



Volume 11 Issue 1 Year 2026 | Page 171-180 ISSN: 2527-9866

Received: 15-12-2025 | Revised: 30-12-2025 | Accepted: 20-02-2026

Sentiment Analysis E-Wallet Application Services Using the SVM and LSTM Methods

Mochammad Dzikri Arya Darmansyah¹, Anik Vega Vitianingsih², Anastasia Lidya Maukar³, SY. Yuliani⁴, Seftin Fitri Ana Wati⁵

^{1,2}Informatics Department, University of Dr. Soetomo, Surabaya, Indonesia, 60118

³Industrial Engineering Department, President University, Bekasi, Indonesia, 17550

⁴Informatics Department, Multimedia Nusantara University, Jakarta, Indonesia, 15810

⁵Information System Department, University of Pembangunan Nasional "Veteran" Jawa Timur, Indonesia, 60294

e-mail: m.dzikri@gmail.com¹, vega@unitomo.ac.id^{2*}, almaukar@president.ac.id³, sy.yuliani@umn.ac.id⁴, seftin.fitri.si@upnjatim.ac.id⁵

*Correspondence: vega@unitomo.ac.id

Abstract: The rapid growth of financial technology services in Indonesia has increased the volume of user reviews, yet their utilization for sentiment-based insights remains limited in the e-wallet sector. This study compares the effectiveness of Support Vector Machine (SVM) and Long Short-Term Memory (LSTM) in classifying the sentiment of 3,185 DANA e-wallet reviews collected from Google Play Store and Instagram. The research process includes text preprocessing, lexicon-based labeling, and feature extraction using TF-IDF for SVM and word embeddings for LSTM. Model evaluation is conducted using a confusion matrix based on accuracy, precision, and recall, without inferential statistical testing. The results show that LSTM outperforms SVM, achieving an accuracy of 86.66%, recall of 81.86%, and precision of 82.09%, while the best SVM variant with an RBF kernel attains an accuracy of 84.93%. This study contributes by identifying key service-related factors influencing user satisfaction and dissatisfaction and by providing practical, sentiment-based insights to support service quality improvement. The novelty lies in the multi-platform analysis of Indonesian e-wallet reviews and the direct comparison of classical machine learning and deep learning approaches without statistical hypothesis testing. These findings confirm the effectiveness of deep learning for sentiment analysis of unstructured Indonesian text.

Keywords: sentiment analysis, deep learning, e-wallet, text mining.

1. Introduction

In Indonesia, fintech has experienced dramatic growth, encouraging innovation in digital payment services, such as electronic wallets (e-wallets) [1]. E-wallets, such as DANA, have transformed transaction behaviour by enabling easy payments, instant transfers, top-up services, and integration with merchants and other digital goods/services [1]. Furthermore, promotions, such as discounts and cashback, are significant factors affecting the average e-wallet adoption rate and usage events across consumer segments [2]. Dana was the third-most-popular digital wallet app after Gopay and OVO during the pandemic [3]. It was launched in 2018 as a mobile phone app for non-cash financial services that allows users to top up their investment balances or purchase everyday items. At Dana, we are the fastest-growing digital wallet in terms of user numbers, with an approximately 4% growth rate even during the COVID-19 pandemic, according to a YouGov survey [3]. Because Dana's app has so many users, one is apt to find both negative and neutral reviews, as well as positive ones, among the terms that were not exclusively digital or social media-related.

The challenge in the current work, however, is that growth in user volume may not always be accompanied by a positive service experience [4]. The case DANA E-Wallet Application (selected based on its leading role and prevalence in the digital wallet market of Indonesia) has led to the highest number of registered users to date [4] this suggests that through its large exposure, DANA naturally elicits more reviews, feedback and complaints via digital forums such as Google Play Store and Instagram compared to other platforms which therefore categorizes it as the richest source of cover data in terms being most comprehensive; manipulated; informative dataset for mining [4]. Commonly reported grievances are "stuck" transactions or missing balances after transfers/top-ups, failed refunds, authentication problems like not receiving SMS OTPs (One Time Password), signs of potential security breaches/data leaks according to their perception, slow/unhelpful responses from Customer Service (both bots and manual agents), ambiguous promotional comms/cashback mechanisms leading to confusion and user angst. Such cases impact customer satisfaction, the service provider's trustworthiness, and possible user churn [5]. In this study, we examine how user attitude toward the DANA e-wallet application is distributed across service elements and how SVM and LSTM perform in categorizing these sentiments, given the aforementioned background.

Inspired by a previous study on user reviews of an e-wallet application on the Google Play Store [6], the researchers analyzed sentiment in the DANA application using a dataset of experiences and an SVM, achieving 80.81% accuracy for positive sentiment and 84.06% for negative sentiment. Meanwhile, an earlier study showed that a sentiment analysis application based on a Naive Bayes algorithm, trained on Play Store user reviews of the Dana app, reported 88% accuracy [7]. Results like the above could be improved by more effectively leveraging Deep Learning, which inherently adapts to variations, complexities, and sequential contexts in user review text. Thus, motivated by the literature gap, this study aimed to propose an e-wallet user review sentiment analysis using Support Vector Machine (SVM) and Long Short-Term Memory (LSTM). SVM can efficiently separate sentiment classes by finding the best hyperplane and kernel functions, as it can handle non-linear data distributions [8]. Meanwhile, the Long Short-Term Memory (LSTM) model was adopted in this study because it can capture contextual patterns and sequential information in texts [9]. Both approaches were employed to compare performance across sentiment polarity classification, thereby enabling better analysis of e-wallet user perception.

The purpose of the study was to implement SVM and LSTM techniques to identify sentiment-based applications for the e-wallet service (Dana) on Google Play and Instagram. To construct this dataset, web scraping will be used to collect user reviews across both digital platforms. The dataset is collected from user reviews on the Google Play Store and Instagram for the DANA e-wallet application. This research is unlike previous work, which mainly employs a single digital platform, i.e., the Google Play Store or Instagram, for sentiment analysis of Dana. This study enables a deeper understanding of user opinions across different communication forms by integrating data from two platforms. Contributions of this research include elucidating critical service factors affecting (dis)satisfaction that are often overlooked and generating sentiment-driven insights to guide e-wallet service providers in improving service quality.

2. Literature Review

The processing of natural language that targets to identify and classify a person's opinion/emotion regarding an object as negative, positive or neutral is sentiment analysis (natural language processing) [10]. Several previous works have conducted sentiment analysis of user reviews on the Play Store for e-wallet applications. In another study [6], the SVM method was applied to DANA application reviews, achieving 80.81% positive sentiment accuracy and 84.06% negative sentiment accuracy. Another study [7] also used the same reviews and achieved 88% accuracy with Naive Bayes.

Table 1. Comparison of Previous Research

Papers	Platform	Algorithm	Dataset Source	Gap
[6]	Dana	SVM	Google Play Store	Only using a single algorithm.
[10]	Dana	Naive Bayes	Google Play Store	Not using SVM and LSTM algorithms.

Based on Table 1, which presents a summary of previous research related to sentiment analysis, a research gap can be identified in the following aspect of data sources:

1. **Dataset Source Gap:** The features of the data evaluated were restricted to the setting of mobile application reviews because prior research often employed a single data source, namely, user reviews from the Google Play Store. The study's findings cannot always be applied to other data sources, such as social media or various review platforms, which have distinct linguistic patterns, degrees of formality, and emotional intensity.
2. **Algorithm Gap:** Some research employs a single classification algorithm, such as Support Vector Machines or Naive Bayes, without comparing approaches, thereby failing to provide a complete view of each algorithm's relative performance.

Based on previous research gaps, this study uses two data sources, both focused and relevant to the DANA context: user reviews from the Google Play Store and Instagram comments. It applies two classification methods, Support Vector Machine and LSTM, to produce a more representative sentiment analysis aligned with the characteristics of application reviews, while also contributing to an understanding of the influence of data sources and classification methods on sentiment analysis.

3. Methods

In this research, the system begins with the input stage, where users collect E-Wallet application review data from the Instagram and Google Play Store using scraping techniques and upload the scraped data to the system, along with supporting data such as an Indonesian slang dictionary and an Inset lexicon. The raw data will undergo several pre-processing steps, including database storage, text cleaning with slang normalization, case folding, tokenization, removal of custom stopwords, stemming, and lemmatization, to ensure standardized data. After pre-processing, the system uses hybrid labelling, combining lexicon-based techniques with manual sample validation to classify reviews into three sentiment categories: positive, negative, and neutral.

The labelled data is then processed using two separate feature engineering approaches: TF-IDF with N-grams for SVM model preparation and word embeddings for LSTM model preparation. The overall workflow of the proposed system is illustrated in Figure 1. To obtain user reviews of the E-Wallet program on the Google Play Store and Instagram, data were collected via scraping. Google Colab was used for the scraping procedure. Three sentiment categories, positive, negative, and neutral, were created from the gathered data. In total, 3,185 reviews were collected, comprising 1,931 from the Google Play Store and 1,254 from Instagram.

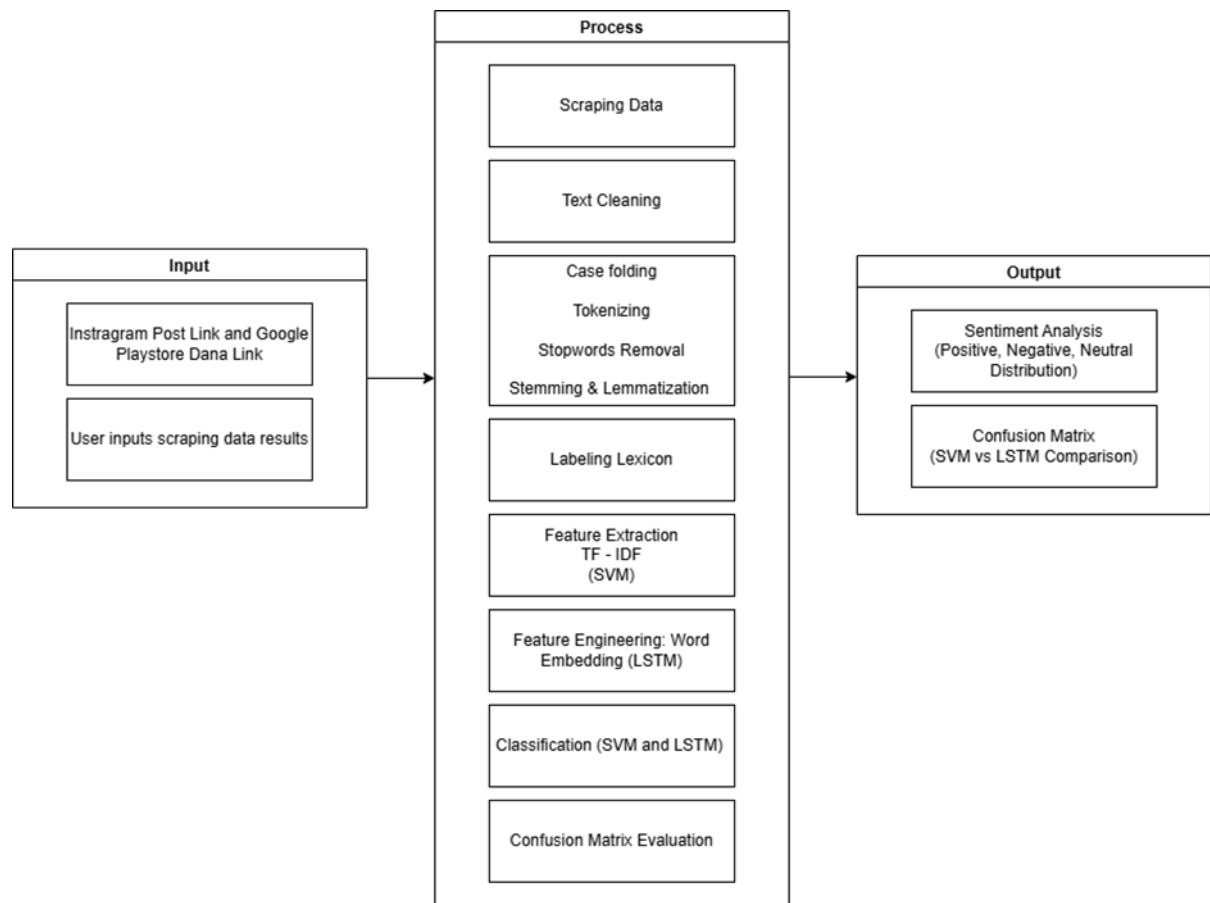


Figure 1. Research Methodology

A. Data Pre-processing

The initial stage in processing text data involves applying several techniques to clean and transform raw text into a format that can be effectively analyzed [11]. The main objectives of text pre-processing are to remove unnecessary information, standardize data, and reduce data complexity [12].

1. **Data Cleaning.** This stage is a procedure for eliminating interference caused by inconsistent or irrelevant data.
2. **Case Folding.** This process standardizes text by converting all characters to a uniform case, thereby minimizing data redundancy and improving computational efficiency during classification.
3. **Tokenization.** The stage of selecting text units to be analyzed and separating them based on analysis units, which can be words, word groups, or phrases, using Bag-Of-Words (BOW), which is the grammar and text sequence in documents without consideration.
4. **Stopword.** This step helps reduce noise and improve efficiency in analyzing meaningless text, such as prepositions. These words are called stop words.
5. **Stemming.** This technique transforms words into their base form by stripping away prefixes and suffixes.

B. Data Labelling

Lexicon-based sentiment analysis is a method that uses a dictionary or collection of words (a lexicon) with assigned weights to assess the sentiment of a text [13]. An Indonesian sentiment dictionary, InSet (Indonesian Sentiment Lexicon), was used in a lexicon-based method for data labelling [14]. Every word in the comment is given a polarity value (positive, negative, or neutral),

and the category of each comment is determined by adding up all of its scores. Positive sentiment is defined as a total score > 0 , neutral sentiment as a total score $= 0$, and negative emotion as a total score < 0 [14]. After labelling, the dataset was split into training and test sets at an 80:20 ratio.

C. TF-IDF

TF-IDF (Term Frequency Inverse Document Frequency) is a word weighting technique based on statistical word occurrence and document importance [15]. Inverse Document Frequency (IDF) indicates how infrequently a word appears in all documents, whereas Frequency Term (TF) indicates how frequently a term appears in a document. TF-IDF is a word weight that is produced by multiplying TF and IDF to assess a word's significance in a document [15]. The subsequent formulas of this methodological equation are articulated in (1), (2), and (3) [15].

$$W_{tf_{t,d}} = \{1 + tf_{t,d}, \text{ if } tf_{t,d} > 0\} \quad (1)$$

$$idf_t = \left(\frac{N}{df(t)} \right) \quad (2)$$

$$W_{t,d} = W_{tf_{t,d}} \times idf_t \quad (3)$$

D. Support Vector Machine

A machine learning technique used for both classification and regression tasks [23]. SVMs choose the best hyperplane to maximize the margin, or distance, between the classes [16]. This research uses three SVM kernels for sentiment classification, namely polynomial, Radial Basis Function (RBF), and Sigmoid. The RBF kernel computes the similarity between vectors in the feature space with a γ parameter that determines how sensitive the output is to changes in the input. The polynomial kernel measures the relationship between vectors with a given degree [16]. The sigmoid kernel uses the sigmoid activation function to capture non-linear interactions. This kernel, which links the kernel function to neural network techniques, is widely used in neural network models. The formulas for the Polynomial, RBF, and Sigmoid kernels are shown in Equations (4), (5), and (6) [17].

$$\text{Polynomial Kernel: } (x, y) = (x, y + c)^d \quad (4)$$

$$\text{RBF Kernel: } K(x, x') = \exp(-\gamma(x, y)^2) \quad (5)$$

$$\text{Sigmoid Kernel: } K(x, x_k) = \tanh[kx_k^T x + \theta] \quad (6)$$

In this research, the Support Vector Machine (SVM) models were implemented using polynomial, Radial Basis Function (RBF), and sigmoid kernels. The learning framework provided default hyperparameter settings for the SVM models, so no additional hyperparameter tuning or cross-validation was performed. A hold-out strategy with a 80:20 split was used to divide the dataset into training and test sets.

E. LSTM

A type of neural architecture that belongs to the recurrent neural network (RNN) family. By storing information for a predetermined period and avoiding typical forgetting of previous details, LSTM is specifically designed to handle sequential data, such as text, sound, or time series [18]. Word embedding techniques convert textual data into numerical representations before the LSTM model processes it. Word embedding aims to produce vector representations that capture the semantic information (meaning) of words in a corpus. The cosine similarity between word vector

representations can be used to quantify semantic similarity [19]. The LSTM model's distinctive structure, consisting of three primary gates (input, forget, and output gates) that control information flow, enables it to learn intricate patterns in sequential data while retaining crucial knowledge over time [18]. The subsequent formulas that describe each component of the LSTM mechanism are presented in Equations (7), (8), (9), (10), (11), and (12) [20].

$$\text{Input Gate: } it = (W_i \times [x_t + h_{t-1}]) + b_i \quad (7)$$

$$\text{Forget Gate: } ft = \sigma(W_f \times [x_t + h_{t-1}] + b_f) \quad (8)$$

$$\text{Candidate Cell: } ct = \tanh(WC \times [x_t + h_{t-1}] + b_C) \quad (9)$$

$$\text{Output Cell: } ot = \sigma(W_o \times [x_t + h_{t-1}] + b_o) \quad (10)$$

$$\text{Cell State Update: } Ct = ft \times C_{t-1} + it \times Ct \quad (11)$$

$$\text{Hidden State: } h_t = ot \times \tanh(Ct) \quad (12)$$

Before being processed by the network, textual data was converted into numerical sequences for the LSTM model using word embedding techniques. An 80:20 ratio was used to divide the dataset into training and testing sets. The Adam optimizer and categorical cross-entropy loss function were used to train the LSTM model over 10 epochs with a batch size of 16. The output layer used a softmax activation function.

G. Confusion Matrix

In machine learning, a confusion matrix is a key tool for evaluating the performance of classification models by comparing predicted outputs with the actual labels [21]. The performance of classification models is assessed using the Confusion Matrix, which provides metrics such as Accuracy, Precision, and Recall. Accuracy calculates the percentage of correct predictions by dividing the total test results by the number of True Positives and True Negatives. Precision measures how well the model predicts positive classifications by comparing the number of true optimistic predictions (True Positives) to the total number of optimistic forecasts. The recall metric evaluates the model's ability to identify each true positive by reducing the percentage of True Positives by the total number of positive data points. Equations (13), (14), and (15) provide the subsequent formulas for the evaluation metrics: recall, accuracy, and precision [20].

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (13)$$

$$\text{Precision} = \frac{TP}{TP + FP} \quad (14)$$

$$\text{Recall} = \frac{TP}{TP + FM} \times 100\% \quad (15)$$

4. Results and Discussion

A. Data Scraping

The automated process of utilizing software to extract data from websites is known as web scraping. Using this technology, structured data that isn't immediately machine-processable, such as CSV, can be extracted from HTML [22]. This data covers diverse perspectives, concerns, and compliments from users regarding the features, performance, ease of use, and experience with e-wallet applications such as DANA. To make the analysis process easier, the data was gathered via

web scraping methods and saved in *.csv format. This dataset is the primary material for the sentiment labelling and classification process using machine learning methods. Here is an example of data scraping results in Table 2.

Table 2. Data Scraping

Text
@Baru ada telp dari dana, diangkat malah mati.... Nggak jelas bener MIN baca DM,kenapa akun dana saya keluar terus Terbaik ❤️❤️ duh danaku ngga bisa dibuka..

B. Data Pre-processing

After the scraping process completes the basic raw review data, the pre-processing phase begins with cleansing unwanted characters (e.g., symbols, punctuation, numbers, and text). Then case folding is applied, lowering all letters to a single form of data. The next step is tokenization, which splits sentences into individual words. Then, common words (stopwords) that do not contribute much to the analysis are removed. It also has to remove inflexion from words (words are represented uniformly). The steps of feature pre-processing are listed in Table 3.

Table 3. Data Pre-processing

Preprocessing Datasets	Result
Raw Data	@Tim2 yang gak kompeten dan malas diganti saja yang loyalis responsif
Data Cleaning	Tim2 yang gak kompeten dan malas diganti saja yang loyalis responsif
Case Folding	tim yang gak kompeten dan malas diganti saja yang loyalis responsif
Tokenization	['tim', 'yang', 'gak', 'kompeten', 'dan', 'malas', 'diganti', 'saja', 'yang', 'loyalis', 'responsif', 'dan', 'kerja', 'maksimal']
Stopword removal	['tim', 'gak', 'kompeten', 'malas', 'diganti', 'loyalis', 'responsif', 'kerja', 'maksimal']
Stemming	['tim', 'gak', 'kompeten', 'malas', 'ganti', 'loyalis', 'responsif', 'kerja', 'maksimal']

C. Data Labelling

For previously pre-processed text collections (cleaned, case-folded, tokenized, filtered, and stemmed), we add the examples to our training set using Lexicon-Based labelling. This approach judges each word in the document by checking whether it appears in a sentiment-related word list from an appropriate sentiment lexicon. The Lexicon-Based steps are shown in Table 4.

Table 4. Lexicon Based

Clean Dataset	Score	Sentiment
kapas tunggu proses	-1	Negative
minggu dana manual cashout transfer akun dana baru uang	0	Neutral
semangat cuy sikat promo dana	7	Positive

The sentiment distribution is visualized using several diagrams after the evaluation stage, making the entire analysis easier to understand. Following pre-processing and tagging, a total of 1,254 Instagram user comments and 1.931 Google Play reviews were successfully transformed into clean, structured data that could be analyzed using lexicon-based methods.

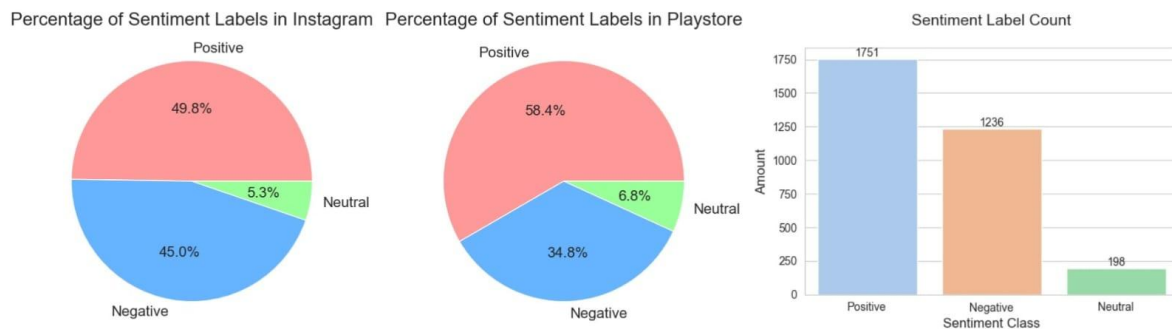


Figure 2. Lexicon-Based Visualization Result

According to Figure 2, neutral sentiment has the fewest data points (198), while positive sentiment has the highest (1751), followed by negative sentiment (1236). In Figure 2, neutral sentiment has the fewest data points (198), while positive sentiment has the highest (1,751), followed by negative sentiment (1,236). This indicates that positive sentiment predominates in the dataset's reviews. Suggests that positive sentiment predominates in the dataset's reviews.

D. TF-IDF

A pre-processing of the text data is computed to obtain the occurrence frequency of each word in a review for the TF in the TF-IDF weighting stage. To emphasize more informative terms, IDF is then calculated to indicate how rare a word is across the entire document. These two numbers are then multiplied to form a TF-IDF weight, which is a real number in the one-line-per-word format. The outcome of this weighting is a feature matrix, which is fed to the classification model (SVM), allowing it to recognize sentiment patterns by considering the relevance of each word to the document.

For illustrative purposes, the word frequency (TF) of each document, the number of documents containing the word (DF), and the word's IDF value, representing the word's relevance across the corpus, are presented in the table to illustrate the TF-IDF computation. Terms that appear more frequently in the document will have higher TF, and terms that are rare across all documents will have higher IDF values according to their informativeness. Table 5 shows the weights assigned to each word in determining its importance at different stages of the analysis, taking both TD and IDF values into account.

Table 5. TF-IDF

Term	TF					DF	IDF
	D1	D2	D3	D4	D5		
aku		0,5782		0,2147		5	0,7482
uang	0,1265		0,6221		0,2143	2	0,8242
rusak	0,2554			0,3256		3	0,5178

F. Models Evaluation

Model evaluation is the last phase of the procedure. The model's performance is then visualized in a confusion matrix, which shows the number of correct and incorrect predictions for each sentiment class. This matrix makes it easier to see how the system distinguishes between positive, negative, and neutral categories. The performance metrics obtained with the SVM kernels and the LSTM model are presented in Table 6.

Table 6. Model Performance

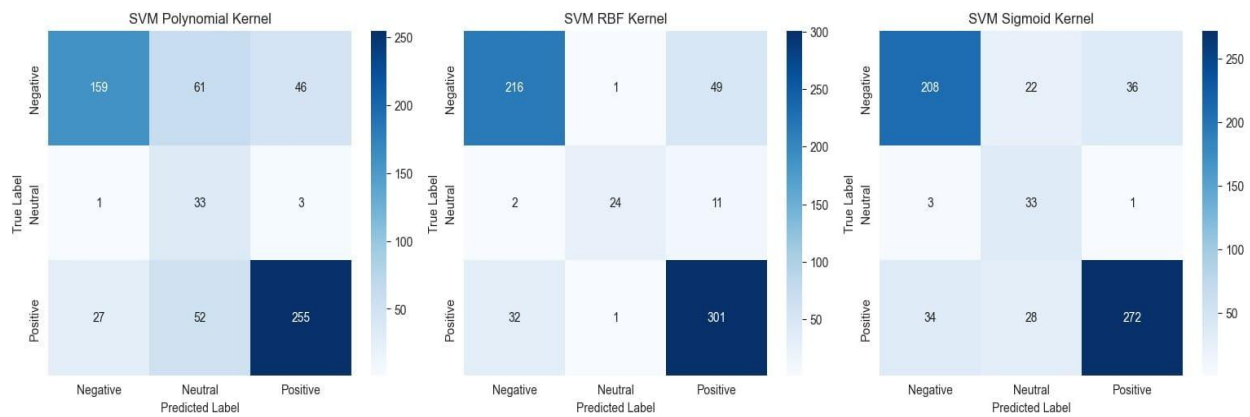
Model		Metric		
		Accuracy	Precision	Recall
SVM	Polynomial	70.17%	75.10%	63.84%
	RBF	84.93%	78.73%	87.36%
	Sigmoid	80.53%	82.94%	70.89%
LSTM		86.66%	81.86%	82.09%

The evaluation results compare how well the LSTM model and various SVM kernel variations perform sentiment classification. The LSTM model achieved 86.66% accuracy, 81.86% recall, and 82.09% precision, performing best and demonstrating that it can recognize linguistic patterns more effectively by comprehending the sequential relationships between words. LSTM performs well at balancing positive data and reducing classification errors, as evidenced by its high recall and precision. The SVM model's evaluation results reveal differences in kernel performance. The Polynomial SVM performs worst, with 70.17% accuracy and 63.84% precision, suggesting that this kernel is less effective at identifying sentiment patterns.

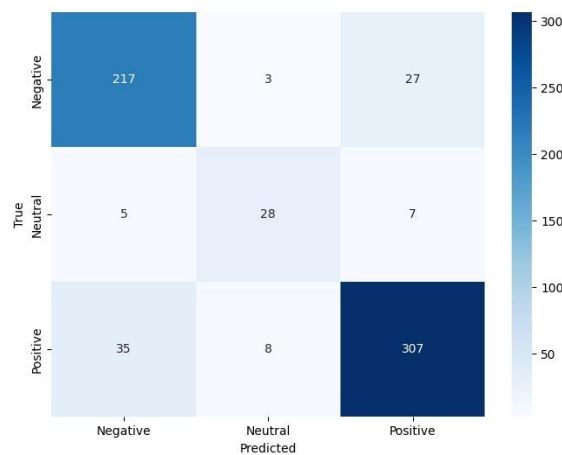
The best-performing kernel was the RBF SVM, which achieved 84.93% accuracy and 87.36% precision. Although its precision was lower, Sigmoid SVM achieved the highest recall of 82.94% and accuracy of 80.53% among SVM kernels. Compared with other SVM kernels, SVM RBF showed the most reliable and accurate performance overall. Because it can learn long-term context from text, LSTM is shown in Figure 3(b) to outperform the other two models generally. SVM RBF, which offers the best performance balance among SVM models, comes in second. These results imply that deep learning-based models perform better on Indonesian text data than standard kernel-based methods.

The confusion matrices provide a comprehensive summary of how each SVM kernel predicts positive, neutral, and negative labels, as shown in Figure 3. An error analysis examined the misclassified samples generated by both models. The findings demonstrate the ambiguity of neutral sentiment in user-generated text by showing that both SVM and LSTM frequently misclassified neutral evaluations as either positive or negative, with an accuracy of 86.66%, recall of 81.86%, and precision of 82.09%. The Long Short-Term Memory (LSTM) model performed best in this research for sentiment categorization. These outcomes demonstrate that the LSTM performs better on sequential text data than other models. Although the LSTM model was more accurate than the SVM, the performance difference was not significant. To further confirm the robustness of the performance disparities, future research should include statistical significance testing, such as paired t-tests or McNemar's test.

The findings of this research align with those of [6], which found that SVM achieved an accuracy above 80% in classifying the sentiment of DANA user evaluations. However, by adding another data source from Instagram and demonstrating that LSTM performs better when addressing informal and context-dependent expressions, this research extends previous work. Then, the performance in this research is marginally lower than the results previously reported [7], which used Naive Bayes to attain greater accuracy. These models may be due to the inclusion of multi-platform data with more varied linguistic features.



(a) SVM Model Evaluation



(b) LSTM Model Evaluation

Figure 3. (a) SVM Model Evaluation (b) LSTM Model Evaluation

5. Conclusions

This research effectively evaluated the efficacy of LSTM and SVM techniques for sentiment classification of 3,185 reviews of the DANA e-wallet application. Without using inferential statistical tests, model performance was evaluated using a confusion matrix to compute classification metrics, specifically accuracy, precision, and recall. According to the test results, the LSTM model outperforms all SVM model versions, with an accuracy of 86.66% and a recall of 81.86%. These results show that LSTM is superior at capturing word relationships and sequential context in review texts. Furthermore, with an accuracy of 84.93%, the SVM model with a Radial Basis Function (RBF) kernel was the best-performing SVM model among those examined; however, its performance was still inferior to that of LSTM. These findings demonstrate that the deep learning method typically yields more accurate sentiment classification in unstructured Indonesian text. The study's practical conclusions can serve as a foundation for creating data-driven sentiment analysis systems that facilitate efficient decision-making and enhance service quality. The use of lexicon-based sentiment classification techniques, which are currently unable to adequately capture the linguistic context in detail, including irony, sarcasm, and implicit emotions in review texts, is one of the study's drawbacks. To improve the precision and depth of sentiment analysis, further research is advised to develop more robust sentiment labelling methods, leverage a broader range of data sources, and explore cutting-edge strategies.

References

- [1] O. Feriyanto, "Peran Fintech Dalam Meningkatkan Akses Keuangan di Era Digital," *Gemilang J. Manaj. dan Akuntansi*, vol. 4, no. 3, pp. 99–114, 2024, doi: <https://doi.org/10.56910/gemilang.v4i3.1573>.
- [2] I. D. Filany and N. R. C. Putri, "The Influence Of Using E-Wallets And Discounts On Impulse Buying In The Marketplace," *3 RD Int. Conf. Econ. BUSINESS, Manag. Res.*, vol. 3, pp. 544–565, 2024.
- [3] F. A. Larasati, D. E. Ratnawati, and B. T. Hanggara, "Analisis Sentimen Ulasan Aplikasi Dana dengan Metode Random Forest," *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 6, no. 9, pp. 4305–4313, 2022.
- [4] W. Yusditar, "Analisis Pengaruh Kualitas Layanan Dan Promosi Terhadap Kepuasan Pelanggan Pada Generasi Z Dalam Industri E-Commerce," vol. 02, no. 02, pp. 63–72, 2024, doi: <https://doi.org/10.51622/kafebis.v2i2.2671>.
- [5] R. Saputri and M. E. Baining, "Pengaruh Kualitas Sistem Dan Kualitas Pelayanan Terhadap Manfaat Bersih Aplikasi Mobile Banking Dengan Variabel Intervening Kepuasan Pengguna," vol. 17, no. 1, pp. 126–138, 2024, doi: <https://doi.org/10.51903/e-bisnis.v17i1.1801>.
- [6] W. E. Saputro, H. Yuana, and W. D. Puspitasari, "Analisis Sentimen Pengguna Dompot Digital Dana Pada Kolom Komentar Google Play Store Dengan Metode Klasifikasi Support Vector Machine," vol. 7, no. 2, pp. 1151–1156, 2023.
- [7] Sriani, Armanysah, and M. A. Rambe, "Analisis Sentimen Pengguna Aplikasi Dana Menggunakan Metode Naïve Bayes," *JIFOTECH (JOURNAL Inf. Technol.)*, vol. 4, no. 2, pp. 174–182, 2024.
- [8] S. D. Parameswari, M. Lubis, S. Suakanto, Y. Z. Ramadhan, N. Amanah, and R. A. Dila, "Studi Perbandingan Naï ve Bayes dan Support Vector Machine (SVM) dalam Analisis Sentimen Pengguna Metaverse," vol. 4, no. 3, pp. 1059–1065, 2025, doi: <https://doi.org/10.55826/jtmit.v4i3.1122>.
- [9] M. F. B. Hannan Asrawi, Sriwinar, "Implementasi Fitur Ekstraksi Glove Pada Teks Klasifikasi Menggunakan Algoritma Long-Short Term Memory," vol. 07, no. 03, pp. 708–717, 2025, doi: <https://doi.org/10.54209/jatilima.v7i03.1694>.
- [10] S. Andini, R. Kurniawan, and S. Anwar, "Analisis Sentimen Pengguna X Mengenai Opini Masyarakat Tentang Dinasti Politik Menggunakan Algoritma Naïve Bayes," *JITET (Jurnal Inform. dan Tek. Elektro Ter.)*, vol. 13, no. 2, pp. 665–671, 2025, doi: <http://dx.doi.org/10.23960/jitet.v13i2.6299>.
- [11] M. Nursalman, J. Kusnendar, and U. F. Fadhila, "Implementation of K-Nearest Neighbor with Cosine Similarity for Classification Abstract International Journal of Computer Science," *2018 Int. Conf. Inf. Technol. Syst. Innov. ICITSI 2018 - Proc.*, no. February, pp. 43–48, 2018, doi: [10.1109/ICITSI.2018.8696072](https://doi.org/10.1109/ICITSI.2018.8696072).
- [12] A. Murugan, H. Chelsey, and N. Thomas, *Practical Text Analytics*, vol. 2. 2019.
- [13] Muhammad Fernanda Naufal Fathoni, Eva Yulia Puspaningrum, and Andreas Nugroho Sihananto, "Perbandingan Performa Labeling Lexicon InSet dan VADER pada Analisa Sentimen Rohingya di Aplikasi X dengan SVM," *Modem J. Inform. dan Sains Teknol.*, vol. 2, no. 3, pp. 62–76, 2024, doi: [10.62951/modem.v2i3.112](https://doi.org/10.62951/modem.v2i3.112).
- [14] Z. Fatah and R. A. Ningsih, "Analisis Sentimen Komentar YouTube terhadap Tragedi Demo 25 Agustus," *JAMASTIKA*, vol. 4, no. 2, pp. 217–224, 2025.
- [15] R. Wang and Y. Shi, "Research on application of article recommendation algorithm based on Word2Vec and Tfidf," in *2022 IEEE International Conference on Electrical Engineering, Big Data and Algorithms (EEBDA)*, IEEE, 2022, pp. 454–457. doi: [10.1109/EEBDA53927.2022.9744824](https://doi.org/10.1109/EEBDA53927.2022.9744824).
- [16] G. A. B. Baliputra, S. Kacung, and B. Santoso, "Sentiment Analysis of Rohingya Refugees in Aceh using Support Vector Machine (SVM) and Multinomial Logistic Regression," *istemasi J. Sist. Inf.*, vol. 14, no. 3, pp. 1332–1343, 2025, doi: <https://doi.org/10.32520/stmsi.v14i3.5159>.
- [17] V. Anggraeni, S. Putri, A. V. Vitianingsih, R. Hamidan, A. L. Maukar, and N. T. Pratitis, "Sentiment Analysis On Social Media Instagram Of Depression Issues Using Nb Method," *J. INOVTEK POLBENG - INFORMATICS Ser.*, vol. 9, no. 2, pp. 802–813, 2024.
- [18] S. J. Pipin and H. Kurniawan, "Analisis Sentimen Kebijakan MBKM Berdasarkan Opini Masyarakat di Twitter Menggunakan LSTM," *J. SIFO Mikroskil*, vol. 23, no. 2, pp. 197–208, 2022, doi: <https://doi.org/10.55601/jsm.v23i2.900>.
- [19] M. A. El-affendi, B. N. Al-tamimi, F. Albalwy, and A. Hussain, "Arabic Sentiment Analysis Based

- on Word Embeddings and Deep Learning," *MDPI J.*, vol. 12, no. 6, pp. 1–21, 2023, doi: doi.org/10.3390/computers12060126.
- [20] M. Ilona, J. Bria, A. V. Vitianingsih, A. L. Maukar, and S. Y. Yuliani, "Sentiment Analysis of User Reviews on Maxim Application Using the Long Short-Term Memory (LSTM) Methods," vol. 5, no. 3, pp. 941–952, 2025.
- [21] A. Syahri and F. Muttakin, "Sentimen Analisis Pada Ulasan Aplikasi Ajaib Di Google Play Store Dengan Algoritma Support Vector Machine," *J. INOVTEK POLBENG - SERI Inform.*, vol. 9, no. 1, pp. 257–271, 2024, doi: <https://doi.org/10.35314/isi.v9i1.4047>.
- [22] M. A. Khder, "Web Scraping or Web Crawling : State of Art , Techniques , Approaches and Application," *Int. J. Adv. Soft Comput.*, vol. 13, no. 3, pp. 145–168, 2021, doi: 10.15849/IJASCA.211128.11.
- [23] Tedyyana, A., Ghazali, O., & W. Purbo, O. (2024). Enhancing intrusion detection system using rectified linear unit function in pigeon inspired optimization algorithm. *IAES International Journal of Artificial Intelligence (IJ-AI)*, 13(2), 1526-1534. doi:<http://doi.org/10.11591/ijai.v13.i2.pp1526-1534>