

Sentiment Analysis of the TikTok Tokopedia Seller Center Application Using Support Vector Machine (SVM) and Naive Bayes Algorithms

Faddilla Aulia Dara *

Information System Study Program, Universitas Mercu Buana Yogyakarta, Sleman Regency, Special Region of Yogyakarta, Indonesia.

Corresponding Email: 221220075@student.mercubuana-yogya.ac.id.

Irfan Pratama

Information System Study Program, Universitas Mercu Buana Yogyakarta, Sleman Regency, Special Region of Yogyakarta, Indonesia.

Email: irfanp@mercubuana-yogya.ac.id.

Received: December 12, 2024; Accepted: January 20, 2025; Published: April 1, 2025.

Abstract: The TikTok Tokopedia Seller Center application is a collaboration between TikTok and Tokopedia designed to help sellers manage their stores and boost sales. Despite offering various features, complaints about poor user experience often appear in reviews on the Google Play Store. This study aims to analyze user sentiment towards the TikTok Tokopedia Seller Center application using a dataset of 2,000 reviews, using the Support Vector Machine (SVM) and Naive Bayes algorithms to classify positive, negative, and neutral sentiments. In addition, this study also attempts to compare the effectiveness of these algorithms in sentiment analysis and evaluate the performance of two weighting methods: TF-IDF and Term Presence. The dataset used was taken by scraping review data on the Google Play Store in Python, as many as 2000 user review datasets. This study found 1,171 negative sentiments, 735 positive sentiments, and 94 neutral sentiments. The results showed that the accuracy of SVM (0.81 and 0.78) was higher than Naive Bayes (0.69 and 0.75). It is hoped that this research can help potential users to find user sentiment towards the application and provide valuable information for application developers to understand user needs and expectations so that developers can improve application features more appropriately and effectively.

Keywords: Google Play Store; Naive Bayes; Sentiment Analysis; Support Vector Machine; TikTok Tokopedia Seller Center.

1. Introduction

In the current era of information technology advancement, more and more business platforms in the digital world are adopting a new form of combining social media and e-commerce called social e-commerce. Social e-commerce is a platform based on social networks, promoting commodity transactions through user-generated content, social interactions, and other means to socialize e-commerce, thereby helping to develop e-commerce and create a more interactive and efficient business experience [1]. Foreign platforms that have used this business system and are growing rapidly in their countries include the Red or XiaoHongShu platform in China, Line Shopping in Japan, Amazon Live in America, and others. In Indonesia, the application that implements the social e-commerce system is the TikTok Tokopedia Seller Center application, which is a collaborative application between TikTok and Tokopedia, where a unique integration between the TikTok short video social media platform and Tokopedia's e-commerce infrastructure to reach more audiences through a content-based approach and social interaction [2].

The TikTok Tokopedia Seller Center application emerged to provide an easy-to-use platform and help sellers reach a broader market. This application offers complete features, such as store management, payment features, product advertising, sales reports, and other important features. With these features, sellers can maximize marketing strategies and improve customer interactions. However, with the various advantages of the application, sellers must face several challenges, such as customer complaints about specific features, inadequate service, and the overall user experience with the application [4]. The user can write down their experience while using the application on the Google Play Store comment section, where each user can provide good or bad ratings and reviews of an application [5]. Therefore, sentiment analysis is a very appropriate method to understand and extract meaning from user reviews in more depth. Sentiment analysis is a technique used to identify and classify positive, negative, and neutral sentiments in user review texts. With sentiment analysis, application developers can more easily understand user perceptions of the application and make targeted improvements [6]. The data used are user reviews on the Google Play Store, which will be valuable information for improving application performance. The review on the Google Play Store contains information about user satisfaction. These reviews allow researchers to analyze sentiment towards the TikTok Tokopedia Seller Center application [20].

Sentiment analysis usually uses Natural Language Processing (NLP) and machine learning algorithms to classify sentiment in reviews. This study will use two popular algorithms in sentiment analysis, namely the Naive Bayes method and the Support Vector Machine (SVM). Naive Bayes is a learning algorithm often used for data classification. The Naive Bayes algorithm aims to determine the opinion on a particular text by calculating the probability of it falling into a particular sentiment category, such as positive or negative [7]. Meanwhile, Support Vector Machine (SVM) is an algorithm for classification and regression. This algorithm works well in complex and linearly non-separable cases [29]. This study also aims to compare the performance of the SVM and Naive Bayes algorithms with the InSet lexicon label and compare the performance of the vectorization method using TF-IDF and Term Presence to find out which is better in the sentiment analysis of this application [8].

Many previous studies related to sentiment analysis have been conducted [9]. One of them is sentiment analysis of reviews towards electric cars using Naive Bayes Classifier and Support Vector Machine Algorithm [17]. The results of the study showed that the accuracy of the SVM algorithm was 90%, while the Naive Bayes algorithm had an accuracy of 88%. Thus, SVM has a better ability to classify than Naive Bayes in this study. In other studies, namely sentiment analysis of user reviews on Covid-19 information applications using Naive Bayes Classifier, Support Vector Machine, and K-Nearest Neighbor [16]. The results of the study stated that SVM has 76.5% accuracy, followed by NB 3%, and KNN has 59.1% accuracy. Another study was the analysis of user reviews for the Peduli Lindungi application on Google Play using the Support Vector Machine and Naive Bayes Algorithm Based on Particle Swarm Optimization [10]. The accuracy results of the two algorithms showed that the highest accuracy value is the PSO-based SVM algorithm compared to the PSO-NB algorithm.

In another study on the comparison of Support Vector Machine and Naive Bayes on Twitter Data Sentiment Analysis [11]. The result of the comparison of the accuracy of the SVM kernel RBF with NB, the highest accuracy value is obtained from the SVM kernel RBF method, which is 88. Research on sentiment analysis on E-Sports for education curriculum using Naive Bayes and Support Vector Machine [12] obtained test results with an accuracy value of 70.32% for the Naive Bayes method, an accuracy value of 66.92% for the Support Vector Machine method using the SMOTE. Research on Naive Bayes Classifier method analysis and Support Vector Machine (SVM) Student Graduation Prediction [13]. The results with the AUC (Area Under Curve) value with Naive Bayes is 0.919 and for the SVM model is 0.083.

Although there have been many studies on sentiment analysis of various data, there have not been many studies that discuss the direct comparison of the performance between the Support Vector Machine and Naive Bayes classification algorithms along with the comparison of the effectiveness between TF-IDF and Term Presence. Therefore, this study attempts to fill this gap by analyzing the sentiment of the TikTok Tokopedia

Seller Center application using the Support Vector Machine and Naive Bayes algorithms [8]. In addition, this study is expected to help better understand user sentiment towards the TikTok Tokopedia Seller Center application by implementing the Support Vector Machine (SVM) and Naive Bayes algorithm and aims to provide accurate data processing [9]. That way, the results of this sentiment analysis can help potential users to find out user sentiment towards the application and provide valuable information for application developers to understand user needs and expectations so that developers can improve application features more precisely and effectively.

2. Related Work

Sentiment analysis has become an indispensable tool in understanding user perceptions and experiences across various digital platforms, including social media and e-commerce applications. Several studies have explored the effectiveness of different algorithms and techniques in sentiment analysis, particularly focusing on the TikTok Tokopedia Seller Center application. This section reviews relevant literature to provide a comprehensive understanding of the methodologies and findings in the field. Hidayat and Nastiti (2024) conducted a comparative study of pre-trained IndoBERT-base and IndoBERT-lite models in classifying sentiments from TikTok Tokopedia Seller Center reviews. Their results indicated that IndoBERT-base outperformed IndoBERT-lite in terms of accuracy and F1-score. This suggests that more sophisticated models, despite their higher computational requirements, are better equipped to handle the complexity and nuances of user-generated content [2]. This finding aligns with the broader literature on the effectiveness of deep learning models in natural language processing tasks, where pre-trained models have shown significant improvements in various benchmarks. Tanniewa, Hamrul, and Sarina (2023) implemented the Support Vector Machine (SVM) algorithm to analyze sentiments of users of the TikTok Shop Seller Center application. They reported that SVM achieved a high accuracy rate, particularly in distinguishing between positive and negative sentiments. The study emphasized the robustness of SVM in handling non-linear relationships and its effectiveness in sentiment analysis. This is consistent with the findings of Muttaqin and Kharisudin (2021), who conducted sentiment analysis on Gojek app reviews using SVM and KNN. They found that SVM had a higher accuracy compared to KNN, further validating the effectiveness of SVM in sentiment analysis tasks [5].

Wijaya *et al.* (2024) applied both Naive Bayes and K-Nearest Neighbor (KNN) algorithms to analyze the sentiments of TikTok Shop Seller Center reviews. They found that Naive Bayes performed slightly better than KNN, especially in terms of precision and recall, indicating its suitability for short and straightforward reviews [4]. This finding is supported by the work of Indriyani, Fauzi, and Faisal (2023), who performed sentiment analysis on TikTok app reviews using both Naive Bayes and SVM algorithms. They concluded that SVM provided more accurate results, particularly in classifying nuanced and lengthy reviews, while Naive Bayes was effective for shorter and more straightforward reviews [6]. Atmajaya, Febrianti, and Darwis (2023) compared the performance of SVM and Naive Bayes in analyzing sentiments of ChatGPT on Twitter. They found that SVM outperformed Naive Bayes in terms of accuracy and F1-score, highlighting the strength of SVM in dealing with large and diverse datasets [7]. Similarly, Salma and Silfianti (2021) conducted sentiment analysis on user reviews of COVID-19 information applications using Naive Bayes Classifier, SVM, and K-Nearest Neighbors. They found that SVM achieved the highest accuracy of 76.5%, further confirming the superior performance of SVM in sentiment analysis tasks [16].

Prajaitan Febrianti (2024) explored the use of TF-IDF Vectorizer, COUNTVECTORIZER, and HASHINGVECTORIZER with parameter optimization in machine learning for sentiment analysis of the 2024 Indonesian election. Their results showed that TF-IDF Vectorizer, when optimized, significantly improved the accuracy of sentiment classification [18]. This highlights the importance of text vectorization techniques in enhancing the performance of machine learning models in sentiment analysis. Oktavia Praneswara and Cahyono (2023) conducted sentiment analysis on TikTok Shop Seller Center reviews using the Naive Bayes algorithm. They found that the algorithm effectively classified user sentiments, providing valuable insights into user experiences and areas for improvement [14]. Ainunnisa and Sulastris (2023) performed sentiment analysis on the TikTok application using SVM, Logistic Regression, and Naive Bayes. They reported that SVM and Logistic Regression achieved higher accuracy compared to Naive Bayes, particularly in handling complex and nuanced reviews [15]. Suryani *et al.* (2023) conducted sentiment analysis on electric cars using both Naive Bayes Classifier and SVM. They found that SVM provided more accurate results, especially in classifying positive and negative sentiments, while Naive Bayes was effective for neutral sentiments [17]. This study underscores the importance of using multiple algorithms to capture a comprehensive range of sentiments and improve overall classification accuracy. Hasibuan *et al.* (2024) analyzed the sentiment of TikTok Shop features using Naive Bayes and K-Nearest Neighbor (KNN). They emphasized the importance of feature engineering and optimization in improving the performance of sentiment analysis models. Their results showed that optimized features significantly enhanced the accuracy of sentiment classification [19]. This aligns with the broader

literature on the importance of feature selection and preprocessing in machine learning tasks, where well-engineered features can lead to more accurate and reliable models.

Musfiroh *et al.* (2021) conducted sentiment analysis on online learning experiences in Indonesia using the InSet lexicon. They found that lexicon-based approaches, particularly InSet, were effective in identifying sentiment in Indonesian text data. This study highlights the importance of domain-specific lexicons in sentiment analysis, especially for languages with unique linguistic characteristics [24]. Triyanti (2023) analyzed the sentiment of users of a language learning application using the Naive Bayes algorithm. They found that the algorithm effectively classified user sentiments, providing valuable insights into user satisfaction and areas for improvement [25]. This study underscores the importance of sentiment analysis in understanding user experiences and improving product development. Fazal and Andraini (2022) compared the performance of SVM and Naive Bayes on Twitter data. They found that SVM outperformed Naive Bayes in terms of accuracy and F1-score, particularly in handling large and diverse datasets [27]. This finding is consistent with the broader literature on the effectiveness of SVM in sentiment analysis tasks, where its ability to handle non-linear relationships and large datasets makes it a preferred choice. Pasaribu and Sriani (2023) conducted sentiment analysis on Shopee application user reviews using the Naive Bayes algorithm. They found that the algorithm effectively classified user sentiments, providing valuable insights into user experiences and areas for improvement [29]. This study highlights the importance of application-specific sentiment analysis in understanding user perceptions and improving service quality. Bhuiyan *et al.* (2024) conducted a comparative study of machine learning models in sentiment analysis of customer feedback in the banking sector. They found that SVM and Random Forest outperformed other algorithms in terms of accuracy and F1-score, particularly in handling complex and diverse datasets [30]. This study underscores the importance of selecting appropriate algorithms based on the characteristics of the dataset and the specific requirements of the sentiment analysis task. Santoso and Wibowo (2022) conducted sentiment analysis on Buzzbreak app reviews using the Naive Bayes Classifier on the Google Play Store. They found that the algorithm effectively classified user sentiments, providing valuable insights into user experiences and areas for improvement [33]. This study highlights the importance of lexicon-based approaches in sentiment analysis, particularly for domain-specific applications. The literature on sentiment analysis of the TikTok Tokopedia Seller Center application and other digital platforms provides valuable insights into the effectiveness of various algorithms and techniques. SVM and Naive Bayes are among the most commonly used algorithms, with SVM generally outperforming Naive Bayes in handling complex and nuanced reviews. Text vectorization techniques, feature engineering, and domain-specific lexicons play crucial roles in enhancing the performance of sentiment analysis models. Future research should focus on developing more sophisticated models and techniques to further improve the accuracy and reliability of sentiment analysis in various applications.

3. Research Method

3.1 Research Instrument

This research utilized Python, leveraging libraries relevant to text processing needs. In terms of data sources, the study gathered comments from the TikTok Tokopedia Seller Center application on the Google Play Store through a scraping mechanism.

3.2 Sentiment Analysis Workflow

Figure 1 illustrates the stages of the research. The stages of this research consist of a literature study, data collection, data preprocessing, classification modeling, and finally, evaluation of the classification model.

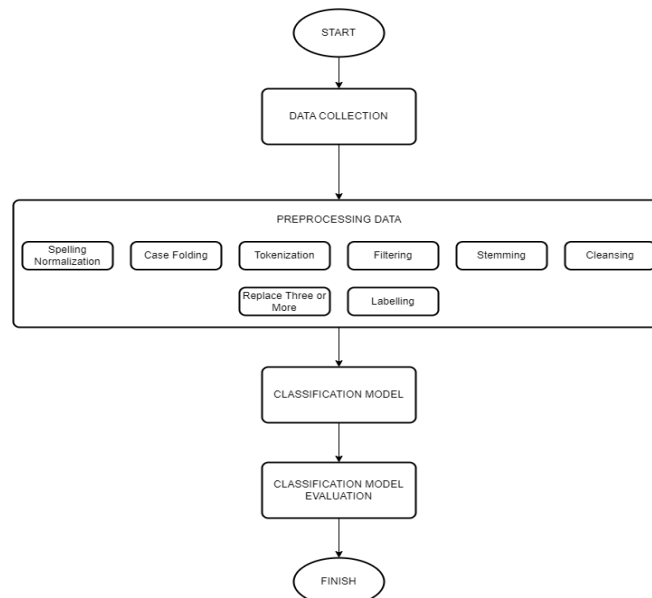


Figure 1. Sentiment Analysis Workflow

1) Data Collection

The TikTok Tokopedia Seller Center app review dataset was collected using web scraping in Python. The review dataset was taken from user reviews of the TikTok Tokopedia Seller Center app on the Google Play Store. The data scraping process focused on the 2000 most relevant review data.

2) Data Preprocessing

At this stage, the raw and dirty dataset must be cleaned and prepared before use. Preprocessing is an effective method for cleaning and organizing initially unstructured data into structured and meaningful data [6]. The data preprocessing stage involves several processes:

3) Spelling Normalization

Spelling normalization involves changing or improving non-standard words into standard ones. The review data often contains non-standard, abbreviated, and misspelled words. Therefore, the review data containing non-standard words must be standardized to ensure the classification process runs smoothly [14].

4) Case Folding

Case folding involves changing all letters in a text to lowercase. This step is essential to ensure consistency in processing text data, avoiding differences caused by upper and lower case letters [26].

5) Tokenization

Tokenization is the process of splitting sentences into word units or tokens separated by spaces or punctuation, such as commas. This step is crucial for further analysis and text processing [25].

6) Filtering

Filtering involves removing unnecessary characters, such as numbers, symbols, emojis, and punctuation. The process includes reading each data line and detecting whether there are letters, numbers, symbols, emojis, and punctuation. All numbers, symbols, emojis, and text in the form of links and punctuation that are detected are deleted [25].

7) Stemming

Stemming is the process of reducing words to their root form by removing prefixes or suffixes. For example, the sentence "For flashlight sellers, please make releasing accounts more than 3 times more manageable. If you want to release them through any means, there is no solution at all" will be processed to remove conjunctions, pronouns, and prepositions, leaving only essential words [28].

8) Cleansing

Cleansing involves detecting and repairing corrupt or inaccurate data. This stage is crucial because dirty data can affect the accuracy of the analysis or predictions. Before cleaning the data, the dataset will be reviewed to understand its description and information related to the dataset [27].

9) Labeling

Initial labeling is done to mark data with appropriate sentiment, such as positive, negative, and neutral labels [15]. In this study, the labeling process will use the Lexicon InSet to identify the initial sentiment based on the lexicon weight of the sentence. Lexicon InSet is an Indonesian language lexicon used in microblogs. It contains around 10,000 words, with 3,612 words having positive sentiment and 6,612 words having negative sentiment. Each word in the dictionary has a weight ranging from -5 (very negative) to 5 (very positive) [17].

The sentiment score is produced by mapping all the word tokens from the previous step to the dictionary. The polarity score calculation process adds up all the word weights detected by the system. The review data will then be classified into sentiment types using the algorithm [18]. The following is a general description of the sentence classification algorithm into sentiment:

If sentiment score > 0, then Positive Sentiment (1)

If sentiment score < 0, then Negative Sentiment (2)

If sentiment score = 0, then Neutral (3)

The classification of review sentences into positive, negative, and neutral sentiments is determined by the number of polarity scores (sentiment scores) obtained. Review sentences are classified as positive if the polarity score is greater than 0, negative if the polarity score is less than 0, and neutral if the polarity score is equal to 0 [18].

3.3 Classification Model

At this stage, classification modeling will be carried out using the Support Vector Machine (SVM) and Naive Bayes algorithms. According to Agarwal *et al.* (2011), Naïve Bayes works well on short and straightforward reviews but struggles with longer and more nuanced feedback due to its assumptions about word independence. Meanwhile, according to Cristianini and Shawe-Taylor (2000), SVM is very effective in separating positive and negative sentiments, even in datasets where sentiment boundaries are not clearly defined. SVM's robustness in handling non-linear relationships and the use of kernel functions make it a powerful tool for sentiment classification. The user review data includes both short and concise texts as well as long and complex texts, making the use of these two algorithms suitable for comparison to provide more comprehensive results in sentiment analysis [33]. Before proceeding to the classification model process, word weighting will be performed using TF-IDF and Term Presence to optimize the performance of the algorithms in classifying sentiment, resulting in more accurate and reliable sentiment classification. This process increases accuracy and efficiency in sentiment analysis [16][29]. TF-IDF is a weighting method consisting of two values derived from two different weighting methods: Term Frequency (TF) and Inverse Document Frequency (IDF). If a word appears frequently in the document, the output with the term will produce a high TF-IDF value, while words that rarely appear in the document will produce a low value. By using TF-IDF, important words will have a high value, and vice versa. Conversely, if a word in the reference word list is found in the weighted data, the value of the word in the feature vector will be given a value of 1 regardless of the number of occurrences of the word [19].

$$\text{idf}(t_i, d_j) = \log \frac{|D|}{\#d(t_i)} \quad (2)$$

$$\text{tf idf}(t_i, d_j) = \text{tf}(t_i, d_j) \times \text{idf}(t_i, d_j) \quad (3)$$

Term Presence is a weighting method in a text document that looks at the existence of a list of words or terms in a document. If a feature contained in the reference feature list is found in the assessed document, it will be scored 1 in the feature vector regardless of the number of occurrences of the feature. However, if the feature is not found in the document, it is given a score of 0 in the feature space [32]. TF-IDF and Term Presence were chosen as weighting methods because TF-IDF can help reduce noise from words that are too common but have no meaning, and it provides high weighting values to words that have meaning. Term Presence provides a simple but effective weighting value to identify keywords in sentiment and helps reduce bias from words that are repeated more than once in one review. In the weighting process, TF-IDF will be assisted by the cosine similarity method to run more optimally. The comparison between the dataset managed by TF-IDF and the presence of terms at the classification model evaluation stage will determine the most effective weighting method. This study uses a training dataset for classification modeling, with 20% of the data used for training and 80% for testing, from a total of 2000 datasets.

3.4 Classification Model Evaluation

At this stage, the performance of the classification model will be assessed using matrices such as accuracy, precision, recall, and F1-score in the confusion matrix [7].

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (4)$$

$$\text{Recall} = \frac{TP}{TP + FN} \quad (5)$$

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

$$\text{F1 - score} = 2 \times \frac{\text{Recall} \times \text{Precision}}{\text{Recall} + \text{Precision}} \quad (7)$$

These metrics will provide a comprehensive evaluation of the classification model's performance, helping to identify the most effective algorithm and weighting method for sentiment analysis.

4. Result and Discussion

4.1 Results

4.1.1 Data Collection

At the data collection stage, the researcher collects data that will be used as the main material for this study. Data is taken by scraping TikTok Tokopedia Seller Center application review data on the Google Play Store application with Python. The data scraping process begins with the data scraping installation process on Python, then calling numpy, pandas, and others. Then, determine the sorting option to retrieve review data that will be recalled in the following order: link from the TikTok Tokopedia Seller Center application page. Then, the reviews will be taken in the form of Indonesian language reviews; in the Indonesian region, the most relevant reviews of 2000 reviews were selected. Next, call the previously selected data in the review column. After that, the columns that will be taken are only the content column containing reviews and the score column containing ratings. Then, the selected columns are renamed and saved in CSV format. The following is the result of data scraping in table 1.

Table 1. Result of Data Collection

	Review	Rating
0	suka karena ladang bisnis ku dsini, tapi tolong pencairan dana lebih di percepatan soalnya uangnya buat modal kembali dan semoga tokoku bisa mendapatkan subsidi ongkir full jadi tidak harus dibebankan kepelangganku terima kasih	s
1	Tolonglah kasih pilihan tarik saldo bisa pakai rekening bank beda nama pemilik, seperti marketplace pada umumnya. Ini aplikasi jualan untuk meningkatkan umkm masyarakat umum loh, jadi sifatnya bukan lagi privat marketplace. Harusnya tidak ada pembatasan dalam hal ini	1
2	Untuk pencairan dana dari konsumen ke akun tiktok sungguh lama banget jadi untuk yang bermodal kecil sungguh sangat merugikan karna uang yang harusnya dimodalkan kembali malah lama untuk diputar kembali. Sangat tidak rekomendasi untuk pedagang kecil apabila ingin berjualan di tiktok Sangat sangat merugikan	1

4.1.2 Data Preprocessing

At this stage, after the dataset is obtained. This raw dataset will be cleaned or prepared first before being analyzed. In this data preprocessing, several processes will be carried out, namely spelling normalization, case folding, tokenization, filtering, stemming, cleansing, and there is an additional replace three or more to remove repeated letters in words. The following is Table 2 of the data preprocessing results.

Table 2. Result of Data Preprocessing

	Before Data Preprocessing	After Data Preprocessing
0	Saya sangat kecewa dengan tiktok seller center, akun saya tiba tiba di blokir dan saya tidak bisa mengakses nya lagi, padahal masih banyak saldo yang belum saya tarik. Saya sudah berusaha menghubungi cs tiktok tapi jawabannya tidak membantu sama sekali. Dana saya tidak bisa cair.	kecewa tiktok seller center akun blokir akses nya saldo tarik usaha hubung cs tiktok jawab bantu dana cair
1	Aplikasi bagus... Memudahkan.. Namun sayang penarikan Hasil jualannya kurang cepet ka.. tolong perbaiki.. kalo bisa 1 hari setelah paket di terima langsung bisa di tarik ka.... Saya selaku modal kecil kewalahan kalo order tinggi minim modal.. modal kecil buat di putar lagi ka..	aplikasi bagus mudah sayang tari hasil jual cepet ka tolong baik kalo paket terima langsung tarik ka modal kewalahan kalo order minim modal modal putar ka
2	Min setiap saya verifikasi pemilik selalu gagal, foto ktp selalu terdeteksi indikasi modifikasi gambar. Padahal itu ktp asli saya dan valid Cara mengatasinya bagaimana?	min verifikasi milik gagal foto ktp deteksi indikasi modifikasi gambar ktp asli valid atas

Table 2 shows a little of the results of the entire data preprocessing process. There are differences before and after data pre-processing, namely all letters are changed to lowercase, some words with affixes are lost and become basic words, punctuation is removed, and so on. The results of data preprocessing can be saved in CSV format. After that, the initial labeling process using the InSet lexicon begins. The data that has been cleaned previously will be labeled using the InSet lexicon into three types of sentiment, namely positive,

negative, and neutral. Table 3 shows the results of the labeling stage, and Table 4 shows the number of sentiment results.

Table 3. Labeling Results

	Review After Data Preprocessing	Polarity Score	Sentiment
0	min verifikasi milik gagal foto ktp deteksi indikasi modifikasi gambar ktp asli valid atas	-1	Negative
1	hati hati daftar gaes lihat nama identitas sesuai bank kali daftar aplikasi huruf tanda baca salah proses cair dana hasil jual alam lebih aplikasi rekomendasi banget jual	0	Negative
2	barang ga pic up ekspedisi kirim lambat seller langgar ga masuk akal klo kaya gitu ngapain opsi pic up mending otomatis drop of	-1	Negative

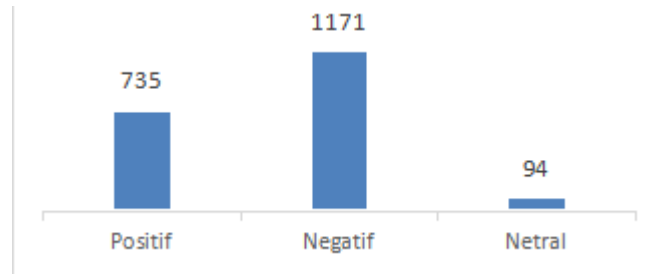


Figure 2. Number of Review Sentiment

Based on the sentiment results obtained, it is known that the most common sentiment is negative sentiment with a total of 1171, while positive sentiment is 735, and neutral sentiment is 94.

4.1.3 Classification Model and Classification Model Evaluation

At this stage, classification will be carried out using the SVM and Naive Bayes classification models. However, before entering the classification model process, a weighting process will be carried out on each word in the review first using TF-IDF and Term Presence so that when the classification model is carried out the results obtained are more accurate. Before starting the weighting process, data preprocessing will be carried out again so that the data can be given more accurate weights, then it can be continued with the first weighting process, namely the TF-IDF method assisted by cosine-similarity, then continued with the Term Presence weighting method. After the weighting process is complete, the next step is the SVM and Naive Bayes classification modeling process which will be trained first with a comparison of training data and testing data of 20:80 from the 2000 datasets used. The following are the results of the classification modeling process using TF-IDF.

Table 4. Result of TF-IDF Weighting

Review	TF-IDF Score
gak atur pilih ekspedisi pakai j t pickup nya larang atur gak kyk belah gak atur gk menye tingkat lgi kualitas ekspedisi kirim pilih ekspedisi nya	0,047419
smoga beli terima kasih	0,000000
guna tiktok shop kirim tanggal kirim kirim ya logistik nya detail nya ya kirim barang nya	0,101551

Next, an evaluation of the classification model will be carried out using a confusion matrix to determine the efficiency and performance of the two algorithms, which will later be known as the accuracy, precision, recall, and f1-score values from a total of 2000 datasets. At this stage, it will also be known what weighting method helps optimize the performance of the classification model. The following are the results of the evaluation of the classification model using a confusion matrix.

Table 5. Result of the Evaluation Process with TF-IDF

Methods	Sentiment	Accuracy	Precision	Recall	f1-score
SVM	Negative	0,82	0,80	0,95	0,87
	Negative		0,00	0,00	0,00
	Negative		0,85	0,71	0,77
Naive Bayes	Negative	0,69	0,66	0,99	0,79
	Negative		0,00	0,00	0,00
	Negative		0,94	0,32	0,48

Table 5 shows a comparison of the results of the evaluation of the classification model using TF-IDF. The SVM accuracy value of 0.82 is higher than Naive Bayes, with an accuracy value of 0.69. In the 'Negative' section, the SVM f1-score value has the highest value of 0.87, while Naive Bayes gets a value of 0.79. Both models show low results in the 'Neutral' section, with each recall value getting 0.00. Likewise, in the 'Positive' section, the SVM f1-score value is higher at 0.77 compared to Naive Bayes at 0.48.

Table 6. Result of the Evaluation Process with Term Presence

Methods	Sentiment	Accuracy	Precision	Recall	f1-score
SVM	Negative	0,78	0,81	0,86	0,84
	Negative		0,00	0,00	0,00
	Negative		0,79	0,74	0,77
Naive Bayes	Negative	0,75	0,77	0,84	0,80
	Negative		0,00	0,00	0,00
	Negative		0,72	0,71	0,73

Meanwhile, table 6 shows a comparison of the results of the classification model evaluation using Term Presence. The results of the SVM accuracy value were obtained at 0.78, higher than Naive Bayes with an accuracy value of 0.75. In the 'Negative' section, the SVM f1-score value has the highest value of 0.84, while Naive Bayes gets a value of 0.80. Both models still show low results in the 'Neutral' section, with each recall value getting 0.00. Likewise, in the 'Positive' section, the SVM f1-score value is higher at 0.77 compared to Naive Bayes, which is 0.73. Based on the evaluation results of the evaluation of the classification model on Support Vector Machine (SVM) and Naive Bayes, the accuracy value of the SVM model is higher using TF-IDF and Term Presence compared to the Naive Bayes model. This can happen due to several factors such as data imbalance where it is known that the labeling results with the Inset lexicon show that most of the sentiments obtained are negative sentiments compared to other sentiments. In addition, other factors include the variation in the length of the review text and the more complex review sentence structures, making the SVM algorithm more effective for use in this study compared to Naive Bayes. Then the TF-IDF weighting method is much better at optimizing the performance of the classification model compared to the Term Presence method. This can happen because of TF-IDF's ability to capture more complex and relevant information from reviews, compared to Term Presence with its simple binary representation.

4.2 Discussion

In the data collection phase, the researcher gathered data from the Google Play Store by scraping reviews of the TikTok Tokopedia Seller Center application using Python. The process involved setting up the necessary libraries, such as numpy and pandas, and configuring the data scraping to select 2000 relevant Indonesian language reviews. The selected data included the review content and the rating, which were then cleaned and saved in CSV format. This dataset served as the primary material for the study (Table 1). Following the data collection, the preprocessing stage was crucial to clean and prepare the dataset for analysis. The preprocessing steps included normalizing spelling, converting text to lowercase (case folding), tokenizing, filtering, stemming, and removing repeated letters in words. These processes ensured that the data was in a suitable format for further analysis. The results of the preprocessing are shown in Table 2, where the cleaned data is compared to the original reviews. The preprocessing significantly simplifies the text, making it more manageable for sentiment analysis. After preprocessing, the data was labeled using the InSet lexicon, which categorized the reviews into three sentiment types: positive, negative, and neutral. The labeling results are presented in Table 3, and the distribution of sentiments is summarized in Table 4. The majority of the reviews (1171) were labeled as negative, followed by positive (735) and neutral (94) sentiments. This imbalance in sentiment distribution is a critical factor to consider in the subsequent classification models. The classification models used in this study were Support Vector Machine (SVM) and Naive Bayes. Before classification, a weighting process was conducted using two methods: TF-IDF (Term Frequency-Inverse Document Frequency) and Term Presence. The weighting process aims to assign more accurate weights to the words in the reviews, thus improving the classification accuracy. The results of the classification model evaluation are presented in Tables 5 and 6. Using the TF-IDF method, the SVM model achieved a higher accuracy (0.82) compared to Naive Bayes (0.69). In the negative sentiment category, SVM also showed a higher f1-score (0.87) compared to Naive Bayes (0.79). For positive sentiment, SVM again outperformed Naive Bayes with an f1-score of 0.77 versus 0.48. However, both models performed poorly in the neutral sentiment category, with recall values of 0.00 for both. When using the Term Presence method, the SVM model again demonstrated superior performance with an accuracy of 0.78 compared to Naive Bayes (0.75). In the negative sentiment category, SVM had a higher f1-score (0.84) compared to Naive Bayes (0.80). For positive sentiment, SVM also showed a higher f1-score (0.77) compared to Naive Bayes (0.73). Similar to the TF-IDF method, both models performed poorly in the neutral sentiment category.

The superior performance of SVM can be attributed to its ability to handle non-linear relationships and complex data structures, which are common in text data. The imbalance in the sentiment distribution, where negative sentiments are more prevalent, also favors SVM, as it is more robust in handling imbalanced datasets [7]. The TF-IDF method, which captures the importance of words based on their frequency and rarity, further enhances the performance of the classification models, particularly in identifying relevant features in the reviews [18]. However, the poor performance in the neutral sentiment category suggests that the model may struggle with identifying reviews that do not express a clear positive or negative sentiment. This could be due to the limited number of neutral reviews in the dataset, which makes it challenging for the models to learn the characteristics of neutral sentiment. Future research could explore techniques such as oversampling the neutral reviews or using more sophisticated models to improve the classification of neutral sentiment. The SVM model with TF-IDF weighting is the most effective for classifying the sentiment of reviews from the TikTok Tokopedia Seller Center application. The high accuracy and f1-scores in the negative and positive sentiment categories indicate that this model can reliably identify user sentiments. The insights gained from this study can help TikTok Tokopedia Seller Center to understand user feedback and make necessary improvements, particularly in areas such as faster fund disbursement and more flexible account verification processes.

5. Conclusion and Recommendations

Based on the research conducted on the TikTok Tokopedia Seller Center application reviews, it is evident that the sentiment of user reviews is predominantly negative. The study results indicate that out of 2000 data points, negative reviews accounted for 1171, positive reviews for 735, and neutral reviews for only 94. Additionally, the evaluation of the classification models using the Support Vector Machine (SVM) and Naive Bayes algorithms revealed that the SVM algorithm is more effective in performing sentiment analysis compared to the Naive Bayes algorithm. The SVM algorithm achieved accuracy values of 0.81 with TF-IDF and 0.78 with Term Presence, whereas the Naive Bayes algorithm achieved accuracy values of 0.69 with TF-IDF and 0.75 with Term Presence. This discrepancy can be attributed to several factors, such as data imbalance, where the InSet lexicon predominantly labeled the sentiments as negative. Other factors include the variability in review text lengths and the complexity of review sentence structures, which make the SVM algorithm more effective for this study compared to Naive Bayes. Similarly, the TF-IDF weighting method outperformed the Term Presence method, likely due to TF-IDF's ability to capture more complex and relevant information from the reviews.

The most effective algorithm in this study is the Support Vector Machine, and the most effective weighting method is TF-IDF. The predominantly negative sentiment of user reviews indicates that users are dissatisfied with their experience using the application. Therefore, developers need to improve the application to enhance user experience and introduce features that are useful and easy to use. For future improvements, the data preprocessing stage should be more thorough, as several reviews still contain non-standard words and abbreviations that are not listed in the dictionary, leading to a less optimal dataset. Additionally, increasing the number of reviews used in the study could improve the generalization of the model and reduce potential biases. It is also recommended to consider more advanced classification models, such as deep learning, to further enhance sentiment classification accuracy, especially for neutral sentiments. Continuous evaluation of user sentiment after implementing improvements is essential to assess the impact of these changes on user satisfaction. This study aims to provide valuable insights for developers to enhance the quality of the TikTok Tokopedia Seller Center application..

References

- [1] Tian, W., Xiao, Y., & Xu, L. (2021). What XiaoHongShu users care about: An analysis of online review comments. Retrieved from <https://www.kuchuan.com>
- [2] Hidayat, W. A., & Nastiti, V. R. S. (2024). Perbandingan kinerja pre-trained IndoBERT-base dan IndoBERT-lite pada klasifikasi sentimen ulasan Tiktok Tokopedia Seller Center dengan model IndoBERT. *Jurnal Sistem Informasi (JSII)*, 11(2), 13-20. <https://doi.org/10.30656/jsii.v11i2.9168>
- [3] Tanniewa, A. M., Hamrul, H., & Sarina. (2023). Implementasi algoritma Support Vector Machine terhadap analisis sentimen penggunaan Aplikasi Tiktok Shop Seller Center. In *Seminar Nasional Teknologi Informasi dan Komputer*. <https://jurnal.ftkom.uncp.ac.id/index.php/semantik/article/view/41>

- [4] Wijaya, M., T. L. P. A. S. (2024). Penerapan algoritma Naive Bayes dan KNN dalam menganalisis sentimen Aplikasi Tiktok Shop Seller Center berdasarkan review Google Playstore. *Scientific Student Journal for Information, Technology and Science*, 5(2). <https://doi.org/10.30865/mahasiswa.v5i2.1014>
- [5] Muttaqin, M. N., & Kharisudin, I. (2021). Analisis sentimen aplikasi Gojek menggunakan Support Vector Machine dan K Nearest Neighbor. *UNNES Journal of Mathematics*, 10(2), 22-27. <https://doi.org/10.15294/ujm.v10i2.48474>
- [6] Indriyani, F. A., Fauzi, A., & Faisal, S. (2023). Analisis sentimen aplikasi TikTok menggunakan algoritma naïve bayes dan support vector machine. *TEKNOSAINS: Jurnal Sains, Teknologi dan Informatika*, 10(2), 176-184. <https://doi.org/10.37373/tekno.v10i2.419>
- [7] Atmajaya, D., Febrianti, A., & Darwis, H. (2023). Metode SVM dan Naive Bayes untuk analisis sentimen ChatGPT di Twitter. *The Indonesian Journal of Computer Science*, 12(4). <https://doi.org/10.33022/ijcs.v12i4.3341>
- [8] Saputra, R., & Hasan, F. N. (2024). Analisis sentimen terhadap program makan siang & susu gratis menggunakan algoritma Naive Bayes. *Jurnal Teknologi dan Sistem Informasi Bisnis*, 6(3), 411-419. <https://doi.org/10.47233/jteksis.v6i3.1378>
- [9] Lubis, A. Y., & Setyawan, M. Y. H. (2024). Analisis sentimen terhadap aplikasi Pospay menggunakan algoritma Support Vector Machine dan Naive Bayes. *Jurnal Teknologi dan Sistem Informasi Bisnis*, 6(3), 514-521. <https://doi.org/10.47233/jteksis.v6i3.1310>
- [10] Mustopa, A., Hermanto, A., Pratama, E. B., Hendini, A., & Risdiansyah, D. (2020). Analysis of user reviews for the Peduli Lindungi application on Google Play using the Support Vector Machine and Naive Bayes Algorithm Based on Particle Swarm Optimization. In *2020 5th International Conference on Informatics and Computing (ICIC)*. <https://doi.org/10.1109/ICIC50835.2020.9288655>
- [11] Styawati, S., Isnain, A. R., Hendrastuty, N., & Andraini, L. (2021). Comparison of Support Vector Machine and Naïve Bayes on Twitter data sentiment analysis. *Jurnal Informatika: Jurnal Pengembangan IT (JPIT)*, 6(1), 56-60. <https://doi.org/10.30591/jpit.v6i1.3245>
- [12] Ardianto, R., Rivanie, T., Alkhalifi, Y., Septia Nugraha, F., & Gata, W. (2020). Sentiment analysis on e-sports for education curriculum using Naive Bayes and Support Vector Machine. *Jurnal Ilmu Komputer dan Informasi (Journal of Computer Science and Information)*, 13(2). <https://doi.org/10.21609/jiki.v13i2.885>
- [13] Mukodimah, S., Muslihudin, M., Rohmadi Mustofa, D., Susianto, D., Bakti Nusantara, I., & Pringsewu, S. (2022). Naive Bayes classifier method analysis and Support Vector Machine (SVM) student graduation prediction. *Neuroquantology*, 20(12), 3522-3533. <https://doi.org/10.14704/NQ.2022.20.12.NQ77360>
- [14] Praneswara, A. O., & Cahyono, N. (2023). Analisis sentimen ulasan Aplikasi TikTok Shop Seller Center di Google Playstore menggunakan Algoritma Naive Bayes. *The Indonesian Journal of Computer Science*, 12(6). <https://doi.org/10.33022/ijcs.v12i6.3473>
- [15] Ainunnisa, I. R., & Sulastri, S. (2023). Analisis sentimen Aplikasi Tiktok dengan Metode Support Vector Machine (SVM), Logistic Regression dan Naive Bayes. *Jurnal Teknologi Sistem Informasi dan Aplikasi*, 6(3), 423-430. <https://doi.org/10.32493/jtsi.v6i3.31076>
- [16] Salma, A., & Silfianti, W. (2021). Sentiment analysis of user review on COVID-19 information applications using Naïve Bayes Classifier, Support Vector Machine, and K-Nearest Neighbors. *International Research Journal of Advanced Engineering and Science*, 6(4), 158-162. Retrieved from <http://irjaes.com/wp-content/uploads/2021/11/IRJAES-V6N4P61Y21.pdf>
- [17] Suryani, S., Fayyad, M. F., Savra, D. T., Kurniawan, V., & Estanto, B. H. (2023). Sentiment analysis of towards electric cars using Naive Bayes Classifier and Support Vector Machine Algorithm. *Public Research Journal of Engineering, Data Technology and Computer Science*, 1(1), 1-9. <https://doi.org/10.57152/predatecs.v1i1.814>

- [18] Panjaitan, F. (2024). Perbandingan penggunaan Tfidfvectorizer, Countvectorizer, dan Hashingvectorizer dengan optimalisasi parameter pada Machine Learning untuk analisis sentimen Pemilu 2024. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 8(4), 7413-7419.
- [19] Hasibuan, S. S., Angraini, A., Saputra, E., & Megawati, M. (2024). Sentimen analisis terhadap fitur TikTok Shop menggunakan Naive Bayes dan K-Nearest Neighbor. *Jurnal Media Informatika Budidarma*, 8(1), 303. <https://doi.org/10.30865/mib.v8i1.7238>
- [20] Fahmi, R. N., Nursyifa, N., & Primajaya, A. (2021). Analisis sentimen pengguna Twitter terhadap kasus penembakan Laskar FPI oleh Polri dengan metode Naive Bayes Classifier. *JIKO (Jurnal Informatika dan Komputer)*, 5(2), 61-66. <http://dx.doi.org/10.26798/jiko.v5i2.437>
- [21] Nugraha, P., Matheos Sarimole, F., & Studi Sistem, P. (2025). Analisis sentimen kepuasan publik terhadap masa kepemimpinan Shin Tae Yong menggunakan algoritma Naive Bayes. *Jurnal Teknologi Informasi dan Komunikasi*, 9(1), 3020. <https://doi.org/10.35870/jtik.v9i1.3020>
- [22] Gifari, O. I., Adha, M., Hendrawan, I. R., & Durrand, F. F. S. (2022). Analisis sentimen review film menggunakan TF-IDF dan Support Vector Machine. *Journal of Information Technology*, 2(1), 36-40.
- [23] Permana, H., Chrisnanto, Y. H., Ashaury Informatika, H., Jenderal Achmad Yani Cimahi Jl Terusan Jend Sudirman, U., Cimahi Sel, K., Cimahi, K., & Barat, J. (2023). Analisis sentimen terhadap bakal calon presiden 2024 dengan algoritma multinomial Naive Bayes dan Oversampling SMOTE. *Jurnal Mahasiswa Teknik Informatika*, 7(5). <https://ejournal.itn.ac.id/index.php/jati/article/view/7309>
- [24] Musfiroh, D., Khaira, U., Utomo, P. E. P., & Suratno, T. (2021). Analisis sentimen terhadap perkuliahan daring di Indonesia dari Twitter dataset menggunakan InSet Lexicon. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 1(1), 24-33. <https://doi.org/10.57152/malcom.v1i1.20>
- [25] Triyanti. (2023). Analisis sentimen pengguna aplikasi belajar bahasa menggunakan algoritma Naive Bayes. Unpublished thesis, Universitas Teknokrat Indonesia. Retrieved from <http://repository.teknokrat.ac.id/4834/1/skripsi19311170.pdf>
- [26] Salma, A., & Silfianti, W. (2021). Sentiment analysis of user reviews on COVID-19 information applications using Naive Bayes Classifier, Support Vector Machine, and K-Nearest Neighbor. *International Research Journal of Advanced Engineering and Science*, 6(4), 158-162. Retrieved from <http://irjaes.com/wp-content/uploads/2021/11/IRJAES-V6N4P61Y21.pdf>
- [27] Fazal, R. (2022). Membandingkan Support Vector Machines dan Naive Bayes pada analisis sentimen data Twitter. *Jurnal Portal Data*, 2(10).
- [28] Kamilla, A. C., Priyani, N., Priskila, R., & Pranatawijaya, V. H. (2024). Analisis sentimen film Agak Laen dengan kecerdasan buatan: Text mining metode Naive Bayes Classifier. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 8(3), 2923-2928.
- [29] Pasaribu, N. A., & Sriani. (2023). The Shopee application user reviews sentiment analysis employing Naive Bayes Algorithm. *International Journal Software Engineering and Computer Science (IJECS)*, 3(3), 194-204. <https://doi.org/10.35870/ijsecs.v3i3.1699>
- [30] Apriani, E., Oktavianalisti, F., Monasari, L. D. H., Winarni, I., & Hanif, I. F. (2024). Analisis sentimen penggunaan TikTok sebagai media pembelajaran menggunakan Algoritma Naive Bayes Classifier. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 4(3), 1160-1168. <https://doi.org/10.57152/malcom.v4i3.1482>
- [31] Wahyuni, W. (2022). Analisis sentimen terhadap opini feminisme menggunakan Metode Naive Bayes. *Jurnal Informatika Ekonomi Bisnis*, 4(4), 148-153. <https://doi.org/10.37034/infec.v4i4.162>
- [32] Bhuiyan, R. J., Akter, S., Uddin, A., Shak, M. S., Islam, M. R., Rishad, S. M. S. I., Sultana, F., & Hasan-Or-Rashid, Md. (2024). Sentiment analysis of customer feedback in the banking sector: A comparative study of machine learning models. *The American Journal of Engineering and Technology*, 6(10), 54-66. <https://doi.org/10.37547/tajet/Volume06Issue10-07>

- [33] Santoso, D. P., & Wibowo, W. (2022). Analisis sentimen ulasan Aplikasi Buzzbreak menggunakan Metode Naïve Bayes Classifier pada Situs Google Play Store. *Jurnal Sains dan Seni ITS*, 11(2), D190-D196. https://ejurnal.its.ac.id/index.php/sains_seni/article/download/72534/7063.