

CONSUMER SENTIMENT ANALYSIS IN THE MARKET THROUGH ONLINE PRODUCT REVIEWS USING PHOBERT-BILSTM ON TIKI DATA

Duc Khuat Van, Duy Hang Nguyen

Posts and Telecommunications Institute of Technology

Abstract: Nowadays, with the rapid growth of social media and the surge in e-commerce, customers' expression of opinions through online product reviews has become an indispensable part of the landscape. Sentiment analysis plays a pivotal role in automatically extracting subjective information, particularly customers' emotions and opinions, from these reviews. Recognizing the importance of this, this article introduces a comprehensive system that starts with data collection from the Tiki e-commerce website. Subsequently, it employs two predictive models: the BERT-BiLSTM model, which combines Bidirectional Encoder Representations from Transformers (BERT) with Bidirectional Long Short-Term Memory (BiLSTM), and the PhoBERT-BiLSTM model, an advanced variant that integrates PhoBERT embeddings with BiLSTM. These models aim to accurately forecast emotional trends within user product reviews. The experimental results demonstrate that both models achieve remarkable performance metrics, with the PhoBERT-BiLSTM model achieving higher accuracy, precision, recall, and F1-score of 91.95%, 92.16%, 91.95%, and 91.99%, respectively, compared to the BERT-BiLSTM model. Consequently, the findings of this study provide a precise means of predicting consumers' emotional trends, particularly concerning specific product lines.

Keywords: Sentiment Analysis, Online Product Review Classification, E-commerce, Deep Learning.

1. INTRODUCTION

The increasing popularity of social media and smart mobile devices has transformed the way we communicate online. Users nowadays not only consume information from providers but also generate and share their own opinions, often through posting textual reviews about products or their experiences. These opinions significantly influence consumers' future purchasing decisions and have implications for businesses and individuals. Today's consumers commonly use search tools to find product reviews or usage experiences before making a purchase.

However, online information is often complex and lacks organization, making the process of searching for reviews a time-consuming and multifaceted task. This poses a challenge for consumers [1]. Sentiment analysis, often referred to as opinion mining, is the process of extracting customer opinions to evaluate products. It falls within the field of machine learning. In the current context of increasing online data, it is considered crucial as user-generated text is abundant online. Sentiment analysis pertains to researching users' thoughts and perceptions of a product. The importance of sentiment analysis or opinion mining is growing each day as data continues to expand. Machines need to be reliable and effective in interpreting and understanding human emotions and feelings [2]. Recently, there has been significant attention to predicting user emotions. To analyze the content of social media, Yoo proposed a system for predicting user emotions. To represent text data, the work utilized a two-dimensional representation of word2vec. The emotion analysis model was constructed using Convolutional Neural Networks. Validated using the Sentiment140 dataset, containing 800,000 positive and 800,000 negative documents, the proposed model outperforms basic methods, specifically Naïve Bayes, Support Vector Machine (SVM), and Random Forest (RF) [3]. In another article, sentiment analysis (SA) was conducted on Amazon product review data. The RFSVM method, a combination of RF and SVM, was employed to leverage the capabilities of both classifiers. Performance evaluation metrics including accuracy, recall, F-Measure, and precision were utilized to compare the proposed method with basic approaches, specifically RF and SVM. When using a dataset consisting of 500 positive and 500 negative samples, the test results demonstrated that RFSVM outperformed the basic methods in all three performance metrics [4]. Furthermore, in a study, multi-class and binary classification for Amazon mobile phones was conducted using a supervised machine learning algorithm. Bi-BiLSTM with GloVe embeddings and joint learning embeddings were applied in this research. Additionally, the BERT model was utilized. The BERT model achieved outstanding results in both multi-class and binary classification, with accuracy rates of 94% and 98%, respectively. On the other hand, BiLSTM with joint learning embeddings also yielded excellent results, with an accuracy of 93% for multi-class classification and 97% for binary classification [5]. Moreover, Based on statements from investors and consumers published on the

Corresponding author: Duc Khuat Van,

Email: duckv@ptit.edu.vn

Manuscript received: 04/2024, revised: 05/2024, accepted: 06/2024.

Chinese internet, this article applies the BERT, LSTM, and BERT-LSTM models to predict emotional trends. Among these models, BERT-LSTM achieves the highest accuracy and recall, demonstrating its predictive capabilities for both overall and core samples. The results highlight the advantages of the combined BERT-BiLSTM prediction model in sentiment analysis. It can accurately forecast the emotional trends of internet users during major events, providing technical support for energy market decision-making. Notably, the BERT-BiLSTM model outperforms the BERT (0.8559 and 0.5576) and LSTM models (0.7775 and 0.0747) with accuracy and recall scores of 0.8620 and 0.7078, respectively. This research can accurately predict the emotional trends of internet users during social events, aiding in capturing energy market trends [6]. In this paper, we have constructed a diagnostic model for the specific problem of user sentiment analysis, with our primary contributions listed as follows:

- Development of a data collection system related to quality reviews and product service in the clothing category from the TIKI e-commerce platform.
- Utilization of deep learning (BERT-BiLSTM and PhoBERT-BiLSTM) to evaluate the diagnostic performance of our model. The combined approach of the PhoBERT and BiLSMT model achieved high diagnostic efficiency with an accuracy of 91.95

The rest of the work is presented as follows. Section II introduces developing a data collection system. Section III describes the data preprocessing of our work. Section IV presents the methodology, result, and discussion introduced in Section V. The last section, Section VI, is the conclusion.

II. DEVELOPING A DATA COLLECTION SYSTEM

This research focuses on the process of collecting information about products and user reviews on the Tiki.vn ecommerce platform. Our system is divided into three key components. The first part involves gathering a list of Product IDs for all products in the 'clothing' category on Tiki.vn. Subsequently, in the second part, we utilize APIs to extract detailed information about each product, including product names, prices, and other relevant details. The third part is dedicated to collecting comment data from the list of Product IDs. We continue to use APIs to collect user comments and store them in our system. This data collection system allows us to aggregate essential information about products and user feedback on Tiki.vn, serving as a valuable database for future research and analysis. After completing the online data collection process from the Tiki data gathering system, we acquired a dataset containing product reviews, which encompassed 12,662 reviews within the clothing category. This dataset comprises productID, idcomments, reviewText, reviewerID, Sumary, overall, reviewTime, unixReviewTime, and customername, as described in Table 1. We selected the review text for experimentation from the 'reviewText' field.

TABLE I: Properties of the dataset

Property	Description
productID	The product identifier code
idcomments	A unique code for the comment

reviewText	Detailed content of the review or comment about the product.
reviewerID	A unique identifier for the person writing the Comment
Sumary	Summary of the review's content.
overall	The overall rating of the product
reviewTime	Time of the review (Raw time)
unixReviewTime	Time of the review (Unix time)
customer_name	The name of the customer

III. DATA PREPROCESSING

Typically, raw data contains a lot of unnecessary information. These factors directly impact the performance of machine learning or deep learning models. Therefore, we have removed elements such as punctuation, spelling errors, typographical mistakes, and unwanted characters from our dataset.

A. Punctuation Marks Removal

The removal of punctuation marks from the data is performed as they do not carry meaningful information for classification tasks, simplifying the process.

B. Eliminating Unnecessary Characters and Symbols.

The removal of unnecessary characters and symbols was executed as part of our dataset preprocessing, aiming to enhance model performance. As online product reviews were gathered, the dataset inevitably incorporated a variety of superfluous characters, including '@', '=', '&', as well as emoticons like smiley and angry faces, vv. These characters and symbols, due to their lack of relevance, possessed the potential to introduce confusion into our models. To ensure optimal model performance, the dataset was standardized through the elimination of these characters and symbols.

C. Editing spelling errors, typographical mistakes, and standardizing text to lowercase

The role of correcting spelling errors, typographical mistakes, and standardizing text to lowercase is deemed crucial in the processing of natural language and deep learning. Ensuring that spelling errors are rectified within the input text guarantees that no spelling inaccuracies are present, thus assuring a correct understanding of the context. Simultaneously, unifying typographical formats and converting text to lowercase aids in eliminating inconsistencies in text formatting, facilitating the model's ease of learning and context comprehension. These operations are carried out to provide clean and uniform input data for deep learning models, thereby enhancing their performance in comprehending and processing natural language, spanning from error correction to precise semantic predictions.

D. Labeling the data

To establish the fundamental truth, three labels were assigned to the dataset, specifically encompassing three states: positive, negative, and neutral, assigned to each review text based on its overall rating (ranging from 1 to 5). In this process, reviews with overall ratings of 1–2 were classified as negative and assigned the label 1, whereas reviews with overall ratings of 4–5 were

considered positive and assigned the label 0. The remaining reviews were designated as neutral and assigned the label 2 [7].

TABLE II: Describes the labeling process

Sentence	Sentiment	Label
"Ao đẹp lắm, chất lượng tuyệt vời".	Positive	0
Móng mặc xấu	Negative	1
Với giá tiền đó thì ok	Neutral	2

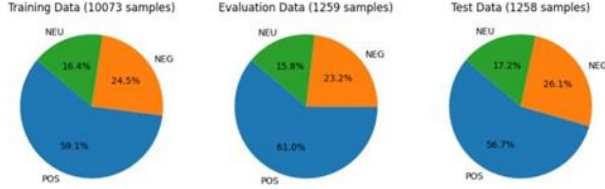


Fig. 1: Dataset Split

E. Dataset Split

Dataset Split is considered a crucial task because the proportions of the partitions impact the evaluation of deep learning model performance. Our data was divided into training, validation, and testing datasets. In the sentiment analysis case, our training, validation, and testing datasets contain 8814 (70%), 1259 (10%), and 2516(20%) instances, respectively.

IV. METHODOLOGY

Upon the completion of the data collection and data pre-processing phases, the construction of a predictive model was initiated using data from three distinct datasets, namely the training, validation, and test datasets. The predictive model was executed through a three-step process shown in Fig 2 as follows: Step 1 - Model Fitting Construction. Step 2 - Model Evaluation Construction. Step 3 - Model Estimation Construction.

A. BERT

A significant breakthrough in the field of natural language processing (NLP) is represented by BERT. Pre-trained language models utilize word context, learning word occurrences and representations from unlabeled training data. An exemplary instance is BERT, one of the leading deep-learning models for sentiment analysis. BERT is pre-trained to consider word context from both directions. It extracts richer contextual features from a sequence, potentially improving performance for various NLP tasks. BERT is trained using a masked language model (MLM), where a randomly selected word in a sentence is masked and replaced with the [MASK] token. The model then predicts the masked word using context from both sides. In addition to MLM, BERT has another objective, predicting the next sentence, and enabling fine-tuning for specific tasks without major model adjustments. The model can perform sentence pair classification to determine the semantic relationship between them. The utilization of the Transformer architecture enables BERT to comprehensively understand language by focusing on the semantics and context of words from the very beginning. Individual words and sentences are effectively represented,

consolidating language knowledge from extensive datasets. This facilitates the surpassing of the limitations associated with traditional Word Embedding techniques, as BERT comprehends multifaceted contexts, learns from new data, and enhances performance across numerous NLP applications. BERT has revolutionized our comprehension and processing of natural language [8-9-10]. The overall architecture of the BERT model is shown in Fig. 3.

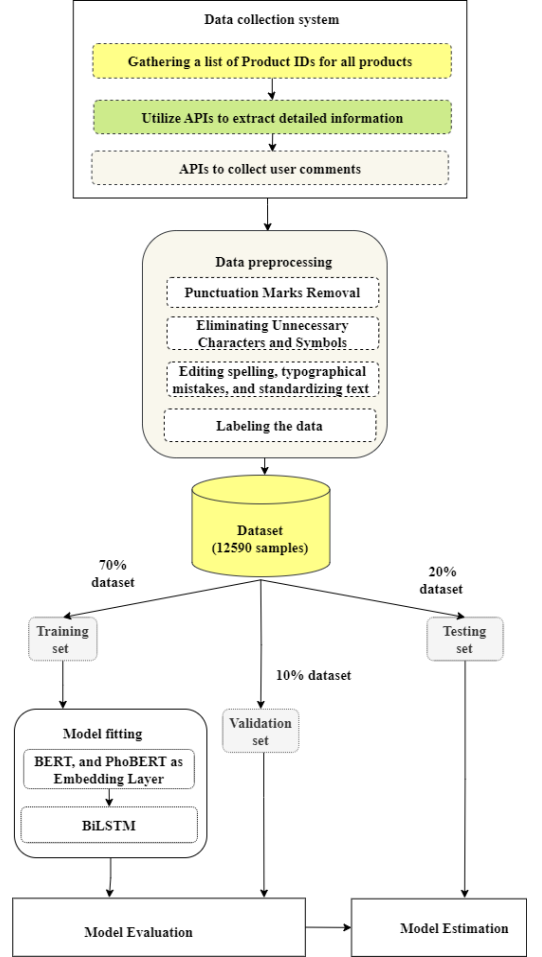


Fig.2: Methodology

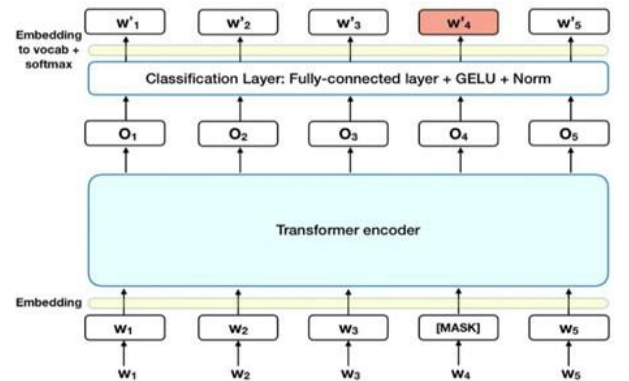


Fig. 3: The overall architecture of BERT

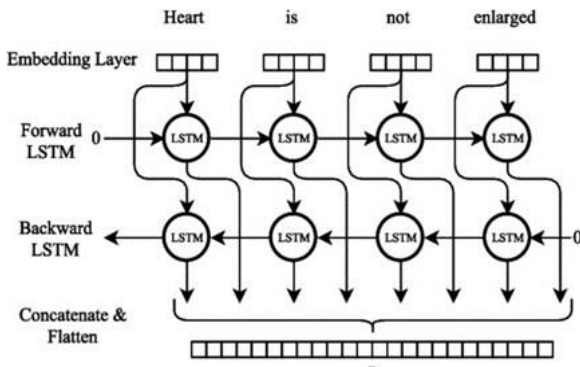


Fig.4: Overall architecture of BiLSTM

B. PhoBERT

PhoBERT is a state-of-the-art pre-trained language model tailored specifically for the Vietnamese language. Developed by the VinAI Research team, PhoBERT is based on the BERT (Bidirectional Encoder Representations from Transformers) architecture, which allows it to understand the context of words in a bidirectional manner. This model is designed to handle the unique complexities of Vietnamese, a tonal language with rich morphology. By training on a large corpus of Vietnamese text, PhoBERT significantly enhances the performance of various natural language processing tasks such as sentiment analysis, named entity recognition, and machine translation, making it an invaluable tool for advancing Vietnamese language processing [17,18].

C. BiLSTM

The issues of gradient explosion and vanishing gradient in RNN are addressed by the bidirectional BiLSTM model, known as BiLSTM. An enhanced version of RNN referred to as BiLSTM, is utilized for handling sequences of varying lengths and for mitigating the problem of information loss in recurrent neural networks. BiLSTM is composed of key components, including the cell state, temporary cell state, hidden state, forget gate, memory gate, and output gate. The operation of BiLSTM is divided into three primary stages: the forget stage, the selective memory stage, and the output stage. During the forget stage, the forget gate controls the retention of important information while the discarding of unimportant data. The selective memory stage is dedicated to "memorizing" essential information and filtering out non-essential data. In the end, the output stage determines which information is designated as the output [11-12-13]. The processing of sequence data is executed efficiently by the BiLSTM model, as depicted in Fig 4.

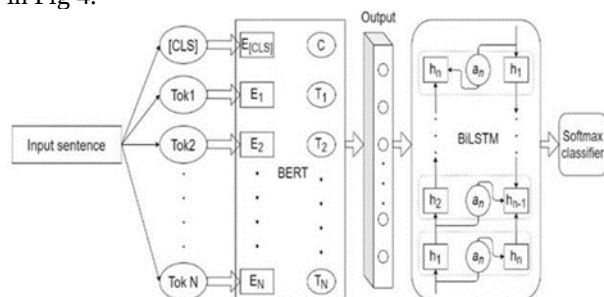


Fig.5: The overall architecture of BERT-BiLSTM

D. BERT-BiLSTM

Based on BERT and BiLSTM, a combined approach called BERT-BiLSTM has been employed to construct a model and predict consumer sentiment based on online product reviews. The structure of this method is depicted below. BERT-BiLSTM utilizes BERT as an upstream source of knowledge and BiLSTM as a downstream source of knowledge. BERT is capable of learning statistical characteristics of nearby words, while BiLSTM can capture contextual information. This aligns with the way human language operates, as fundamental grammar often relies on statistical patterns and specific meanings within context. Consequently, BERT-BiLSTM has the ability to analyze consumer sentiment based on their product reviews [6,14,15,16]. Fig. 5 shows the overall architecture of the BERT-BiLSTM model

E. phoBERT-BiLSTM

The combination of PhoBERT and BiLSTM models offers a powerful approach for customer sentiment analysis in product reviews. PhoBERT, a pre-trained language model specifically designed for Vietnamese, effectively captures the contextual meaning of words in a bidirectional manner. When integrated with a BiLSTM, which excels at learning long-term dependencies in sequential data, this hybrid model enhances the accuracy of sentiment classification. By leveraging PhoBERT's contextual understanding and BiLSTM's sequential processing capabilities, the model can more accurately classify customer sentiments, providing valuable insights into consumer opinions and improving the overall understanding of customer feedback in Vietnamese [19].

F. Model Fitting

After the data preprocessing process, three essential datasets are obtained, including the training dataset, the validation dataset, and the test dataset. During this phase, the training dataset is provided to the BERT model to conduct training and extract preliminary word vectors. Subsequently, these preliminary word vectors are input into the BiLSTM model, which is a variation of the LSTM network, to execute further analysis and processing. Ultimately, the output results from the BiLSTM model are passed through a Softmax layer to generate the final matrix of sentiment orientation classification.

G. Model Evaluation

The validation dataset is employed after the model fitting step to adjust the model and select crucial hyperparameters, such as learning rate, the number of LSTM layers, embedding size, and so on. The primary objective of using the validation dataset is to examine the model's performance on data it has not been trained on and assess parameter changes. This serves to enhance the model and ensure its effective operation on real-world data when applied to specific tasks.

H. Model Estimation

Following the conclusion of the model evaluation phase, the test dataset is utilized to assess the overall performance. Based on BERT and BiLSTM, a combined approach called BERT-BiLSTM has been employed to construct a model and predict consumer sentiment based on online product reviews. The structure of this method is

depicted below. BERT-BiLSTM utilizes BERT as an upstream source of knowledge and BiLSTM as a downstream source of knowledge. BERT is capable of learning statistical characteristics of nearby words, while BiLSTM can capture contextual information. This aligns with the way human language operates, as fundamental grammar often relies on statistical patterns and specific meanings within context. Consequently, BERT-BiLSTM has the ability to analyze consumer sentiment based on their product reviews [6- 14-15-16]. Fig. 5 shows the overall architecture of the BERT- BiLSTM model. The test dataset is an independent dataset that the model has not encountered during the training and prior evaluation processes. The primary objective is to examine the model's overall capability in predicting, classifying, or executing specific tasks on new and unfamiliar data. The results from the test dataset serve to evaluate the model's general proficiency and ensure that it operates effectively when applied to real-world scenarios.

I. Evaluation Metrics

To validate the classification performance of our models, we have applied several evaluation matrices namely Accuracy, Precision, Recall, and F1 Score. We have generated the confusion matrix for both of the models and calculated the value of the matrices by using the following equations.

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (1)$$

$$\text{Precision} = \frac{TP}{TP + FP} \quad (2)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (3)$$

$$\text{F1-score} = 2 \cdot \frac{\text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}} \quad (4)$$

Where: TP= True Positive TN= True Negative FP= False Positive FN= False Negative.

V. RESULT AND DISCUSSION

In this phase, an analysis was conducted to evaluate the performance of the proposed models. Result values were recorded with up to two decimal places. The analysis process was carried out in three steps. Firstly, learning curves were generated to analyze the training and validation performance, aiding in the identification of potential issues related to underfitting or overfitting. Subsequently, predictions were made on the test dataset to assess the model's real-world effectiveness. To create the learning curve, a chart representing the loss function's changes across each epoch was drawn by us. Fig 6 and Fig 8 illustrate the loss function curve during the training process of the sentiment analysis model. it can be easily observed that the training and validation losses converged significantly and did not increase further after the eighth epoch. This is a clear indication that our model was effectively learned. After training the PhoBERT-BiLSTM and BERT-BiLSTM models on the dataset, the results obtained are described in Figures 7 and 9, respectively. Both figures indicate that the best weights during training, which yield the highest accuracy and better generalization (avoiding overfitting), are achieved

at the 8th epoch and the 4th epoch, respectively. Since we only save the best model based on the validation loss, the weights of the models corresponding to the 8th and 4th epochs achieve accuracies of 92.95% and 90.62%, respectively. After selecting the best models in the 8th and 3rd epochs respectively for the PhoBERT-BiLSTM and BERT-BiLSTM models, they were evaluated on the test dataset. The results demonstrate that our models achieved high accuracy and performed well on critical evaluation metrics. Detailed results are presented in Table 3. The evaluation results on the test dataset are highly encouraging. Specifically, the PhoBERT-BiLSTM model achieved impressive accuracy, recall, precision, and F1- score of 91.95%, 91.95%, 92.16%, and 91.99% respectively, while the BERT-BiLSTM model achieved notable accuracy, recall, precision, and F1-score of 90.62%, 90.63%, 90.65%, and 90.63% respectively. This indicates the models' capability in effectively classify and predict consumer sentiments in the market. Furthermore, the results reflect the model's generalization ability as there is no evidence of overfitting. The balance of high accuracy, recall, precision, and F1-score demonstrates the efficient performance of the models on the test dataset.

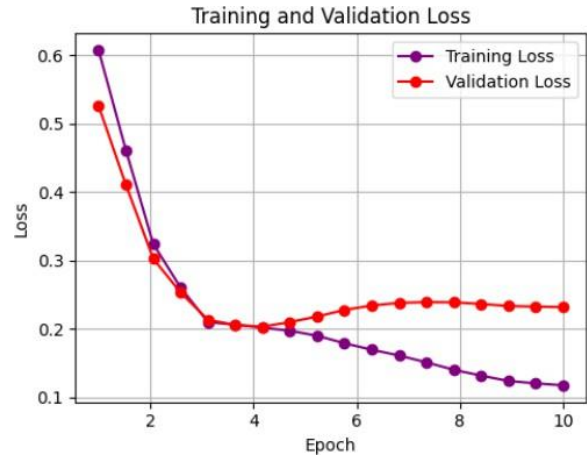


Fig.6: Validation Vs. Training Loss Graph For The PhoBERT- BiLSTM Model

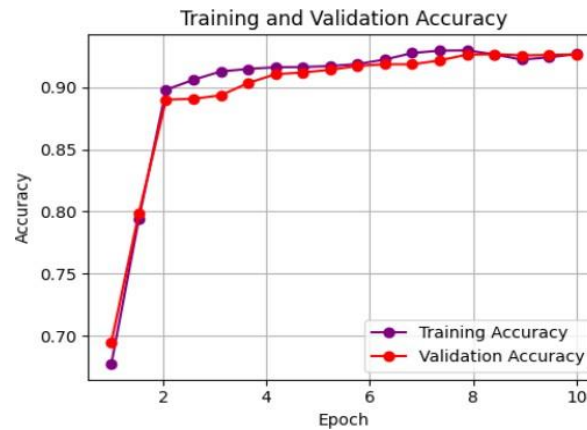


Fig.7: Validation Vs. Training Accuracy Graph For The PhoBERT- BiLSTM Model

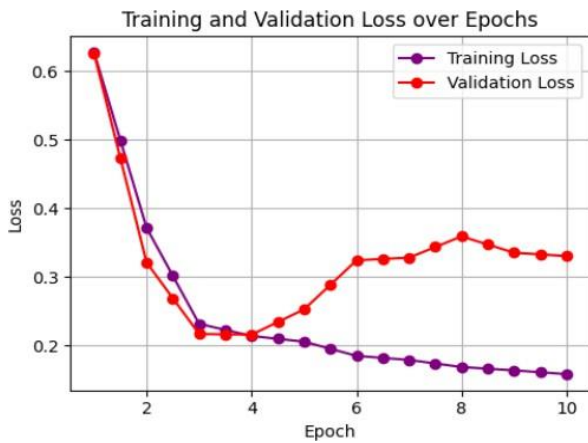


Fig.8: Validation Vs. Training Loss Graph For The BERT-BiLSTM Model

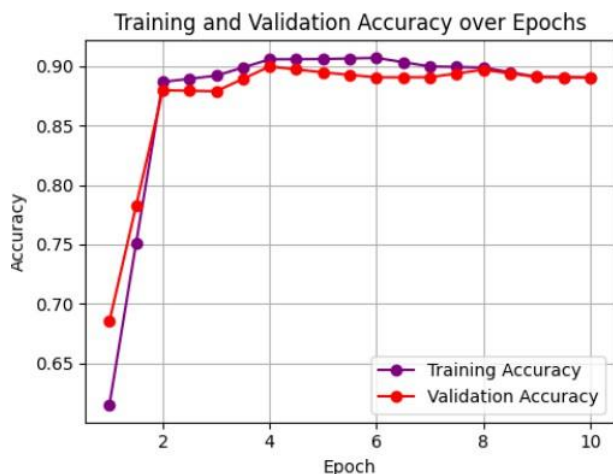


Fig.9: Validation Vs. Training Accuracy Graph For The BERT- BiLSTM Model

TABLE III: Performance of the BERT-BiLSTM and PhoBERT-BiLSTM model on the testing dataset

Model	Acc (%)	Precision (%)	Recall (%)	F1-score (%)
BERT-BiLSTM	90.62	90.65	90.63	90.63
PhoBERT-BiLSTM	91.95	92.16	91.95	91.99

VI. CONCLUSION

In this study, a comprehensive method is proposed, encompassing data collection to the specific application of the deep learning model for predicting the quality of online products and classifying them into three categories based on customer reviews. A data collection system was established from the TIKI e-commerce platform. Subsequently, multiple data pre-processing steps were carried out to clean and prepare the data, feature extraction was performed to generate numerical features from text data. Following that, the combination of BERT-BiLSTM and PhoBERT-BiLSTM models were constructed to make predictions, using accuracy, precision, recall, and F1-score as evaluation metrics. The experimental results demonstrate that the combination of these models achieves remarkable performance, particularly the PhoBERT-BiLSTM model, which attains

an accuracy of 91.95%. In the future, we hope to integrate more data sources and compare them with various state-of-the-art NLP techniques.

ACKNOWLEDGMENT

VINIF partly supports this work. Khuat Van Duc was funded by the Master, Ph.D. Scholarship Programme of Vingroup Innovation Foundation (VINIF), code VINIF.2023.ThS.037.

REFERENCES

- [1] Lin, H.-C. K., Wang, T.-H., Lin, G.-C., Cheng, S.- C., Chen, H.-R., "&" Huang, Y.-M. (2020). Applying sentiment analysis to automatically classify consumer comments concerning marketing 4Cs aspects. *Applied Soft Computing*, 97, 106755. doi:10.1016/j.asoc.2020.106755
- [2] Zhao, W., Guan, Z., Chen, L., He, X., Cai, D., Wang, B., & Wang, Q. (2018). Weakly-Supervised Deep Embedding for Product Review Sentiment Analysis. *IEEE Transactions on Knowledge and Data Engineering*, 30(1), 185–197. doi:10.1109/tkde.2017.2756658
- [3] Yoo SY, Song JI, Jeong OR. Social media contents based sentiment analysis and prediction system. *Expert Syst Appl*. 2018;105:102–11.
- [4] Al Amrani Y, Lazaar M, El Kadiri KE. Random forest and support vector machine based hybrid approach to sentiment analysis. *Procedia Comput Sci*. 2018;127:511–20.
- [5] AlQahtani, Arwa S. M., Product Sentiment Analysis for Amazon Re- views (2021). *International Journal of Computer Science & Information Technology (IJCSIT)* Vol 13, No 3, June 2021, Available at SSRN: <https://ssrn.com/abstract=3886135>.
- [6] Cai, R., Qin, B., Chen, Y., Zhang, L., Yang, R., Chen, S., & Wang, W. (2020). Sentiment Analysis About Investors and Consumers in Energy Market Based on BERT-BiLSTM. *IEEE Access*, 8, 171408–171415. doi:10.1109/access.2020.3024750
- [7] Rintyarna, B. S., Samo, R., & Fatichah, C. (2019). Evaluating the performance of sentence-level features and domain-sensitive features of product reviews on supervised sentiment analysis tasks. *Journal of Big Data*, 6(1). doi:10.1186/s40537- 019-0246-8
- [8] Xie, S., Cao, J., Wu, Z., Liu, K., Tao, X., & Xie, H. (2020). Sentiment Analysis of Chinese E- commerce Reviews Based on BERT. 2020 IEEE 18th International Conference on Industrial Informatics (INDIN). doi:10.1109/indin45582.2020.94421
- [9] A. Vaswani, N. Shazeer, N. Parmar, J. Uszkoreit, L. Jones, A. N. Gomez, Ł. Kaiser, and I. Polosukhin, "Attention is all you need," in *Advances in neural information processing systems*, 2017, pp. 5998–6008
- [10] Z. Zhang, X. Han, Z. Liu, X. Jiang, M. Sun, and Q. Liu, "Ernie: Enhanced language representation with informative entities," *arXiv preprint arXiv:1905.07129*, 2019.
- [11] S. Hochreiter, "The vanishing gradient problem during learning recurrent neural nets and problem solutions," *Int. J. Uncertainty, Fuzziness Knowl.- Based Syst.*, vol. 06, no. 02, pp. 107–116, Apr. 1998
- [12] Xu, G., Meng, Y., Qiu, X., Yu, Z., & Wu, X. (2019). Sentiment Analysis of Comment Texts Based on BiLSTM. *IEEE Access*, 7, 51522–51532. doi:10.1109/access.2019.2909919
- [13] Long, F., Zhou, K., & Ou, W. (2019). Sentiment Analysis of Text Based on Bidirectional LSTM With Multi-Head Attention. *IEEE Access*, 7, 141960–141969. doi:10.1109/access.2019.2942614
- [14] Ge H, Zheng S, Wang Q. Based BERT-BiLSTM- ATT model of commodity commentary on The emotional tendency analysis. In 2021 IEEE 4th International Conference on Big Data and Artificial Intelligence 2021 (pp. 130-133). IEEE.

- [15] Mouthami, K., S. Anandamurugan, and S. Ayyasamy. "BERT-BiLSTM- BiGRU-CRF: Ensemble Multi Models Learning for Product Review Sentiment Analysis." 2022 6th International Conference on Electronics, Communication and Aerospace Technology. IEEE, 2022.
- [16] Ge H, Zheng S, Wang Q. Based BERT-BiLSTM- ATT model of commodity commentary on The emotional tendency analysis. In 2021 IEEE 4th International Conference on Big Data and Artificial Intelligence (pp. 130-133). IEEE.
- [17] Truong, T. L., Le, H. L., & Le-Dang, T. P. (2020, November). Senti- ment analysis implementing BERT-based pre-trained language model for Vietnamese. In 2020 7th NAFOSTED Conference on Information and Computer Science (NICS) (pp. 362-367). IEEE.
- [18] Tran, H. V., Bui, V. T., Do, D. T., & Nguyen, V. V. (2021, November). Combining PhoBERT and SentiWordNet for Vietnamese Sentiment Analysis. In 2021 13th International Conference on Knowledge and Systems Engineering (KSE) (pp. 1-5). IEEE.
- [19] [Hung, B.T. and Huy, T.Q., 2022. Named Entity Recognition Based on Combining Pretrained Transformer Model and Deep Learning. In Artificial Intelligence and Sustainable Computing: Proceedings of ICSISCT 2021 (pp. 311-320). Singapore: Springer Nature Singapore.

PHÂN TÍCH TÂM LÝ NGƯỜI TIÊU DÙNG TRÊN THỊ TRƯỜNG THÔNG QUA ĐÁNH GIÁ SẢN PHẨM TRỰC TUYẾN BẰNG CÁCH SỬ DỤNG PHOBERT-BILSTM TRÊN DỮ LIỆU TIKI

Tóm tắt: Hiện nay, với sự phát triển nhanh chóng của mạng xã hội và sự bùng nổ của thương mại điện tử, việc khách hàng thể hiện ý kiến thông qua đánh giá sản phẩm trực tuyến đã trở thành một phần không thể thiếu của cảnh quan. Phân tích cảm xúc đóng vai trò then chốt trong việc tự động rút trích thông tin chủ quan, đặc biệt là cảm xúc và ý kiến của khách hàng từ những đánh giá này. Nhận thức về tầm quan trọng của vấn đề này, bài viết này giới thiệu một hệ thống toàn diện bắt đầu từ việc thu thập dữ liệu từ trang web thương mại điện tử Tiki. Sau đó, nó áp dụng hai mô hình dự đoán: mô hình BERT-BiLSTM, kết hợp Biểu diễn Bộ mã hóa song hướng từ Transformers (BERT) với Bộ nhớ Ngắn hạn Dài - Hồng trạng thái hai chiều (BiLSTM), và mô hình PhoBERT-BiLSTM, một biến thể tiên tiến tích hợp các nhúng PhoBERT với BiLSTM. Những mô hình này nhằm mục đích dự đoán chính xác xu hướng cảm xúc trong đánh giá sản phẩm của người dùng. Kết quả thử nghiệm cho thấy cả hai mô hình đều đạt được các chỉ số hiệu suất đáng chú ý, với mô hình PhoBERT-BiLSTM đạt được độ chính xác, độ chính xác, thu hồi và điểm F1 cao hơn lần lượt là 91,95%, 92,16%, 91,95% và 91,99% so với mô hình BERT-BiLSTM. Do đó, các kết quả của nghiên cứu này cung cấp một phương tiện chính xác để dự đoán xu hướng cảm xúc của người tiêu dùng, đặc biệt là đối với các dòng sản phẩm cụ thể.

Từ khóa: Phân tích cảm xúc, Phân loại đánh giá sản phẩm trực tuyến, Thương mại điện tử, Học sâu.



Duc Khuat Van, Received a Bachelor of Information Technology from the Posts and Telecommunications Institute of Technology (PTIT) in 2022, and a Master's degree in Computer

Science from PTIT. Currently, he is a lecturer in the Data Engineering department at PTIT. His research areas include Artificial Intelligence (AI), Bioinformatics, Network Optimization, and Sentiment Analysis.

Email: duckv@ptit.edu.vn



Duy Hang Nguyen, received the B.E degree from the Posts and Telecommunications Institute of Technology (PTIT) in 2020. Now she is an assistant lecturer at Faculty Telecommunication 1 of PTIT. Her research interests include predictive analytics, the application of AI, signal processing, high-speed optical communications, and AI for photonics.

Email: Duyth@ptit.edu.vn