

OPTIMIZATION OF SENTIMENT ANALYSIS ALGORITHM ON YOUTUBE MUSIC APPLICATION WITH COMPARISON OF NAIVE BAYES AND SUPPORT VECTOR MACHINE

Shafa Khairunnisa Azzahra¹⁾, Abdul Barir Hakim²⁾

¹²System Informations Department, Universitas Ary Ginanjar

email: abdulbarir.h@uag.ac.id

ABSTRACT

This study conducted a sentiment analysis of a new application created by YouTube, namely the YouTube Music application. Sentiment analysis is used to understand user opinions about a service or product. In conducting sentiment analysis, researchers compared the Naïve Bayes and Support Vector Machine algorithm methods, which were then optimized using Particle Swarm Optimization for both algorithms. Research data was collected through web scraping from the Google Play Store which contained reviews from YouTube Music application users. Each review was labeled with positive and negative sentiment based on the context and emotions contained therein. The results showed that the Naïve Bayes algorithm model had a higher accuracy rate of 87.17% and the Support Vector Machine had an accuracy of 91.80%. The Particle Swarm Optimization method successfully optimized the evaluation process using the Confusion Matrix, with the initial Naïve Bayes accuracy of 87.17% to 91.80%, the initial accuracy of the Support Vector Machine of 85.20% to 85.51%. The results of sentiment analysis using Naïve Bayes with Particle Swarm Optimization on the YouTube Music application show that users responded positively with a total of 5,967 positive sentiments and 1,657 negative sentiments.

Keywords: Sentiment Analysis, YouTube Music, Naïve Bayes, Support Vector Machine, Particle Swarm Optimization

INTRODUCTION

YouTube has become the fastest growing and most popular social media platform worldwide. Through this site, users are presented with various information in a reliable and easily accessible video format. YouTube gives its users the freedom to search for and watch videos directly, as well as contribute by uploading videos to YouTube servers and sharing them all over the world. As a leading video sharing site, YouTube allows users to upload, watch, and share video clips free of charge [1]. In fact, media and television companies such as CNN, CNBC, TRANSTV, and KOMPAS, as well as various organizations and institutions, have utilized YouTube Channels to share their content.

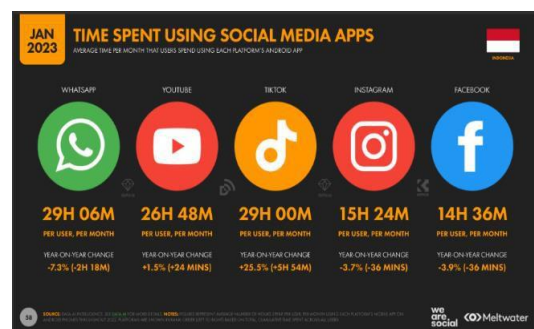


Figure 1 Data on Social Media Usage Activities in Indonesia

Figure 1 displays data on social media user activity in Indonesia in January 2023, which reveals that users in Indonesia spend a significant amount of time on YouTube, with an average of 26 hours and 48 minutes per month.

On November 7, 2019, YouTube officially launched its music streaming service, YouTube Music, simultaneously in seven Asian countries, including Indonesia. This service is not only

available in the form of an application, but can also be accessed via a web player (CNN Indonesia, 2019). YouTube Music offers a variety of licensed songs and albums, complete with playlist and radio features. Uniquely, unlike other music applications, YouTube Music also provides video content in addition to audio, so users can enjoy a richer music experience [2].

YouTube Music is a music streaming service brought by YouTube, formerly known as YouTube Red. The platform allows users to browse and watch music videos and songs based on playlists, genres, and recommendations. In addition to the free version, users can get an ad-free experience by subscribing to YouTube Music Premium. In addition to being ad-free, this subscription also gives access to features like background audio playback and offline playback by downloading songs. The standard subscription price for YouTube Music Premium is Rp. 59,000 per month [13].

The presence of an application often triggers various responses from the public, ranging from satisfaction to disappointment with its features. The impact of these responses can be significant for YouTube's business, namely affecting brand image, user satisfaction levels and in the long term, loyalty and use of the application. Each application has unique advantages and disadvantages. One way for users to express their opinions, both positive and negative, is through reviews on the Google Play Store. These reviews can be analyzed using text mining process. Text mining is a data analysis technique that allows to explore and extract valuable information from text data sources, especially documents [8]. Based in this the reviews on the Google Play Store is some kind of unstructured document that can be analyzed using text mining. One of the main application of text mining techniques

is sentiment analysis, which can extract public opinion, in the form of reviews in Google Play Store about Youtube Music Application, to measure the sentiment contained in the reviews with the aim of evaluating the extend to which the application is accepted by the public.

Sentiment analysis is a method used to collect information from various platforms, including the Google Play Store, with the aim of predicting user moods, analyzing and describing their emotions automatically. This method aims to identify patterns of emotions, both positive and negative, expressed by users [3]. Sentiment analysis is an automated process of processing text data with the aim of providing accurate information. The main purpose of sentiment analysis is to extract public opinion about a topic or issue contained in unstructured text data. Currently, many people use sentiment analysis in their daily activities, because people's views play an important role in making decisions about products or services. [7]. So Sentiment analysis is selected in this study. Additionally, beside Sentiment Analysis, this study uses a feature selection method, namely Particle Swarm Optimization, to improve the accuracy of Naïve Bayes classification. In a study by [4] it was revealed that the Naïve Bayes method optimized with Particle Swarm Optimization (PSO) achieved the best accuracy of 78.33%. These results show an increase in accuracy of 4.66% compared to Naïve Bayes without PSO optimization which only achieved an accuracy of 73.67% using 1000 datasets. Feature selection optimization aims to improve classification accuracy. The application of the Particle Swarm Optimization algorithm to the Naïve Bayes algorithm significantly increases accuracy, especially in sentiment analysis.

Research that has been conducted [5] shows that the use of a classification model

with the SVM method optimized through PSO provides satisfactory performance in analyzing reviews related to Google Classroom. Before applying PSO, the SVM method achieved an accuracy of 79%. However, after using PSO for feature selection, the accuracy increased to 83%. This proves that PSO is an effective feature selection method for SVM in sentiment analysis, as evidenced by the 4% increase in accuracy in this study. Feature selection and parameter tuning in Support Vector Machine have a significant impact on the results of classification accuracy. To overcome this optimization problem, the Particle Swarm Optimization algorithm is often applied as a solution. PSO not only solves the optimization challenge, but also effectively addresses the feature selection problem, improving the overall performance of the classification model [6].

Based on previous research Naïve Bayes and SVM was used in conducting Sentiment Analysis. Naïve Bayes, first proposed by Thomas Bayes in the 18th century, is a classification algorithm that incorporates simple principles of Bayesian theory. This algorithm is known as the Naïve Bayes classifier, which is known for its accuracy and efficiency when applied to large datasets. This method is often used in the context of machine learning because it is able to provide a high level of accuracy with a relatively simple calculation process [9]. While the other algorithm SVM, which introduced by Vapnik in 1992, is a machine learning method that relies on the principle of Structural Risk Minimization [10]. This method attempts to separate the data space using linear or nonlinear classification to distinguish between different classes. The main concept of SVM is the use of hyperplane as a separator of two data classes, which acts as a boundary between positive and negative classes, thus allowing for more accurate and effective predictions [3].

According to [11] Particle Swarm Optimization (PSO) is an optimization method that aims to determine process parameters to achieve optimal response values. PSO mimics the social behavior of flocks of birds or fish in their natural habitat. Although each individual or particle acts independently using its own intelligence, their behavior is still influenced by the collective dynamics of the group. When one particle or bird finds an optimal path to a food source, other members of the group will follow that path, even if they are far from that location. Each individual or particle is treated as a point in a certain dimensional space, with two main factors determining their status in the search space: the position and velocity of the particle.

Based on previous research showing that the use of Particle Swarm Optimization (PSO) can improve accuracy of Naïve Bayes and SVM. This study compares several algorithmic approaches to determine the best accuracy in reviewing sentiment analysis on the YouTube Music application. The methods compared include Naïve Bayes without PSO, Naïve Bayes based on PSO, Support Vector Machine without PSO, and Support Vector Machine based on PSO. This study aims to determine the most effective way to improve the accuracy of user review sentiment analysis.

RESEARCH METHODOLOGY

In this study, the problem-solving design was created based on CRISP-DM. CRISP-DM stands for Cross-Industry Standard Process for Data Mining. It is a widely used process model for data mining and data science projects, originally developed in the late 1990s by a consortium including Daimler-Benz, SPSS, and NCR. Despite its age, it remains highly relevant and is still adopted in both academia and industry for guiding data-driven problem solving as described in [15], [16].

As stated in the framework This study is divided into 6 phases, as follows:

1) Business Understanding

At this stage, an understanding of the research object is carried out. The problem raised is that YouTube created a new application, namely the YouTube Music application. Therefore, the researcher created a solution by creating an algorithm classification model that can later be used to classify public sentiment regarding the application. In this study, the researcher used user review data for the YouTube Music application from January-December 2023 as the object of this research which was taken from the Google Play Store. At this stage, user comments on the YouTube Music application will be categorized into positive and negative sentiment categories.

2) Data Understanding

At this stage is the stage of understanding the data that will be used as the material to be studied. The steps taken at this stage are to collect review data from YouTube Music application users in January-December 2023 as research objects taken from the Google Play Store using the Data Miner tool. In this study, the data used is in the form of reviews found on the Google Play Store on the YouTube Music application, where the domain address of this website is as follows:

https://play.google.com/store/apps/details?id=com.google.android.apps.youtube.music&pcampaignid=web_share.

Data retrieval was carried out using the scraping technique using the Data Miner tool which is an extension of Google Chrome.

3) Data Preparation

At this stage, the raw data obtained will be converted into qualitative data, then cleaned until it becomes data that is ready to be processed. At this stage, the collected review data will go through several stages of text processing:

a. Cleaning

First, the data will be cleaned, namely cleaning the reviews from unimportant characters such as hashtags, mentions, URLs, symbols, etc.

b. Labeling

YouTube Music app reviews are divided into two categories: positive and negative. Positive reviews include praise and satisfaction, while negative reviews contain complaints and dissatisfaction.

c. Preprocessing

In this step, data that has been labelled are preprocessed using stemming, tokenize, filter stopwords and filter token by length.

Stemming is a process in text processing that aims to reduce words to their basic form or root word. Tokenize will split words in a sentence into individual parts of the word itself. Filter Stopword will remove words that contain conjunctions. Filter Token (by length) will remove words that are too short or too long. Words with less than four letters and words with more than twenty-five letters will be removed. After that the resulting data will be split into training data and testing data. The Training data will be used in Modeling, and the Testing Data will be used to create Confusion Matrix and then to calculate the accuracy according to the result of the Confusion Matrix.

4) Modeling

At this stage, apply the appropriate modeling technique by preparing the algorithm to be used. In this study, the

method chosen is Naïve Bayes and Support Vector Machine which are then optimized using feature selection using the Particle Swarm Optimization algorithm. These methods will be applied to the Training Data to create the model.

5) Evaluation

In the evaluation stage, this study will try to apply the model from the Training Data, to predict the label of the Testing Data and create confusion matrix from the result of the prediction. This confusion matrix can be used to calculate Accuracy. From Accuracy the best method can be selected. The process of calculating

performance using the following confusion matrix is an important step in evaluating the performance of a classification model. Confusion Matrix is a matrix-shaped metric that provides a clear picture of how accurate each class is with the selected algorithm. By using a confusion matrix, we can accurately determine the number of true positives, true negatives, false positives, and false negatives, which will later help us evaluate the overall performance of the model in more detail and depth. The formula for calculating accuracy in forming a classification model is as shown in Table 1.

Table 1 Confusion Matrix

		<i>Actual Class</i>	
		<i>Positive</i>	<i>Negative</i>
<i>Predicted Class</i>	<i>Positive</i>	<i>True Positive (TP)</i>	<i>False Positive (FP)</i>
	<i>Negative</i>	<i>False Negative (FN)</i>	<i>True Negative (TN)</i>

The formula for calculating accuracy is based from this Confusion Matrix. The formula is as followed:

$$Accuracy = \frac{TP + TN}{TP + FP + FN + TN}$$

6) Deployment

At this stage, the knowledge and information obtained are presented so that they can be used to provide knowledge in the form of sentiment analysis research to further researchers using this algorithm method. And can be a recommendation for YouTube Music related to the results of the sentiment analysis obtained, and the results of the sentiment analysis can be used as a decision material for YouTube Music for strategies in improving quality.

RESULTS AND ANALYSIS

In accordance with the research method, the results of this study are explained based on the CRISP-DM phases.

1) Business Understanding

In this stage, understanding the problems faced in the research is carried out. The author will find out which level of accuracy is better between the Naïve Bayes method without Particle Swarm Optimization, Naïve Bayes using Particle Swarm Optimization, Support Vector Machine without Particle Swarm Optimization and Support Vector Machine using Particle Swarm Optimization. In the process of comparing accuracy and evaluation in this study, a confusion matrix is used, then the average accuracy value is calculated and the results are obtained. This test is carried out to obtain the best accuracy in the testing process. After obtaining a model with higher accuracy, it

can be said that the model is a better model to use in conducting sentiment analysis.

The sample results of data scraping using the Data Miner tool can be seen in table 2 below.

2) Data Understanding

Table 2 Data Scraping Results

No	User	Date	Review
1	Agus Setyanto	31 Desember	Oke
2	Taufik	31 Desember	Apa ini
3	thoyib A Y	31 Desember	Ku kira adanya youtube musik menikmati musik di youtube tanpa harus stay di aplikasi. Ternyata sama aja di aplikasi aplikasinya mungkin menyenangkan dan tanpa harus
...			
10.892	Firman Hidayat	1 Januari 2023	Mantab Bisa Dengar Semua Lagu

The results of scraping YouTube Music application review data in January-December 2023 on the Google Play Store amounted to 10,892 reviews. The reviews used were only reviews in Indonesian and the data generated was raw data, the review data still contained symbols, hashtags, mentions. Therefore, the next process is data cleaning to remove symbols, hashtags, mentions that are still in the reviews.

3) Data Preparation

At this stage, the raw data that has been obtained will be made into quality data,

then cleaning will be carried out so that the raw data becomes data that is ready to be processed in modelling process.

a. Cleaning

In this step, the data is cleaned from unimportant characters such as hashtags, mentions, URLs, symbols, and any other character that will not be used in classification. Table 3 shows the example results of removing the Hashtag from the reviews. In this step, every hashtags are searched and removed from the reviews.

Table 3 Cleaning Hashtag Example Results

<i>Replace Hastag</i>	
Before	After
Mantap, aplikasi Super Keren..Youtube Music is #1 tiada tandingannya. sangat compatible disemua perangkat Android 5.0 keatas. Mantap..	Mantap, aplikasi Super Keren..Youtube Music is 1 tiada tandingannya. sangat compatible disemua perangkat Android 5.0 keatas. Mantap..
1 Kata #Mantap	1 Kata Mantap

The next step is cleaning the reviews from Mention. The example of this process step is as shown in Table 4. Every

mentions in the reviews are searched and deleted from the reviews.

Table 4 Replace Mention Results

<i>Replace Mention</i>	
Before	After
j@ringan nya lelet	jringan nya lelet

After removing mention from the reviews, the next step is removing unused symbols from the reviews. Using symbol

removal steps all the symbol in the reviews can be removed. The example of this can be viewed in Table 5.

Table 5 Replace Simbol Results

<i>Replace Simbol</i>	
Before	After
Kenapa aplikasinya gk bisa terbuka padahal udah premium, udah upgrade juga?!! Tolong kasi solusi saya harus apa!!	Kenapa aplikasinya gk bisa terbuka padahal udah premium udah upgrade juga Tolong kasi solusi saya harus apa

The next step is trimming the review to remove excess spacing. In this process,

each excess spacing can be detected and removed.

Table 6 Trim Results

<i>Trim</i>	
Before	After
Woyy yang koment banyak iklan , ya kalo gamau banyak iklan makanya premium dong yaelah murah x 1 bulan mah	Woyy yang koment banyak iklan ya kalo gamau banyak iklan makanya premium dong yaelah murah x 1 bulan mah

The results of the Trim process contained in the review have been removed using the Trim operator. The results of the review data that has used Trim can be seen in Table 6. After the data scraping process, the data produced 10,892 reviews, and after cleaning the review data became 7,615 reviews.

b. Labeling

Based on the previous stage, namely data cleaning, the data that has been cleaned is 7,615. The next process is data labeling, data labeling is done manually in Excel by labeling positive or negative sentiments on existing data reviews. Manual labeling is due to time constraints in the study, labeling is done based on the context of the review, the context of the review is considered positive when users feel helped by the features, services or are

happy with their experience. The context of the review is considered negative when users feel disturbed, disappointed or have problems with features or services. Data labeling is done by testing sentiment labeling using 1,523 data. Table 2 shows the example of labelling each review.

Table 7 Manual Review Labeling Example

Review	Sentiment
wah update sekarang parah sih masa gak bisa download lagu padahal kalo mau nyetel lagu yg gak didownload data yg kepake kan gede jadi gak bisa denger lagu jam jam sekarang apa apa mesti premium	Negatif
saya senang dengan aplikasi ini semua lagu kenang dan nostalgia ada semua	Positif
bagus dan suara nya jernih sekali terimakasih	Positif
maaf min saya pindah ke grup belah lebih asik dari pada youtube music udah banyak iklan dan gak bisa main di latar belakang lagi tuh gak kaya grup belah walau banyak kurang nya tapi ada lebih nya bisa main di latar belakang jadi gak mesti idup terus layar hp saya	Negatif

c. Data Preprocessing

Data obtained from Google Play Store with data scraping techniques and 1,523 reviews have been labeled, the data cannot be used because the data still contains noise or errors. To fix this, it is necessary to carry out the Preprocessing stage on the 1,523 review data that have been labeled.

The following are the steps taken in data Preprocessing, namely Stemming, Tokenizing, Filter Stopword, and Filter Token (by Length), and finally splitting the data into Training Data and Testing Data.

The first step in preprocessing the data is Stemming, which is finding the root or basic word for each word in the data.

Table 8 Stemming Process Results

<i>Stemming</i>	
Before	After
Tidak bisa membeli premium gagal terus aplikasi sampah	tidak bisa beli premium gagal terus aplikasi sampah

Words that have previously undergone the Stemming process are still intact words after going through the Stemming process,

the word becomes a basic word. The Stemming results can be seen in Table 8.

Table 9 Tokenize Process Results

<i>Tokenize</i>	
Before	After
aplikasi yang sangat rekomendasi dengar music jadi lebih gampang dan mudah top banget	Aplikasi, Yang, Sangat, Rekomendasi, Dengar, Music, Jadi, Lebih, Gampang, Dan, Mudah, Top, Banget

Sentences that were previously Tokenized are still intact sentences after going through the Tokenizing process, the sentence is divided into words per word.

The results of Tokenizing can be seen in table 9 above.

Sentences that were previously Tokenized are still intact sentences after

going through the Tokenizing process, the sentence is divided into words per word.

The results of Tokenizing can be seen in Table 9 above.

Table 10 List Stopword

No.	List Stopword
1	ada
2	adalah
3	adanya
4	adapun
...	
758	yang

The next process is removing the stopwords. Stopwords are common words in a language that are usually filtered out or removed during text processing, because these words are considered to carry little meaningful information and are often not useful for text analysis. The stopwords list used in this research is downloaded from the Kaggle site at the

following address:
<https://www.kaggle.com/datasets/oswinrh/indonesian-stoplist>. The list of stopwords that have been downloaded is 758 stopwords and can be used immediately. You can see the list of Indonesian stopwords that have been downloaded in Table 10.

Table 11 Stopword Filter Process Results

Filter Stopword	
Before	After
musiknya banyak sekali dan lengkap enak di dengar dan lagu lawas banyak pula	musiknya banyak lengkap enak dengar lagu lawas banyak

Sentences that were previously processed by the Stopword Filter are still intact sentences after going through the Stopword Filter process, words that are

conjunctions are deleted. The results of the Stopword Filter can be seen in table 11 above.

Table 12 Token Filter Process Results

Filter Token	
Before	After
terus tingkat inovasi karena aplikasi ini sangat manfaat khusus untuk mudah cari lagu	terus tingkat inovasi karena aplikasi sangat manfaat khusus untuk mudah cari lagu

Sentences that were previously processed by the Token Filter are still intact sentences after going through the Token Filter process, words that are less than three letters or more than twenty-five letters will be deleted. The results of the Token Filter can be seen in table 12 above.

After labeling and preprocessing the review data, the next process is to divide the data into training data and test data. A

total of 1,523 data that have passed preprocessing are divided into training data and test data. The proportion in dividing the data is 80% for Training Data and 20% for Testing Data. Random Sampling (shuffled sampling) method is used in dividing the data to ensure that the training and testing data are not biased based on the order of the data. In this data division, there are 1,218

training data and 304 test data with the following calculations:

$$\text{Number of Training Data} = \text{Total Dataset} \times \frac{\text{Training Data Ratio}}{100}$$

$$\text{Number of Training Data} = 1523 \times \frac{80}{100}$$

$$\text{Number of Training Data} = 1218$$

$$\text{Number of Test Data} = \text{Total Dataset} \times \frac{\text{Test Data Ratio}}{100}$$

$$\text{Number of Test Data} = 1523 \times \frac{20}{100}$$

$$\text{Number of Test Data} = 305$$

So after calculation there are 1218 training data and 304 test data. Then the process of dividing the data to training and test data must be done according to these numbers.

4) Modeling

In this phase the creation of models was conducted. In this study, the methods used to create the models are the Naïve Bayes classification method and the Support Vector Machine classification method. Additionally these two methods will be equipped with PSO. So that these two methods can be compared with or without PSO. So there are 4 models created in this phase, namely Naïve Bayes Model, Support Vector Machine Model, Naïve Bayes with PSO Model and Support Vector Machine with PSO Model. These

models are created by applying each method with Training Data.

5) Evaluation

In this phase the performance of each model was compared with one another to get the the model with the best performance. Each of the four models created in previous phase then used to predict the label of Testing Data. Then the resulting prediction was summarized using Confusion Matrix. Then from the Confusion Matrix the performance can be compared to get the best accuracy. Model with the accuracy will be used in the next phase to do Sentiment Analysis for the whole 10,892 reviews.

Table 13 shows the Confusion Matrix and Accuracy of Naïve Bayes without PSO

Table 13 Naive Bayes Confusion Matrix Results.

<i>Naïve Bayes</i>			
Accuracy: 87,17%			
	True Negatif	True Positif	Class Precision
Pred. Negatif	41 (TN)	12 (FN)	77,36%
Pred. Positif	27 (FP)	224 (TP)	89,24%
Class Recall	60,29%	94,92%	

Based on the results of the accuracy calculation using the Confusion Matrix formula, it shows that the accuracy level

using the Naïve Bayes algorithm without PSO is 87.17%. After this the next model which is Naïve Bayes with PSO is

evaluated. Table 14 shows the Confusion Matrix and Accuracy of Naïve Bayes with PSO.

Table 14 Naive Bayes Optimization Results with PSO

Naïve Bayes + Particle Swarm Optimization			
Accuracy: 91,80% +/-4,10% (micro average: 91,78%)			
	True Negatif	True Positif	Class Precision
Pred. Negati	50 (TN)	7 (FN)	87,72%
Pred. Positif	18 (FP)	229 (TP)	92,71%
Class Recall	73,53%	97,03%	

Based on the results of the accuracy calculation using the Confusion Matrix formula, it shows that the level of accuracy using the Naïve Bayes algorithm with PSO

is 91.80%. The next model to be processed is SVM without PSO. Table 15 shows the Confusion Matrix and Accuracy of SVM without PSO.

Table 15 Confusion Matrix Results with SVM

Support Vector Machine			
Accuracy: 85,20%			
	True Negatif	True Positif	Class Precision
Pred. Negati	24 (TN)	1 (FN)	96,00%
Pred. Positif	44 (FP)	235 (TP)	84,23%
Class Recall	35,29%	99,58%	

Based on the results of the accuracy calculation using the Confusion Matrix formula, it shows that the level of accuracy using the Support Vector

Machine algorithm is 85.20%. The last model to be processed is SVM with PSO. Table 16 shows the Confusion Matrix and Accuracy of SVM with PSO.

Table 16 Support Vector Machine Optimization Results with PSO

Support Vector Machine + Particle Swarm Optimization			
Accuracy: 85,51% +/-4,47% (micro average: 85,53%)			
	True Negatif	True Positif	Class Precision
Pred. Negati	25 (TN)	1 (FN)	96,15%
Pred. Positif	43 (FP)	235 (TP)	84,53%
Class Recall	36,76%	99,58%	

Based on the results of the accuracy calculation using the Confusion Matrix formula, it shows that the level of accuracy using the Support Vector Machine algorithm with PSO is 85.51%.

These for performances of each models then compared with each other to get the best performance according to the accuracy of each model. The comparison is shown in table 4.17.

Table 17 Comparison of Accuracy Levels

	Confusion Matrix
<i>Naïve Bayes</i>	87,17%
<i>Naïve Bayes + PSO</i>	91,80%
<i>Support Vector Machine</i>	85,20%
<i>Support Vector Machine + PSO</i>	85,51%.

Based on these results the best method to do sentiment analysis in this study is by using Naïve Bayes + PSO because with this method the accuracy achieved gets the highest value.

6) Deployment

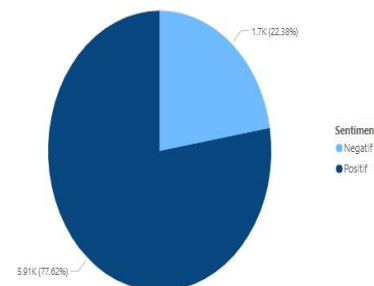
From the evaluation results, it is known that the Naïve Bayes algorithm method with Particle Swarm Optimization is better than the Support Vector Machine algorithm. At this Deployment phase, Naïve Bayes algorithm with Particle Swarm Optimization is used to do sentiment analysis. The initial stage was to create a model using labeled data of 1,523 reviews, then the model was applied to all cleaned review data totalling 7,615 reviews.

From the prediction process using the Naïve Bayes algorithm with Particle Swarm Optimization. The results of the positive sentiment analysis prediction were 5,957 reviews and the results of the negative sentiment analysis prediction were 1,657 reviews. The results of the sentiment analysis using Naïve Bayes can be seen in table 18.

Table 18 Sentiment Analysis Prediction Results

Prediction s (Sentiment)	P 5957 reviews	Negativ 1657 reviews
--------------------------------	----------------------	----------------------------

Based on the results of sentiment analysis predictions using the Naïve Bayes algorithm with Particle Swarm Optimization, the results obtained are more positive sentiments than negative sentiment analysis. The results are visualized in the form of a pie chart as in Figure 2.



Index	Nominal value	Absolute count	Fraction
1	Positif	5957	0.782
2	Negatif	1657	0.218

Figure 2 Sentiment Analysis Results

From the pie chart visualization, the results of the overall sentiment analysis were obtained, where for positive sentiment there were 5957 reviews or 78.2% of the total 7615 reviews. So this shows that the public response to the YouTube Music application is very positive because the high percentage of positive reviews indicates that users are generally satisfied with the services and features provided by the YouTube Music application. For negative sentiment, there were 1657 reviews or 21.8% of the total 7615 reviews. So this shows that although the majority of responses to the YouTube Music application are positive, there are still a number of users who are dissatisfied or have problems with the application. This percentage of negative reviews reflects aspects that need to be fixed or improved by developers to meet user expectations.

Visualization is done using Word Cloud. The visualization results show that there are many words included in the positive sentiment category for the YouTube Music application. The results of the positive sentiment Word Cloud can be seen in Figure 3.

From the visualization, the word “Bagus” appears with the largest size, indicating that this word appears most often in the dataset with positive sentiment. The size of a word in a word cloud directly correlates with its frequency in the dataset because the larger the word, the more the word appears. The frequency of occurrence of the most frequently occurring words is summarized in table 19 which displays the top five words.

Word	Occurrence in Positive Class	Example in Review
Bagus	1750	“Aplikasi ini lengkap banget banyak musik pilihan yg bagus ” “ Bagus banget makasih developer telah buat aplikasi ini aku suka banget udah itu aja”
Aplikasi	980	“ Aplikasi musik yang asik buat di dengar” “Keren aplikasi enak dan mudah dgunakan”
Musik	759	“ Musik selalu terupdate dan aplikasinya ringan” “Aplikasi yang pas untuk dengerin musik banyak rekomendasi pilihan musik ”
Mantap	724	“ Mantap banget fiturnya enak mudah di gunakan” “Aplikasi simple banyak playlist lagu2 keren kita bisa tau apa yg sedang trend saat ini lalu aplikasi ini mantan ”
Lagu	533	“Sangat senang bisa pilih lagu lagu favorit untuk di dengar dengan fitur fitur yang mudah dimengerti” “Banyak pilihan lagu dan sangat mudah digunakan aplikasi anti ribet deh pokok”



From the visualization, the word “Advertisement” appears with the largest size, indicating that this word appears most often in the dataset with negative sentiment. The size of a word in the word cloud directly correlates with its frequency in the dataset because the larger the word,

the more the word appears. The frequency of occurrence of the most frequently occurring words is summarized in table 20 which displays the top five words.

Table 20 Occurrence of Words in the Negative Class

Word	Occurrence in Positive Class	Example in Review
Iklan	421	“Terlalu sering banyak iklan gak jelas dan sering henti sendiri musik” “Sangat buruk iklan selalu ada sangat mengganggu”
Aplikasi	396	“ Aplikasi gak jelas orang mau simpan playlist malah muncul navigasi tidak sedia mau buka channel akun malah ada tulis jadi salah coba lagi aplikasi gak jelas mending spotify” “Lah lah kok update yg ini nambah ribet dalah skip ganti aplikasi lain dah gak seru lg aplikasi”
musik	389	“Dulu musik bisa didownload dan dengar <i>offline</i> sekarang harus premium musik mati jika aplikasi tutup” “ Musik suka mati sendiri saat di ganti tab atau layar mati atau layar hidup mau <i>free</i> atau <i>premium</i> sama aja kendala dah gitu putar <i>offline</i> musik yang <i>premium</i> gak bisa putar wkwkwkwk kocak”
Lagu	387	“Gak bisa dengerin sambil main sosmed tiap kali setel lagu dari sini di kembali musik nya otomatis mati aplikasi jelek” “Saya lagi dengar lagu mati sendiri alias di jeda kenapa harus di jeda segala”
Youtube	361	“Gak jauh beda sama youtube gabisa dimatiin layar” “Jujur aja kalau dibandingin sama spotify joox ini gak ada apa apa nya liat dari segi iklan walaupun sama sama bukan langgan iklan banyak ini parah masa tiap 1 lagu iklan trs ngga bisa diminimize pula ya masa mau dengerin musik harus stay di situ trs udah kalo pas hp dikantongin suka ngga sengaja pencet mending lewat youtube biasa aja bisa diminimize iklan juga gak terlalu banyak”

The conclusion from the negative sentiment word cloud results table is that users often complain about problems with ads, poor quality, and unsatisfactory music listening experience on the YouTube Music application. These things need to be fixed by YouTube Music developers so that the application's image can be improved in the eyes of users.

CONCLUSION

Conclusions obtained from the research that has been conducted are as follows:

1. Naïve Bayes shows a better level of accuracy than the Support Vector Machine algorithm model. The average accuracy of the Naïve Bayes

model is 87.17%, while the Support Vector Machine is 85.20%.

2. The Particle Swarm Optimization method successfully optimized the evaluation process using the Confusion Matrix, with the initial Naïve Bayes accuracy of 87.17% to 91.80%, the initial Support Vector Machine accuracy of 85.20% to 85.51%.
3. The results of the classification sentiment analysis of the YouTube Music application show that users give more positive responses than negative responses. With the number of positive sentiments reaching 5,957 and negative sentiments as many as 1,657, it can be concluded that the YouTube Music application is well received by the public. This reflects user satisfaction with the various features and services offered by this application.
4. Positive reviews show that a good user experience is very important. Developers should continue to pay attention to user feedback and make necessary adjustments to ensure a satisfying experience. With these recommendations, YouTube Music developers can continue to improve and enhance their application, maintaining user satisfaction and strengthening a positive image in the eyes of the public.
5. Negative sentiment reviews show a lot of complaints about the app's performance and features not working properly. Developers need to perform regular maintenance and app updates to fix bugs, improve stability and ensure all features are working properly. By implementing these recommendations, YouTube Music developers can reduce user complaints, improve their experience

and can increase user satisfaction and loyalty towards the app.

The following suggestions need to be done to improve this research:

1. The data labeling process in this study is still subjective, because it uses personal assumptions to label. For further researchers, it is recommended that in labeling, several people are involved to do the labeling or to do the labeling based on the rating or stars in the application review.
2. Sentiment Aspect Analysis: In addition to general sentiment classification (positive and negative), perform sentiment aspect analysis to identify opinions on various app features, such as UI, performance, and lyrics features.

REFERENCES

- [1] Mastanora, R. (2018). Dampak Tontonan Video Youtube pada Perkembangan Kreativitas Anak Usia Dini. In Refika Mastanora: Vol. I (Issue 2).
- [2] Listiyani, D. (2019, November 6). YouTube Music Akhirnya Rilis di Indonesia, Ini Fitur yang dibawa. <https://www.inews.id/Techno/Intern-et/Youtube-Music-Akhirnya-Rilis-Di-Indonesia-Ini-Fitur-Yang-Dibawa>
- [3] Kevin, V., Que, S. S., Iriani, A., & Purnomo, H. D. (2020). Analisis Sentimen Transportasi Online Menggunakan Support Vector Machine Berbasis Particle Swarm Optimization. In *Jurnal Nasional Teknik Elektro dan Teknologi Informasi* | (Vol. 9, Issue 2).
- [4] Putra, T. D., Utami, E., & Kurniawan, M. P. (2023). Analisis Sentimen Pemilu 2024 dengan Naive Bayes Berbasis Particle Swarm Optimization (PSO). *EXPLORE – Volume 13 No 1 Tahun 2023*

- [5] Mursianto, A. G., Widiyanto, D., & Tri Wahyono, B. (n.d.). Analisis Sentimen Ulasan Pengguna Pada Aplikasi Google Classroom Menggunakan Metode SVM Dan Seleksi Fitur PSO. JURNAL INFORMATIK Edisi Ke, 18, 2022.
- [6] Handayani, R. N. (2021). Optimasi Algoritma *Support Vector Machine* untuk Analisis Sentimen pada Ulasan Produk Tokopedia Menggunakan PSO. Media Informatika, Vol. 20, (2), 97–108.
- [7] Sujjada, A., Nurfazri Novianti, J., & Griha Tofik Isa, I. (2023). Analisis Sentimen Terhadap Review Bank Digital Pada Google Play Store Menggunakan Metode *Support Vector Machine (SVM)* (Vol. 9, Issue 2). Jurnal Rekayasa Teknologi Nusa Putra Vol. 9., No. 2, Agustus 2023.
- [8] Pakpahan, D., Widyastuti, H. (2014). Aplikasi Opinion Mining dengan Algoritma *Naïve Bayes* untuk Menilai Berita Online. Jurnal Integrasi, 6(1), 1–10.
- [9] Handayani, F., Feddy, D., & Pribadi, S. (2015). Implementasi Algoritma *Naive Bayes Classifier* dalam Pengklasifikasian Teks Otomatis Pengaduan dan Pelaporan Masyarakat melalui Layanan *Call Center* 110. Jurnal Teknik Elektro Vol. 7 No. 1, Juni 2015.
- [10] Monika Parapat, I., & Tanzil Furqon, M. (2018). *Penerapan Metode Support Vector Machine (SVM) Pada Klasifikasi Penyimpangan Tumbuh Kembang Anak* (Vol. 2, Issue 10).
- [11] Sateria, A., Saputra, I. D., & Dharta, Y. (2018). Penggunaan Metode *Particle Swarm Optimization (PSO)* Pada Optimasi Multirespon Gaya Tekan dan Momen Torsi Penggurdian Material Komposit *Glass Fiber Reinforce Polymer (GFRP)* yang Ditumpuk dengan Material Stainless Steel (SS). Jurnal Manutech Vol.10, No.1, hal 1-57.
- [12] Rininda, G., Santi, I. H., & Kirom, S. (2023). Penerapan SVM dalam Analisis Sentimen Pada Edlink Menggunakan Pengujian *Confusion Matrix*. In *Jurnal Mahasiswa Teknik Informatika* (Vol. 7, Issue 5).
- [13] Hariyanto. (2020, August 9). YouTube Music, Aplikasi Musik yang mudah mencari lagu hits. <https://ajaib.co.id/Youtube-Music-Aplikasi-Musik-Yang-Mudah-Mencari-Lagu-Hits/>.
- [14] Nahjan, M. R., Heryana, N., & Voutama, A. (2023). Implementasi Rapidminer dengan Metode Clustering K-Means untuk Analisa Penjualan Pada Toko Oj *Cell*. In *Jurnal Mahasiswa Teknik Informatika* (Vol. 7, Issue 1).
- [15] P. Chapman, J. Clinton, R. Kerber, T. Khabaza, T. Reinartz, C. Shearer, and R. Wirth, "CRISP-DM 1.0: Step-by-step data mining guide," The CRISP-DM Consortium, 2000. [Online]. Available: <https://www.the-modeling-agency.com/crisp-dm.pdf>
- [16] R. Wirth and J. Hipp, "CRISP-DM: Towards a standard process model for data mining," in Proc. 4th Int. Conf. Practical Applications of Knowledge Discovery and Data Mining, London, U.K., 2000, pp. 29–39.