

Performance Optimization of Gojek Reviews Sentiment Classification Using the SVM Method with Hyperparameter Tuning

Avrillia Andhini^{1,*}, Nailah Hafsa¹, Shanda Nurhalizah¹, Agus Perdana Widarto², Anjar Wanto²

¹ Information Systems Study Program, Sekolah Tinggi Ilmu Komputer Tunas Bangsa, Pematangsiantar, Indonesia
Jl. Kartini, Proklamasi, Kec. Siantar Barat, Sumatera Utara, 21143, Indonesia

² Master of Informatics Study Program, Sekolah Tinggi Ilmu Komputer Tunas Bangsa, Pematangsiantar, Indonesia
Jl. Kartini, Proklamasi, Kec. Siantar Barat, Sumatera Utara, 21143, Indonesia

Email: ^{1,*} avrilliaandhini494@gmail.com, ² naylahafsa72@gmail.com, ³ shandanurhalizah99@gmail.com,
⁴ agusperdana@amiktunasbangsa.ac.id, ⁵ anjarwanto@gmail.com

Correspondence Author Email: avrilliaandhini494@gmail.com

Submitted: 01/06/2026; Accepted: 12/06/2026; Published: 30/06/2026

Abstract—This study aims to optimize multiclass sentiment classification on Indonesian-language Gojek application reviews using a linear Support Vector Machine (SVM) combined with GridSearchCV-based hyperparameter tuning. The increasing volume of user-generated reviews on digital transportation platforms presents significant challenges for manual sentiment analysis due to the noisy and unstructured characteristics of textual data. The dataset used in this study consists of approximately 5,000 Indonesian-language Google Play Store reviews categorized into negative, neutral, and positive sentiments. The preprocessing stages include text cleaning, case folding, tokenization, stopword removal, and stemming using the Sastrawi library, while feature extraction is performed using TF-IDF unigram representation. The classification model employs a linear SVM optimized through 5-fold cross-validation on the cost parameter (C). Experimental results show that the optimized SVM achieved an accuracy of 63.40%, outperforming the baseline model with an accuracy of 61.29%. The findings indicate that hyperparameter optimization improves model generalization and classification stability for sparse high-dimensional Indonesian text data. Furthermore, this study highlights the challenges of multiclass sentiment classification in informal Indonesian-language reviews, particularly for neutral sentiment categories with semantic overlap. Overall, the proposed approach demonstrates the effectiveness of optimized linear SVM for Indonesian-language sentiment analysis in online transportation application reviews.

Keywords: Sentiment Analysis; Support Vector Machine; Hyperparameter Tuning; GridSearchCV; Gojek Reviews

1. INTRODUCTION

The digital transformation of the Industry 4.0 era has accelerated the development of online transportation in Indonesia through fast, flexible, and efficient smartphone-based services (Elgeldawi et al., 2021; Putri Angraini Aziz et al., 2025). This growth is supported by high internet and mobile device usage, driving the development of platforms like Gojek that offer app-based transportation, food delivery, digital payments, and logistics services (Waliyul Arinni & Manshur Ali Suyanto, 2023; Furqon, 2023). The increasing number of Gojek users has also led to a rise in user reviews on the Google Play Store and App Store (Amri et al., 2024). User reviews serve as a key indicator in reflecting customer satisfaction, service quality, app performance, as well as various technical and payment issues experienced by users (Nugroho & Sudarno, 2025; Rahman et al., 2024). User opinions not only influence a company's image but also shape potential users' decisions when choosing digital services; therefore, companies require an automated opinion analysis system to support service evaluation and strategic decision-making (Moslehpour et al., 2022; Ismail et al., 2025).

The increasing number of user reviews has made manual analysis less effective, as it requires significant time and human resources (Chhetri et al., 2025). The complexity of reviews is further heightened by the use of informal language, abbreviations, typos, mixed languages, and variations in emotional expressions within unstructured text data (Bae et al., 2023). These conditions have driven the adoption of computational approaches such as sentiment analysis in text mining and natural language processing to automatically and accurately classify user opinions (Bharadwaj, n.d.; Kaur & Sharma, 2023). Sentiment analysis is used to categorize opinions into positive, negative, or neutral categories (Firizkiansah et al., 2025). Furthermore, advancements in machine learning through algorithms such as Naive Bayes, Decision Tree, Logistic Regression, Random Forest, KNN, and SVM have improved text classification accuracy compared to manual or rule-based methods (Ulinuha et al., n.d.; Yuniar & Kismiantini, 2023; Lakshmi et al., 2025; Abia & Johnson, 2024; Tufail et al., 2023; Fajri & Primajaya, 2023).

The ability of Support Vector Machines (SVM) to form optimal hyperplanes on high-dimensional data makes them widely used in text classification (Arifin et al., n.d.; Wang & Hu, n.d.). This method is effective on sparse TF-IDF data and produces more accurate and stable sentiment classification results compared to Naive Bayes (Sueno, 2020; Flower, 2023; Syahputra et al., 2022). The combination of TF-IDF and SVM has also been shown to improve the quality of sentiment classification in digital app reviews (Prabowo et al., 2025). However, SVM performance is highly influenced by the selection of hyperparameters such as the kernel, cost (C), and gamma, which affect the model's accuracy, precision, recall, and F1-score (Salman & Al-Jawher, 2024; Fariesya Suhaila Md Sazihan et al., 2025; Dewi et al., 2023; Adrian & Laksana, 2025; Pannakkong et al., 2022). Hyperparameter tuning is necessary to optimize the performance of machine learning models (Guido et al., 2023). One method is Grid Search, which can systematically test parameter combinations and has been shown to improve accuracy, precision, recall, and the stability of SVM models compared to default parameters (Chong & Shah, n.d.; Reddy et al., 2024). In text classification, data

quality is significantly influenced by *preprocessing* steps such as *case folding*, *cleaning*, *tokenization*, *stopword removal*, and *stemming* to reduce *noise* in user reviews (Khanduja & Kaur, 2023; Suresh Kumar et al., 2024; Fatmawati, 2024). In Indonesian text, the use of the *Sastrawi library* has proven effective in improving the consistency of root words and feature quality (Das et al., 2023; Mufti et al., n.d.). Additionally, the TF-IDF method is widely used in feature extraction because it can generate effective numerical representations for *machine learning* algorithms (Ghatora et al., 2024; Semary et al., 2024). TF-IDF is used to calculate the importance of a word in a document based on the word's frequency of occurrence in the document and the entire corpus (Srivastava et al., n.d.). TF-IDF effectively generates numerical representations for machine learning, and its combination with a linear Support Vector Machine demonstrates high performance in high-dimensional text classification (Alemerien et al., 2024; Cahyani & Saraswati, 2023).

There remains a *research gap* in sentiment analysis research because most previous studies have focused solely on comparing classification algorithms without conducting in-depth *hyperparameter* optimization (Tukino & Fifi, n.d.; Khoiruddin et al., n.d.). In fact, parameter configuration has a significant impact on the performance of *machine learning* models (Altalhan et al., 2025). Some studies also still use default parameters for SVM, resulting in suboptimal classification outcomes (Maulana et al., 2023). Additionally, performance evaluation is generally limited to *accuracy* without comprehensively considering *precision*, *recall*, and *F1-score* (Maisat Eka Darmawan & Fauzan Dianta, 2023). Research on reviews of online transportation apps remains limited because the majority of previous studies used datasets from Twitter and marketplaces (AlBadani et al., 2022). In fact, reviews of transportation apps have complex characteristics, such as opinions regarding service, payment, and direct transaction experiences (Putri Angraini Aziz et al., 2025).

Additionally, the prevalence of *typos*, informal language, and mixed languages can increase data *noise*, while some previous studies have not optimally applied text *preprocessing* (Br Sinulingga & Sitorus, 2024; Rihastuti et al., 2025). Another limitation of previous research lies in the lack of focus on optimizing linear Support Vector Machines in the sparse feature space resulting from TF-IDF transformation (Bima Aditya et al., 2025). Some studies merely compared several machine learning algorithms without conducting parameter sensitivity analysis regarding the characteristics of high-dimensional text data (Aria Hidayatullah & Marcellina, n.d.). However, TF-IDF-processed data possesses sparse matrix and high-dimensional feature space characteristics that require specific parameter configurations for the model to operate optimally (Tariq et al., 2025). Furthermore, some previous studies have focused solely on improving accuracy without considering the stability of model performance in Indonesian text classification (Wibowo et al., 2024).

This study focuses on optimizing sentiment classification of Gojek reviews using a linear Support Vector Machine with Grid Search on Indonesian-language text data. The research stages include text preprocessing, TF-IDF feature extraction, model parameter optimization, and evaluation using accuracy, precision, recall, F1-score, and a confusion matrix. The selection of the linear kernel is based on its suitability for the sparse TF-IDF feature space and its lower computational complexity compared to nonlinear kernels (Karyawati et al., 2023). The main novelty of this study lies in the focus on optimizing the linear Support Vector Machine in the domain of Indonesian-language online transportation app reviews using CRISP-DM-based preprocessing and systematic hyperparameter tuning. This study not only focuses on the implementation of classification algorithms but also evaluates the impact of parameter optimization on the stability of model performance in the sparse Indonesian text feature space. Thus, this study is expected to provide an academic contribution to the development of Indonesian-language sentiment analysis while offering practical benefits to digital transportation companies in automatically and more accurately understanding user opinions (Andriati et al., 2025).

2. RESEARCH METHODS

2.1 Dataset

This study uses a dataset of Gojek app reviews obtained from Kaggle in CSV format. The dataset consists of approximately 5,000 Indonesian-language Google Play Store user reviews with two main attributes: review text and rating. The study employs multiclass sentiment classification with three sentiment categories: negative, neutral, and positive. Ratings of 1–2 are categorized as negative sentiment, a rating of 3 as neutral sentiment, and ratings of 4–5 as positive sentiment. The formulation of sentiment labels is as follows:

$$S(r) = \begin{cases} \text{Negative}, & r \leq 2; \\ \text{Neutral}, & r = 3; \\ \text{Positive}, & r \geq 4 \end{cases} \quad (1)$$

$$\text{Negative}, \text{ } r \leq 2; \quad (2)$$

$$\text{Neutral}, \text{ } r = 3; \quad (3)$$

$$\text{Positive}, \text{ } r \geq 4 \quad (4)$$

Next, the categorical sentiment labels were transformed into numerical representations using the *LabelEncoder* method from the Scikit-learn library, where 0 represents Negative, 1 represents Neutral, and 2 represents Positive (Matteucci et al., 2023). This encoding process was essential to enable machine learning algorithms to process

categorical sentiment classes in a structured numerical format. By converting textual labels into numerical values, the dataset became more suitable for computational analysis and classification modeling. Furthermore, this transformation also ensured consistency in data representation during the training and evaluation stages of the model. The overall characteristics and distribution of the research dataset are presented in Table 1.

Table 1. Research Dataset Characteristics

Component	Description
Dataset Name	Gojek Review Dataset
Dataset Source	Kaggle
Data Format	CSV
Data Volume	Approx. 5,000 reviews
Language	Indonesian
Key Features	reviews, ratings
Platform	Google Play Store
Data Type	Unstructured text data
Classification Method	Multiclass Classification
Sentiment Classes	Negative, Neutral, Positive

Table 1 shows the characteristics of the dataset used in this study, which consists of approximately 5,000 Indonesian-language user reviews of the Gojek app obtained from the Google Play Store via the Kaggle platform. The dataset has two main attributes reviews and ratings which are used in the process of multiclass sentiment labeling.

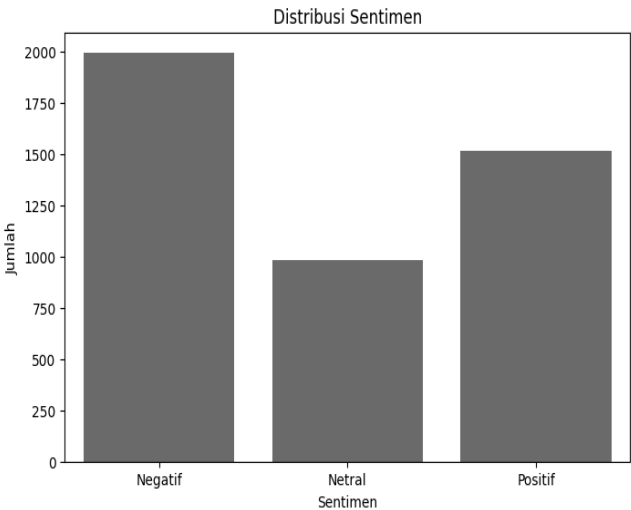


Figure 1. Sentiment Class Distribution

Figure 1 shows the sentiment class distribution in the research dataset. The negative class has the largest number of data points, while the neutral class has the fewest. This indicates an imbalance in the data distribution that could potentially affect model performance, particularly in recognizing neutral sentiment patterns in multiclass classification (Thölke et al., 2022). The sentiment class distribution in the dataset is shown in Table 2.

Table 2. Sentiment Class Distribution

Sentiment Class	Label	Number of Data Points
Negative	0	1995
Neutral	1	986
Positive	2	1519

Based on Table 2, the sentiment class distribution shows that the neutral class has fewer data points compared to the negative and positive classes, resulting in an imbalanced data distribution (Hasanin et al., 2019). This condition has the potential to affect model performance, particularly in recognizing neutral sentiment patterns during the multiclass classification process.

$$Dataset = 0.8 \times Training + 0.2 \times Testing \tag{5}$$

This equation illustrates the division of the dataset into training and testing data with an 80:20 ratio. Eighty percent of the data is used as training data to train the machine learning model, while the remaining 20% is used as testing data to evaluate the model's performance. This division aims to test the model using data it has never encountered before, thereby making the evaluation results more objective and accurate.

2.2 Research Framework

This study consists of several main stages, namely dataset collection, text preprocessing, feature extraction using TF-IDF, classification using Support Vector Machines (SVM), hyperparameter tuning using GridSearchCV, and model performance evaluation. The entire experimental process was conducted using Python on Google Colaboratory. The research framework used in this study is shown in Figure 2.

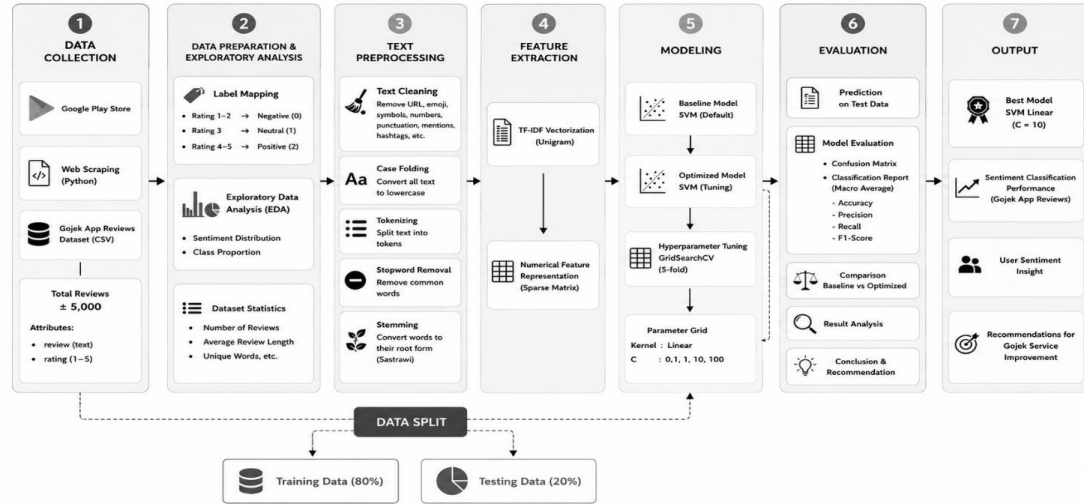


Figure 2. Research Framework

Figure 2 illustrates the overall research workflow, beginning with dataset collection, text preprocessing, TF-IDF feature extraction, sentiment classification using Support Vector Machines (SVM), hyperparameter optimization using GridSearchCV, and model evaluation through a confusion matrix and classification report. The preprocessing stage was conducted to improve the quality and consistency of textual data by reducing noise contained in user reviews. This stage consisted of several sequential processes, namely text cleaning, case folding, tokenization, stopword removal, and stemming using the Sastrawi library. During text cleaning, irrelevant elements such as URLs, emojis, punctuation marks, symbols, hashtags, numbers, and special characters were removed using the Regular Expression (Regex) method to minimize noise and improve textual consistency. Subsequently, case folding converted all text into lowercase to ensure lexical uniformity, while tokenization using the NLTK library segmented sentences into individual word tokens. Furthermore, stopword removal was applied to eliminate non-informative words with minimal contribution to sentiment analysis, followed by stemming to transform inflected words into their root forms and reduce morphological variations within the dataset.

After the preprocessing stage, the textual data were transformed into numerical feature representations using the unigram TF-IDF approach with the parameters $max_features = 5000$, $ngram_range = (1,1)$, $min_df = 2$, and $max_df = 0.95$. This approach was selected because TF-IDF effectively represents word importance within documents while maintaining computational efficiency for sparse and high-dimensional text classification tasks. In addition, the unigram representation was employed to capture dominant lexical patterns in Indonesian-language user reviews while reducing feature complexity and improving model interpretability. The resulting TF-IDF features were subsequently utilized as input for the SVM classification model to identify sentiment patterns within the dataset. Furthermore, the generated feature representations enabled the classification model to distinguish sentiment categories more effectively by emphasizing informative terms with higher relevance scores. The TF calculation used in this study can be expressed as follows:

$$TF(t, d) = \frac{f_{t,d}}{\sum_k f_{k,d}} \quad (6)$$

$$IDF(t) = \log \frac{N}{df(t)} \quad (7)$$

$$TFIDF(t, d) = TF(t, d) \times IDF(t) \quad (8)$$

The resulting TF-IDF features were subsequently used as input for the Support Vector Machine (SVM) classification model. In this study, two SVM approaches were implemented, namely the baseline SVM and the optimized SVM model. SVM aims to identify the optimal hyperplane that maximizes the margin between classes, which can be mathematically expressed as follows:

$$w \cdot x + b = 0 \quad (9)$$

A linear kernel was selected because TF-IDF representations produce sparse and high-dimensional feature spaces that are more effectively handled using linear decision boundaries. Compared to nonlinear kernels, linear SVM

offers lower computational complexity while maintaining stable classification performance in large-scale text mining tasks. The hyperparameter tuning configuration is shown in Table 3.

Table 3. Hyperparameter Tuning Parameters

Parameter	Value
Kernel	Linear
C	0.1, 1, 10, 100
max_features	5000
ngram_range	(1,1)
min_df	2
max_df	0.95
Tuning Method	GridSearchCV
Cross-Validation	5-Fold

Table 3 shows the parameters used in the hyperparameter tuning process using GridSearchCV. This study used a linear kernel with several variations of the cost parameter (C) to obtain the best model configuration based on cross-validation results. GridSearchCV tested all parameter combinations to identify the best model based on cross-validation accuracy. Model performance was evaluated using a confusion matrix and a classification report, with metrics including accuracy, precision, recall, and F1-score. The use of several evaluation metrics aims to provide a more comprehensive analysis of model performance in multiclass sentiment classification.

$$Accuracy = \frac{TP + TN}{TP + TN + FP + FN} \quad (10)$$

$$Precision = \frac{TP}{TP + FP} \quad (11)$$

$$Recall = \frac{TP}{TP + FN} \quad (12)$$

$$F1 = 2 \times \frac{Precision \times Recall}{Precision + Recall} \quad (13)$$

The confusion matrix is used as an evaluation tool to analyze the distribution of prediction results for each sentiment class in a more comprehensive manner. Through this matrix, the model's ability to classify data correctly can be identified by comparing actual labels with predicted labels. In addition, the confusion matrix helps reveal prediction errors occurring in each class, thereby providing deeper insights into the strengths and weaknesses of the multiclass sentiment classification model.

2.3 Support Vector Machine Classification Model

2.3.1 Baseline Support Vector Machine

Support Vector Machine (SVM) works by constructing an optimal hyperplane that separates data classes with the maximum margin in the feature space. In this study, the baseline model uses Scikit-learn's default parameters and serves as a reference before the hyperparameter optimization process is conducted. The representation of the SVM hyperplane can be expressed using the following equation:

$$w \cdot x + b = 0 \quad (14)$$

In this equation, w represents the weight vector, x denotes the feature vector of the analyzed data, and b refers to the bias term used to adjust the position of the hyperplane within the feature space. The primary objective of SVM is to maximize the margin between classes in order to achieve better classification performance. The maximum-margin concept in SVM is represented using the following optimization function:

$$\min \frac{1}{2} \|w\|^2 \quad (15)$$

The objective function of Support Vector Machine (SVM) is formulated to maximize the margin between classes while minimizing the complexity of the classification model. The optimization problem is expressed as follows:

$$\min \frac{1}{2} \|w\|^2 \quad (16)$$

2.3.2 Hyperparameter Tuning Using GridSearchCV

To improve the classification performance, this study employed hyperparameter optimization using the GridSearchCV method. This approach systematically evaluates multiple parameter combinations through a cross-validation mechanism to identify the most effective model configuration. A 5-fold cross-validation strategy was applied during the tuning process to ensure a more reliable and robust model evaluation. Furthermore, the hyperparameter optimization process was conducted not only to enhance classification accuracy but also to minimize the risk of overfitting and improve the model's generalization capability when handling unseen data.

By determining the optimal parameter settings, the resulting model is expected to achieve more stable and consistent performance across different data distributions. The optimization process focused on tuning the cost parameter (C) in the linear Support Vector Machine (SVM) model. Various parameter values were systematically

evaluated using the GridSearchCV method to determine the configuration that provides the best classification results. The evaluation considered both predictive accuracy and the model’s ability to generalize effectively to test data. The detailed parameter tuning configuration for the optimized SVM model is presented in Table 4.

Table 4. Hyperparameter Tuning

Parameter	Value
Kernel	Linear
C	0.1, 1, 10, 100
Tuning Method	GridSearchCV
Cross-Validation	5-Fold

Table 4 shows the parameter combinations used in the Support Vector Machine optimization process using GridSearchCV. The tuning process was conducted to obtain optimal parameters capable of improving sentiment classification performance on Indonesian-language text data. GridSearchCV works by testing all parameter combinations to obtain the best model configuration based on cross-validation accuracy (Guido et al., 2023). To provide a visual illustration of the effect of hyperparameter tuning on sentiment data separation, this study presents a feature space visualization using Principal Component Analysis (PCA). The visualization was performed on the baseline SVM and optimized SVM after the dimension reduction process on the TF-IDF representation.

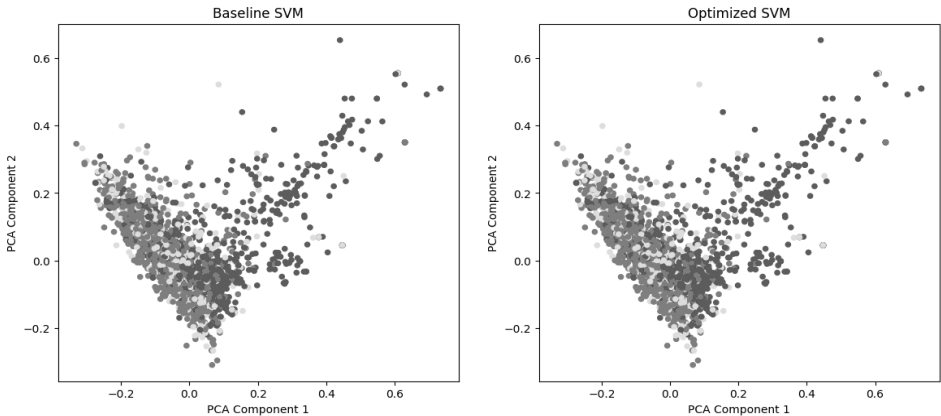


Figure 3. Visualization of Baseline and Optimized SVM Using PCA

Figure 3 shows the visualization results of sentiment data distribution using Principal Component Analysis (PCA) for the baseline SVM and optimized SVM. This visualization aims to illustrate the model’s ability to separate the distributions across sentiment classes after the classification process is performed. In the baseline SVM, the data distribution across classes still shows a fairly high level of overlap, so the separation between sentiment categories is not yet optimal. After applying hyperparameter tuning using GridSearchCV, the optimized SVM demonstrates a more distinct and structured class distribution. These results indicate that optimizing the cost parameter (C) contributes to enhancing the model’s ability to form a more optimal hyperplane, particularly when dealing with sparse, high-dimensional TF-IDF data representations.

3. RESULTS AND DISCUSSION

3.1 Data Preprocessing Results

The preprocessing stage was conducted to improve the quality of text data prior to classification using the Support Vector Machine (SVM). Gojek user reviews collected from the Google Play Store generally contain textual noise such as emojis, symbols, typographical errors, URLs, numbers, and informal language that may affect feature representation quality. Therefore, preprocessing was applied to generate cleaner and more structured text data through several stages, including text cleaning, case folding, tokenization, stopwords removal, and stemming using the Sastrawi library. Text cleaning was performed to remove irrelevant characters such as emojis, punctuation marks, symbols, and URLs, while case folding converted all text into lowercase format to maintain lexical consistency. Tokenization segmented sentences into individual word tokens, stopwords removal eliminated non-informative terms, and stemming transformed inflected words into their root forms to improve feature consistency during the classification process. An example of the preprocessing results is shown in Table 5.

Table 5. Example of Data Preprocessing Results

Original Review	Preprocessing Results
“My GoPay account is blocked!!!🤡” “Bunch of trash drivers!!!”	“My GoPay account is blocked” “Bunch of trash drivers”

Original Review	Preprocessing Results
“Just downloaded Gojek and got a new phone, then topped up”	“Just downloaded Gojek on my new phone and topped up”

Based on Table 5, the preprocessing steps successfully reduced textual noise such as emojis, punctuation, and irrelevant words while improving word consistency through text cleaning and stemming. After preprocessing, feature extraction was performed using TF-IDF unigram with parameters `max_features = 5000`, `min_df = 2`, and `max_df = 0.95`. The TF-IDF transformation generated a high-dimensional sparse text representation used as the primary input for the SVM classification model, allowing each document to be represented as a numerical vector based on the importance of words within the dataset.

3.2 SVM Baseline Classification Results

The initial classification stage was performed using a baseline Support Vector Machine with default parameters from the *Scikit-learn* library. The baseline model was used as an initial reference before hyperparameter optimization. The baseline SVM achieved an accuracy of 61.29%, indicating that linear SVM is capable of performing multiclass sentiment classification on sparse Indonesian-language review data with relatively stable performance. This result confirms the suitability of linear hyperplanes for high-dimensional TF-IDF feature representations commonly found in text classification problems. The results of the baseline SVM confusion matrix are shown in Figure 4.

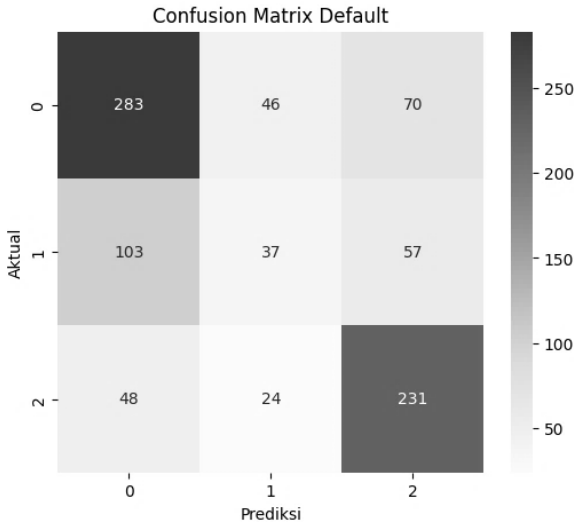


Figure 4. Baseline SVM Confusion Matrix

Figure 4 shows that the baseline SVM is able to classify negative and positive classes quite well, but still has difficulty recognizing the neutral class. Most neutral reviews were misclassified as either positive or negative sentiments due to semantic overlap between sentiment categories. Neutral opinions often contain implicit sentiment cues, mixed contextual expressions, or ambiguous linguistic patterns that are difficult to capture using unigram TF-IDF representations. The baseline SVM classification report results are shown in Table 6.

Table 6. Baseline SVM Classification Report

Class	Precision	Recall	F1-Score
Negative	65%	71%	68%
Neutral	35%	19%	24%
Positive	65%	76%	70%
Accuracy	61.29%		

Table 6 shows that the baseline SVM is able to classify negative and positive sentiments quite well, but still has difficulty recognizing the neutral class. The low recall in the neutral class indicates semantic overlap between sentiment classes as well as an imbalance in data distribution. Additionally, the use of unigram TF-IDF has not been fully capable of capturing the semantic context in neutral opinions, so some data have feature representations similar to both positive and negative sentiments.

3.3 Results of SVM Hyperparameter Tuning

To improve the classification model’s performance, this study applied hyperparameter tuning using the GridSearchCV method with 5-fold cross-validation. The tuning process was performed on the cost (C) parameter of the linear Support Vector Machine. Based on the tuning results, the optimal parameters obtained are:

- a. Kernel = Linear
- b. C = 10

The increase in the cost parameter (C) improves the model’s ability to construct a more discriminative hyperplane by penalizing classification errors more strictly during the optimization process. As a result, the optimized SVM achieves better separation between sentiment classes within the sparse TF-IDF feature space. The confusion matrix results for the optimized SVM are shown in Figure 5.

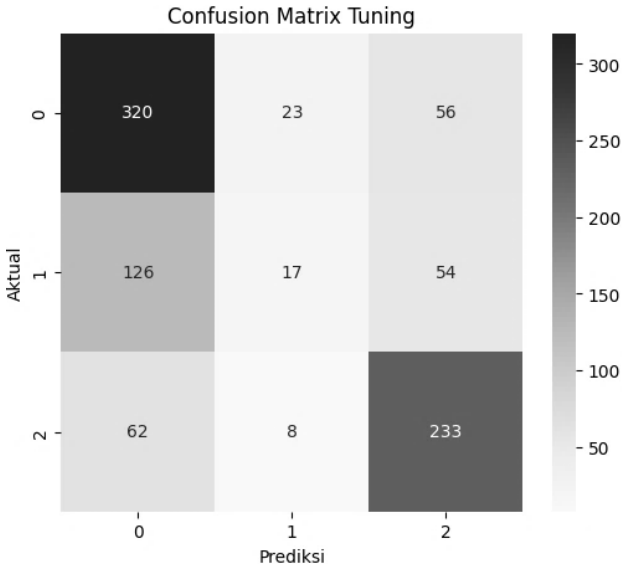


Figure 5. Confusion Matrix of the Optimized SVM

Figure 5 shows the confusion matrix of the optimized SVM following the hyperparameter tuning process using GridSearchCV. The classification results indicate that the model is able to classify negative and positive classes quite well, as evidenced by the dominance of values on the main diagonal of the confusion matrix. This suggests that optimizing the cost parameter (C) helps the model form a more optimal hyperplane to separate sentiment classes. In addition, the tuning process improves the model’s ability to reduce misclassification between sentiment categories, particularly in distinguishing strongly polarized opinions. The optimized model also demonstrates more stable prediction performance across the testing dataset compared to the baseline model.

However, the neutral class achieved lower performance due to semantic ambiguity and overlap between sentiment categories. Neutral reviews often contain mixed opinions or contextual expressions that are difficult to classify consistently. In addition, the imbalance of neutral samples may limit the classifier’s ability to learn representative patterns. Nevertheless, the optimized SVM outperformed the baseline model, indicating that hyperparameter tuning improved the model’s generalization capability. The classification report for the optimized SVM is presented in Table 7.

Table 7. Classification Report of the Optimized SVM

Class	Precision	Recall	F1-Score
Negative	63%	80%	71%
Neutral	35%	9%	14%
Positive	68%	77%	72%
Accuracy	63.40%		

Based on Table 7, the Optimized SVM yields improved performance compared to the baseline model, particularly in terms of negative class recall and positive class F1-score. The evaluation results indicate that hyperparameter tuning enhances the model’s ability to classify negative and positive sentiments, with negative class recall reaching 80% and positive class F1-score at 72%. This performance improvement indicates that optimizing the cost parameter (C) helps the model form a more optimal hyperplane in the high-dimensional vector space, thereby reducing classification errors between dominant classes. As a result, the model can distinguish negative and positive sentiment patterns more effectively and produce more consistent classification results across the testing dataset.

However, performance on the neutral class remains relatively low, with a recall of 9%. The confusion matrix results show that most neutral data is still predicted as either negative or positive sentiment due to semantic ambiguity and data distribution imbalance. Additionally, the use of informal language, typos, abbreviations, and mixed languages in Gojek user reviews increases the complexity of text classification, so that the unigram TF-IDF representation is not yet fully capable of optimally capturing the semantic context of neutral opinions. Overall, the findings demonstrate that the combination of TF-IDF and optimized linear SVM is effective for multiclass sentiment classification on Indonesian-language app reviews. The application of GridSearchCV-based hyperparameter tuning contributes positively to improving model generalization and classification stability, especially for dominant sentiment classes in high-dimensional text data.

3.4 Comparison of Baseline and Optimized SVM

A comparison of baseline and optimized SVM performance was conducted to evaluate the effect of hyperparameter tuning on improving the sentiment classification performance of Gojek reviews. The evaluation was performed by comparing the accuracy values between the baseline SVM and the optimized SVM after the tuning process using GridSearchCV. The comparison of the accuracy of the two models is shown in Figure 6.

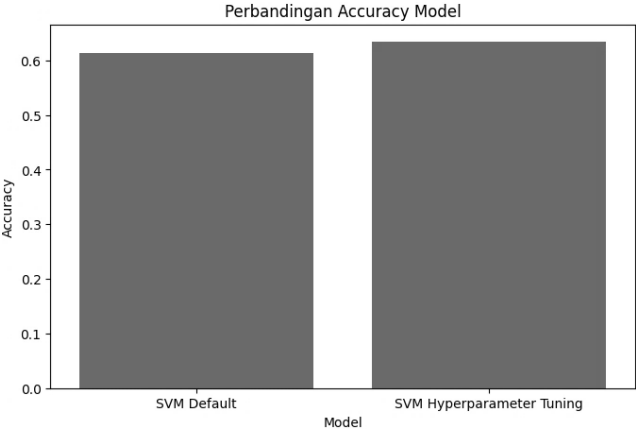


Figure 6. Model Accuracy Comparison

Based on Figure 6, the optimized SVM achieved a higher accuracy than the baseline SVM. These results indicate that hyperparameter tuning using GridSearchCV can improve the performance of Gojek review sentiment classification. Optimizing the cost parameter (C) helps the model form a more optimal hyperplane, thereby improving the separation between sentiment classes. Thus, the tuning process contributes positively to the model’s generalization ability and stability on Indonesian-language text data. The comparison of model accuracy results is shown in Table 8.

Table 8. Comparison of Baseline and Optimized SVM Accuracy

Model	Accuracy
Baseline SVM	61.29%
Optimized SVM	63.40%

Table 8 shows that the Optimized SVM achieved an accuracy of 63.40%, outperforming the baseline SVM’s 61.29% with an improvement of 2.11%. These findings indicate that hyperparameter tuning using GridSearchCV enhances the model’s generalization capability in multiclass sentiment classification. In addition to improving accuracy, the optimized SVM also increased precision and F1-score for the negative and positive classes. Parameter optimization enables the model to construct a more discriminative hyperplane, thereby reducing classification errors in majority sentiment classes. Although the performance improvement is relatively moderate, the results confirm that appropriate parameter selection positively contributes to classification stability and predictive performance in sparse Indonesian TF-IDF-based sentiment classification.

3.5 Discussion

The experimental results demonstrate that the integration of text preprocessing, TF-IDF feature extraction, and linear Support Vector Machine classification provides an effective approach for multiclass sentiment analysis on Indonesian-language Gojek reviews. The preprocessing stages significantly reduced textual noise and improved feature consistency, particularly for user-generated reviews containing informal language, abbreviations, typographical errors, and mixed-language expressions. Furthermore, the TF-IDF representation successfully transformed unstructured textual data into sparse numerical vectors suitable for high-dimensional text classification tasks. The findings also confirm that linear SVM is highly suitable for sparse TF-IDF feature spaces because linear hyperplanes can efficiently separate sentiment classes while maintaining relatively low computational complexity. In addition, hyperparameter optimization using GridSearchCV positively improved model performance by enabling the classifier to construct a more discriminative hyperplane and achieve better generalization capability for unseen testing data.

Despite the improvement in classification performance, the neutral sentiment category remained challenging due to semantic ambiguity and contextual overlap between sentiment classes. Neutral reviews frequently contain mixed opinions and implicit emotional expressions that are difficult to capture using unigram TF-IDF representations. In addition, the imbalanced class distribution reduced the classifier’s ability to learn representative patterns for minority sentiment categories. Another challenge identified in this study is the linguistic complexity of Indonesian-language user-generated reviews, including slang expressions, inconsistent spelling, and code-mixing between languages, which limits the effectiveness of lexical-based feature extraction methods in capturing deeper semantic meaning. Nevertheless, the findings confirm that optimized linear SVM remains an effective and computationally efficient approach for multiclass sentiment classification on sparse Indonesian-language review datasets. Future

research is encouraged to explore contextual embedding approaches such as Word2Vec, FastText, IndoBERT, or transformer-based architectures combined with data balancing techniques to improve semantic understanding and minority-class classification performance.

4. CONCLUSION

This study demonstrates that the integration of text preprocessing, TF-IDF feature extraction, and a linear Support Vector Machine is effective for performing multiclass sentiment classification on Indonesian-language Gojek user reviews characterized by noisy and unstructured textual patterns. The implementation of GridSearchCV-based hyperparameter optimization successfully improved the classification performance, increasing the model accuracy from 61.29% in the baseline model to 63.40% in the optimized model. The experimental findings confirm that appropriate hyperparameter configuration plays an important role in improving classification stability, sentiment class separability, and model generalization capability within sparse high-dimensional TF-IDF feature spaces. The optimized linear SVM demonstrated more stable predictive performance, particularly for dominant sentiment classes, indicating that parameter optimization significantly contributes to enhancing machine learning performance in Indonesian-language sentiment analysis tasks. Despite the improvement in overall classification performance, the neutral sentiment category remains challenging due to semantic ambiguity, contextual overlap between sentiment classes, informal linguistic expressions, and imbalanced data distribution. These findings indicate that unigram TF-IDF representations still have limitations in capturing deeper semantic and contextual relationships within Indonesian-language user-generated reviews. This study contributes to the advancement of Indonesian-language sentiment analysis research, particularly in optimizing machine learning approaches for multiclass text classification in online transportation application reviews. In practical applications, the proposed approach can support digital transportation companies in understanding user opinions more efficiently and accurately for service evaluation and strategic decision-making processes. Future research is encouraged to explore contextual embedding methods such as Word2Vec, FastText, IndoBERT, or transformer-based architectures combined with data balancing techniques to improve semantic understanding and minority-class classification performance in Indonesian-language sentiment analysis.

REFERENCES

- Abia, V. M., & Johnson, E. H. (2024). Sentiment Analysis Techniques: A Comparative Study of Logistic Regression, Random Forest, and Naive Bayes on General English and Nigerian Texts. *Journal of Engineering Research and Reports*, 26(9), 123–135. <https://doi.org/10.9734/jerr/2024/v26i91268>
- Adrian, M. I., & Laksana, E. A. (2025). Optimalisasi Parameter Support Vector Machine dengan Algoritma PSO untuk Tugas Klasifikasi Sentimen Ulasan IMDb. *Jurnal Algoritma*, 22(1), 288–299. <https://doi.org/10.33364/algoritma/v.22-1.2306>
- AlBadani, B., Shi, R., & Dong, J. (2022). A Novel Machine Learning Approach for Sentiment Analysis on Twitter Incorporating the Universal Language Model Fine-Tuning and SVM. *Applied System Innovation*, 5(1). <https://doi.org/10.3390/asi5010013>
- Alemerien, K., Al-Ghareeb, A., & Alksasbeh, M. Z. (2024). Sentiment Analysis of Online Reviews: A Machine Learning Based Approach with TF-IDF Vectorization. *Journal of Mobile Multimedia*, 20(5), 1089–1116. <https://doi.org/10.13052/jmm1550-4646.2055>
- Altalhan, M., Algarni, A., & Turki-Hadj Alouane, M. (2025). Imbalanced Data Problem in Machine Learning: A Review. *IEEE Access*, 13, 13686–13699. <https://doi.org/10.1109/ACCESS.2025.3531662>
- Amri, P., Suri, D. M., & Syuhada. (2024). The analysis of ride hailing user characteristics from app reviews. *Jurnal Siasat Bisnis*, 241–262. <https://doi.org/10.20885/jsb.vol28.iss2.art7>
- Andriati, D. A., Laple, R., Putra, S., & Saputra, G. D. (2025). Analisis Sentimen Komentar Pengguna Aplikasi Klik Indomaret Di Google Playstore Menggunakan Metode Support Vector Machine (SVM)-Dea Andini Andriati et.al Analisis Sentimen Komentar Pengguna Aplikasi Klik Indomaret Di Google Playstore Menggunakan Metode Support Vector Machine (SVM). 07. <https://doi.org/10.54209/jatilima.v7i03.1653>
- Aria Hidayatullah, K. M., & Marcellina, D. (n.d.). Analisis Sentimen Pelanggan Pempek Cek Ida Terhadap Penilaian Produk Menggunakan Layanan GoFood Pada Aplikasi Gojek Dengan Klasifikasi Support Vector Machine (SVM).
- Arifin, N., Enri, U., & Sulistiyowati, N. (n.d.). STRING (Satuan Tulisan Riset dan Inovasi Teknologi) Penerapan Algoritma Support Vector Machine (Svm) Dengan Tf-Idf N-Gram Untuk Text Classification.
- Aufar, A. F., Mochamad Alfian Rosid, Eviyanti, A., & Astutik, I. R. I. (2023). Optimizing Text Preprocessing for Accurate Sentiment Analysis on E-Wallet Reviews. *JICTE (Journal of Information and Computer Technology Education)*, 7(2), 42–50. <https://doi.org/10.21070/jicte.v7i2.1650>
- Bae, J., Yu Hung, C., & van Lent, L. (2023). Mobilizing Text As Data. In *European Accounting Review* (Vol. 32, Number 5, pp. 1085–1106). Routledge. <https://doi.org/10.1080/09638180.2023.2218423>
- Bharadwaj, L. (n.d.). Sentiment Analysis in Online Product Reviews: Mining Customer Opinions for Sentiment Classification. In *IJFMR23056090* (Vol. 5, Number 5). Retrieved www.ijfmr.com
- Bima Aditya, A., Samsudin, S., Pandu Rizki, W., Mahendra, M., & Setiawan, A. (2025). Jurnal Informatika Terpadu Analisis Sentimen Ulasan Aplikasi Gojek Menggunakan Support Vector Machine Dan Random Forest. *Jurnal Informatika Terpadu*, 11(2), 134–143. <https://journal.nurulfikri.ac.id/index.php/JIT>

- Br Sinulingga, J. E., & Sitorus, H. C. K. (2024). Analisis Sentimen Opini Masyarakat terhadap Film Horor Indonesia Menggunakan Metode SVM dan TF-IDF. *Jurnal Manajemen Informatika (JAMIKA)*, 14(1), 42–53. <https://doi.org/10.34010/jamika.v14i1.11946>
- Cahyani, S. N., & Saraswati, G. W. (2023). Implementation Of Support Vector Machine Method In Classifying School Library Books With Combination Of Tf-Idf And Word2vec. *Jurnal Teknik Informatika (Jutif)*, 4(6), 1555–1566. <https://doi.org/10.52436/1.jutif.2023.4.6.1536>
- Chhetri, V., Upadhyay, K., Siddique, A. B., & Farooq, U. (2025). *What Users Value and Critique: Large-Scale Analysis of User Feedback on AI-Powered Mobile Apps*. <http://arxiv.org/abs/2506.10785>
- Chong, K., & Shah, N. (n.d.). Comparison of Naive Bayes and SVM Classification in Grid-Search Hyperparameter Tuned and Non-Hyperparameter Tuned Healthcare Stock Market Sentiment Analysis. In *IJACSA International Journal of Advanced Computer Science and Applications* (Vol. 13, Number 12). Retrieved www.ijacsa.thesai.org
- Das, R. K., Islam, M., Hasan, M. M., Razia, S., Hassan, M., & Khushbu, S. A. (2023). Sentiment analysis in multilingual context: Comparative analysis of machine learning and hybrid deep learning models. *Heliyon*, 9(9). <https://doi.org/10.1016/j.heliyon.2023.e20281>
- Dewi, C., Indriawan, F. A., & Christanto, H. J. (2023). Spam classification problems using support vector machine and grid search. *International Journal of Applied Science and Engineering*, 20(4). [https://doi.org/10.6703/IJASE.202312_20\(4\).006](https://doi.org/10.6703/IJASE.202312_20(4).006)
- Elgeldawi, E., Sayed, A., Galal, A. R., & Zaki, A. M. (2021). Hyperparameter tuning for machine learning algorithms used for arabic sentiment analysis. *Informatics*, 8(4). <https://doi.org/10.3390/informatics8040079>
- Fajri, M., & Primajaya, A. (2023). Komparasi Teknik Hyperparameter Optimization pada SVM untuk Permasalahan Klasifikasi dengan Menggunakan Grid Search dan Random Search. In *Journal of Applied Informatics and Computing (JAIC)* (Vol. 7, Number 1). <http://jurnal.polibatam.ac.id/index.php/JAIC>
- Fariesya Suhaila Md Sazihan, N., Mu, N., Abdul Latiff, azzah, Noordini Nik Abd Malik, N., Mashohor, S., & Taha Al-Dhief, F. (2025). *Evaluating the Effectiveness of Parameter Tuning for Support Vector Machine on Voice Pathology Database*. 24(2), 118–124.
- Fatmawati, D. (2024). Sentiment Classification of IT Service Feedback via TF-IDF. *COGITO Smart Journal*, 10(2).
- Firizkiansah, A., Muhammad, A., & Maulana, I. R. (2025). *Optimasi Klasifikasi Data Teks Menggunakan Algoritma Logistic Regression dengan TF-IDF dan SMOTE* (Vol. 2, Number 1).
- Flower, Dr. K. L. (2023). Text Classification from positive and unlabeled examples using Support Vector Machine (SVM). *Interantional Journal Of Scientific Research In Engineering And Management*, 07(11), 1–11. <https://doi.org/10.55041/IJSREM27391>
- Furqon, D. A. (2023). *Science and Technology The Journey Of Gojek: Becoming Indonesia's First Decacorn Startup* (Vol. 7, Number 2). <http://jurnal.uts.ac.id>
- Ghatora, P. S., Hosseini, S. E., Pervez, S., Iqbal, M. J., & Shaukat, N. (2024). Sentiment Analysis of Product Reviews Using Machine Learning and Pre-Trained LLM. *Big Data and Cognitive Computing*, 8(12). <https://doi.org/10.3390/bdcc8120199>
- Guido, R., Groccia, M. C., & Conforti, D. (2023). A hyper-parameter tuning approach for cost-sensitive support vector machine classifiers. *Soft Computing*, 27(18), 12863–12881. <https://doi.org/10.1007/s00500-022-06768-8>
- Hasanin, T., Khoshgoftaar, T. M., Leevy, J. L., & Bauder, R. A. (2019). Severely imbalanced Big Data challenges: investigating data sampling approaches. *Journal of Big Data*, 6(1). <https://doi.org/10.1186/s40537-019-0274-4>
- Ismail, H., Teknologi, I., Bisnis, D., & Malang, A. (2025). Analysis of The Influence of Website Application Quality, Service Quality and Trust on Consumer Decisions to Transact on Gojek Applications in the Pier Industrial Area of Pasuruan District. *Eduvest-Journal of Universal Studies*, 5. <http://eduvest.greenvest.co.id>
- Karyawati, A. E., Wijaya, K. D. Y., Supriana, I. W., & Supriana, I. W. (2023). A Comparison Of Different Kernel Functions Of Svm Classification Method For Spam Detection. *JITK (Jurnal Ilmu Pengetahuan Dan Teknologi Komputer)*, 8(2), 91–97. <https://doi.org/10.33480/jitk.v8i2.2463>
- Kaur, G., & Sharma, A. (2023). A deep learning-based model using hybrid feature extraction approach for consumer sentiment analysis. *Journal of Big Data*, 10(1). <https://doi.org/10.1186/s40537-022-00680-6>
- Khanduja, D. K., & Kaur, S. (2023). The Categorization of Documents Using Support Vector Machines. *International Journal of Scientific Research in Computer Science and Engineering*, 11(6), 1–12. <https://doi.org/10.26438/ijscse.v11i6.112>
- Khoiruddin, Y., Fauzi, A., & Siregar, A. M. (n.d.). *Analisis Sentimen Gojek Indonesia Pada Twitter Menggunakan Algoritme Naïve Bayes Dan Support Vector Machine*.
- Lakshmi, L., Ali, A. B. M., Devi, K. D. S., Rafiq, M., Shernazarov, I., Othman, N. A., & Khan, M. I. (2025). Opinion mining in e-commerce: Evaluating machine learning approaches for sentiment analysis. *Results in Control and Optimization*, 19. <https://doi.org/10.1016/j.rico.2025.100575>
- Maisat Eka Darmawan, Z., & Fauzan Dianta, A. (2023). Implementasi Optimasi Hyperparameter GridSearchCV Pada Sistem Prediksi Serangan Jantung Menggunakan SVM. *Teknologi*, 13(1), 8–15. <https://doi.org/10.26594/teknologi.v13i1.3098>
- Matteucci, F., Arzamasov, V., & Boehm, K. (2023). *A benchmark of categorical encoders for binary classification*. <http://arxiv.org/abs/2307.09191>
- Maulana, A., Inayah Khasnaputri Afifah, Asghafi Mubarrak, Kiagus Rachmat Fauzan, Ardhan Dwintara, & Zen, B. P. (2023). Comparison Of Logistic Regression, Multinomialnb, Svm, And K-Nn Methods On Sentiment Analysis Of Gojek App Reviews On The Google Play Store. *Jurnal Teknik Informatika (Jutif)*, 4(6), 1487–1494. <https://doi.org/10.52436/1.jutif.2023.4.6.863>
- Moslehpour, M., Ismail, T., Purba, B., & Wong, W. K. (2022). What makes go-jek go in indonesia? The influences of social media marketing activities on purchase intention. *Journal of Theoretical and Applied Electronic Commerce Research*, 17(1), 89–103. <https://doi.org/10.3390/jtaer17010005>
- Mufti, N., Aprianoputri, A., & Gustriansyah, R. (n.d.). *Klasifikasi Sentimen Ulasan Aplikasi Gojek Berbasis Decision Tree Dengan Optimasi Grid Search*. <https://doi.org/10.12928/jstic.v14i1.31092>
- Nugroho, Y. A., & Sudarno, S. (2025). Analisis Sentimen Aplikasi Gojek Pada Ulasan Pengguna di Google Play Store Menggunakan Metode Support Vector Machine. *Journal of Information System Research (JOSH)*, 6(4), 2171–2179. <https://doi.org/10.47065/josh.v6i4.7891>

- Pannakkong, W., Thiwa-Anont, K., Singthong, K., Parthanadee, P., & Buddhakulsomsiri, J. (2022). Hyperparameter Tuning of Machine Learning Algorithms Using Response Surface Methodology: A Case Study of ANN, SVM, and DBN. *Mathematical Problems in Engineering*, 2022. <https://doi.org/10.1155/2022/8513719>
- Prabowo, B. A., Hindasyah, A., & Khalid Rivai, A. (2025). Sentiment Analysis Of Pln Mobile Application Services Using Naive Bayes, Support Vector Machine (Svm) And Decision Tree Methods. *Jurnal Riset Informatika*, 7(3), 236–243. <https://doi.org/10.34288/jri.v7i3.378>
- Putri Angraini Aziz, Nur Ilahi, S. B., Sumiarni Moka, & Sajiah, A. M. (2025). Penerapan Hadoop untuk Analisis Sentimen Berbasis Big Data pada Ulasan Aplikasi Transportasi Online. *SATESI: Jurnal Sains Teknologi Dan Sistem Informasi*, 5(1), 51–60. <https://doi.org/10.54259/satesi.v5i1.4051>
- Rahman, R. A., Pranatawijaya, V. H., & Sari, N. N. K. (2024). *Analisis Sentimen Berbasis Aspek pada Ulasan Aplikasi Gojek* (Vol. 4, Number 1).
- Reddy, K. B., Aditya, K. J., & Arivazhagan, D. (n.d.). Margin Maximization of Text Classification based on Support Vector Machine. In *International Journal for Research in Applied Science & Engineering Technology (IJRASET)* (Vol. 11). Retrieved www.ijraset.com
- Rihastuti, S., Rosyidi, A., Singopuran, N., Kartasura, K., Sukoharjo, K., & Tengah, J. (2025). Perbandingan Kinerja Support Vector Machine Dan Random Forest Untuk Klasifikasi Sentimen Pengguna Aplikasi Gojek Dengan Optimasi SMOTE. *Jurnal Algoritme*, 6(1), 131–141. <https://doi.org/10.35957/algoritme.v5i3.13463>
- Salman, A. H., & Al-Jawher, W. A. M. (2024). Performance Comparison of Support Vector Machines, AdaBoost, and Random Forest for Sentiment Text Analysis and Classification. *Journal Port Science Research*, 7(3), 300–311. <https://doi.org/10.36371/port.2024.3.8>
- Semary, N. A., Ahmed, W., Amin, K., Plawiak, P., & Hammad, M. (2024). Enhancing machine learning-based sentiment analysis through feature extraction techniques. *PLoS ONE*, 19(2 February). <https://doi.org/10.1371/journal.pone.0294968>
- Srivastava, R., Bharti, P. K., & Verma, P. (n.d.). Comparative Analysis of Lexicon and Machine Learning Approach for Sentiment Analysis. In *IJACSA International Journal of Advanced Computer Science and Applications* (Vol. 13, Number 3). Retrieved www.ijacsa.thesai.org
- Sueno, H. T. (2020). Multi-class Document Classification using Support Vector Machine (SVM) Based on Improved Naïve Bayes Vectorization Technique. *International Journal of Advanced Trends in Computer Science and Engineering*, 9(3), 3937–3944. <https://doi.org/10.30534/ijatcse/2020/216932020>
- Suresh Kumar, K., Radha Mani, A. S., Ananth Kumar, T., Jalili, A., Gheisari, M., Malik, Y., Chen, H. C., & Jahangir Moshayedi, A. (2024). Sentiment Analysis of Short Texts Using SVMs and VSMs-Based Multiclass Semantic Classification. *Applied Artificial Intelligence*, 38(1). <https://doi.org/10.1080/08839514.2024.2321555>
- Syahputra, R., Yanris, G. J., & Irmayani, D. (2022). SVM and Naïve Bayes Algorithm Comparison for User Sentiment Analysis on Twitter. *Sinkron*, 7(2), 671–678. <https://doi.org/10.33395/sinkron.v7i2.11430>
- Tariq, H., Majeed, M., & Ahmad, M. (2025). Optimizing SVM Performance through Combinatorial Hyperparameter Tuning and Model Selection. *International Journal Bioautomation*, 29(2), 117–144. <https://doi.org/10.7546/ijba.2025.29.2.000981>
- Thölke, P., Mantilla-Ramos, Y.-J., Abdelhedi, H., Maschke, C., Dehgan, A., Harel, Y., Kemtur, A., Berrada, L. M., Sahraoui, M., Young, T., Bellemare Pépin, A., El Khantour, C., Landry, M., Pascarella, A., Hadid, V., Combrisson, E., O’Byrne, J., & Jerbi, K. (2022). *Class imbalance should not throw you off balance: Choosing the right classifiers and performance metrics for brain decoding with imbalanced data*. <https://doi.org/10.1101/2022.07.18.500262>
- Tufail, S., Riggs, H., Tariq, M., & Sarwat, A. I. (2023). Advancements and Challenges in Machine Learning: A Comprehensive Review of Models, Libraries, Applications, and Algorithms. In *Electronics (Switzerland)* (Vol. 12, Number 8). MDPI. <https://doi.org/10.3390/electronics12081789>
- Tukino, & Fifi. (n.d.). *Penerapan Support Vector Machine Untuk Analisis Sentimen Pada Layanan Ojek Online*. Retrieved <http://journal.aptikomkepri.org/index.php/JDDAT>
- Ulinuha, A., Majid, E., & Nuari, R. (n.d.). *Performance Comparison Of Bert Metrics And Classical Machine Learning Models (Svm, Naive Bayes) For Sentiment Analysis Perbandingan Kinerja Metrik BERT Dan Model Machine Learning Klasik (SVM, Naive Bayes) Untuk Analisis Sentimen*. 10(2), 2025.
- Waliyul Arinni, R., & Manshur Ali Suyanto, A. (2023). Positioning Analysis of Online Transportation Companies in Indonesia Based on Marketing Mix Aspects. *International Journal of Scientific Research and Management (IJSRM)*, 11(11), 5505–5515. <https://doi.org/10.18535/ijrm/v11i11.em13>
- Wang, H., & Hu, D. (n.d.). *Comparison of SVM and LS-SVM for Regression*.
- Wibowo, J. A., Mawardi, V. C., & Sutrisno, T. (2024). Visualisasi Word Cloud Hasil Analisis Sentimen Berbasis Fitur Layanan Aplikasi Gojek Dengan Support Vector Machine. *Jurnal Serina Sains, Teknik Dan Kedokteran*, 2(1), 61–70. <https://doi.org/10.24912/jsstk.v2i1.32058>
- Yuniar, P., & Kismiantini. (2023). Analisis Sentimen Ulasan pada Gojek Menggunakan Metode Naive Bayes. *Statistika*, 23(2), 164–175. <https://doi.org/10.29313/statistika.v23i2.2353>