

Sentiment Analysis of Mobile Provider Application Reviews Using Naive Bayes Algorithm and Support Vector Machine

Tiara Sari Ningsih^{1)*}, Teguh Iman Hermanto²⁾, Imam Ma'ruf Nugroho³⁾

¹⁾²⁾³⁾ Sekolah Tinggi Teknologi Wastukencana Purwakarta, Indonesia

¹⁾tiarasarin00@wastukencana.ac.id, ²⁾teguhiman@wastukencana.ac.id, ³⁾imam.ma@wastukencana.ac.id

Submitted : Feb 6, 2024 | Accepted : Feb 12, 2024 | Published : Apr 1, 2024

Abstract: A cellular provider is the type of support that is required to carry out telecommunications activities in order for them to go smoothly. The ease and efficiency with which communication can proceed will depend on the cellular operator used. Prospective users frequently consult reviews posted by past users of the mobile provider application while deciding which mobile carrier to utilize. The Google Play Store is one place to look for reviews of programs from mobile providers. Sentiment analysis is required because potential application users find it difficult to decide which cellular carrier application to utilize due to the abundance of evaluations. The aim of this study is to examine user reviews of applications from cellular providers and compare the accuracy levels of the two algorithms—Naïve Bayes Classification (NBC) and Support Vector Machine (SVM)—that are employed. The research object is focused on the three most popular applications in Indonesia according to the goodstate website, namely Telkomsel, IM3 and XL Axiata. After testing using the Naïve Bayes Clasification method, the accuracy value obtained in the MyTelkomsel application is 75%, MyIM3 80% and MyXL 72%. While the Support Vector Machine method obtained an accuracy value in the MyTelkomsel 77%, MyIM3 80% and MyXL 76% applications.

Keywords: Cellular Provider; Google Play Store; Naïve Bayes Classification; Sentiment Analysis; Support Vector Machine

INTRODUCTION

Based on the survey results of the Indonesian Internet Service Providers Association (APJII) in June 2020, it shows that internet users in Indonesia increased by 8.9%, reaching 196 million, and will continue to grow, develop, and expand network infrastructure in Indonesia (Hakim et al., 2021). The growing number of internet users is a sizable market opportunity for internet service provider companies (Hakim et al., 2021). By creating good internet services and meeting the needs of the community, various internet service provider companies are competing to offer the best services for the community (Garcia & Berton, 2021). The use of the internet, which ranks first, is the most frequently used service to communicate (Ananda & Pristyanto, 2021). In carrying out these telecommunication activities, support is needed so that activities can run smoothly, namely in the form of a cellular provider. Choosing the right cellular provider will affect the extent to which communication activities can run smoothly and efficiently.

According to goodstats data for 2023, some of the cellular provider companies that have the largest market in Indonesia are Telkomsel, Indosat, and XL Axiata. In the first quarter of 2023, Telkomsel, the largest cellular provider in Indonesia, recorded 156.8 million subscribers, followed by Indosat with 98.5 million subscribers and XL Axiata with 57.9 million subscribers. Choosing the right cellular provider will affect the extent to which communication activities can run smoothly and efficiently. To choose a cellular provider to use, prospective users often rely on reviews left by previous users of cellular provider applications, and one source of information for finding cellular provider application reviews is the Google Play Store.

Google Play Store is a digital content provider service that offers a variety of online product stores featuring games, apps, books, music, and movies in various categories. This service is owned by Google (Herlinawati et al., 2020). In the Google Play Store, there are several features, one of which is the review feature for users of available applications or services. However, the number of reviews that are too many and diverse makes it difficult for prospective application users to make decisions about the applications or services they want to use. In this case,

*name of corresponding author



This is anCreative Commons License This work is licensed under a Creative Commons Attribution-NonCommercial 4.0 International License.

sentiment analysis has become an important approach to understanding user opinions and perceptions through reviews. Sentiment analysis is the activity of analyzing a person's opinion, attitude, or emotion about a particular product, topic, or issue so that it can be known that it includes positive and negative sentiments (Rejeki & Vina, 2023). Using sentiment analysis, we can classify these reviews into positive, negative, or neutral categories, thus providing valuable insights into how users perceive and rate an app or service.

Based on the background and previous research, in this study researchers will use two algorithms, namely Naïve Bayes Clasification (NBC) and Support Vector Machine (SVM), in analyzing cellular provider applications on the Google Play Store. The Naïve Bayes Clasification method was first put forward by a scientist from England, Thomas Bayes. Naïve Bayes Clasification is a method whose classification is done using probability and statistics (Maulidah et al., 2021). Support Vector Machine (SVM) is a classification method based on the concept of hyperplane separation (Rivanie et al., 2021). By using these two methods, it is expected to help users or service providers evaluate the quality of internet services provided and see a comparison of the accuracy level between the two algorithms used.

LITERATURE REVIEW

The study conducted by Sulton et al. previously examined sentiment analysis of myindihome user reviews using support vector machine and naïve bayes classification methods. The dataset samples were taken from the Google Play store user reviews. The average total accuracy of the SVM method was 86.54% higher than that of the NBC method, which had an average total accuracy of 84.69%. (Hakim et al., 2021). Both algorithms are able to predict values well, but the Support Vector Machine algorithm is better at it.

Further research, conducted by Fadhilah and Yoga, entitled Twitter User Sentiment Analysis of Internet Provider Services Based on the test results, the SVM algorithm using the linear kernel and the RBF kernel has good evaluation performance results in terms of accuracy, precision, and recall, which are relatively the same. So it can be said that the SVM algorithm uses both RBF and linear kernels (Ananda & Pristyanto, 2021).

Research conducted by Aan et al. With the title Comparison of Naïve Bayes and Support Vector Machine for Indihome Customer Review Classification shows that the application of the Linear Kernel Support Vector Machine algorithm is better than the Naïve Bayes Classifier algorithm with an accuracy value of 82.11%, precision 76.44%, recall 88.01%, and AUC value 0.909 (Rohanah et al., 2021). In Indihome Customer Review Classification, the Support Vector Machine algorithm performs better, it can be concluded.

The next research conducted by Dimas et al., entitled Sentiment Analysis of the MyXL Application Using the Support Vector Machine Method Based on User Reviews on the Google Play Store, shows that the linear kernel is better than the use of the RBF kernel, with an average accuracy value of 88%, prediction of 88%, recall of 88%, and f-measure of 88% (Diandra Audiansyah et al., 2022). It can be concluded that the Support Vector Machine method performs well in classification in this study because the accuracy value has a good value.

The next research conducted by Said et al., with the title Text Mining Classification of Public Opinion Towards Providers in Indonesia, aims to be able to identify sentiments that come from the public opinion of provider card users, which is a type of unstructured data. This study performs classification using the Naïve Bayes Classifier (NBC) algorithm using 10 experiments. The results show that the best accuracy value in this study is 68.85% (Rizaldi et al., 2021).

The aforementioned research has been conducted with great skill; in certain investigations, the average with a higher accuracy value is the Support Vector Machine. The authors intend to compare two algorithms the Support Vector Machine and Naive Bayes Classification algorithms for cellular carrier applications in order to carry out additional research in light of their findings. The Google Play store user reviews of applications from cellular providers give the dataset that will be used.

METHOD

The following are the stages of research that will be carried out in this study: The stages of this research include literature study, data collection, text preprocessing, feature extraction, classification, evaluation, and visualization.

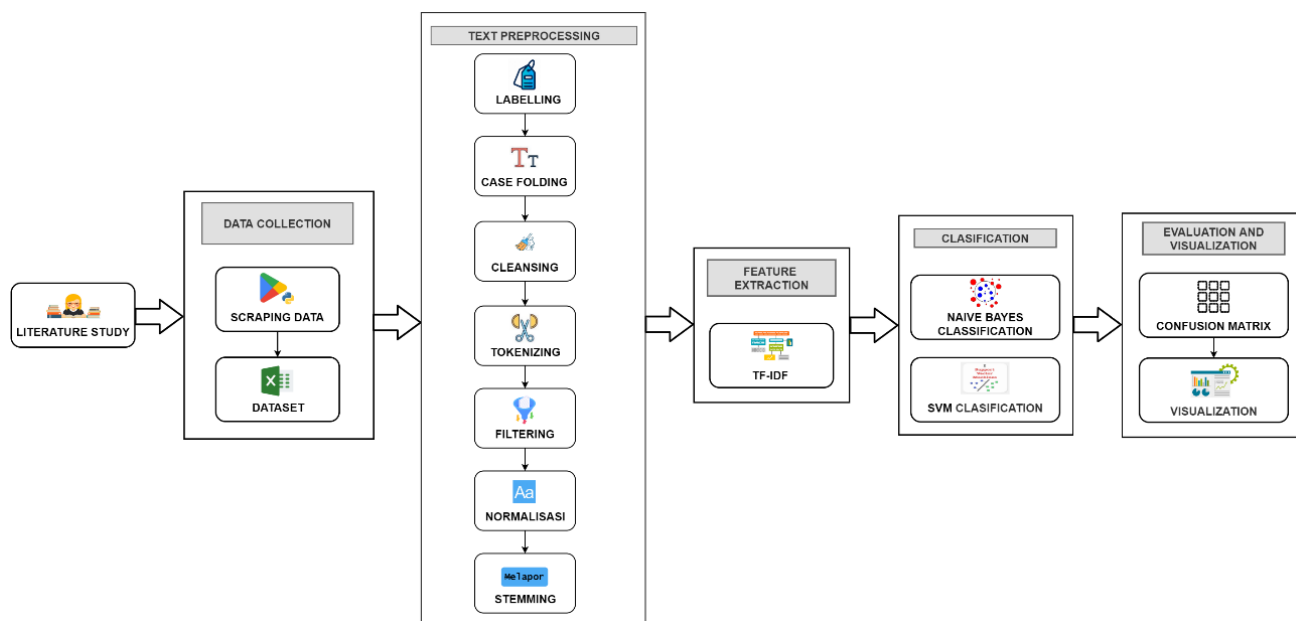


fig. 1 Phases of Research

Literature Study

The literature study stage involves the process of collecting data from various sources, such as journals and articles, which will be used as reference materials that support research. In this study, the authors sought data or literature materials from journals or articles related to sentiment analysis, classification, mobile providers, Naive Bayes classification, and support vector machines.

Data Collection

Data was collected from the top 3 mobile service provider applications in Indonesia based on the GoodState website. These applications include MyTelkomsel, MyXL, and MyIM3. From the scraping results, we managed to collect as much as 7,000 data points from each application for further analysis or research purposes.

Text Preprocessing

Text preprocessing is the initial stage of text mining preparation, where text data with noise will be reduced so that it can be processed further. Text preprocessing is efficient for reducing high-dimensional features and noise so that it will reduce data processing time and increase accuracy (Rizaldi et al., 2021). The first stage in text preprocessing is the labeling method used to determine whether a text has a positive, negative, or neutral opinion (sentiment) (Santoso et al., 2022). The second stage, namely case folding, functions to convert uppercase letters into lowercase letters (Rohanah et al., 2021). The third cleaning stage is to clean the document from components that have no relationship with the information in the document, such as characters or symbols, numbers, emoticons, and URL links (Putra et al., 2020). The fourth stage, tokenization, is done to break down text that was originally a sentence into words (Ardiani et al., 2020). The fifth stage, filtering, is done to eliminate words that have no meaning or have no effect on classification results (Rohanah et al., 2021). Next, the sixth stage, namely normalization, is carried out to change words that are not standardized into standard words and are ready to be processed (Putra et al., 2020). The last stage is stemming the process of returning various word formations to basic word formations by removing affixes (Putra et al., 2020).

Feature Extraction

The feature extraction stage attempts to prepare the data for categorization. In this research, words are weighted in relation to documents using TF-IDF (Term Frequency-Inverse Document Frequency). The TF-IDF approach combines the Inverse Document Frequency (IDF) and Term Frequency (TF) approaches. The number of words or terms in a text is known as term frequency (TF), and the frequency of occurrence of those words or terms across documents is known as inverse document frequency (IDF) (Safryda Putri & Ridwan, 2023).

Classification

Classification is a step that will be applied in testing datasets that have completed several previous stages. The stages that have been completed are data scraping, text preprocessing, and feature extraction. Then the next stage

*name of corresponding author



This is anCreative Commons License This work is licensed under a Creative Commons Attribution-NonCommercial 4.0 International License.

is to classify and test accuracy. In the classification process, the data is divided into training data and testing data. Training data is used to develop classification models with the Naïve Bayes Classifier and Support Vector Machine methods. While testing data is used to obtain classification results.

Evaluation dan Visualization

In the evaluation stage, the classification results of the review data are compared using a confusion matrix. Confusion Matrix A confusion matrix is a process used in analyzing the accuracy of classification models made to identify data from different classes. By measuring the level of accuracy, the performance of a classification model can be known (Afrillia et al., 2022). In addition, to make the data easier to understand by readers, at this stage, data visualization is carried out using word clouds to display words based on their frequency of occurrence, both in uploads with positive and negative sentiments. A word cloud is a graphical representation of a document created by plotting words that often appear in documents in two-dimensional space and collecting them like a cloud (Agustina et al., 2021).

RESULT

Data Collection

Data scraping was conducted from November 8 to 15, 2023, on the Google Play Store using the Python programming language version 3. Data was taken from the three most popular mobile service applications in Indonesia, according to the GoodState website. These applications include MyTelkomsel, MyIM3, and MyXL. From the scraping results, we managed to collect 7,000 data points from each application.

Text Preprocessing

In the text preprocessing stage, it is divided into several stages, starting with labeling, case folding, cleaning, tokenizing, filtering, normalization, and stemming. The results of the text preprocessing stage are as follows:

1. Labeling: this stage is done to classify reviews into positive and negative reviews. Labeling is done based on rating; if the rating is 1 or 2, then the review goes into the negative class, and if the rating is 4 or 5, then the review goes into the positive class. Rating 3 is removed because it is ambiguous. The labeling implementation process can be seen in Table 1.

Table 1. Labeling Process

Review	Rating	Label
SETELAH UPDATE KENAPA MAKIN LEMOT?! mau buka mytelkomsel pun ga bisa. Tolonglah ubah jadi lebih simpel lagi tampilannya agar tidak berat	1	Negatif
Simpel, praktis dan gak ribet. Aplikasi myIM3 semakin jauh lebih baik dan menyenangkan	4	Positif
aplikasi yg sangat membantu buat saya pengguna XL dan selalu dapat bonus paket data menarik	5	Positif

2. Case folding: is performed to convert all text into lowercase. The implementation process at the case folding stage can be seen in Table 2.

Table 2. Case Folding Process

Before Case Folding	After Case Foling
SETELAH UPDATE KENAPA MAKIN LEMOT?! mau buka mytelkomsel pun ga bisa. Tolonglah ubah jadi lebih simpel lagi tampilannya agar tidak berat	setelah update kenapa makin lemot?! mau buka mytelkomsel pun ga bisa. tolonglah ubah jadi lebih simpel lagi tampilannya agar tidak berat

3. Cleaning: this stage is done to clean the data from any kind of noise or interference, such as special characters, URLs, or irrelevant punctuation. The cleaning implementation process can be seen in Table 3.

Table 3. Cleaning Process

Before Cleaning	After Cleaning
setelah update kenapa makin lemot?! mau buka mytelkomsel pun ga bisa. tolonglah ubah jadi lebih simpel lagi tampilannya agar tidak berat	setelah update kenapa makin lemot mau buka mytelkomsel pun ga bisa tolonglah ubah jadi lebih simpel lagi tampilannya agar tidak berat

4. Tokenizing: this stage is done to separate the sentence into words. The tokenizing implementation process can be seen in Table 4.

*name of corresponding author



This is anCreative Commons License This work is licensed under a Creative Commons Attribution-NonCommercial 4.0 International License.

Table 4. Tokenizing Process

Before Tokenizing	After Tokenizing
setelah update kenapa makin lemot mau buka mytelkomsel pun ga bisa tolonglah ubah jadi lebih simpel lagi tampilannya agar tidak berat	['setelah', 'update', 'kenapa', 'makin', 'lemot', 'mau', 'buka', 'mytelkomsel', 'pun', 'ga', 'bisa', 'tolonglah', 'ubah', 'jadi', 'lebih', 'simpel', 'lagi', 'tampilannya', 'agar', 'tidak', 'berat']

5. Filtering is performed to filter and remove irrelevant or common words, such as stopwords. The filtering implementation process can be seen in Table 5.

Table 5. Filtering Process

Before Filtering	After Filtering
['setelah', 'update', 'kenapa', 'makin', 'lemot', 'mau', 'buka', 'mytelkomsel', 'pun', 'ga', 'bisa', 'tolonglah', 'ubah', 'jadi', 'lebih', 'simpel', 'lagi', 'tampilannya', 'agar', 'tidak', 'berat']	['update', 'lemot', 'buka', 'mytelkomsel', 'tolonglah', 'ubah', 'simpel', 'tampilannya', 'berat']

6. Normalization: this stage is carried out to convert words into their basic form. The implementation process at the normalization stage can be seen in Table 6.

Table 6. Normalization Process

Before Normalization	After Normalization
['update', 'lemot', 'buka', 'mytelkomsel', 'tolonglah', 'ubah', 'simpel', 'tampilannya', 'berat']	['update', 'lemot', 'buka', 'mytelkomsel', 'tolonglah', 'ubah', 'simpel', 'tampilannya', 'berat']

7. Stemming is done to change the attached words into basic words. The stemming implementation process can be seen in Table 7.

Table 7. Stemming Process

Before Stemming	After Stemming
['update', 'lemot', 'buka', 'mytelkomsel', 'tolonglah', 'ubah', 'simpel', 'tampilannya', 'berat']	['update', 'lot', 'buka', 'mytelkomsel', 'tolong', 'ubah', 'simpel', 'tampil', 'berat']

Feature Extraction

At this stage, the word-weighting process is carried out with the help of TF-IDF. Word weighting is considered important; if a word appears more often in a document, then the contribution value will be greater, but if the word often appears in several documents, it will have a smaller contribution. The following are the stages for word weighting using TF-IDF:

1. Take review samples from the text preprocessing results. Review samples for the TF-IDF process can be seen in Table 8.

Table 8. Sample Review

Sample Review	
D1	MyIM3 ternyata sangat bagus
D2	bagus aplikasi nya saya sering dapat kuota gratis
D3	aplikasi yang mudah dan praktis untuk membeli kuota atau pulsa.

2. The next step is to find the frequency of occurrence (TF) of a word in the document sample. The process of the calculation can be seen in Table 9.

Table 9. Find Term Frequency

Term	(D1)	(D2)	(D3)
MyIM3	1	0	0
bagus	1	1	0
aplikasi	0	1	1
kuota	0	1	1
gratis	0	1	0
mudah	0	0	1
praktis	0	0	1
pulsa	0	0	1

3. Next, determine document frequency (DF), which is the number of terms (T) that appear in all documents. The process of calculation can be seen in Table 10.

Table 10. Determining Document Frequency (DF)

T	(D1)	(D2)	(D3)	DF
MyIM3	1	0	0	1
bagus	1	1	0	2
aplikasi	0	1	1	2
kuota	0	1	1	2
gratis	0	1	0	1
mudah	0	0	1	1
praktis	0	0	1	1
pulsa	0	0	1	1

4. Then calculate the Inverse Document Frequency (IDF) value by calculating the log value of the result of the number of documents divided by the Document Frequency (DF) value. The process of the calculation can be seen in Table 11 as follows:

Table 11. Calculating IDF Value

T	DF	N/DF	IDF
MyIM3	1	3	$\text{Log } 3 = 0,477$
bagus	2	1,5	$\text{Log } 1,5 = 0,176$
aplikasi	2	1,5	$\text{Log } 1,5 = 0,176$
kuota	2	1,5	$\text{Log } 1,5 = 0,176$
gratis	1	3	$\text{Log } 3 = 0,477$
mudah	1	3	$\text{Log } 3 = 0,477$
praktis	1	3	$\text{Log } 3 = 0,477$
pulsa	1	3	$\text{Log } 3 = 0,477$

5. The next step is to calculate the TF-IDF value. The process can be seen in Table 12.

Table 12. Calculating TF-IDF Value
TF*IDF

(D1)	(D2)	(D3)
0,477	0	0
0,176	0,176	0
0	0,176	0,176
0	0,176	0,176
0	0,477	0
0	0	0,477
0	0	0,477
0	0	0,477

6. Finally, sum up all the TF-IDF values in the document. The results can be seen in Table 13.

Table 13. Final Result TF-IDF

D1	D2	D3
0,653	1,005	1,783

Classification

Two algorithms are used in the classification process: Support Vector Machine (SVM) and Naïve Bayes Classification (NBC). An illustration of the categorization model is shown below.

```
from sklearn.svm import SVC
classifier = SVC(kernel = 'linear', random_state = 0)
classifier.fit(X_train_vect, y_train)
svm_pred = classifier.predict(X_test_vect)
# Classification report
print(classification_report(y_test, svm_pred))

# Confusion Matrix
class_label = ["negative", "positive"]
df_cm = pd.DataFrame(confusion_matrix(y_test, svm_pred), index=class_label, columns=class_label)
sns.heatmap(df_cm, annot=True, fmt='d')
plt.title("Confusion Matrix")
plt.xlabel("Predicted Label")
plt.ylabel("True Label")
plt.show()
```

fig. 2 Naive Bayes Classification Method

```
from sklearn.naive_bayes import MultinomialNB
classifier = MultinomialNB()
classifier.fit(X_train_vect, y_train)
nb_pred = classifier.predict(X_test_vect)

# Classification report
print(classification_report(y_test, nb_pred))

# Confusion Matrix
class_label = ["negative", "positive"]
df_cm = pd.DataFrame(confusion_matrix(y_test, nb_pred), index=class_label, columns=class_label)
sns.heatmap(df_cm, annot=True, fmt='d')
plt.title("Confusion Matrix")
plt.xlabel("Predicted Label")
plt.ylabel("True Label")
plt.show()
```

fig. 3 Support Vector Machine Method

Evaluation and visualization

The results of the classification process will then be evaluated and visualized. The following is an explanation of the evaluation and classification process.

1. Confusion Matrix

In the evaluation stage, the results of data classification by Naïve Bayes Classifier (NBC) and Support Vector Machine (SVM) algorithms are compared using the Confusion Matrix.

a) MyTelkomsel Application

The confusion matrix evaluation stage using Jupyter Notebook with NBC and SVM algorithms using data samples from the MyTelkomsel application produces values for the NBC classification

model, namely accuracy (74%), recall (66%, precision (71%, and F1-Score (69%). As for the SVM classification model, the accuracy is 77%, recall is 75%, precision is 72%, and F1-Score is 73%.

Table 14. MyTelkomsel Data Testing Results

	NBC	SVM
Accuracy	74%	77%
Recall	66%	75%
Precision	71%	72%
F1-Score	69%	73%

Figure 4 and 5 shows the confusion matrix, which is 2 x 2, presenting the positive classification class and the negative classification. As for the results of manual confusion matrix calculations described in tables 15 and 16, And the results of the highest accuracy value obtained by the SVM algorithm are 77%.

Table 15. Accuracy Value of NBC MyTelkomsel

$$\text{Accuracy} = \frac{tp + tn}{tp + fp + fn + tn} \times 100\% = \frac{655 + 389}{655 + 200 + 156 + 389} \times 100\% = \frac{1044}{1400} \times 100\% = 74,57$$

Table 16. Accuracy Value of SVM MyTelkomsel

$$\text{Accuracy} = \frac{tp + tn}{tp + fp + fn + tn} \times 100\% = \frac{639 + 439}{639 + 150 + 172 + 439} \times 100\% = \frac{1078}{1400} \times 100\% = 77$$

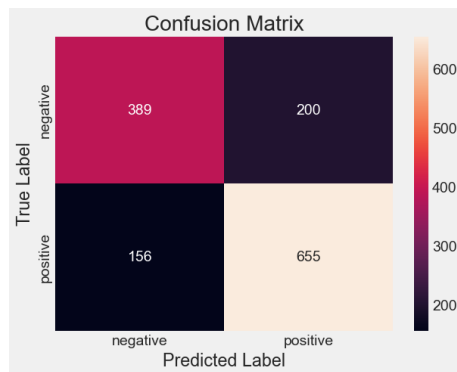


fig. 4 Confusion Matrix MyTelkomsel NBC

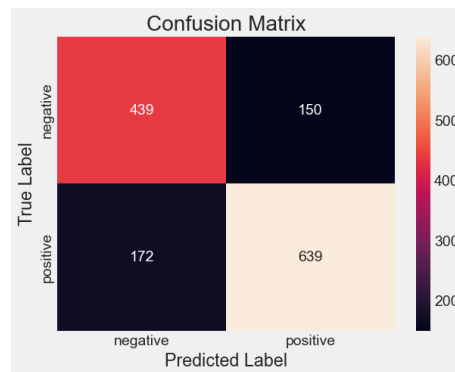


fig. 5 Confusion Matrix MyTelkomsel SVM

b) MyIM3 Application

The confusion matrix evaluation stage using jupyter notebook with NBC and SVM algorithms using data samples from the MyIM3 application produces values for the NBC classification model, namely accuracy 79%, Recall 72%, Precision 78%, and F1-Score 75%. As for the SVM classification model, namely accuracy 79%, Recall 75%, Precision 76%, and F1-Score 75%.

Table 17. MyIM3 Data Testing Results

	NBC	SVM
Accuracy	79%	79%
Recall	72%	75%
Precision	78%	76%
F1-Score	75%	75%

Figure 6 and 7 shows the confusion matrix with a size of 2 x 2, which represents the positive classification class and the negative classification. As for the results of manual confusion matrix calculations explained in tables 18 and 19, And the results of the accuracy value on both algorithms get the same value of 79%.

Table 18. Accuracy Value of NBC MyIM3

$$\text{Accuracy} = \frac{tp + tn}{tp + fp + fn + tn} \times 100\% = \frac{694 + 424}{694 + 165 + 117 + 424} \times 100\% = \frac{1118}{1400} \times 100\% = 79,85$$

Table 19. Accuracy Value of SVM MyIM3

$$\text{Accuracy} = \frac{tp + tn}{tp + fp + fn + tn} \times 100\% = \frac{676 + 439}{676 + 150 + 135 + 439} \times 100\% = \frac{1115}{1400} \times 100\% = 79,64$$

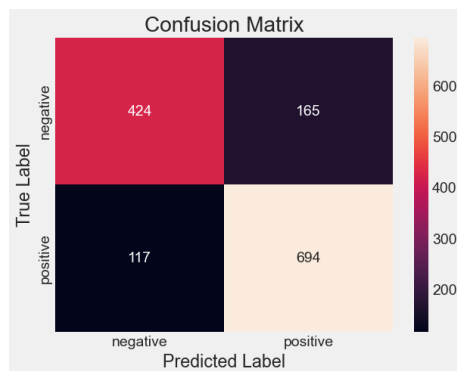


fig. 6 Confusion Matrix MyIM3 NBC

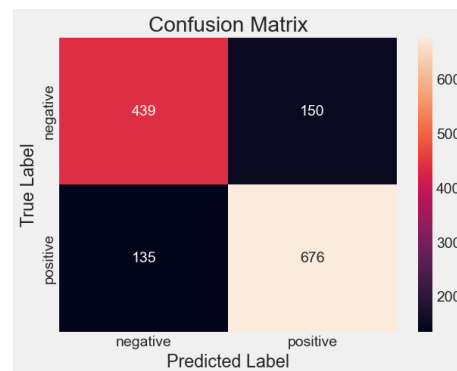


fig. 7 Confusion Matrix MyIM3 SVM

c) MyXL Application

The confusion matrix evaluation stage using jupyter notebook with NBC and SVM algorithms using data samples from the MyXL application produces values for the NBC classification model, namely accuracy 72%, Recall 73%, Precision 65%, and F1-Score 69%. Meanwhile, the SVM classification model is 75% accuracy, 75% Recall, 70% Precision, and 72% F1-Score.

Table 20. MyXL Data Testing Results

	NBC	SVM
Accuracy	72%	75%
Recall	73%	75%
Precision	65%	70%
F1-Score	69%	72%

Figure 8 and 9 shows the confusion matrix, which is 2 x 2, presenting the positive classification and negative classification classes. The results of the manual confusion matrix calculation are explained in tables 21 and 22. And the results of the highest accuracy value obtained by the SVM algorithm are 75%.

Table 21. Accuracy Value of NBC MyXL

$$\text{Accuracy} = \frac{tp + tn}{tp + fp + fn + tn} \times 100\% = \frac{579 + 430}{579 + 158 + 233 + 430} \times 100\% = \frac{1009}{1400} \times 100\% = 72,07$$

Table 22. Accuracy Value of SVM MyXL

$$\text{Accuracy} = \frac{tp + tn}{tp + fp + fn + tn} \times 100\% = \frac{624 + 439}{624 + 149 + 188 + 439} \times 100\% = \frac{1063}{1400} \times 100\% = 75,92$$

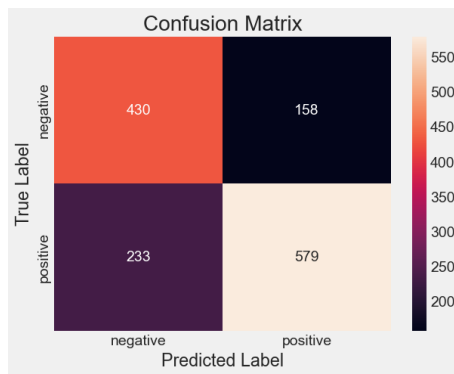


fig. 8 Confusion Matrix MyXL NBC

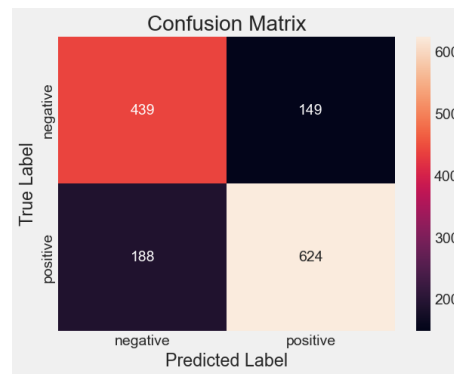


fig. 9 Confusion Matrix MyXL SVM

Based on the results of the confusion matrix calculation above, it can be concluded in Figure 10 that the Support Vector Machine (SVM) algorithm has a good average value for two of the three applications. So it can be proven that the Support Vector Machine (SVM) algorithm is suitable for this research, namely sentiment analysis on mobile provider application reviews on the Google Play Store.

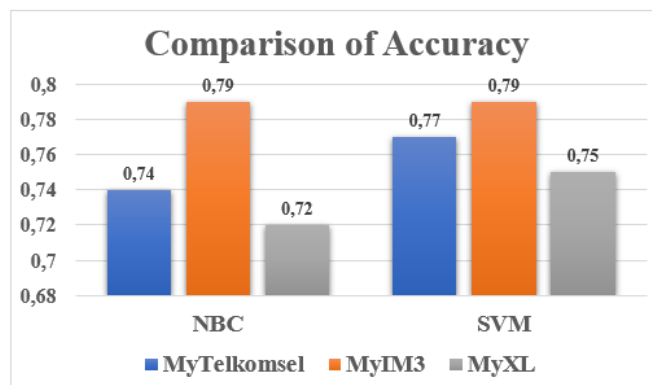


fig. 10 Comparison of NBC and SVM Accuracy Values

2. Visualization

In Figure 11, the sentiment comparison of the MyTelkomsel application shows more positive sentiment, namely, 57.1% of the 4000 reviews are positive, and negative sentiment, as much as 42.9% of the 3000 reviews are negative. In Figure 12, the sentiment comparison of the MyIM3 application shows more positive sentiment, namely, 57.1% of the 4000 reviews are positive, and negative sentiment, as much as 42.9% of the 3000 reviews are negative. In Figure 13, the sentiment comparison of the MyXL application has more positive sentiment, which is 57.1%, or 4000 reviews, positive, and negative sentiment, which is 42.9%, or 3000 reviews, negative.

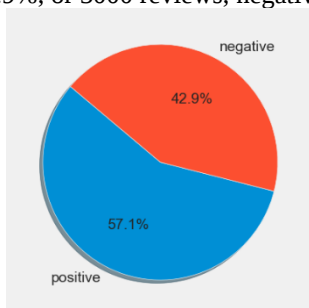


fig. 11 Sentiment MyTelkomsel

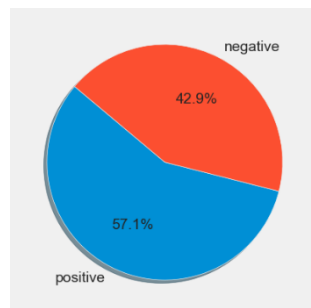


fig. 12 Sentiment MyIM3

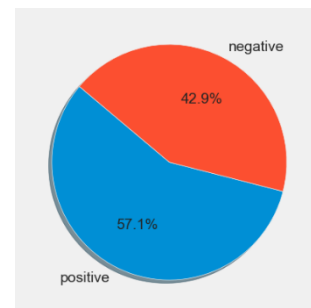


fig. 13 Sentiment MyXL

To make the comparison easier to understand, the author presented the data with WordCloud for positive and negative sentiment towards the three applications. In Figure 14 for wordcloud positive reviews, MyTelkomsel, one of the positive words that often appear, is easy, and Figure 15 for wordcloud negative

[illegible]

fig. 15 Wordcloud Negative MyTelkomsel



fig. 17 Wordcloud Negative MyIM3



fig. 19 Wordcloud Negative MyXL

- Menggunakan Algoritma Support Vector Machine. *MATRIK : Jurnal Manajemen, Teknik Informatika Dan Rekayasa Komputer*, 20(2), 407–416. <https://doi.org/10.30812/matrik.v20i2.1130>
- Ardiani, L., Sujaini, H., & Tursina, T. (2020). Implementasi Sentiment Analysis Tanggapan Masyarakat Terhadap Pembangunan di Kota Pontianak. *Jurnal Sistem Dan Teknologi Informasi (Justin)*, 8(2), 183. <https://doi.org/10.26418/justin.v8i2.36776>
- Diandra Audiansyah, D., Eka Ratnawati, D., & Trias Hanggara, B. (2022). Analisis Sentimen Aplikasi MyXL menggunakan Metode Support Vector Machine berdasarkan Ulasan Pengguna di Google Play Store. *Jurnal Pengembangan Teknologi Informasi Dan Ilmu Komputer*, 6(8), 3987–3994. <http://j-ptiik.ub.ac.id>
- Garcia, K., & Berton, L. (2021). Topic detection and sentiment analysis in Twitter content related to COVID-19 from Brazil and the USA. *Applied Soft Computing*, 101, 107057. <https://doi.org/10.1016/j.asoc.2020.107057>
- Hakim, N. S., Putra, A. J., & Khasanah, A. U. (2021). Sentiment analysis on myindihome user reviews using support vector machine and naïve bayes classifier method. *International Journal of Industrial Optimization*, 2(2), 141. <https://doi.org/10.12928/ijio.v2i2.4449>
- Herlinawati, N., Yuliani, Y., Faizah, S., Gata, W., & Samudi, S. (2020). Analisis Sentimen Zoom Cloud Meetings di Play Store Menggunakan Naïve Bayes dan Support Vector Machine. *CESS (Journal of Computer Engineering, System and Science)*, 5(2), 293. <https://doi.org/10.24114/cess.v5i2.18186>
- Maulidah, N., Supriyadi, R., Utami, D. Y., Hasan, F. N., Fauzi, A., & Christian, A. (2021). Prediksi Penyakit Diabetes Melitus Menggunakan Metode Support Vector Machine dan Naive Bayes. *Indonesian Journal on Software Engineering (IJSE)*, 7(1), 63–68. <https://doi.org/10.31294/ijse.v7i1.10279>
- Putra, M. W. A., Susanti, Erlin, & Herwin. (2020). Analisis Sentimen Dompot Elektronik Pada Media Sosial Twitter Menggunakan Naïve Bayes Classifier. *IT Journal Research and Development (ITJRD)*, 5(1), 72–86.
- Rejeki, F., & Vina, A. (2023). Analisa Sentimen Mengenai Kenaikan Harga Bbm Menggunakan Metode Naïve Bayes Dan Support Vector Machine. *JSai (Journal Scientific and Applied Informatics)*, 6(1), 1–10. <https://doi.org/10.36085/jsai.v6i1.4628>
- Rivanie, T., Pebrianto, R., Hidayat, T., Bayhaqy, A., Gata, W., & Novitasari, H. B. (2021). Analisis Sentimen Terhadap Kinerja Menteri Kesehatan Indonesia Selama Pandemi Covid-19. *Jurnal Informatika*, 21(1), 1–13. <https://doi.org/10.30873/ji.v21i1.2864>
- Rizaldi, S. T., Khairi, A. Al, & Mustakim. (2021). Text Mining Classification Opini Publik Terhadap Provider di Indonesia. *Univesitas Islam Negeri Sultan Syarif Kasim Riau*, 1(3), 2579–5406.
- Rohanah, A., Rianti, D. L., & Sari, B. N. (2021). Perbandingan Naïve Bayes dan Support Vector Machine untuk Klasifikasi Ulasan Pelanggan Indihome. *STRING (Satuan Tulisan Riset Dan Inovasi Teknologi)*, 6(1), 23. <https://doi.org/10.30998/string.v6i1.9232>
- Safryda Putri, D., & Ridwan, T. (2023). Analisis Sentimen Ulasan Aplikasi Pospay Dengan Algoritma Support Vector Machine. *Jurnal Ilmiah Informatika*, 11(01), 32–40. <https://doi.org/10.33884/jif.v11i01.6611>