



Comparative Study of Naive Bayes and SVM for E-Commerce Sentiment Classification on Shopee

Yoga Fradana*, Cahyo Adi Nugraha, Frans Nicko Apriansyah, Tri Mutiara Illahi, Ken Ditha Tania, Allsela Meiriza

Information Systems, Faculty of Computer Science, Universitas Sriwijaya

Email: ^{1,*}yogafradana69@email.com, ²cahyoadinugraha29@gmail.com, ^{3,*}fransnicko22@gmail.com, ⁴trimutiara3@gmail.com, ⁵kenyatania@gmail.com, ⁶allsela_meiriza@yahoo.co.id

(* : fransnicko22@gmail.com)

Submitted: 28/04/2026; Accepted: 12/06/2026; Published: 29/07/2026

Abstract— This study compares the performance of Naive Bayes and Support Vector Machine (SVM) for sentiment classification of men's shirt product reviews on Shopee. A dataset of 500 reviews was collected via web scraping and processed through case folding, tokenizing, stopword removal, and stemming, followed by TF-IDF feature extraction. The data was split at an 80:20 ratio and evaluated using accuracy, precision, recall, and F1-score. The main contribution of this study is demonstrating that despite both algorithms achieving equal overall accuracy of 93%, SVM outperforms Naive Bayes in detecting negative sentiment on a class-imbalanced dataset, with SVM attaining a negative class recall of 0.87 and F1-score of 0.88 compared to 0.80 and 0.87 for Naive Bayes. These findings provide practical guidance for selecting an appropriate classifier in imbalanced e-commerce review classification tasks.

Keywords: Data Mining; Sentiment Analysis; Naive Bayes; Support Vector Machine; Text Mining; Shopee

1. INTRODUCTION

Rapid advancements in digital technology and internet connectivity have become major catalysts for the expansion of the e-commerce industry in Indonesia. Shopee stands out as one of the platforms that has grown substantially, offering a wide range of product categories, including fashion items for men [1]. The high volume of transactions occurring on this platform continuously produces a large accumulation of user-generated reviews. Such reviews serve as a reflection of customers' experiences, assessments, and overall levels of satisfaction with the items they have purchased [2].

Customer reviews represent a contemporary form of electronic word-of-mouth (e-WOM) communication that significantly shapes the purchasing behavior of potential buyers [3]. Nevertheless, the sheer magnitude of review data makes it impractical for sellers or consumers to process this information through manual inspection alone. Consequently, the development of automated methods capable of categorizing review content into positive and negative sentiments has become essential for enabling more effective use of this data [4].

Sentiment analysis constitutes a branch of text mining focused on extracting and classifying subjective opinions within textual content based on their polarity [5]. Among the classification methods commonly applied in this domain are Naive Bayes and Support Vector Machine (SVM). Naive Bayes is widely recognized for its simplicity and computational efficiency in natural language processing tasks [6], whereas SVM operates as a margin-based classifier that demonstrates strong performance when handling high-dimensional and sparse textual representations [7].

A number of prior research efforts have positioned both algorithms as suitable candidates for comparative evaluation in text classification due to their fundamentally distinct computational paradigms [8]. That said, the relative effectiveness of each algorithm may shift considerably based on factors such as the nature of the data, dataset volume, and the linguistic features of the corpus under study [9]. This underscores the importance of conducting dedicated empirical investigations to determine which algorithm performs best under specific conditions, particularly in the context of men's shirt product reviews gathered from the Shopee platform.

Against this backdrop, the present study is designed to systematically evaluate and juxtapose the classification capabilities of Naive Bayes and Support Vector Machine in determining the sentiment polarity of men's shirt product reviews, utilizing a dataset of 500 independently collected data points. The outcomes of this research are intended to offer practical guidance in identifying the most suitable algorithmic approach for sentiment classification tasks involving e-commerce review data.

Based on the background above, this study addresses a specific research gap: although comparative studies on Naive Bayes and SVM exist in the text classification literature, empirical investigation focused on Shopee men's shirt reviews—a domain characterized by class imbalance between positive and negative samples—remains limited. Most prior studies do not explicitly analyze minority-class detection performance, which is precisely the most actionable information for sellers seeking to improve product quality. This study therefore asks: (1) how can text mining be applied to classify sentiments in men's shirt reviews on Shopee; (2) how does Naive Bayes perform on this dataset; (3) how does SVM perform on the same dataset; and (4) which algorithm achieves superior classification performance based on accuracy, precision, recall, and F1-score. The objectives of this study are: (1) to implement a text mining pipeline for preprocessing and classifying sentiments in the review data; (2) to evaluate the performance of Naive Bayes on the dataset; (3) to





evaluate the performance of SVM on the same dataset; and (4) to compare both algorithms to identify which is more suitable for sentiment classification of imbalanced e-commerce review data.

Theoretically, this study contributes to the body of knowledge on machine learning-based sentiment analysis for e-commerce data, and serves as a reference for future researchers investigating similar text classification problems. Practically, the findings offer sellers on Shopee actionable insights to understand consumer perceptions, identify product weaknesses, and improve quality and service based on automated review analysis.

2. RESEARCH METHODOLOGY

2.1 Research Type and Approach

This study adopts a quantitative research design grounded in a comparative experimental framework. The quantitative orientation reflects the nature of the analytical procedure, which involves transforming textual data into numerical representations and assessing classification model performance through established statistical metrics, including accuracy, precision, recall, and F1-score. The comparative dimension of the study centers on empirically evaluating two machine learning classifiers—Naive Bayes and Support Vector Machine (SVM)—as applied to sentiment classification tasks [10].

Thematically, this research operates within the domain of data mining, with a particular concentration on text mining as the primary approach for processing and categorizing written content [4]. The dataset utilized consists of men's shirt product reviews extracted from the Shopee e-commerce platform through an automated scraping process, comprising a total of 500 review entries.

All computational procedures in this research were executed using Python version 3.x as the primary programming language. The key libraries employed encompassed: (1) scikit-learn, utilized for building the Naive Bayes classifier (MultinomialNB), the SVM model (LinearSVC), TF-IDF vectorization (TfidfVectorizer), and computing evaluation metrics; (2) PySastrawi, applied for Indonesian-language stemming operations; (3) pandas, used for structured data handling and management tasks; and (4) matplotlib and seaborn, employed for generating data visualizations and rendering confusion matrices. The entire development and experimentation workflow was conducted within the Google Colaboratory (Colab) cloud-based environment.

2.2 Data Source and Collection

The dataset employed in this research constitutes primary data consisting of customer-written reviews of men's shirt products sourced directly from the Shopee e-commerce platform. The data acquisition was carried out using Selenium-based web scraping, which automated the extraction of publicly available review texts from product pages. The scraping process targeted verified buyer reviews only, excluding promotional or seller-added content. Inclusion criteria required that each review entry contained non-empty text and an associated star rating. No minimum review length was enforced, as short reviews were considered to still carry sentiment value.

In total, 500 customer reviews were gathered. The collected data covers a broad range of consumer experiences related to product quality, sizing, materials, shipping performance, and overall satisfaction. Only review text and star rating were retained; all personally identifiable user information was excluded to uphold ethical research standards. The resulting dataset contains 350 positive and 150 negative samples (70:30 ratio), reflecting a natural class imbalance common in e-commerce review data [11]. To address this imbalance, stratified splitting was applied during data partitioning to ensure proportional class representation in both training and testing sets. Explicit oversampling or undersampling techniques (e.g., SMOTE) were not applied in this study and are acknowledged as a limitation.

Sentiment labels were assigned to each review by leveraging the star rating system provided by Shopee customers. Reviews rated 4 or 5 stars were designated as expressing positive sentiment, while those rated 1 to 3 stars were categorized as negative. This star-rating-based labeling strategy was adopted because numerical ratings serve as a quantifiable and consistent proxy for customer satisfaction, and this approach has been well-established in previous sentiment analysis studies involving e-commerce review datasets [12].

2.3 Research Stages

The procedural flow of this research was structured in a stepwise and methodical manner to ensure that the sentiment classification pipeline is executed with consistency and yields reliable, well-optimized models. The sequence of steps followed throughout the research is described below:

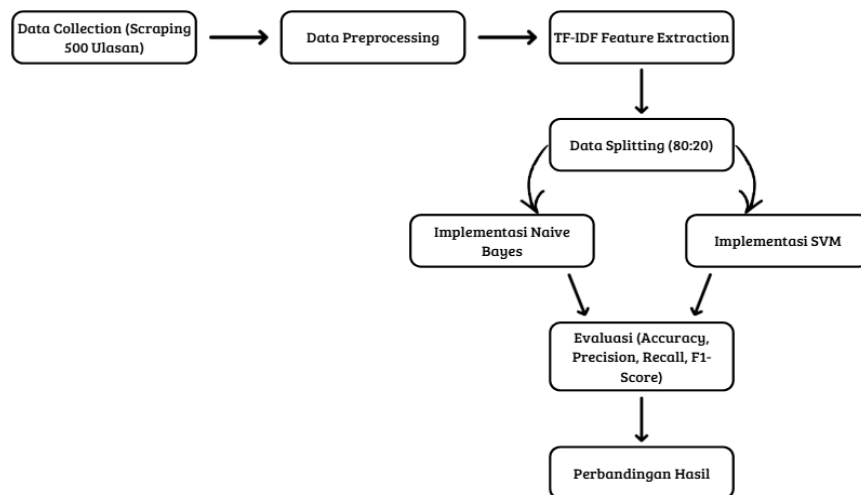


Fig 1. Research Workflow Diagram of Sentiment Classification Using Naive Bayes and SVM

1. **Data Collection:** The initial stage of the research was the collection of men's shirt product review data from the Shopee platform through scraping techniques. A total of 500 reviews were obtained in the form of raw text.
2. **Data Preprocessing:** The preprocessing stage was performed to cleanse and structure the raw text data prior to the classification phase, applying four sequential techniques: Case Folding, Tokenizing, Stopword Removal, and Stemming.
3. **Feature Transformation (TF-IDF):** After preprocessing, the text data is converted into a numerical representation using the Term Frequency-Inverse Document Frequency (TF-IDF) method implemented via scikit-learn's `TfidfVectorizer`. The configuration used `max_features=1000` to cap vocabulary size, with `ngram_range=(1,1)` (unigram only) and `sublinear_tf=True` to apply log normalization to term frequencies. These settings were selected to balance computational efficiency and representational adequacy for the 500-review corpus.
4. **Data Splitting:** The processed dataset was then divided into training data and testing data with a ratio of 80:20 using the stratified split method.
5. **Classification Algorithm Implementation:** The classification process was carried out using Naive Bayes and Support Vector Machine (SVM). For Naive Bayes, the `MultinomialNB` implementation was used with the default smoothing parameter `alpha=1.0` (Laplace smoothing) to handle zero-probability features. For SVM, the `LinearSVC` implementation was used with a linear kernel and the regularization parameter `C=1.0` (default), which was found to be appropriate for this dataset size and feature dimensionality.
6. **Model Evaluation:** The final stage is model performance evaluation using the confusion matrix with metrics: Accuracy, Precision, Recall, and F1-Score.

2.4 Classification Methods

2.4.1 Naive Bayes

Naive Bayes is a probabilistic classifier grounded in Bayes' Theorem, operating under the assumption that all input features contribute independently to class determination. When applied to sentiment analysis tasks, each term present in a document is treated as statistically independent from other terms in forming the overall sentiment prediction. Despite the idealized nature of this independence assumption, Naive Bayes has consistently demonstrated solid performance in text classification problems. Its primary strengths include low computational overhead and robust applicability to datasets of small to moderate size [6].

2.4.2 Support Vector Machine (SVM)

Support Vector Machine (SVM) is a discriminative classification algorithm that operates by identifying an optimal decision boundary, or hyperplane, that maximally separates instances of different classes. In scenarios where text data has been transformed into a numerical vector space through TF-IDF, SVM seeks the widest possible margin between the positive and negative sentiment regions. The algorithm is particularly well-suited for high-dimensional, sparse data structures typical of text corpora [7].

2.5 Model Evaluation Method

To assess the effectiveness of both classifiers, this study applies a model evaluation procedure designed to quantify the performance of the Naive Bayes and Support Vector Machine (SVM) algorithms. The foundation of this evaluation is the confusion matrix, from which key performance indicators are derived and computed.

1. Confusion Matrix

A confusion matrix is a structured tabular tool that facilitates the assessment of a classification model's predictions by cross-referencing them against ground truth labels. In the context of binary sentiment classification (positive and negative), the matrix is composed of four fundamental elements: True Positive (TP), True Negative (TN), False Positive (FP), and False Negative (FN).

2. Accuracy

Accuracy is a measure of the model's overall correctness in classifying data. Accuracy is calculated using the formula:

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

3. Precision

Precision measures the model's accuracy in predicting the positive class:

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

4. Recall

Recall measures the model's ability to detect all data that truly belongs to the positive class:

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

5. F1-Score

F1-Score is the harmonic mean of precision and recall:

$$F1 = 2 \times (\text{Precision} \times \text{Recall}) / (\text{Precision} + \text{Recall}) \quad (4)$$

3. RESULT AND DISCUSSION

3.1 Dataset Description

The dataset employed in this research is drawn from customer reviews of men's shirt products, collected independently through a scraping mechanism applied to the Shopee e-commerce platform. The compiled dataset contains 500 review entries structured across two principal attributes: the review field, which stores the raw customer-written text, and the sentiment field, which holds the corresponding classification label.

An examination of the label distribution reveals that 350 of the reviews (70%) carry a positive sentiment label, while the remaining 150 reviews (30%) are labeled negative. This disparity indicates a class imbalance within the dataset, with positive instances outnumbering negative ones at a 7:3 ratio. Such imbalance is a frequently encountered characteristic in e-commerce review datasets [11], as customers who are satisfied with their purchase tend to submit reviews at higher rates than those who are dissatisfied.

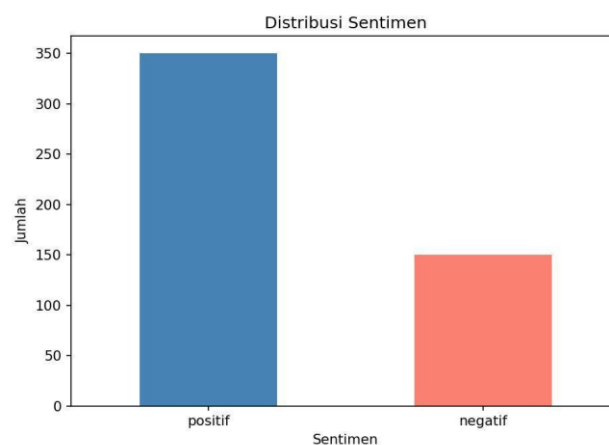


Fig 2. Sentiment Distribution of the Men's Shirt Product Review Dataset on Shopee

3.2 Preprocessing Results

Prior to entering the classification pipeline, the raw text data underwent a preprocessing phase designed to eliminate noise and bring the text into a consistent, standardized format. This phase was structured around four sequential operations: case folding, tokenizing, stopword removal, and stemming.

Case folding involved transforming all textual content to lowercase and stripping out non-alphabetic symbols such as punctuation marks and numeric digits. Tokenizing then decomposed each sentence into its constituent word units.

Stopword removal filtered out high-frequency, low-information words within sentiment analysis contexts—such as “yang”, “di”, and “dan”—leveraging the Sastrawi library. Finally, stemming reduced each remaining word to its morphological root form; for instance, the word “jahitan” was reduced to “jahit” and “terjangkau” was transformed into “jangkau”.

The following are examples of preprocessing results for several reviews in the dataset:

Table 1. Example of Text Review Preprocessing Results

No	Original Text	Preprocessing Result
1	Unexpected, the item matches expectations, honest seller, awesome!	unexpected item match expectation honest seller awesome
2	this flannel is affordable, the quality is not cheap, recommended!	flannel price afford quality not cheap recommended
3	Honestly, the stitching quality is neat, no loose threads sticking out, recommended!	honestly quality stitch neat no loose thread stick out recommended

Upon completion of the preprocessing phase, the cleaned text data was vectorized into a numerical format using the Term Frequency-Inverse Document Frequency (TF-IDF) method, with the `max_features` hyperparameter set to 1000. Through this transformation, 306 distinct features were generated and subsequently fed into the classification models. The resulting feature matrix was then partitioned into 400 training samples (80%) and 100 test samples (20%) via the stratified split procedure.

3.3 Naive Bayes Classification Results

The sentiment classification process using the Naive Bayes algorithm was carried out by training the model on 400 training data items and then testing it on 100 testing data items. The model's prediction results are visualized using the confusion matrix as follows:

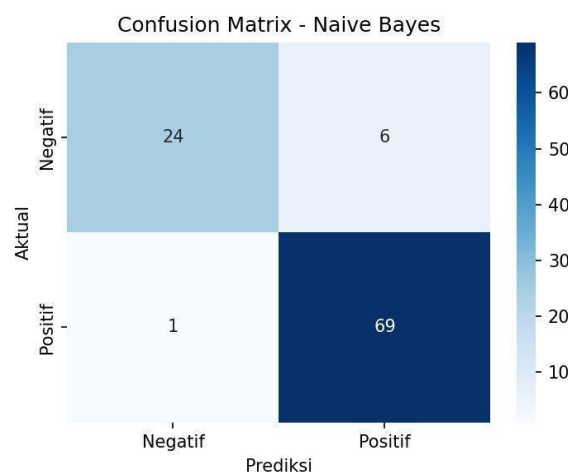


Fig 3. Naive Bayes Confusion Matrix

Based on the confusion matrix, the performance evaluation results of the Naive Bayes algorithm can be summarized in the following table:

Table 2. Naive Bayes Evaluation Results

Class	Precision	Recall	F1-Score	Support
Negative	0.96	0.80	0.87	30
Positive	0.92	0.99	0.95	70
Weighted Avg	0.93	0.93	0.93	100

The Naive Bayes classifier attained an overall accuracy of 93%. Within the positive class, the model demonstrated an exceptionally high recall of 0.99, reflecting its strong capacity to correctly identify nearly all positive review instances. In contrast, recall performance on the negative class was considerably lower at 0.80, implying that roughly 6 negative reviews were misclassified as positive. This shortfall in minority class recognition is attributable to the skewed distribution of the dataset, wherein the scarcity of negative samples constrains the model’s ability to learn robust decision boundaries for that class.

3.4 Support Vector Machine (SVM) Classification Results

The sentiment classification process using the Support Vector Machine (SVM) algorithm was carried out with a linear kernel configuration and trained using the same 400 training data items, then tested on 100 testing data items. The model's prediction results are visualized using the confusion matrix as follows:

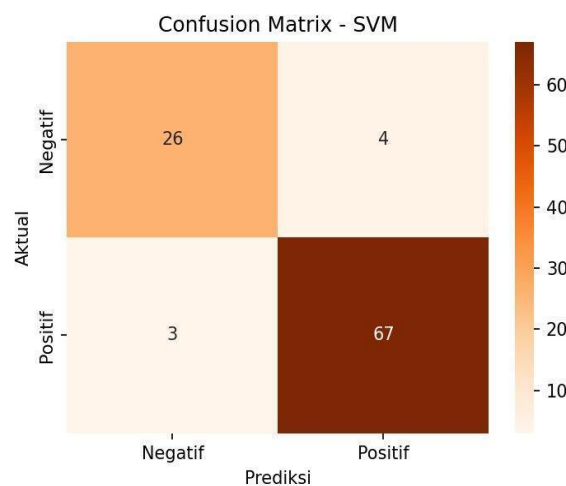


Fig 4. Support Vector Machine Confusion Matrix

Based on the confusion matrix, the performance evaluation results of the SVM algorithm can be summarized in the following table:

Table 3. SVM Evaluation Results

Class	Precision	Recall	F1-Score	Support
Negative	0.90	0.87	0.88	30
Positive	0.94	0.96	0.95	70
Weighted Avg	0.93	0.93	0.93	100

The SVM classifier also achieved an accuracy of 93%. For the negative class, SVM recorded a recall of 0.87, successfully identifying 26 out of 30 negative reviews. When compared to Naive Bayes, SVM exhibited greater sensitivity toward the minority class (negative), though this came at a slight cost to precision—the SVM negative class precision of 0.90 fell marginally below Naive Bayes' value of 0.96. This pattern reflects SVM's tendency to cast a wider net in predicting negative sentiment, which effectively minimizes the number of missed negative reviews.

3.5 Performance Comparison of Naive Bayes and Support Vector Machine

Having completed the classification experiments using both algorithms, a comprehensive comparative analysis of their performance was conducted, evaluating Naive Bayes and Support Vector Machine (SVM) across the full set of metrics: accuracy, precision, recall, and F1-score.

Table 4. Performance Comparison of Naive Bayes and SVM

Metric	Naive Bayes	SVM
Accuracy	93%	93%
Precision (Negative)	0.96	0.90
Precision (Positive)	0.92	0.94
Recall (Negative)	0.80	0.87
Recall (Positive)	0.87	0.96
F1-Score (Negative)	0.87	0.88
F1-Score (Positive)	0.95	0.95
Weighted Avg F1	0.93	0.93

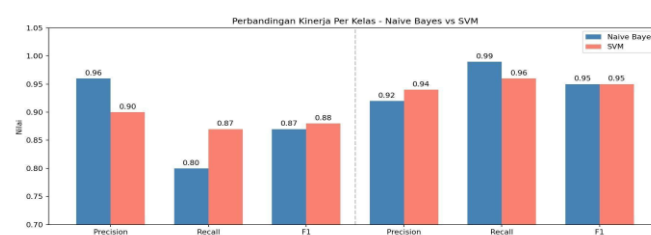


Fig 5. Performance Comparison Graph of Naive Bayes and SVM

Based on the table and graph above, both algorithms produced the same accuracy value of 93%. However, when examined more deeply at the per-class level, there are significant differences between the two algorithms, particularly in the ability to detect the negative class.



With respect to negative class precision, Naive Bayes outperformed SVM, recording a value of 0.96 against SVM's 0.90. This means that when Naive Bayes classifies a review as negative, this judgment tends to be accurate more often. However, this high precision comes with a trade-off: the negative class recall for Naive Bayes reached only 0.80, indicating that approximately 6 negative reviews went undetected and were erroneously labeled as positive. Previous research by Angreyani found that similar constraints in Naive Bayes' ability to handle minority-class samples within e-commerce review data [13].

SVM, on the other hand, achieved a superior negative class recall of 0.87, demonstrating a greater overall capacity to identify negative reviews. This characteristic carries significant practical value, as negative reviews represent the most operationally relevant feedback for sellers seeking to address quality concerns. The 70:30 class imbalance within the dataset further amplified the behavioral differences between the two algorithms in handling the minority class.

Regarding the F1-score for the negative class, SVM attained 0.88 compared to 0.87 for Naive Bayes, reflecting a marginally superior equilibrium between precision and recall in identifying negative sentiment instances. Based on these findings, despite both classifiers yielding identical overall accuracy, SVM emerges as the more appropriate choice for sentiment classification on imbalanced e-commerce review corpora, especially when the priority is maximizing detection of the minority class. Riki similarly demonstrated SVM's advantage over Naive Bayes in managing class-imbalanced sentiment data [9].

In terms of validation methodology, the study employed a stratified train-test partitioning strategy with an 80:20 split ratio. Stratified sampling was deliberately chosen to preserve the proportional distribution of positive and negative class instances across both the training and testing subsets, thereby mitigating evaluation bias that commonly arises when working with class-imbalanced datasets [11]. While cross-validation was not incorporated into the experimental design, the convergence of results observed across both Naive Bayes and SVM provides reasonable confidence that the performance estimates are sufficiently reliable and reproducible.

Concerning the scale of the dataset, a total of 500 reviews may appear modest in comparison with larger corpora employed in some research endeavors. However, this sample size aligns with the minimum volume generally deemed acceptable for machine learning-based sentiment analysis conducted in academic settings [14]. That said, caution is warranted when attempting to generalize the findings to other product categories or digital platforms, and future research would benefit from expanding the dataset accordingly.

4. CONCLUSION

This study compared the performance of Naive Bayes and Support Vector Machine (SVM) in classifying the sentiment of 500 Shopee customer reviews on men's shirt products. The results showed that both algorithms achieved the same overall accuracy of 93%. However, SVM demonstrated better performance in identifying negative sentiment, with a recall of 0.87 and an F1-score of 0.88, compared with recall and F1-score values of 0.80 and 0.87, respectively, for Naive Bayes. Considering that the dataset was imbalanced, comprising 70% positive reviews and 30% negative reviews, SVM can be regarded as the more appropriate model for this classification task. Its higher sensitivity to the minority class is particularly useful for sellers because it supports the early identification of negative customer feedback and facilitates more responsive product and service improvements.

Nevertheless, this study has several limitations. The dataset consisted of only 500 reviews obtained from a single seller and was limited to one product category, namely men's shirts on the Shopee platform. Consequently, the findings may not be directly generalized to other product categories, sellers, or e-commerce platforms. In addition, no specific class-imbalance handling techniques were applied, and model evaluation was conducted using a stratified train-test split rather than cross-validation. Although stratified splitting maintained proportional class representation in the training and testing sets, it may not fully reflect the stability of model performance across different data partitions.

Future research should expand the dataset by including a larger number of reviews from multiple sellers, product categories, and e-commerce platforms. Techniques such as Synthetic Minority Over-sampling Technique (SMOTE), class weighting, or other resampling methods should also be investigated to improve minority-class detection. Furthermore, cross-validation should be applied to obtain a more robust and reliable performance assessment. The use of more advanced language models, such as IndoBERT, may also be explored to determine whether contextual text representations can further improve sentiment classification performance.

REFERENCES

- [1] Puteri, G. A., Trisnawarman, D., & Budiyantera, A, "Comparison of Shopee and Tokopedia user sentiment using Naive," Jurnal Ilmu Komputer dan Sistem Informasi, vol. 13, no. 1, 2025.
- [2] Wibisono, A. C., Nadira, T. S., & Sutabri, T, "Customer sentiment analysis on the Shopee platform using the Naive Bayes method," Nusantara Journal of Multidisciplinary Science, 2024.



- [3] Andriani, N., Ma'arif, M. S., Amalia, I. N., & Rozikin, C, "Sentiment analysis of Shopee application users using Naive Bayes Classifier," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 12, no. 1, 2023. <https://journal.eng.unila.ac.id/index.php/jitet/article/view/3631>
- [4] Fauzi, M. A., & Arifin, A. Z, "Text mining for sentiment analysis of marketplace application reviews using machine learning," *Jurnal RESTI*, vol. 5, no. 3, pp. 512-519, 2021.
- [5] Halim, F., & Wicaksono, S. A, "Implementation of TF-IDF and SVM for sentiment classification of e-commerce product reviews," *Jurnal Teknologi dan Sistem Komputer*, vol. 11, no. 2, pp. 89-97, 2023.
- [6] Pratama, D. A., & Nugroho, A, "Sentiment analysis of marketplace users using the Naive Bayes algorithm based on text mining," *Jurnal Ilmiah Informatika Komputer*, vol. 26, no. 3, pp. 203-212, 2021.
- [7] Rahman, A., & Prasetyo, E, "Comparison of SVM and Naive Bayes methods in sentiment analysis of marketplace reviews," *Jurnal Teknologi Informasi dan Ilmu Komputer (JTIK)*, vol. 9, no. 4, pp. 811-818, 2022.
- [8] Sanjaya, T. P. R., Fauzi, A., & Masruriyah, A. F. N, "Sentiment analysis of reviews on Shopee e-commerce using Naive Bayes and SVM algorithms," *INFOTECH Journal*, vol. 4, no. 1, 2023.
- [9] Riki, A., Dameria, E., Hermawan, A., Junaedi, J., & Kurnia, Y, "Optimizing sentiment classification of e-commerce product reviews: A comparative study of Naive Bayes and SVM with SMO," *JUITA: Jurnal Informatika*, vol. 13, no. 3, pp. 287-295, 2025.
- [10] Jannah, N. Z. B., & Kusnawi, K, "Comparison of Naive Bayes and SVM in sentiment analysis of product reviews on marketplaces," *Sinkron: Jurnal dan Penelitian Teknik Informatika*, vol. 8, no. 2, pp. 727-733, 2024.
- [11] Nabilah, N., Rohimah, S., & Loka, S. K. P, "Analyzing customer sentiment towards marketplace reviews using machine learning algorithms," *IJATIS: Indonesian Journal of Applied Technology and Innovation Science*, vol. 3, no. 1, 2024.
- [12] R. Rismansyah, "Sentiment analysis of Shopee reviews on the Google Play Store using the Naive Bayes algorithm," *JEIS: Jurnal Elektro dan Informatika Swadharma*, vol. 5, no. 1, 2025.
- [13] Angreyani, J., & Pernando, Y, "Sentiment analysis of Shopee e-commerce reviews using the Naive Bayes algorithm," *J-Com: Jurnal Informatika*, vol. 5, no. 1, 2024. <https://jurnal.stmikroyal.ac.id/index.php/j-com/article/view/3570>
- [14] Bahtiar, S. A. H., Dewa, C. K., & Luthfi, A, "Comparison of Naive Bayes and Logistic Regression in sentiment analysis," *Journal of Information Systems and Informatics*, vol. 5, no. 2, 2023. <https://journal-isi.org/index.php/isi/article/view/539>
- [15] Yulianeu, A., & Setiawan, R, "Sentiment analysis on product reviews from Shopee marketplace using the Naive Bayes classifier," *Lontar Komputer*, vol. 13, no. 2, 2022.
- [16] Al Fachrozi, M., & Tania, K. D, "Comparison of Naive Bayes, SVM, K-NN, Decision Tree, and Random Forest in Sentiment Analysis Based on SeaBank Application Aspects," *JURTEKSI*, vol. XII, no. 1, pp. 69-80, 2025.
- [17] Fatihaturrahmah, A., & Tania, K. D, "Review: A Hybrid Approach of Aspect-Based Sentiment Analysis and Knowledge Extraction for Evaluating Security Perceptions in Digital Payment Applications," *Scientific Journal of Informatics*, vol. 12, no. 4, 2025.
- [18] Ardhani, D. A., & Tania, K. D, "Knowledge discovery on e-commerce customer churn using interpretable machine learning: A comparative study of SHAP-based classifiers," *Journal of Applied Informatics and Computing*, vol. 9, no. 5, pp. 2695-2702, 2025.
- [19] Ayuningtiyas, P., Tania, K. D, "Sentiment-Based Knowledge Discovery on the iPusnas Application Using Machine Learning and Deep Learning Methods," *Journal of Applied Informatics and Computing*, vol. 9, no. 5, 2025.
- [20] Depari, A. D. S., Tania, K. D., & Sevtiyuni, P. E, "Application of Machine Learning Methods and SMOTE Technique for Diabetes Prediction," *Jurnal Sistem Komputer dan Informatika*, vol. 7, no. 2, 2025.