study exclusively retains emotion words representing positive and negative sentiments. A portion of the self-constructed emotion lexicon ontology presented in the article is illustrated in Table 2. Based on emotional intensity, we categorize emotion words into five classes: 1, 3, 5, 7, and 9. For positive polarity, the emotional features are determined based on the emotional intensity, while for negative polarity, the emotional features are obtained by multiplying the emotional intensity by -1 . Standardization procedures have been applied to normalizing the emotional features. Ultimately, the formulated expressions for emotional components are represented by Equations (1)–(3).

$$ET(w_i) = \begin{cases} Sw_i, & w_i \in EL \\ 0, & w_i \notin EL \end{cases} \quad (1)$$

$$Sw_i = \frac{\tilde{Q}_i - \mu_Q}{\sigma_Q} \quad (2)$$

$$\tilde{Q}_i = \begin{cases} Q_i, & L_i = 1 \\ -Q_i, & L_i = 2 \end{cases} \quad (3)$$

Table 2. Example of self-built emotion lexicon.

Words	Sentiment Classification	Strength	Polarity
dirty and messy	NN	7	2
breed	NN	5	2
as the acme of perfection	PC	7	1
pollution-free	PH	9	1
conspicuous bag	PB	3	1
be far ahead	PH	9	1
make a call	PH	9	1
blue thin shiitake mushroom	NB	7	2

In the equations, w_i represents a word, Sw_i represents the emotional features that need to be concatenated to the word vector, EL represents the emotion lexicon, Q_i represents emotional intensity, L_i represents emotional polarity, and $\tilde{Q}_i$ represents the distinguished polarity of the emotional intensity. Equation (2) is used for normalization of the sentiment feature values in the sentiment dictionary. We adopt Z-score normalization to transform the sentiment feature values into a distribution with a mean of 0 and a standard deviation of 1, where μ_Q is the mean of the sentiment features and σ_Q is the standard deviation of the sentiment features. The absolute value of the obtained emotional feature indicates the intensity of the emotion associated with the word. The emotional inclination of a word is represented by the positive or negative sign of the feature value; a negative feature value suggests a negative emotional inclination, while a positive value indicates a positive emotional inclination. After obtaining the emotional feature $ET(w_i)$ for each word using the emotion lexicon, it is concatenated with the word vectors e_i obtained from the word

embedding layer to obtain the input for the final semantic extraction layer. The formula for vector concatenation is shown in Equation (4):

$$e'_i = [e_i, ET(w_i)] \quad (4)$$

2.2.3. Semantic Extraction Layer

The semantic extraction layer of our model utilizes the BiLSTM (bidirectional long short-term memory) model. LSTM, which stands for long short-term memory, was introduced by Hochteiter et al. as a type of recurrent neural network (RNN). It addresses the issues of gradient vanishing and exploding that are present in traditional RNN models during training.

The structure of the LSTM module is shown in Figure 4. LSTM consists primarily of the cell state C_t , temporary cell state $\tilde{C}_t$, hidden state h_t , forget gate f_t , memory gate i_t , and output gate o_t . The operation of LSTM can be summarized as follows: at each time step, the forget gate f_t , input gate i_t , and output gate o_t control maintenance of the information in the cell state. This facilitates the transmission of useful information for subsequent time steps while discarding irrelevant information. The computation of these gates involves the hidden state h_{t-1} from the previous time step and the current input x_t . The formulas for calculating these gates are expressed in Equations (5)–(10):

$$f_t = \sigma(W_f \cdot [h_{t-1}, x_t] + b_f) \quad (5)$$

$$i_t = \sigma(W_i \cdot [h_{t-1}, x_t] + b_i) \quad (6)$$

$$\tilde{C}_t = \tanh(W_C \cdot [h_{t-1}, x_t] + b_C) \quad (7)$$

$$C_t = f_t * C_{t-1} + i_t * \tilde{C}_t \quad (8)$$

$$o_t = \sigma(W_o \cdot [h_{t-1}, x_t] + b_o) \quad (9)$$

$$h_t = o_t * \tanh(C_t) \quad (10)$$

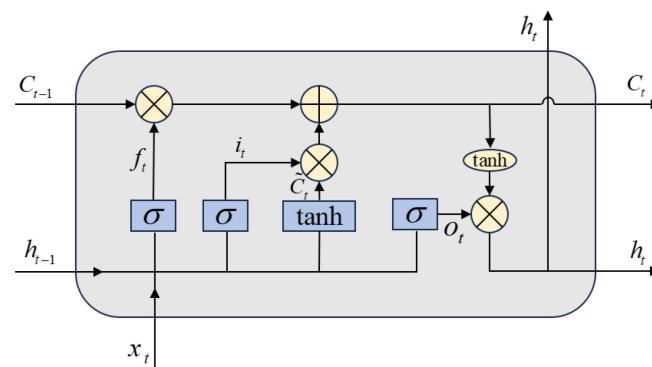


Figure 4. LSTM module structure.

In the equations, x_t represents the input at time step t ; h_{t-1} denotes the hidden state from the previous time step; W_f , W_i , W_C , and W_o represent the weight matrices; b_f , b_i , b_C , and b_o are the bias terms introduced during the training process; and σ represents the sigmoid activation function.

One issue with unidirectional LSTM is its inability to encode information from later time steps. The issue is addressed by BiLSTM, which combines forward LSTM with backward LSTM. Thus, this enables BiLSTM to consider both the forward and backward directions of textual information simultaneously. This bidirectional nature contributes to improving its performance in sentiment analysis tasks, as sentiment analysis is often influenced by various factors in the text.

The model selects the feature vector $E \in R^n$, which is the concatenation of the output of the last layer trained using Roberta-wwm-ext and the feature values extracted from

the emotion lexicon. This feature vector is then added to weight $W_e \in R^{e \times n}$ to obtain $A = (a_1, a_2, a_3 \dots a_n)$ as the input for the semantic extraction layer. The calculation is expressed in Equation (11):

$$a_t = g_1(W_e E_t + b_a) \quad (11)$$

In the equation, $1 \leq n \leq t$, where n is the dimensionality of the feature vector obtained after the sentence has been trained using Roberta-wwm-ext. g_1 is the activation function *Sigmoid*, and b_a is the bias vector.

The vector input to the hidden layer undergoes forward and backward LSTM operations, resulting in forward and backward outputs, respectively. The final output vector v_i of the BiLSTM is obtained by concatenating the hidden state result $\vec{h}_t$ from the last time step of the forward LSTM and the hidden state result $\overleftarrow{h}_t$ from the last time step of the backward LSTM. The calculation is expressed in Equation (8):

$$v_i = \vec{h}_t + \overleftarrow{h}_t \quad (12)$$

2.2.4. Sentiment Classification Layer

The output from the bidirectional LSTM is fed into a fully connected layer, which utilizes the ReLU activation function. Finally, the output of the fully connected layer is passed through the softmax function for classification, yielding the ultimate classification result. The calculation is represented in Equation (13):

$$P(\gamma|H, W_s, b_s) = \text{softmax}(W_s \cdot V + b_s) \quad (13)$$

3. Experiment and Analysis

This section is dedicated to implementing the proposed neural network model to achieve accurate recognition of the sentiment tendencies in texts related to the safety of aquatic products. The motivation behind this effort stems from the crucial role of aquatic food products in food safety. Analyzing user expressions of sentiment toward aquatic food products aids in our understanding of some of the current issues related to the safety of aquatic product quality. However, due to the rapid evolution of the internet, modern text sentiment analysis models face challenges. Therefore, our research constructed a custom lexicon of internet slang expressions to more accurately capture the emotional expressions of internet slang in text. Subsequently, according to meticulous parameter adjustments, we ensured that the model could better adapt to diverse contexts and sentiment tendencies when dealing with internet slang. To comprehensively assess our optimization strategies, multiple evaluation metrics were employed in this paper to thoroughly examine the performance of the model from various dimensions.

3.1. Experimental Environment

The experiments in this study were conducted on a device with an Intel(R) Core(TM) i9-10900K CPU @ 3.70 GHz, 64 GB RAM, and an NVIDIA GeForce RTX 3090 graphics card. The operating system used was Ubuntu 18.04, with Python 3.8.10, CUDA 11.0.221, and PyTorch 1.10.0 used for building the proposed network model and the comparison models used in the experiments.

3.2. Evaluation Metrics

To assess the performance of the model on the dataset, this study selected three commonly used evaluation metrics in classification models: accuracy, recall, and F1 score. Their calculation formulas are shown in Equations (14)–(16):

$$\text{accuracy} = \frac{TP + TN}{TP + FN + FP + TN} \quad (14)$$

$$recall = \frac{TP}{TP + FN} \quad (15)$$

$$F1 = \frac{2TP}{2TP + FP + FN} \quad (16)$$

In the formulas, TP represents true positives, TN represents true negatives, FN represents false negatives, and FP represents false positives. TP is the number of positive samples correctly predicted by the model, TN is the number of negative samples correctly predicted, FN is the number of positive samples incorrectly predicted as negative, and FP is the number of negative samples incorrectly predicted as positive.

3.3. Experimental Parameter Setting

To fully demonstrate the effectiveness of the proposed method in our research, comparative experiments involving various widely used sentiment analysis models were conducted by us. This series of comparative experiments aims to intuitively validate the reliability and proficiency of the proposed method in terms of the network performance. In the experiments, the word embeddings generated by the pre-trained model Roberta-wwm-ext were utilized by us to endow the model with powerful semantic representation capabilities. During the model training process, we opted for the adaptive moment estimation (Adam) algorithm as the optimizer to expedite the convergence of training and enhance the overall training efficiency. Detailed hyperparameter tuning was conducted to obtain the best experimental results. These parameters included the number of training epochs, batch size, learning rate, the number of hidden units in LSTM (HiddenSize), and random dropout of hidden units during training. The objective of this process is to ensure that the model achieves optimal performance on the dataset used in this study.

The term “epoch” refers to the number of times the entire training dataset is used to train the model. In each epoch, the model traverses the entire training dataset, adjusting its weights and parameters using an optimization algorithm (such as gradient descent) to minimize the training loss. Typically, training a model involves multiple epochs to ensure that the model adequately learns the features of the data and improves its performance. However, using too many epochs may lead to overfitting, where the model performs well on the training data but poorly on unseen data. Taking into account the dataset used in this study and the selected model, the number of epochs for the experimental model in this study was set to 6, aiming for a balance between capturing the patterns in the data and avoiding overfitting.

The “batch size (Batchsize)” defines the number of samples used to update the model weights in each training step. When training neural networks, the entire training dataset is typically divided into small batches, with each batch containing a certain number of training samples. Batchsize for the model in this study was set to 16. This decision was based on considering the computational resources available in the experimental environment, as well as the size and characteristics of the dataset used in this study.

The learning rate is used to control the magnitude of the parameter adjustments during each update step of the model. It determines the size of the step taken when updating the parameters along the gradient direction. For the model in this study, the learning rate was set to 1×10^{-5} .

The number of hidden units (Hidden Size) and the random dropout of hidden units (Dropout) are crucial parameters affecting the performance and training process of long short-term memory networks. For the model in this study, Hidden Size was set to 320, and Dropout was set to 0.5.

While BiLSTM models offer advantages in capturing the temporal dependencies in text sequences, they are not without challenges. One of the primary issues associated with BiLSTM models is the risk of overfitting, especially when dealing with small datasets. Overfitting occurs when the model learns the training data too well, including the noise, leading to poor generalization on unseen data. We carefully selected the optimal Dropout parameter according to multiple experiments to mitigate the risk of overfitting. Addition-

ally, BiLSTM models are prone to vanishing and exploding gradient problems, although to a lesser extent than traditional RNNs. These issues can hinder the models' ability to learn long-range dependencies effectively. We addressed these issues by employing gradient clipping and meticulously initializing the model parameters. Furthermore, BiLSTM models can be computationally intensive due to the bidirectional processing of sequences, which may pose challenges in terms of the training time and resource requirements.

3.4. Comparing Model Settings

The method of combining sentiment lexicons with deep learning has been studied for other types of review data. However, similar research has not been conducted on aquatic product review datasets. The sentiment lexicon used in this paper is independently constructed and differs from those used in other studies. Regarding the settings of comparative experiments, this paper uses generic text sentiment analysis methods and divides the experiments into two parts: text word vector extraction and text semantic extraction. The comparative experiments aim to demonstrate the superiority of the proposed word embedding model and the advantages of the text semantic extraction model used in this paper. Therefore, we select commonly used word embedding models such as Word2Vec, BERT, and Roberta-wwm-ext. Additionally, generic text sentiment analysis models with good performance in neutral text sentiment research are chosen, including RNN, textCNN, LSTM networks, and BiLSTM. Different combinations of word embedding models and semantic extraction models are used to obtain different comparative experimental models. Through the design of comparative experiments, we aim to demonstrate the advantages of our method for the aquatic product review data used in this paper. Therefore, in this paper, our model is labeled "our method", and the comparative experimental models are as follows:

1. Word2Vec-RNN [21]: The model utilizes Word2Vec as the word embedding model to obtain word vectors. Subsequently, it employs an RNN for feature training and classification.
2. Word2Vec-LSTM [22]: Word2Vec is utilized by the model as the word embedding model to obtain the word vectors. It employs LSTM for feature training and classification, with LSTM configured as two layers and the number of units in the hidden layer set to 320.
3. Word2Vec-BiLSTM [23]: The model utilizes Word2Vec as the word embedding model to obtain word vectors and represent the text features. Subsequently, it uses BiLSTM to process these features for classification.
4. Bert-RNN [24]: The input text is first encoded using the BERT model to obtain deep semantic representations of the text. Subsequently, these semantic representations serve as inputs to an RNN, which is responsible for processing the serialized semantic information and ultimately outputting the predicted sentiment category.
5. Bert-textCNN [25]: The model first encodes the input text using the BERT model to obtain contextual representation vectors for each word. These vectors are then used as inputs to the textCNN model. The textCNN model, through operations like convolutional and pooling layers, effectively extracts the key features from the text. The extracted features are put into a fully connected layer, and the output of the fully connected layer represents the final sentiment classification result.
6. Bert-LSTM [26]: BERT is first utilized by the model for feature extraction from the text. Subsequently, LSTM is employed to further process these features, and finally, a softmax classifier is used to classify the processed features.
7. Bert-BiLSTM [27]: The model utilizes a pre-trained BERT model by putting the text sequence into BERT to obtain contextual representations for each word. The output from BERT is then fed into a BiLSTM layer to further capture the sequential information in the text. Following the BiLSTM layer, a fully connected layer is added, and the final input to the softmax layer is used to classify the text into different sentiment categories.

8. Roberta-wwm-ext-RNN: The text sequence is put into the pre-trained Roberta_wwm-ext model to obtain contextual representations for each word. Subsequently, an RNN is employed for feature training and classification.
9. Roberta-wwm-ext-textCNN: A model utilizing the pre-trained Roberta-wwm-ext model. The text sequence is put into this model to obtain contextual representations for each word. Subsequently, these vectors serve as inputs to the textCNN model, which extracts features according to operations like convolutional and pooling layers. Finally, the extracted features are put into a fully connected layer to output the ultimate sentiment classification result.
10. Roberta-wwm-ext-LSTM: The pre-trained Roberta-wwm-ext model is utilized by the network to input text sequences, obtaining contextual representations for each word. Subsequently, LSTM is employed to further process these features, and finally, a softmax classifier is used to classify the processed features.
11. Roberta-wwm-ext-BiLSTM: The word embedding model of the network utilizes Roberta-wwm-ext to input the text sequences and obtain contextually relevant representations for each word. Subsequently, BiLSTM is employed to extract the features, and the hyperparameters of the semantic feature extraction model are consistent with those of the proposed method.

3.5. Network Training and Comparative Analysis of the Results

The results of the model established in this paper, as well as the various comparative models set out in this paper, for the dataset are presented in Table 3.

Table 3. Comparison of results between sentiment analysis networks.

Networks	Accuracy (%)	Recall (%)	F1 (%)
Word2vec-RNN	87.60	86.58	87.01
Word2vec-LSTM	88.75	88.48	88.62
Word2vec-BiLSTM	89.07	88.83	88.96
Bert-RNN	92.59	92.25	92.48
Bert-textCNN	91.70	91.23	91.51
Bert-LSTM	92.63	92.27	92.46
Bert-BiLSTM	93.43	92.86	93.09
Roberta-wwm-ext-RNN	93.05	92.32	92.18
Roberta-wwm-ext-textCNN	92.63	92.01	92.44
Roberta-wwm-ext-LSTM	93.27	92.62	92.97
Roberta-wwm-ext-BiLSTM	94.05	93.36	93.72
Our method	95.49	94.53	95.18

We compared our method with model 3 and model 7 to evaluate the effectiveness of our proposed approach using a custom sentiment lexicon and Roberta-wwm-ext as the word embedding model. Table 3 shows that our proposed embedding model outperforms the use of Word2Vec by 6.42 percentage points and improves upon using BERT as the word embedding model by 2.06 percentage points. The text representation based on Word2Vec shows the lowest performance across all metrics. This is mainly attributed to Word2Vec's excellent representation of semantic relationships between words in the text but neglect of the issues related to polysemy in different contexts and long-distance semantic associations.

To validate the effectiveness of bidirectional LSTM in feature extraction, we compared models 8–11. In this comparison, all the models use Roberta-wwm-ext as the word embedding model. In terms of the classification accuracy, bidirectional LSTM outperforms LSTM by 0.68 percentage points and surpasses textCNN by 1.42 percentage points. This demonstrates its superior performance and effectively verifies its ability to address the issues of vanishing and exploding gradients traditionally associated with RNN models during training.

We compare model 11 with the method proposed in this paper to demonstrate the effectiveness of the sentiment dictionary designed and used in this study. Compared to the

model used in this paper, model 11 eliminates the feature of using a sentiment dictionary to enhance the word vectors. The rest of the network structure is consistent with the network structure of this paper. The results in the table show that the accuracy of the model without the sentiment dictionary is 1.44 percentage points lower than that of the model with the sentiment dictionary. In terms of recall and F1 score, using a sentiment dictionary to extract features from some difficult-to-judge online popular words shows higher values. From this, we can conclude that on the same aquatic product review dataset, the model's performance can be improved, and the accuracy of sentiment orientation judgment can be enhanced by the addition of a sentiment dictionary.

To visually observe the relative performance advantages of the proposed model compared to other networks on the dataset, as well as the convergence of the model, the changes in accuracy with respect to the step length for each model on this dataset are presented in Figure 5a,b.

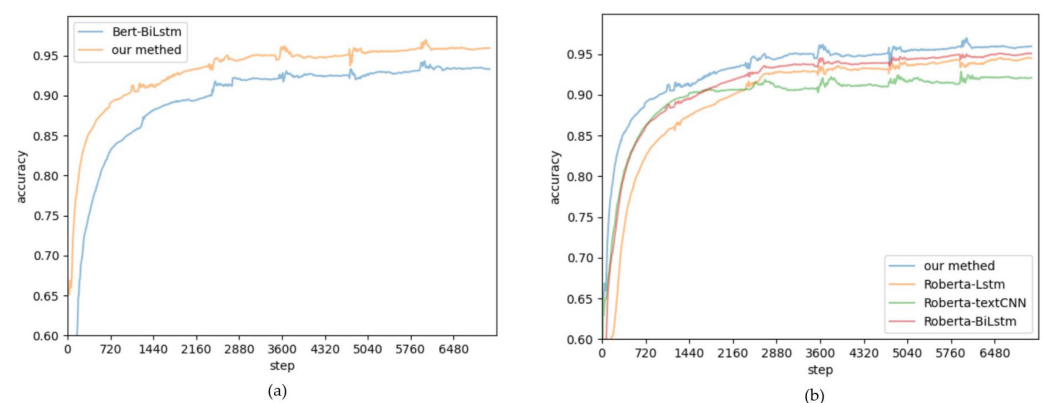


Figure 5. Comparison of training accuracy curves of network models. (a) Training accuracy curves for different word embedding models; (b) training accuracy curves for training different bias analysis models in this article.

From Figure 5a, it can be observed that the proposed word embedding model outperforms BERT. The proposed word embedding model exhibits higher accuracy and a faster growth rate in the first epoch. As the number of training iterations increases, the model's accuracy gradually stabilizes. From Figure 5b, it can be observed that bidirectional LSTM has advantages compared to other text sentiment analysis models. Although textCNN converges quickly, its accuracy is lower compared to the LSTM network. Bidirectional LSTM exhibits higher accuracy and a faster convergence speed than unidirectional LSTM. This is because bidirectional LSTM considers both past and future information simultaneously. A more comprehensive understanding of the contextual relationships in the sequence is provided, and long-term dependencies are captured. From Figure 5, it can be observed that the accuracy of each network exhibits slight oscillations at the beginning of each epoch, before stabilizing. This could be attributed to the choice of hyperparameters such as the learning rate, optimizer momentum, and the randomness in the training data. As training progresses, the model gradually learns the patterns in the data and adjusts its weights more finely, and the network's output becomes more stable.

In addition to these commonly used evaluation indicators, a rigorous validation process was employed to ensure the robustness of our model's performance. A k-fold cross-validation technique was utilized, where the dataset was divided into k equal-sized subsets. In each iteration, one subset was used as the validation set while the remaining subsets were used for training. This process was repeated k times, with each subset serving as the validation set once. This approach allowed us to evaluate the model's performance across different subsets of the data, reducing the likelihood of overfitting and providing a more comprehensive assessment of the model's generalizability.

4. Conclusions

This paper initially examines the prevalent methods in general sentiment analysis. The specific challenges encountered by these methods in the specialized domain of textual sentiment analysis for aquatic product quality and safety are then highlighted. General sentiment analysis algorithms are unable to accurately identify the emotional tendencies of specialized domain vocabulary. Furthermore, the emotional tendencies within online popular vocabulary, derived from the rapid development of the internet, are unable to be accurately recognized using the traditional generic sentiment analysis methods. This study aims to tackle the aforementioned issues. It conducts a thorough comparative analysis of the current sentiment analysis models and selects Roberta-wwm-ext as the base word embedding model. Additionally, bidirectional LSTM is employed as the base semantic extraction model. Building on this foundation, we introduce a sentiment dictionary as an auxiliary tool. This sentiment dictionary is based on the Dalian University of Technology Chinese Sentiment Lexicon Ontology, to which we added 300 online popular vocabulary words and manually annotated their emotional tendencies. A sentiment dictionary is established in this paper to extract the emotional characteristics of vocabulary words. These emotional features are then concatenated with the base word embedding model. As a result, the proposed word embedding model can more accurately extract word vectors, which facilitates semantic feature extraction. We take into account the distribution of textual data on aquatic product quality and safety. Consequently, consumer reviews of aquatic products from real stores on the JD platform are selected as the base data. Various data processing operations were conducted, resulting in the establishment of an aquatic product quality and safety review dataset with 70,000 real comments and emotional labels. In future research, we will incorporate data from more platforms. This dataset can be made available to a wider audience for exploring sentiment analysis in the context of aquatic product quality and safety. The experimental analysis in this paper demonstrates that the model proposed in this paper exhibits clear superiority in performance based on the data presented in this paper. The proposed network in this paper achieves an accuracy rate of 95.49%, showing an improvement of 6.42 percentage points compared to using Word2Vec as the word embedding model. Furthermore, the use of a sentiment dictionary to extract the emotional features resulted in a 1.44 percentage point increase in accuracy compared to not using a sentiment dictionary. The role of the sentiment dictionary in the word embedding model is fully demonstrated by this result. Moreover, this method can provide technical support for calculating the concentration of negative emotions. As a result, a basis for risk assessment according to public opinion related to aquatic product quality and safety can be provided by this method.

Author Contributions: Conceptualization, M.C. and G.F.; methodology, X.T.; software, X.T.; validation, M.C. and G.F.; formal analysis, M.C. and X.T.; resources, M.C.; data curation, M.C. and X.T.; writing—original draft preparation, M.C., X.T. and G.F.; writing—review and editing, M.C., X.T. and G.F.; supervision, M.C. All authors have read and agreed to the published version of the manuscript.

Funding: This research was funded by the Research and Development Planning in Key Areas of Guangdong Province, No. 2021B0202070001.

Institutional Review Board Statement: Not applicable.

Informed Consent Statement: Not applicable.

Data Availability Statement: The experimental data related to this paper can be requested from the authors by email. If a researcher is in need of the dataset, they can email mchen@shou.edu.cn.

Conflicts of Interest: The authors declare no conflicts of interest.

References

- Guo, T.; Zhang, Y.; Ye, M.; Ke, H.; Zhang, L.; Xiao, Q.; Wu, Z. A review of recent advances in quality and safety control of aquatic products in China. *Meat Res.* **2019**, *33*, 67–73.
- Deng, Y.; Qian, Y.; Li, X.; Lian, Y. Public Opinion Risk Composite Index of Agro-Food Quality and Safety. *J. Food Saf. Qual.* **2018**, *9*, 4734–4741.
- Nanli, Z.; Ping, Z.; Weiguo, L.I.; Meng, C. Sentiment Analysis: A Literature Review. In Proceedings of the 2012 International Symposium on Management of Technology (ISMOT), Hangzhou, China, 8–9 November 2012; pp. 572–576.
- Bing, L. *Sentiment Analysis and Opinion Mining (Synthesis Lectures on Human Language Technologies)*; University of Illinois: Chicago, IL, USA, 2012.
- Medhat, W.; Hassan, A.; Korashy, H. Sentiment Analysis Algorithms and Applications: A Survey. *Ain Shams Eng. J.* **2014**, *5*, 1093–1113. [[CrossRef](#)]
- Ahmad, M.; Aftab, S.; Ali, I. Sentiment Analysis of Tweets Using Svm. *Int. J. Comput. Appl.* **2017**, *177*, 25–29. [[CrossRef](#)]
- Kim, Y. Convolutional Neural Networks for Sentence Classification. *arXiv* **2014**, arXiv:1408.5882.
- Hochreiter, S.; Schmidhuber, J. Long Short-Term Memory. *Neural Comput.* **1997**, *9*, 1735–1780. [[CrossRef](#)] [[PubMed](#)]
- Devlin, J.; Chang, M.-W.; Lee, K.; Toutanova, K. BERT: Pre-Training of Deep Bidirectional Transformers for Language Understanding. In Proceedings of the naacL-HLT, Minneapolis, MN, USA, 2–7 June 2019.
- Dosovitskiy, A.; Beyer, L.; Kolesnikov, A.; Weissenborn, D.; Zhai, X.; Unterthiner, T.; Dehghani, M.; Minderer, M.; Heigold, G.; Gelly, S.; et al. An Image Is Worth 16 × 16 Words: Transformers for Image Recognition at Scale. *arXiv* **2021**, arXiv:2010.11929.
- Dong, L.; Xu, S.; Xu, B. Speech-Transformer: A No-Recurrence Sequence-to-Sequence Model for Speech Recognition. In Proceedings of the 2018 IEEE International Conference on Acoustics, Speech and Signal Processing (ICASSP), Calgary, AB, Canada, 15–20 April 2018; pp. 5884–5888.
- Ravichandiran, S. *Getting Started with Google BERT: Build and Train State-of-the-Art Natural Language Processing Models Using BERT*; Packt Publishing Ltd.: Birmingham, UK, 2021.
- Liu, Y.; Ott, M.; Goyal, N.; Du, J.; Joshi, M.; Chen, D.; Levy, O.; Lewis, M.; Zettlemoyer, L.; Stoyanov, V. RoBERTa: A Robustly Optimized BERT Pretraining Approach. *arXiv* **2019**, arXiv:1907.11692.
- Zhou, P.; Shi, W.; Tian, J.; Qi, Z.; Li, B.; Hao, H.; Xu, B. 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), Berlin, Germany, 7–12 August 2016; pp. 207–212.
- Rong, X. Word2vec Parameter Learning Explained. *arXiv* **2016**, arXiv:1411.2738.
- Pennington, J.; Socher, R.; Manning, C.D. Glove: Global Vectors for Word Representation. In Proceedings of the 2014 Conference on Empirical Methods in Natural Language Processing (EMNLP), Doha, Qatar, 25–29 October 2014; pp. 1532–1543.
- Athiwaratkun, B.; Wilson, A.G.; Anandkumar, A. Probabilistic FastText for Multi-Sense Word Embeddings. *arXiv* **2018**, arXiv:1806.02901.
- Vaswani, A.; Shazeer, N.; Parmar, N.; Uszkoreit, J.; Jones, L.; Gomez, A.N.; Kaiser, Ł.; Polosukhin, I. Attention Is All You Need. *arXiv* **2017**, arXiv:1706.03762.
- Xu, Z. RoBERTa-Wwm-Ext Fine-Tuning for Chinese Text Classification. *arXiv* **2021**, arXiv:2103.00492.
- Xu, L.; Lin, H.; Pan, L.; Ren, H.; Chen, J. Constructing the Affective Lexicon Ontology. *J. China Soc. Sci. Tech. Inf.* **2008**, *27*, 180–185.
- Kurniasari, L.; Setyanto, A. Sentiment Analysis Using Recurrent Neural Network. *Proc. J. Phys. Conf. Ser.* **2020**, *1471*, 012018. [[CrossRef](#)]
- Muhammad, P.F.; Kusumaningrum, R.; Wibowo, A. Sentiment Analysis Using Word2vec and Long Short-Term Memory (LSTM) for Indonesian Hotel Reviews. *Procedia Comput. Sci.* **2021**, *179*, 728–735. [[CrossRef](#)]
- Vimali, J.S.; Murugan, S. A Text Based Sentiment Analysis Model Using Bi-Directional Lstm Networks. In Proceedings of the 2021 6th International Conference on Communication and Electronics Systems (ICCES), Coimbatore, India, 8–10 July 2021; pp. 1652–1658.
- Topbaş, A.; Jamil, A.; Hameed, A.A.; Ali, S.M.; Bazai, S.; Shah, S.A. Sentiment Analysis for COVID-19 Tweets Using Recurrent Neural Network (RNN) and Bidirectional Encoder Representations (BERT) Models. In Proceedings of the 2021 International Conference on Computing, Electronic and Electrical Engineering (ICE Cube), Quetta, Pakistan, 26–27 October 2021; pp. 1–6.
- Shan, D.; Li, H. Group Emotion Recognition for Weibo Topics Based on BERT with TextCNN. *Am. J. Inf. Sci. Technol.* **2023**, *7*, 95–100. [[CrossRef](#)]
- Tseng, H.-T.; Zheng, Y.-Z.; Hsieh, C.-C. Sentiment Analysis Using BERT, LSTM, and Cognitive Dictionary. In Proceedings of the 2022 IEEE International Conference on Consumer Electronics-Taiwan, Taipei, Taiwan, 6–8 July 2022; pp. 163–164.
- Cai, R.; Qin, B.; Chen, Y.; Zhang, L.; Yang, R.; Chen, S.; Wang, W. Sentiment Analysis about Investors and Consumers in Energy Market Based on BERT-BiLSTM. *IEEE Access* **2020**, *8*, 171408–171415. [[CrossRef](#)]

Disclaimer/Publisher’s Note: The statements, opinions and data contained in all publications are solely those of the individual author(s) and contributor(s) and not of MDPI and/or the editor(s). MDPI and/or the editor(s) disclaim responsibility for any injury to people or property resulting from any ideas, methods, instructions or products referred to in the content.