



Comparison of Support Vector Machine and Random Forest for Sentiment Analysis of Gojek Reviews

Erlyne Nadhilah Widyaningrum ^{1}, Filzah Syakirah ¹, M. Fathurahman¹*

¹ Statistics Study Program, Mulawarman University, Indonesia

*Corresponding Author: erlynenadhilah@fmipa.unmul.ac.id

Abstract: The development of digital technology has increased the use of online transportation applications in Indonesia, one of which is Gojek. User reviews on the Gojek application in the Google Play Store reflect the level of user satisfaction and dissatisfaction with the services provided. Therefore, sentiment analysis is needed to classify these reviews into positive and negative categories. This study aims to get the result and performance of Support Vector Machine (SVM) and random forest classification methods in sentiment analysis of Gojek user reviews. The data used are user reviews from January 2026 obtained through a scraping technique. The analysis stages include text processing, word weighting using TF-IDF, and hyperparameter optimization using the random search method. The results show that the SVM method achieved an accuracy of 86,78%, precision of 87,96%, recall of 86,23%, and F1-score of 87,09%. Meanwhile, the random forest method achieved an accuracy of 84,75%, precision of 82,87%, recall of 88,85%, and F1-score of 85,76%. Based on these results, the SVM method demonstrates superior overall performance compared to Random Forest in classifying sentiment of Gojek user reviews.

Keywords: Gojek, Random Forest, SVM, Sentiment Analysis, User Reviews

Introduction

The rapid advancement of information and communication technology has significantly transformed various aspects of human life, including the transportation sector. The emergence of online transportation applications has provided practical and efficient solutions for mobility, especially for people who require fast, flexible, and accessible transportation services. In Indonesia, online transportation services have become an essential part of daily activities because they not only offer transportation services but also support various digital services such as food delivery, logistics, and electronic payment systems. The increasing use of smartphones and internet access has further accelerated the growth of online transportation platforms in Indonesian society.

One of the most widely used digital service platforms in Indonesia is Gojek. Gojek provides various services, including transportation, food delivery, shopping, courier services, and digital payment through GoPay. As one of the largest digital platforms in Indonesia, Gojek continuously receives feedback from users regarding service quality and application performance. User reviews available on the Google Play Store contain opinions, experiences, criticisms, and suggestions from users after using the application. These reviews can reflect the level of user satisfaction and public perception toward the services provided by the application [1]. Therefore, user reviews can be utilized as an important source of information for evaluating service quality and improving application performance.

On January 13, 2026, Gojek experienced several system disruptions, including problems in location detection, difficulties in accessing services, and transaction failures in the GoPay feature. These disruptions caused many users to submit complaints and negative reviews through the Google Play Store platform [2]. This phenomenon demonstrates that user reviews can serve as real-time indicators of application performance and service quality. Positive reviews generally indicate user satisfaction, while negative reviews reflect dissatisfaction or problems experienced by users. Because of the large amount of review data generated continuously, manual analysis becomes inefficient and time-consuming. Therefore, an automated sentiment analysis approach is needed to identify and classify user opinions effectively.

Sentiment analysis is one of the applications of Natural Language Processing (NLP) that is used to analyze and classify opinions, emotions, and attitudes expressed in textual data into certain sentiment categories, such as positive and negative sentiments [3]. Sentiment analysis has been widely applied in various fields, including product review analysis, social media monitoring, customer satisfaction evaluation, and public opinion analysis. In this study, sentiment analysis was conducted on Gojek user reviews collected from the Google Play Store. The analysis aims to identify public sentiment patterns toward the Gojek application and evaluate the effectiveness of machine learning methods in classifying review sentiments.

The classification process in this study applies supervised machine learning methods, namely Support Vector Machine (SVM) and Random Forest. SVM is widely used in text classification because of its ability to handle high-dimensional and nonlinear data effectively [4]. In this study, SVM uses the Radial Basis Function (RBF) kernel, which is capable of modeling complex relationships in textual data. Meanwhile, Random Forest is an ensemble learning method that combines multiple decision trees to improve classification performance and reduce overfitting. Random Forest is also effective for handling large and complex datasets with stable classification performance [5].

Before the classification process, textual review data must be transformed into numerical representations that can be processed by machine learning algorithms. Therefore, this study applies the Term Frequency–Inverse Document Frequency (TF-IDF) method for feature extraction. TF-IDF is commonly used in text mining because it is capable of measuring the importance of words in a document while reducing the influence of commonly occurring words. Previous studies have shown that the combination of TF-IDF with SVM and Random Forest can produce good classification performance in sentiment analysis tasks [6].

In addition to selecting appropriate classification methods, hyperparameter optimization is also important to improve model performance. Conventional optimization methods such as Grid Search often require high computational cost because they evaluate all possible parameter combinations. Therefore, this study utilizes Random Search optimization, which is more computationally efficient while still capable of finding optimal hyperparameter combinations [7]. Random Search is applied to optimize the parameters of the SVM and Random Forest models in order to improve sentiment classification performance.

Based on these considerations, this study aims to compare the performance of SVM and Random Forest methods in sentiment analysis of Gojek user reviews using TF-IDF feature extraction and Random Search optimization. The results of this study are expected to provide insights into public sentiment toward the Gojek application and identify the most effective classification method for sentiment analysis of user reviews.

Methodology

1. Sentiment Analysis

Sentiment analysis is one of the most popular applications of Natural Language Processing (NLP). This analysis focuses on extracting sentiments from text content to identify whether an opinion tends to be positive or negative [8]. Also known as opinion mining, sentiment analysis is widely used to understand public opinions toward certain issues, products, services, or policies.

Currently, the need to analyze public opinion continues to increase in various fields. In business sectors, sentiment analysis can support product evaluation and development, while in government sectors, it can assist in improving policies and public services [9]. The data commonly used in sentiment analysis include articles, tweets, product reviews, and social media comments. The main objective of this analysis is to classify text polarity into positive or negative categories [10].

2. Term Frequency-Inverse Document Frequency

Term Frequency–Inverse Document Frequency (TF-IDF) is a feature extraction method commonly used in text mining and sentiment analysis to convert textual data into numerical representations. TF-IDF measures the importance of a word in a document by considering both the frequency of the word in a specific document and its occurrence across all documents in the dataset. Words that frequently appear in a document but rarely appear in other documents will have higher weights [6].

The Term Frequency (TF) value represents the frequency of a term in a document and can be calculated using the following formula:

$$TF(t, d) = \frac{f_{t,d}}{N_d} \quad (1)$$

where $f_{t,d}$ represents the number of occurrences of term t in document d .

The Inverse Document Frequency (IDF) value is used to measure the importance of a term across all documents and is calculated as follows:

$$\text{IDF}(t) = \log\left(\frac{N}{\text{DF}(t)}\right) \quad (2)$$

where N is the total number of documents and $\text{DF}(t)$ is the number of documents containing term t .

The TF-IDF weight (W) is obtained by multiplying the TF and IDF values, as shown in the following formula:

$$W(t, d) = \frac{f_{t,d}}{N_d} \times \left[\log\left(\frac{N}{\text{DF}(t)}\right) \right] \quad (3)$$

The resulting TF-IDF weights are then used as input features for the SVM and Random Forest classification models in the sentiment analysis process.

3. Support Vector Machine

Support Vector Machine (SVM) is a supervised machine learning method commonly used for classification tasks, including sentiment analysis. SVM was developed by Boser, Guyon, and Vapnik in 1992 and works by finding an optimal hyperplane that separates data into different classes with the maximum margin [11]. The data points located closest to the hyperplane are called support vectors, which play an important role in determining the classification boundary.

In SVM, the optimal hyperplane is determined by maximizing the margin between classes. For linearly separable data, the hyperplane can be expressed as follows:

$$\mathbf{w}^T \mathbf{x}_i + b = 0 \quad (4)$$

where w represents the weight vector and b represents the bias value. SVM aims to obtain the best decision boundary to improve classification performance and generalization ability.

However, not all datasets can be separated linearly. To address this problem, SVM applies the kernel trick to transform data into a higher-dimensional feature space. One of the most widely used kernels is the Radial Basis Function (RBF) kernel because of its ability to capture non-linear relationships in text data [12]. The RBF kernel function is defined as follows:

$$K(\mathbf{x}_i, \mathbf{x}_j) = \exp\left(-\gamma \|\mathbf{x}_i - \mathbf{x}_j\|^2\right) \quad (5)$$

where γ controls the influence of each data point on the classification boundary.

The process of determining data classification can be defined through the following decision function:

$$f(x) = \text{sign}\left(\sum_{i=1}^r \alpha_i y_i K(\mathbf{x}_i, \mathbf{x}_u) + b\right) \quad (6)$$

The performance of SVM is influenced by several hyperparameters, particularly the cost parameter (C) and gamma (γ). The parameter C controls the trade-off between maximizing the margin and minimizing classification errors. A large C value produces a more complex model with lower tolerance for misclassification, while a smaller C value generates a simpler and more generalized model [13]. Therefore, selecting optimal hyperparameter values is important to achieve better classification performance.

SVM is widely applied in sentiment analysis because it performs effectively on high-dimensional and sparse text data. Previous studies have shown that SVM often achieves better performance than several conventional classification methods in text classification tasks. Its ability to produce stable and accurate classification results makes SVM suitable for sentiment analysis of user reviews and opinion mining datasets.

4. Random Forest

Random Forest is a method that can be used for classification tasks. This method was first introduced by Breiman in 2001 [14]. According to Ainur et al. (2025), Random Forest is effective in handling large-scale data and is flexible for various types of data, including text data [15]. By constructing multiple decision trees combined through ensemble learning, this method is able to produce more accurate predictions and effectively reduce overfitting [16]. These advantages make Random Forest widely used, particularly in classification problems.

The Random Forest method consists of a collection of structured trees that generate decisions based on majority voting. The basic technique of this method is derived from the Decision Tree algorithm. A Decision Tree is a method that separates a group of data using a tree structure [17]. The structure of a Decision Tree consists of three main parts: the root node as the starting point, branch nodes, and leaf nodes as the final output. In the context of classification, Random Forest constructs multiple decision trees and combines the prediction results from each tree to obtain more accurate predictions [18]. According to [19] the general working process of Random Forest can be illustrated in Figure 1.

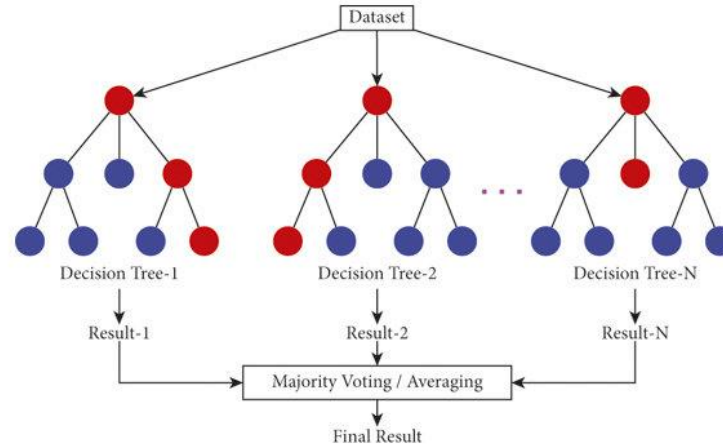


Figure 1: Illustration of the Random Forest Classification Process

Random Forest applies bootstrap aggregating (bagging) and random feature selection to improve model stability and reduce overfitting [20]. In this method, multiple bootstrapped datasets are generated to train different decision trees, while random feature selection is applied at each node to increase diversity among trees and reduce correlation between models. The final prediction is obtained through majority voting from all decision trees [21].

According to [19], the Random Forest classification process consists of bootstrap sampling and decision tree construction. During tree construction, node splitting is performed by selecting the best feature based on the Gini Index. A higher Gini Index indicates more diverse data, while a lower value indicates more homogeneous data. The formula used to calculate the Gini Index is as follows:

$$Gini = 1 - \sum_{l=1}^c \pi_l^2 \quad (7)$$

5. Random Search

Random Search is a hyperparameter optimization method that selects parameter combinations randomly from a predefined search space [22]. Compared to exhaustive search methods, Random Search is more efficient in terms of computation time while still being effective in finding optimal model performance, especially for high-dimensional data [23]. In this study, Random Search was used to optimize the hyperparameters of SVM and Random Forest models.

The main concept of Random Search is to randomly sample hyperparameter values from a specified range and evaluate the performance of each combination. Unlike Grid Search, which tests all possible parameter combinations, Random Search explores only a subset of the search space, making the optimization process faster and more computationally efficient. This approach is particularly useful when dealing with large hyperparameter spaces where exhaustive search requires high computational cost and longer processing time.

Random Search has been widely applied in machine learning because it is capable of finding near-optimal hyperparameter combinations with relatively fewer iterations. Previous studies have shown that Random Search can improve model performance by identifying parameter configurations that produce better classification accuracy and generalization ability [24]. In

addition, the randomness of the search process allows the method to explore diverse parameter combinations and avoid focusing only on limited regions of the search space.

6. Gojek

Gojek is a startup company founded by Nadiem Makarim in 2010. Currently, its services are available in almost all cities across Indonesia. Initially, Gojek only operated in the ride-hailing sector, but as the company rapidly expanded, it broadened its focus to other services such as GoRide, GoCar, GoFood, GoMart, and many others [25]. As an online-based platform providing various services, Gojek has brought significant impacts to society in terms of mobility, shopping, and daily activities. On the Google Play Store, the Gojek application has surpassed more than 100 million downloads with a rating of 4.7. User experiences while using the services provided by the application are reflected in their reviews. These reviews can serve as valuable evaluation material for the company to improve and enhance the quality of services offered by the Gojek application [26].

Results and Discussions

This section presents the results of hyperparameter optimization, sentiment classification, and evaluation of the Support Vector Machine (SVM) and Random Forest models on Gojek user review data. The experiments were conducted using three train–test split proportions, namely 70:30, 80:20, and 90:10.

1. Text Processing

The data obtained from scraping Gojek user reviews on the Google Play Store is inherently raw and unstructured. This is because the extracted text contains various forms of noise, such as punctuation marks, numbers, emojis, and non-standard words, making it unsuitable for direct classification. Therefore, a text processing stage is essential to clean the dataset and enhance its quality. The text processing workflow implemented in this study consists of five consecutive stages: case folding, cleaning, tokenizing, stopwords removal, and stemming.

To provide a clear illustration of how the raw data is transformed at each stage, the processing of a sample user review instance is demonstrated below:

Table 1. Sample Results of the Text Processing Stages

Review	Case Folding	Cleaning	Tokenizing	Stopword Removal	Stemming
Pemesanan makanan dalam aplikasi ini sangat jelek, sudah menunggu lama, pas dapat driver, malah drivernya yang cancel, sudah begitu harus menunggu lama jika ingin membatalkan pesanan. buruk sekali	pemesanan makanan dalam aplikasi ini sangat jelek, sudah menunggu lama, pas dapat driver, malah drivernya yang cancel, sudah begitu harus menunggu lama jika ingin membatalkan pesanan. buruk sekali	pemesanan makanan dalam aplikasi ini sangat jelek, sudah menunggu lama, pas dapat driver, malah drivernya yang cancel, sudah begitu harus menunggu lama jika ingin membatalkan pesanan. buruk sekali	['pemesanan', 'makanan', 'dalam', 'aplikasi', 'ini', 'sangat', 'jelek', 'sudah', 'menunggu', 'lama', 'pas', 'dapat', 'driver', 'malah', 'drivernya', 'yang', 'cancel', 'sudah', 'begitu', 'harus', 'menunggu', 'lama', 'jika', 'ingin', 'membatalkan', 'pesanan', 'buruk', 'sekali', 'pemesanan',	['pemesanan', 'makanan', 'aplikasi', 'sangat', 'jelek', 'tunggu', 'lama', 'pas', 'driver', 'malah', 'drivernya', 'cancel', 'begitu', 'menunggu', 'lama', 'ingin', 'membatalkan', 'pesanan', 'buruk', 'sekali']	['mesan', 'makan', 'aplikasi', 'sangat', 'jelek', 'tunggu', 'lama', 'pas', 'driver', 'malah', 'drivernya', 'cancel', 'begitu', 'tunggu', 'lama', 'ingin', 'batal', 'pesan', 'buruk', 'sekali']

'buruk',
'sekali']

2. Term Frequency- Inverse Document Frequency Score

The Term Frequency-Inverse Document Frequency (TF-IDF) score serves as the feature weighting mechanism utilized in the subsequent classification modeling phase. Following the text preprocessing stage, the TF-IDF calculation is performed to transform the qualitative textual data into a structured numerical representation. This process involves executing distinct computations for Term Frequency (TF) and Inverse Document Frequency (IDF), which are subsequently integrated to yield the final weight. The compiled results of the TF-IDF scoring process are presented in Table 2 below:

Table 2. TF-IDF Weighting Results for Review Terms

No.	Kata	$W(t, d)$						
		1	2	3	...	2947	2948	2949
1	aamiin	0	0	0		0	0	0
2	abai	0	0	0		0	0	0
3	abang	0	0	0		0	0	0
4	abis	0	0	0		0	0	0
:	:	:	:	:		:	:	:
4412	zonk	0	0	0		0	0	0
4413	zoom	0	0	0		0	0	0

Based on the results presented in Table 2, it can be observed that the preprocessed dataset comprises a total of 2,949 distinct review documents (d), yielding a vocabulary size of 4,413 unique terms (t) after the text cleaning phase. Consequently, this transforms the unstructured text into a feature matrix of dimensions 4413×2949 . The matrix exhibits a high degree of sparsity, as indicated by the predominance of zero values (0). This sparsity occurs because each specific term (such as "aamiin", "abai", or "zoom") only appears in a limited subset of the total reviews. A non-zero weight, $W(t, d)$, signifies not only the presence of a word within a particular review but also reflects its relative importance and uniqueness across the entire corpus, which serves as the fundamental input for the machine learning classifiers.

3. Word Cloud



Figure 2. User Reviews of Gojek Applications Word Cloud

Figure 2 shows that the largest words in the word cloud are “driver” and “gojek.” This indicates that most users frequently used the words “driver” and “gojek” when reviewing the Gojek application on Google Play Store. Therefore, it can be interpreted that users

mainly discussed drivers and the services provided by the Gojek application itself in their reviews

4. Optimization of SVM Hyperparameter

The results of SVM hyperparameter optimization using the Random Search method on three train-test split proportions are presented in Table 1. The optimization process determined the optimal values of the cost parameter (C) and gamma (γ) for the RBF kernel. In this study, the hyperparameter search range for C was set from 0,01 to 100, while the gamma parameter ranged from 0.001 to 1.

Table 3. Results of Optimization of Hyperparameter SVM

Data Proportions	Hyperparameter SVM	Optimal Value	Mean CV Accuracy	Individual Accuracy
70:30	C	21,36832907	0,8653	0,861
	γ	0,004335282		
80:20	C	5,456725486	0,8703	0,8678
	γ	0,020914981		
90:10	C	5,456725486	0,87	0,8407
	γ	0,020914981		

Based on Table 3, the 80:20 split proportion produced the best performance with a cross-validation accuracy of 0.8703 and model accuracy of 0.8678. Meanwhile, the 70:30 proportion achieved a model accuracy of 0.8610, and the 90:10 proportion obtained the lowest model accuracy of 0.8407.

The results indicate that the selection of hyperparameter values and data split proportions influenced the classification performance of the SVM model. In this study, the 80:20 split proportion provided the most optimal performance for sentiment classification of Gojek user reviews. Based on the best split proportion, the dataset was divided into training and testing data with an 80:20 ratio. The distribution of positive and negative sentiment data used in the classification process is presented in Table 4.

Table 4. Distribution of Training and Testing Data

Sentiment	Training Data	Test Data	Total
Positive	1220	305	1525
Negative	1139	285	1424
Total	2359	590	2949

Table 4 shows that the dataset consisted of 2949 review data, including 1525 positive reviews and 1424 negative reviews. The training data contained 2359 reviews, while 590 reviews were used as testing data. The distribution between positive and negative sentiments was relatively balanced, indicating that the dataset was suitable for the sentiment classification process and reducing the risk of classification bias toward a particular class.

5. Optimization of Random Forest Hyperparameter

The results of Random Forest hyperparameter optimization using the Random Search method on three train-test split proportions are presented in Table 4. The optimization process focused on determining the optimal values of $n_estimators$ and max_depth . In this study, the hyperparameter search range for $n_estimators$ was set from 10 to 200, while the max_depth parameter ranged from 5 to 40.

Table 5. Results of Optimization of Hyperparameter Random Forest

Data Proportions	Hyperparameter Random Forest	Optimal Value	Mean CV Accuracy	Individual Accuracy
70:30	$n_estimators$	90	0,8498	0,8463
	max_depth	39		
80:20	$n_estimators$	113	0,8521	0,8475
	max_depth	40		
90:10	$n_estimators$	59	0,8602	0,8271

max_depth	40
-----------	----

Based on Table 5, the 80:20 split proportion achieved the best model performance with a cross-validation accuracy of 0.8521 and individual accuracy of 0.8475. The 70:30 proportion produced an individual accuracy of 0.8463, while the 90:10 proportion obtained the lowest individual accuracy of 0.8271 despite having the highest mean cross-validation accuracy value of 0.8602.

These results indicate that the train–test split proportion and hyperparameter configuration influenced the performance of the Random Forest model. In this study, the 80:20 split proportion provided the most stable and optimal performance for sentiment classification of Gojek user reviews. Therefore, the Random Forest classification process used the same distribution of training and testing data as presented in Table 2.

6. SVM Classification

The sentiment classification process using the Support Vector Machine (SVM) method was carried out using the optimal hyperparameters obtained from the Random Search optimization stage. The classification aimed to classify Gojek user reviews into two sentiment categories, namely positive and negative sentiments. Based on the optimization results for the 80:20 data split (where $C = 5.456725486$ and $\gamma = 0.020914981$), the specific decision function $f(x)$ used to determine the sentiment class of each new review document is formulated on (6). To illustrate how the trained model classifies a review, a sample calculation was performed using the TF-IDF feature vector of the first training data instance. By substituting the optimized weights, support vectors, and the constant bias the mathematical expansion of the decision function for this specific instance is calculated as follows:

$$\begin{aligned}
 f(x) &= \sum_{i=1}^r \alpha_i y_i K(x_i, x_t) + b \\
 &= \sum_{i=1}^{2359} \alpha_i y_i K(x_i, x_t) - 0,556933 \\
 &= ((1,358185)(-1) K(x_1 \cdot x_1) + \dots + (1,07 \times 10^{-5})(-1) K(x_{2359} \cdot x_1)) - 0,556933 \\
 &= -0,44305 - 0,556933 \\
 &= -0,999983576.
 \end{aligned}$$

The negative output indicates that the first sample is correctly mapped by the hyperplane. Following this classification process across the entire test set, the overall performance of the SVM model was evaluated using a confusion matrix to measure the classification results on the testing data. The confusion matrix results of the SVM classification model are presented in Table 6.

Table 6. Confusion Matrix of SVM Classification Results

		Prediction		Total
		Negative	Positive	
Actual	Negative	249	36	285
	Positive	42	263	305
Total		291	299	590

Table 6 presents the confusion matrix results of the SVM classification model. Based on the table, the model correctly classified 249 negative reviews and 263 positive reviews. Meanwhile, 36 negative reviews were misclassified as positive, and 42 positive reviews were misclassified as negative.

The total number of correctly classified data was 512 out of 590 testing data, indicating that the SVM model performed well in distinguishing positive and negative sentiments in Gojek user reviews. The results demonstrate that the SVM model with the RBF kernel was effective in handling sentiment classification on text-based data.

7. Random Forest Classification

The sentiment classification process using the Random Forest method was carried out using the optimal hyperparameters obtained from the Random Search optimization stage. The classification aimed to classify Gojek user reviews into two sentiment categories, namely positive and negative sentiments. During the construction of the decision trees, node splitting was determined by selecting the best feature that minimized the data impurity, measured via the Gini Index. To illustrate this mechanism within the first decision tree, the Gini split calculation for the word feature 'akhir' was performed as follows:

$$\begin{aligned}
 Gini_{split}('akhir') &= \sum_{i=1}^2 \left(\frac{n_i}{n} \right) \times Gini(S_i) \\
 &= \left(\left(\frac{n_1}{n} \right) \times Gini(S_1) \right) + \left(\left(\frac{n_2}{n} \right) \times Gini(S_2) \right) \\
 &= \left(\left[\left(\frac{2182}{2215} \right) \times 0,498 \right] + \left[\left(\frac{33}{2215} \right) \times 0,3343 \right] \right) \\
 &= 0,4906 + 0,0046 \\
 &= 0,4956.
 \end{aligned}$$

The feature that yields the lowest Gini split value is selected as the optimal splitting node for the tree layout. Following the construction of all optimized ensemble trees ($n_estimators = 113$), the overall performance of the Random Forest model was evaluated using a confusion matrix to measure the classification results on the testing data. The confusion matrix results of the Random Forest classification model are presented in Table 7.

Table 7. Confusion Matrix of Random Forest Classification Results

		Prediction		Total
		Negative	Positive	
Actual	Negative	229	56	285
	Positive	34	271	305
Total		291	265	590

Table 7 presents the confusion matrix results of the Random Forest classification model. Based on the table, the model correctly classified 229 negative reviews and 271 positive reviews. Meanwhile, 56 negative reviews were misclassified as positive, and 34 positive reviews were misclassified as negative.

The total number of correctly classified data was 500 out of 590 testing data, indicating that the Random Forest model performed well in distinguishing positive and negative sentiments in Gojek user reviews. The results demonstrate that the Random Forest model was effective in handling sentiment classification on text-based data.

8. Evaluation Results

After the classification process was completed, the performance of the SVM and Random Forest models was evaluated using several evaluation metrics, namely accuracy, precision, recall, and F1-score. These metrics were used to measure the effectiveness of each model in classifying positive and negative sentiments from Gojek user reviews. The evaluation results of both classification models are presented in Table 8.

Table 8. Evaluation Metrics of Classification Results

Methods	Accuracy	Pesisi	Recall	F1-Score
SVM	0,8678	0,8796	0,8623	0,8709
Random Forest	0,8475	0,8287	0,8885	0,8576

Table 8 presents the evaluation results of the SVM and Random Forest models using accuracy, precision, recall, and F1-score metrics. Based on the results, the SVM model showed better overall performance than the Random Forest model in classifying sentiments from Gojek user reviews during January 2026. This can be seen from the higher accuracy, precision, and F1-

score values achieved by SVM, indicating that the model was more optimal in recognizing sentiment patterns within the review data.

Although SVM achieved better overall performance, the Random Forest model also produced relatively good results with evaluation values that were not significantly different from SVM. In addition, Random Forest obtained a higher recall value, indicating that the model was better at correctly identifying positive sentiment data. This result shows that Random Forest had a strong capability in recognizing positive user satisfaction, thereby reducing the possibility of positive reviews being misclassified as negative sentiments.

However, overall, the SVM model provided more stable and balanced classification results for both positive and negative sentiments. Therefore, based on the evaluation results, SVM is considered more suitable for sentiment classification of Gojek user reviews because it produced more accurate, consistent, and balanced performance compared to Random Forest.

Conclusion

Based on the results of data analysis and discussion, it can be concluded that both SVM and Random Forest methods were able to perform sentiment classification on Gojek user reviews effectively. The SVM model correctly classified 512 out of 590 testing data, consisting of 249 negative reviews and 263 positive reviews. Meanwhile, the Random Forest model correctly classified 500 out of 590 testing data, consisting of 229 negative reviews and 271 positive reviews.

The evaluation results also showed that both methods achieved good classification performance. The SVM model obtained an accuracy of 86.78%, precision of 87.96%, recall of 86.23%, and F1-score of 87.09%. Meanwhile, the Random Forest model achieved an accuracy of 84.75%, precision of 82.87%, recall of 88.85%, and F1-score of 85.76%. Although Random Forest produced a higher recall value, the SVM model provided better overall performance with higher accuracy, precision, and F1-score values. Therefore, SVM is considered more suitable and effective for sentiment classification of Gojek user reviews in this study.

Acknowledgments

-

References

- [1] D. Nugraha and D. Gustian, "Analisis Sentimen Penggunaan Aplikasi Transportasi Online Pada Ulasan Google Play Store Menggunakan Algoritma Svm (Support Vector Machine)," in *Seminar Nasional Sistem Informasi dan Manajemen Informatika Universitas Nusa Putra*, 2023, vol. 3, pp. 99–106.
- [2] Alwiyah, "Gojek Down di 13 Januari 2026? Banyak User Keluhkan Dampak Layanan yang Bermasalah.," 2026. <https://narasi.tv/Read/Narasi-Daily/Gojek-down-Di-13-Januari-2026-Banyak-User-Keluhkan-Dampak-Layanan-Yang-Bermasalah>.
- [3] V. Machairas, A. Tsangrassoulis, and K. Axarli, "Algorithms for optimization of building design: A review," *Renew. Sustain. Energy Rev.*, vol. 31, no. 1364, pp. 101–112, 2014, doi: 10.1016/j.rser.2013.11.036.
- [4] W. Lestari and S. Sumarlinda, "Implementation of k-nearest neighbor (knn) and suport vector machine (svm) for clasification cardiovascular disease," *Int. J. Multi Sci.*, vol. 2, no. 10, pp. 30–36, 2022.
- [5] I. S. Aisah, B. Irawan, and T. Suprpti, "Algoritma Support Vector Machine (Svm) Untuk Analisis Sentimen Ulasan Aplikasi Al Qur'an Digital," *JATI (Jurnal Mhs. Tek. Inform.*, vol. 7, no. 6, pp. 3759–3765, 2023.
- [6] A. I. W. AIW, M. W. P. MWP, B. K. S. BKY, V. C. VCR, and M. D. A. MDA, "Implementasi Algoritma Model Random Forest, SVM Dan Naive Bayes Untuk Sentimen Analisis Aplikasi Gojek Di Playstore," *Evolusi J. Sains dan Manaj.*, vol. 13, no. 2, pp. 53–64, 2025.
- [7] M. Fajri and A. Primajaya, "Komparasi Teknik Hyperparameter Optimization pada SVM untuk Permasalahan Klasifikasi dengan Menggunakan Grid Search dan Random Search," *J. Appl. Informatics Comput.*, vol. 7, no. 1, pp. 10–15, 2023.
- [8] A. Halim and A. Safuwan, "Analisis Sentimen Opini Warganet Twitter Terhadap Tes

- Screening Genose Pendeteksi Virus Covid-19 Menggunakan Metode Naïve Bayes Berbasis Particle Swarm Optimization," *J. Inform. Teknol. dan Sains*, vol. 5, no. 1, pp. 170–178, 2023.
- [9] 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.
- [10] A. Adiansyah, "Analisis Sentimen Pada Ulasan Aplikasi Home Credit Dengan Metode SVM dan K-NN," *J. Komput. Antart.*, vol. 1, no. 4, pp. 174–181, 2023.
- [11] V. Vapnik, *The nature of statistical learning theory*. Springer science & business media, 2013.
- [12] Z. Arifin, D. F. Rahman, B. S. Rintyarna, and D. Daryanto, "Penerapan Algoritma Support Vector Machine Berbasis Kernel Radial Basis Function dalam Klasifikasi Sel Kanker," *BIOS J. Teknol. Inf. dan Rekayasa Komput.*, vol. 4, no. 2, pp. 100–106, 2023.
- [13] A. Kowalczyk, "Support vector machines succinctly," *Syncfusion Inc*, vol. 317, 2017.
- [14] S. Mahmuda, "Implementasi metode Random Forest pada kategori konten kanal YouTube," *Deleted J.*, 2024.
- [15] G. Ainur, N. Latifah, and F. Nugraha, "Analisis sentimen publik terhadap program makan bergizi gratis menggunakan metode random forest pada platform x," *J. SITECH Sist. Inf. dan Teknol*, vol. 8, no. 1, pp. 75–84, 2025.
- [16] R. Z. Firdaus, S. H. Wijoyo, and W. Purnomo, "Analisis Sentimen Berbasis Aspek Ulasan Pengguna Aplikasi Alfagift Menggunakan Metode Random Forest dan Pemodelan Topik Latent Dirichlet Allocation," *J. Pengemb. Teknol. Inf. dan Ilmu Komput.*, vol. 9, no. 2, 2025.
- [17] G. A. M. Ashfania, T. Prahasto, A. Widodo, and T. Warsokusumo, "Penggunaan Algoritma Random Forest untuk Klasifikasi berbasis Kinerja Efisiensi Energi pada Sistem Pembangkit Daya," *Rotasi*, vol. 24, no. 3, pp. 14–21, 2023.
- [18] E. P. Febtiawan, L. A. Syamsul, I. Akbar, and A. S. Rachman, "Forecasting Produksi Energi Photovoltaic Menggunakan Algoritma Random Forest Classification," *J. Inf. Syst. Res.*, vol. 5, no. 4, pp. 1053–1062, 2024.
- [19] J. Pranata, S. Agustian, and E. Haerani, "Penggunaan Model Bahasa indoBERT pada Metode Random Forest Untuk Klasifikasi Sentimen Dengan Dataset Terbatas," *vol*, vol. 6, pp. 1668–1676, 2024.
- [20] H. Hidayat, A. Sunyoto, and H. Al Fatta, "Klasifikasi Penyakit Jantung Menggunakan Random Forest Classifier," *J. SISKOM-KB (Sistem Komput. dan Kecerdasan Buatan)*, vol. 7, no. 1, pp. 31–40, 2023.
- [21] E. R. B. Sebayang, Y. H. Chrisnanto, and M. Melina, "Klasifikasi Data Kesehatan Mental di Industri Teknologi Menggunakan Algoritma Random Forest," *IJESPG (International J. Eng. Econ. Soc. Polit. Gov.)*, vol. 1, no. 3, pp. 237–253, 2023.
- [22] T. A. E. Putri, T. Widiari, and R. Santoso, "Penerapan Tuning Hyperparameter Randomsearchcv Pada Adaptive Boosting Untuk Prediksi Kelangsungan Hidup Pasien Gagal Jantung," *J. Gaussian*, vol. 11, no. 3, pp. 397–406, 2023.
- [23] M. A. Abubakar, M. Muliadi, A. Farmadi, R. Herteno, and R. Ramadhani, "Random forest dengan random search terhadap ketidakseimbangan kelas pada prediksi gagal jantung," *J. Inform.*, vol. 10, no. 1, pp. 13–18, 2023.
- [24] I. A. Firdaus, "Deteksi Infeksi Mycoplasma Pneumoniae Pneumonia menggunakan Komparasi Algoritma Klasifikasi Machine Learning," *J. Inf. Technol. Comput. Sci.*, vol. 7, no. 1, pp. 35–42, 2022.
- [25] E. Setiawan, W. W. Pamungkas, and M. Mansuri, "Pengujian Algoritma Deep Learning Pada Analisis Terhadap Layanan Gojek Menggunakan Data Media Sosial Twitter," *J. Penerapan Ilmu-Ilmu Komput.*, vol. 10, no. 2, pp. 53–67, 2024.
- [26] J. A. Wibowo, V. C. Mawardi, and T. Sutrisno, "Penerapan Support Vector Machine untuk Analisis Sentimen Fitur Layanan pada Ulasan Gojek," *J. Ilmu Komput. dan Sist. Inf.*, vol. 12, no. 1, 2024.