

Support Vector Machine Based Machine Learning for Sentiment Analysis of User Reviews of the Bibit Application on Google Play Store

Ega Shela Marsiani¹, Fauzan Natsir^{2*}, Redo Abeputra Sihombing³, Millati Izzatillah⁴, Rajiansyah⁵

¹Department of Informatic Engineering, Universitas Indraprasta PGRI, South Jakarta (Indonesia)

²Department of Computer Science and Systems Engineering, Wroclaw University of Science and Technology (Poland)

*Corresponding author: fauzan.natsir@gmail.com



Received 29 October 2025 | Accepted 26 November 2025 | Early Access 30 November 2025

ABSTRACT

The increasing use of financial technology (fintech) applications has significantly transformed investment behaviour in Indonesia. Bibit, one of the leading fintech investment platforms, receives a large number of user reviews on the Google Play Store, reflecting user perceptions and satisfaction levels. However, despite the growing volume of reviews, systematic sentiment analysis remains limited, making it difficult for developers to fully understand user needs and experiences. To address this issue, an artificial intelligence based approach is required to efficiently and objectively extract and analyse user opinions. This study applies a Support Vector Machine (SVM) based machine learning model to conduct sentiment analysis on user reviews of the Bibit application. The research methodology includes data collection, text preprocessing, TF-IDF feature extraction, and sentiment classification into positive, negative, and neutral categories. Of the total collected data, 7,801 reviews (79.99%) were used as training data and 1,561 reviews (20.01%) as test data following the commonly used 80:20 ratio in machine learning. The objective of this study is to identify the dominant sentiment expressed by users and evaluate the performance of the SVM algorithm in classifying sentiment. The experimental results show that the SVM model achieved high accuracy and effectively captured user opinions, providing valuable insights for developers to enhance application quality and improve user engagement on fintech platforms.

KEYWORDS

Sentiment Analysis, Support Vector Machine (SVM), Seed Application, User Reviews.

I. INTRODUCTION

THE development of financial technology applications (*financial technology* or *fintech*) has changed investment patterns and user interaction with financial services in Indonesia. The mobile app allows users to invest, trade, and manage financial products more easily and efficiently [1]. User reviews on app stores (such as Bibit on the Google Play Store) are a valuable source of information that reflects the sentiment, level of satisfaction, and user experience of the services provided. Automating sentiment extraction from such reviews can help developers and stakeholders improve service quality and user experience [2].

Among the various *machine learning* approaches to sentiment classification, the Support Vector Machine (SVM) has proven to be effective for text-based classification tasks. Its ability to handle high-dimensional feature spaces and generate strong *decision boundaries* makes it an ideal method for sentiment analysis of user reviews. This study aims to apply an SVM-based approach in sentiment analysis to user reviews of the Bibit application on the Google Play Store, with the aim of identifying the dominant perception of users and evaluating the performance of the SVM algorithm in the classification of sentiment in this context [3].

The research background on sentiment analysis using the Support Vector Machine (SVM) method is driven by the rapid growth of digital applications and the increasing number of user reviews which are an important source of data for developers to understand user satisfaction and complaints [4]. Various previous literature studies have proven the effectiveness of SVM in sentiment classification on app review data on the Google Play Store, such as on Gojek apps, ChatGPT, Google Maps, and other financial apps including Ajaib and several mutual fund apps. The SVM model has been reported to have high accuracy, both compared to other algorithms such as Naïve Bayes and Random Forest, and is able to handle high-dimensional text data, especially in unbalanced data conditions [5].

However, based on previous studies, there have not been many studies that specifically examine sentiment analysis in the Bibit application with a comprehensive SVM approach, where existing studies use other methods such as BERT for multi-aspect analysis or limit it to two categories of sentiment without considering the dimensions of neutral sentiment and pre-processing features that are relevant for Bahasa [6]. In addition, the novelty of this study lies in the optimization of the analysis pipeline starting from scraping, preprocessing, lexicon-based sentiment labeling, to detailed evaluation

through confusion matrix and classification report for Bibit review, something that has not been fully explored in previous studies that focus more on other investment applications or different analysis scopes [7].

The gap analysis identified is the need for a machine learning model that is tested and optimized specifically in the context of the Bibit application and the use of three sentiment labels (positive, neutral, negative) in order to accommodate the diversity of user expressions more accurately, in addition to the implementation of a complete pipeline from data scraping to performance evaluation and visualization of analysis results. The need for systematic mapping of Bibit user opinions is very important for application developers in making decisions to improve features, marketing strategies, and improve the quality of services based on actual data, so that this research is a strategic effort to contribute scientific novelty and applicative benefits to the world of fintech and similar application development in Indonesia [8].

The novelty of this research in three main contributions: (1) the development of a complete and optimized sentiment-analysis pipeline specifically for the Bibit investment application starting from automated scraping, language specific text preprocessing, lexicon-based labeling, TF-IDF feature extraction, and multi-class SVM classification; (2) the use of three sentiment categories (positive, neutral, negative), which provides a more granular representation of user perceptions compared to prior studies that used only binary labels; and (3) a comprehensive evaluation using accuracy, precision, recall, F1-score, confusion matrix interpretation, and performance visualization, which has not been thoroughly explored in previous research on Bibit reviews.

II. RELATED WORK

Sentiment analysis has been extensively researched in a variety of domains involving user-generated text, such as social media, product reviews, and digital apps. A systematic study of the application of sentiment analysis using SVM shows that this algorithm is one of the most widely used supervised methods for text classification tasks. The study confirmed the effectiveness of SVM in detecting sentiment polarity in user reviews and social media texts. For example, Ahmad et al. conducted a systematic literature review of SVM-based sentiment analysis in the period 2012–2017 [9].

In the context of fintech, several studies have applied sentiment analysis techniques to reviews of users of digital financial services. A study combined SVM with the Particle Swarm Optimization (PSO) method and obtained promising results in the classification of fintech application user sentiment [10]. Another study compared SVM and Naïve Bayes' performance in analyzing user preferences for old and new fintech technologies, and showed that SVM achieved an accuracy of about 87%, higher than NB [11].

The latest advances also explore sentiment analysis on mobile apps, including reviews related to fintech and digital banking. For example, a study in 2025 evaluated machine learning and large language model's approach to Arabic-language digital banking reviews, and found that traditional

machine learning models still outperform LLMs in those domains [12]. Another study in 2023 conducted an aspect-based sentiment analysis on reviews of the Flip digital wallet app in Indonesia using SVM, and obtained high accuracy for the aspects of speed, security, and transaction fees [13].

Overall, although SVM has been widely used and validated in various sentiment analysis contexts, its specific application to user reviews of Bibit's investment application is still rarely discussed. This research contributes by filling the gap through the application of an SVM-based sentiment analysis pipeline that includes the pre-processing stages of text, feature extraction, and classification of Bibit's review data on the Google Play Store.

III. METHODOLOGY

A. Overview

This research was carried out through several stages of the process starting from data collection, then continued with cleaning and preparing data so that it is ready for use. After the data is prepared, labeling is carried out to categorize sentiment, then the data is divided into training data and test data [14]. The model is then built using deep learning algorithms and evaluated to measure its performance. The complete flow of the research methodology can be seen in figure 1 below.

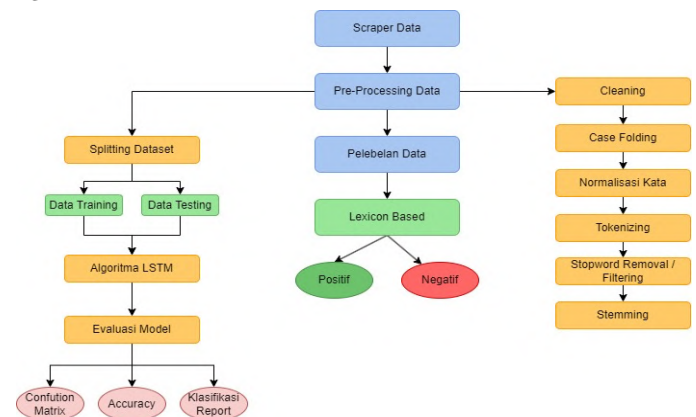


Fig. 1. Methodology

The research methodology starts with a data scraper to collect text data from a specified source. The collected data then goes into the pre-processing stage which consists of several cleaning steps, namely cleaning to eliminate unnecessary characters, case folding to standardize letters, normalizing words to correct spelling, tokenizing to break text into words, stop word removal to remove common words, and stemming to change words to basic forms. After the data is clean, labeling is carried out using a lexicon-based approach that produces three categories: positive, neutral, and negative. The labeled data is then divided into two parts through splitting datasets, namely training data (80%) to train the model and testing data (20%) to evaluate the model's performance. The model used is the Support Vector Machine (SVM) algorithm with linear kernel, which is then evaluated using four metrics: confusion matrix to see the details of true and false predictions, accuracy to measure the level of

correctness, precision to measure the exactness of positive predictions, recall to measure the model's ability to find all positive instances, and F1-score as the harmonic mean of precision and recall to provide a balanced performance measure across all sentiment categories.

B. Materials

The main research is a collection of user reviews of the Bibit investment application obtained from the Google Play Store. The dataset was collected using the Google Play Scraper library on the Python programming language, with the attributes retrieved including review text, rating, and publication date. Reviews that are spammy, duplicate, or use languages other than English are removed to maintain data quality. The study also used several supporting software and libraries such as Python 3.11, Scikit-learn, Pandas, and Sastrawi for the preprocessing and model training process. The Support Vector Machine (SVM) algorithm is applied for sentiment classification, while the TF-IDF (Term Frequency–Inverse Document Frequency) method is used to convert text into numerical representations.

C. Participants

The study analyzed public reviews it was means no human participants were directly involved. Participants in this context refer to users who provide reviews of the Bibit application on the Google Play Store. In total, around 5,000–10,000 reviews were collected, which included positive, negative, and neutral sentiment. These reviews reflect real user experiences and are relevant for sentiment analysis in the context of fintech applications.

D. Tasks and Methods

The context of this research refers to the natural activities of users, namely providing reviews and assessments of the Bibit application on the Google Play Store. No specific experimental tasks are assigned, as the data is obtained from the user's real interaction with the application.

The researcher carried out the main stages as follows:

1. Text Preprocessing: including case folding, tokenization, stop word removal, stemming with Literature, and spelling normalization.
2. Feature Extraction: uses the TF-IDF method to convert text into numerical vectors based on the weight of important words.
3. Sentiment Classification: implement SVM with a linear kernel to group reviews into three sentiment classes: positive, neutral, and negative.
4. Model Evaluation: using the Accuracy, Precision, Recall, and F1-Score metrics, as well as the Confusion Matrix to assess the model's performance.
5. SVM Model Configuration and Hyperparameter Tuning

The Support Vector Machine implementation in this study utilized the linear kernel due to its computational efficiency and effectiveness in high-dimensional text classification tasks. The SVM model was configured with the following technical specifications:

- a. Kernel Function: Linear kernel ($K(x, x') = x \cdot x'$)
- b. Regularization Parameter (C): 1.0 (default)
- c. Multi-class Strategy: One-vs-Rest (OvR) approach

- d. Class Weight: Balanced (to handle class imbalance)
- e. Maximum Iterations: 1000
- f. Convergence Tolerance: 1e-3

Hyperparameter Tuning to optimize model performance, we conducted hyperparameter tuning using 5-fold cross-validation on the training set. The C parameter (regularization strength) was tested across values [0.1, 1.0, 10, 100]. The grid search results showed that C=1.0 provided the optimal balance between model complexity and generalization capability, achieving the highest cross-validation accuracy. Multi-class Classification Strategy is inherently a binary classifier, we employed the One-vs-Rest (OvR) strategy to handle the three class sentiment classification problem (positive, neutral, negative). In this approach, three separate binary classifiers were trained:

1. Positive vs. (Neutral + Negative)
2. Neutral vs. (Positive + Negative)
3. Negative vs. (Positive + Neutral)

During prediction, the classifier that produces the highest confidence score determines the final class label. This strategy was chosen over One-vs-One due to its computational efficiency and straightforward probability estimation for multi-class problems. Mathematical equations are used to describe the process of classification and performance measurement in the Support Vector Machine (SVM) model. The SVM model looks for an optimal hyperplane that can separate data from two different classes (positive and negative). The basic equation of a hyperplane can be expressed as:

$$w \cdot x + b = 0 \quad (1)$$

- w : Vector Weight,
- x : Feature Vector (review text extraction feature),
- b : bias.

The decision function used to determine the data class is:

$$f(x) = \text{sign}(w \cdot x + b) \quad (2)$$

When the results $f(x) > 0$, then the review is categorized as positive sentiment, whereas if $f(x) < 0$, It is categorized as negative sentiment. To obtain an optimal hyperplane, the SVM optimization function is defined as:

$$\min_{w,b} \frac{1}{2} ||w||^2 \quad (3)$$

with limitations:

$$y_i(w \cdot x_i + b) \geq 1, i = 1, 2, \dots, n \quad (4)$$

where y_i is a review class label to- i (positive = +1, negative = -1). To measure the performance of the model, the following metrics are used:

$$\text{Accuracy} = \frac{TP+TN}{TP+TN+FP+FN} \quad (5)$$

$$\text{Precision} = \frac{TP}{TP+FP} \quad (6)$$

$$\text{Recall} = \frac{TP}{TP+FN} \quad (7)$$

$$F1 = 2 \times \frac{\text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}} \quad (8)$$

IV. RESULT AND ANALYSIS

A. Data Collection and Preprocessing

The figure 2 shows the results of the process of scraping user reviews of the Bibit application from the Google Play Store using the Python programming language. This scraping process is the initial stage in the research methodology, serving to collect raw data that will be used in the preprocessing stage of the text before sentiment analysis is carried out. The data is then cleaned, normalized, and converted into the appropriate format for processing by the Support Vector Machine (SVM) model.

The text preprocessing consisted of the following sequential steps:

1. Case Normalization: All text was converted to lowercase to ensure uniformity and reduce vocabulary size (e.g., "Bagus" → "bagus", "APLIKASI" → "aplikasi").
2. Cleaning and Noise Removal. Removed special characters, punctuation marks, and symbols (e.g., @, #, !, ?, -, etc.). Eliminated numbers and alphanumeric characters that do not contribute to sentiment. Stripped URLs, mentions, and hashtags commonly found in user reviews. Removed extra whitespaces and newline characters.
3. Indonesian Text Normalization. Slang and colloquial term standardization: Converted informal Indonesian internet slang to standard forms using a predefined dictionary (e.g., "gak" → "tidak", "banget" → "sangat", "emang" → "memang", "udah" → "sudah"). Typo correction: Fixed common misspellings in user reviews (e.g., "aplkasi" → "aplikasi", "bgus" → "bagus"). Repeated character normalization: Reduced repeated characters used for emphasis (e.g., "baguuuus" → "bagus", "jelek!!!" → "jelek").
4. Tokenization: Split the normalized text into individual words (tokens) using whitespace as delimiters, creating a list of terms for each review.
5. Stopword Removal: Eliminated common Indonesian stopwords that do not carry significant sentiment information, such as "yang", "dan", "di", "ke", "dari", "untuk", "pada", "adalah", and other function words. A custom stopwords list containing 758 Indonesian stopwords from the Sastrawi library was utilized.
6. Stemming: Applied the Sastrawi stemming algorithm to reduce words to their root forms, removing affixes (prefixes, suffixes, infixes) specific to Indonesian morphology. For example:
 - a. "menggunakan" → "guna"
 - b. "mempermudah" → "mudah"
 - c. "terlambat" → "lambat"
 - d. "aplikasinya" → "aplikasi"

This stemming process is crucial for Indonesian language processing as it handles complex affixation patterns including prefixes (me-, di-, ke-, pe-), suffixes (-kan, -an, -i), and their combinations.

After completing the preprocessing pipeline, each review was transformed from raw user-generated text into a clean, normalized token sequence ready for feature extraction using TF-IDF vectorization. For example:

- a. Before preprocessing: "Aplikasinya bagus banget!!! Memudahkan investasi saya".
- b. After preprocessing: "aplikasi bagus sangat mudah investasi".

This systematic preprocessing approach ensures that the text data is standardized and optimized for machine learning analysis while preserving the essential sentiment-bearing terms.

Review ID	Username	Rating	Review Text	Date
0 b589778-5a51-4601-849d-25a4936a8fc4	Pengguna Google	5	good	2025-10-23 02:08:28
1 2c5a7c1b-65c7-4797-8e6c-d0897bc50302	Pengguna Google	1	sudah tau gk bisa bukan di alih ke yang lain ...	2025-10-22 16:03:54
2 766e0e9a-73d2-4d32-b62f-03834bceae15	Pengguna Google	5	ok	2025-10-22 14:19:59
3 8de0dca9-e6a8-4744-84fb-6dcf8525bec4	Pengguna Google	4	sangat keren	2025-10-22 13:07:10
4 8d18415d-4303-40d9-bf1d-c8517362ae58	Pengguna Google	5	baru belajar trading cukup mudah memahaminya	2025-10-22 13:29:13
5 4cd5e03-42c7-47b0-9503-2c535e79ba07	Pengguna Google	2	maaf sebelumnya knp apk nya tidak bisa di buka...	2025-10-22 11:28:32
6 e555894a-1d61-4013-9a62-122b4e09ed05	Pengguna Google	5	bibit is ok for me	2025-10-22 10:46:09
7 58a88fc2-441c-4a76-909b-42a0fb465dbb	Pengguna Google	5	aplikasi terbaik sepanjang masa	2025-10-22 03:53:29
8 ba141ed0-1e0d-4059-bb2a-a8a3203d8a98	Pengguna Google	5	mantap terpercaya	2025-10-22 03:14:02
9 a50af5e0-57d9-4841-8c03-c42e0d92ae0	Pengguna Google	5	mantap deh buat tabungan masa depan	2025-10-22 03:04:35
10 bad8329-c449-4b77-9ef1-e9119527f691	Pengguna Google	5	sangat mempermudah untuk dapat berinvestasi se...	2025-10-21 22:57:51
11 69e355d3-495a-4f5a-994f-015efee8e5cc	Pengguna Google	5	Apakah bisa mengadakan tabungan emas? Kemudahan...	2025-10-21 18:56:49
12 01855ac3-2e2a-4570-885d-b9a4004f1e02	Pengguna Google	1	goblog gabisa login, tolol, malah stak di log...	2025-10-21 15:00:39
13 0713c1fb-1e62-4a60-b203-16c1b91f350d	Pengguna Google	4	apk bagus, tapi saran tambahkan aset digital...	2025-10-21 14:25:44
14 658a6c30-37c7-4019-be46-6425142ac89c	Pengguna Google	5	keren	2025-10-21 13:35:04

Fig. 2. Data Scraping Google Play apps and reviews.

B. Data Splitting

The dataset was divided into training (6,240 samples, 79.99%) and Number of training data and test data stratified random sampling to maintain proportional sentiment distribution in both subsets. This 80:20 split is strategically chosen for this dataset size, as 6,240 training samples provide sufficient statistical power for the SVM algorithm to learn robust decision boundaries between sentiment classes, while 1,561 test samples offer reliable performance estimates with reasonable confidence intervals. The stratified approach ensures each sentiment class maintains its original proportion in both sets, preventing the model from being biased toward the majority class during training and ensuring that test performance metrics accurately reflect real-world distribution, which is critical given the moderate class imbalance (39.57% positive, 33.91% neutral, 26.52% negative).

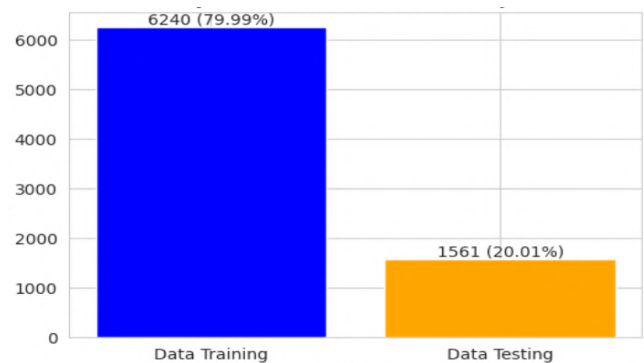


Fig 3. Splitting Training Data and Test Data

The visualization of the bars in the image shows a clear comparison of the amount of data, where blue represents the training data and the orange color represents the test data. This stage is an important part of the research methodology because it determines the balance between the training process and the evaluation of the model, especially in the implementation of the Support Vector Machine (SVM) algorithm used in sentiment analysis.

(11-point F1 gap from positive) because Indonesian users often express criticism indirectly through polite suggestions or mixed-sentiment statements, and the neutral class performs intermediately as it serves as an ambiguous boundary zone between sentiments. The precision-recall patterns reveal critical insights: positive class over-prediction (recall 0.91 > precision 0.89) suggests some neutral reviews are captured as positive, while negative class conservatism (precision 0.81 > recall 0.78) means 22% of actual complaints are missed—a significant concern for automated monitoring systems. The macro-average F1-score of 0.84 is competitive for Indonesian sentiment analysis with lexicon-based labeling, though the 16-point performance gap between best and worst classes indicates room for improvement in handling linguistic nuances, particularly for detecting indirectly expressed negative sentiment that could lead to underestimation of user complaints in practical application.

TABLE I. Sentiment Class

Sentiment Class	Precision	Recall	F1-Score
Positive	0.89	0.91	0.90
Neutral	0.84	0.82	0.83
Negative	0.81	0.78	0.79
Average (Macro)	0.85	0.84	0.84

The positive sentiment class dominated with 3,087 reviews (39.57%), followed by neutral sentiment with 2,645 reviews (33.91%), and negative sentiment with 2,069 reviews (26.52%). These results show that most users give positive responses to the performance of the Bibit application, indicating a high level of user satisfaction.

F. Confusion Matrix and Error Analysis

Figure 7 shows a confusion matrix that provides a detailed overview of the model's classification performance in each sentiment category. The SVM model correctly classified 366 negative reviews, 469 neutral reviews, and 590 positive reviews, resulting in 1,425 correct predictions out of 1,561 test samples (91.3% accuracy). The matrix reveals 136 misclassifications, with the most significant confusion occurring between negative and neutral classes (47 cases).

Negatif	366	47	1
	32	469	28
	1	27	590
	Negatif	Netral	Positif

Fig 7. Confusion Matrix

The model rarely confuses opposite sentiments (only 2 cases total between positive and negative), indicating effective separation of clearly contrasting opinions. The neutral class experiences the most confusion overall (60 misclassifications), as it serves as the boundary between positive and negative

sentiments. These patterns suggest that the model struggles primarily with indirect or politely worded criticism common in Bahasa Indonesia, where complaints are often expressed as neutral-sounding suggestions or factual statements rather than explicit negative expressions.

G. Classification Report

Figure 8 shows the complete classification report that summarizes the SVM model evaluation metrics for each sentiment class, with an overall accuracy of 91.3%. The positive class achieves the strongest performance (precision: 0.953, recall: 0.955, F1-score: 0.954), followed by the negative class (precision: 0.917, recall: 0.884, F1-score: 0.900), while the neutral class shows the lowest performance (precision: 0.864, recall: 0.887, F1-score: 0.875). The macro average F1-score of 0.910 and weighted average of 0.913 demonstrate strong overall performance with minimal class imbalance bias. The main challenge identified is the model's tendency to misclassify negative sentiment as neutral (47 cases), primarily due to indirect criticism and mixed-sentiment expressions typical in Indonesian language, such as reviews using contrastive conjunctions like "tetapi" (but) or polite phrasing like "mungkin perlu ditingkatkan" (maybe needs improvement). Despite these linguistic challenges, the model provides reliable automated sentiment analysis for fintech application reviews, though future improvements using contextual embeddings like IndoBERT could enhance classification of ambiguous cases.

Classification Report for SVM:				
	precision	recall	f1-score	support
Negatif	0.917	0.884	0.900	414.000
Netral	0.864	0.887	0.875	529.000
Positif	0.953	0.955	0.954	618.000
accuracy	0.913	0.913	0.913	0.913
macro avg	0.911	0.908	0.910	1561.000
weighted avg	0.913	0.913	0.913	1561.000

Fig 8. Classification Report for SVM

The dominance of positive sentiment indicates that most users perceive Bibit's interface, investment automation features, and ease of use positively. However, the proportion of negative reviews (26.52%) shows that technical issues, late updates, or transactional delays still affect user perception. The misclassification observed in the confusion matrix particularly between neutral and negative sentiment suggests the presence of ambiguous expressions typical in Bahasa Indonesia, such as indirect criticism or mixed sentiment within a single review. This highlights the need for deeper linguistic modeling. The relatively strong performance of the SVM model demonstrates that traditional machine-learning approaches remain highly competitive for domain-specific sentiment tasks, even when compared to more complex deep-learning models.

V. CONCLUSION

This research successfully applied the Support Vector Machine (SVM) method to conduct sentiment analysis on user reviews of the Bibit application from the Google Play Store. The SVM model classified reviews into three categories positive, neutral, and negative based on text features processed using Text Mining and TF-IDF approaches. The

model achieved the highest precision and recall in the positive class, while performance in neutral and negative classes was slightly lower, likely due to semantic similarities in Indonesian text and ambiguous expressions. The SVM algorithm effectively separated data by optimal margins between sentiment classes, demonstrating its suitability for analyzing user opinions in fintech applications. The research contributes to the development of automated public opinion monitoring systems in Indonesia's growing fintech and digital investment sector. For future research, deep learning approaches such as Bidirectional LSTM or BERT are recommended to improve classification accuracy and natural language understanding. Additionally, collecting data from diverse platforms (Twitter or investment forums) could enrich the analysis and provide a more comprehensive view of public perception toward digital investment applications in Indonesia.

REFERENCES

- [1] T. Salma and F. Natsir, "Penerapan Analytical Hierarchy Process (AHP) Dalam Sistem Pendukung Keputusan Berbasis Web Untuk Pemilihan Ketua OSIS," *J. Inform. SIMANTIK*, vol. 10, no. 2, pp. 8–15, 2025.
- [2] A. P. Tirtopangarsa and W. Maharani, "Sentiment analysis of depression detection on twitter social media users using the K-nearest neighbor method," in *Seminar Nasional Informatika (SEMNASIF)*, 2021, pp. 247–258.
- [3] T. Triyadi, R. A. Sihombing, and F. Natsir, "Perancangan Aplikasi Analisis Sentimen terhadap Opini Penghapusan Skripsi pada Twitter menggunakan Metode Navie Bayes," *Elconika J. Tek. Elektro*, vol. 2, no. 1 SE-Articles, pp. 1–14, Dec. 2023, doi: 10.33752/elconika.v2i1.5518.
- [4] E. B. Susanto, Paminto Agung Christianto, Mohammad Reza Maulana, and Satriedi Wahyu Binabar, "Analisis Kinerja Algoritma Naive Bayes Pada Dataset Sentimen Masyarakat Aplikasi NEWSAKPOLE Samsat Jawa Tengah," *J. CoSciTech (Computer Sci. Inf. Technol.)*, vol. 3, no. 3, pp. 234–241, 2022, doi: 10.37859/coscitech.v3i3.4343.
- [5] A. Impron et al., *Enterprise Resource Planning: Implementasi dan Manajemen*. Penerbit Widina, 2025.
- [6] N. S. Aji, F. Natsir, and S. Istianah, "Penentuan Penjualan Barang Berdasarkan Pengelompokan Produk dengan K-Means Clustering Metode CRISP-DM Pada CV Sembako Dina," *J. Zetroem*, vol. 5, no. 2, pp. 119–126, 2023, doi: 10.36526/ztr.v5i2.3041.
- [7] A. Sudiarjo, M. Praseptiawan, N. G. Setyoningrum, H. M. Drajat, and F. Natsir, "Analysis of Image Improvement and Edge Identification Methods in Watermelon Image," *Innov. Innov. Res. Informatics*, vol. 6, no. 1, pp. 35–40, 2024.
- [8] W. Andriyani et al., *Perangkat Lunak Data Mining*. Penerbit Widina, 2024.
- [9] T. Nurdin, M. Alexandri, W. Sumadinata, and R. Arifianti, "Sentiment Analysis of User Preference for Old Vs New Fintech Technology Using SVM and NB Algorithms," *Manag. Syst. Prod. Eng.*, vol. 31, pp. 373–380, 2023, doi: 10.2478/mspe-2023-0041.
- [10] M. Ahmad, S. Aftab, M. S. Bashir, and N. Hameed, "Sentiment Analysis using SVM: A Systematic Literature Review," *Int. J. Adv. Comput. Sci. Appl.*, vol. 9, 2018, [Online]. Available: <https://api.semanticscholar.org/CorpusID:3695786>
- [11] R. Alawaji and A. Aloraini, "Sentiment Analysis of Digital Banking Reviews Using Machine Learning and Large Language Models," *Electronics*, vol. 14, no. 11, 2025, doi: 10.3390/electronics14112125.
- [12] Warjiyono, S. Aji, Fandhilah, N. Hidayatun, H. Faqih, and Liesnaningsih, "The Sentiment Analysis of Fintech Users Using Support Vector Machine and Particle Swarm Optimization Method," in *2019 7th International Conference on Cyber and IT Service Management (CITSM)*, 2019, pp. 1–5. doi: 10.1109/CITSM47753.2019.8965348.
- [13] N. Hidayati, F. Hamami, and R. Y. Fa'rifah, "Aspect-Based Sentiment Analysis On FLIP Application Reviews (Play Store) Using Support Vector Machine (SVM) Algorithm," *J. Informatics Telecommun. Eng.*, vol. 7, no. 1, pp. 183–197, 2023.
- [14] I. N. T. I. N. Ikhsan, R. Rudiman, and F. Y. Fendy, "Analisis Sentimen Pada Ulasan Aplikasi Google Maps Terhadap Pelayanan Badan Penyelenggara Jaminan Sosial (BPJS) Kesehatan Samarinda Menggunakan Metode K-Nearest Neighbor Dengan Fitur Ekstraksi TF-IDF," *J. Teknol. Inf. J. Keilmuan dan Apl. Bid. Tek. Inform.*, vol. 18, no. 2, pp. 100–116, 2024.



Ega Shela Marsiani

Ega Shela Marsiani, holds a master's degree in computer science and is currently a lecturer in the Informatics Study Program, Faculty of Engineering and Computer Science, Universitas Indraprasta PGRI, Jakarta, Indonesia. Her research focuses on decision support systems, artificial intelligence, data mining, and information system development. She has conducted research on employing AHP, TOPSIS, and machine learning algorithms for data analysis and classification in several fields of information technology.



Fauzan Natsir

He is a lecturer and researcher in the field of Informatics at Department Informatic Engineering, Universitas Indraprasta PGRI, South Jakarta. Many research and papers have been published in various national and international journals. His research interests include Digital Forensic, Software Engineering, Decision Support System, Data Mining, Knowledge Management System, IoT (Internet of Things), Machine Learning, AI (Artificial Intelligence). He also became a lecturer under the Junior Web Developer and Office Applications scheme at BPPTIK, Ministry of Communication and Information Technology.



Redo Abeputra Sihombing

He was born in Pekanbaru in 1991. He completed his undergraduate degree (S1) in 2014 at Budi Luhur University, then continued his master's degree (S2) in 2017 at Budi Luhur University. Currently, he is registered as a permanent lecturer at Indraprasta PGRI University in the Informatics Engineering Study Program. In addition to his academic dedication in education, he is also active as a technology practitioner, working as a developer specializing in Software Engineering, IoT (Internet of Things), Machine Learning, AI (Artificial Intelligence), and other technologies. His commitment to technology development and education keeps him active in various research and scientific publications in the field of information technology and computing. The author can be contacted through LinkedIn "Redo Abeputra Sihombing", Instagram @redoabe, or email at redoabe@gmail.com.



Millati Izzatillah

She is a certified lecturer serving as a Lektor in the Informatics Engineering Department at Universitas Indraprasta PGRI. Her research interests focus on expert systems, information systems, and software engineering, with a particular emphasis on developing innovative and practical technology-based solutions. She has actively participated in academic and professional forums as a speaker, sharing insights on digital transformation, education, and applied informatics. In addition to her academic responsibilities, she is also a

BNSP-certified content creator who integrates technology and creativity to promote digital literacy. Through her work, she aims to bridge the gap between theoretical research and its real-world applications in education and technology. She can be found on social media as @izzatimilla, where she shares insights on education, technology, and digital creativity.



Rajiansyah

He is both an academic and a researcher with a strong commitment to advancing knowledge in the field of computer science and informatics. In his role as a lecturer at Universitas Mulawarman, he focuses on teaching and mentoring students, particularly in areas related to probability and statistics, programming, and applied computer science. His doctoral research at Wroclaw University of Science and Technology builds upon his expertise by exploring advanced methods in artificial intelligence, decision theory, and optimization models to address complex real-world problems. By combining his teaching experience in Indonesia with his research training in Poland, Rajiansyah embodies a global academic perspective, integrating theory, practice, and innovation. His passion lies not only in contributing to scientific progress but also in shaping the next generation of computer scientists and engineers.