

ANALYSIS OF TWITTER USER SENTIMENTS ON INDEPENDENT CURRICULUM INDONESIA

Fahmi Cholid¹, Suparman², Ngatma'in³, Insani Wahyu Mubarak⁴

Department of Magister Mathematics Education, Faculty of Teacher Training and Education, Universitas Ahmad Dahlan, Yogyakarta, Indonesia^{1,2}, Department of Bachelor of Education in Indonesian Language and Literature, Faculty of Teacher Training and Education, Universitas Muhammadiyah Surabaya^{3,4}

fahmi2107050021@webmail.uad.ac.id¹, suparman@ppd.uad.ac.id²,
dirjopenewu@gmail.com³, insaniwm@um-surabaya.ac.id⁴,

ABSTRAK

Pendidikan merupakan aspek yang sangat penting dalam berbagai kehidupan, hal ini tidak lepas dari besarnya peran dan dampak positif yang ditimbulkan dari majunya suatu sistem pendidikan. Pendidikan merupakan kunci utama bagi suatu negara untuk unggul dalam persaingan global. Pendidikan selalu berkaitan dengan kurikulum. Kurikulum merupakan alat yang digunakan untuk mencapai tujuan pendidikan sehingga dapat dikatakan bahwa kurikulum merupakan acuan dalam proses penyelenggaraan pendidikan di Indonesia. Kurikulum pendidikan Indonesia telah mengalami perubahan atau revisi setidaknya sebanyak 10 kali, yaitu pada tahun 1952, 1964, 1968, 1975, 1984, 1994, 2004, 2006, 2013. Kurikulum terbaru di Indonesia yaitu kurikulum merdeka merupakan masa dimana guru dan siswa dapat atau memiliki kebebasan dalam berpikir dan juga bebas dalam beban pikiran sehingga dapat mengembangkan potensi pendidikannya. Penelitian ini bertujuan untuk mengklasifikasikan sentimen pengguna Twitter terhadap kebijakan pemerintah mengenai kurikulum mandiri ke dalam sentimen positif dan sentimen negatif. Metode yang digunakan Naïve Bayes Classifier (NBC) dan Support Vector Machines (SVM). Hasil penelitian menunjukkan bahwa persentase sentimen positif sebesar 47% atau 451 tweet sedangkan sentimen negatif sebesar 53% atau 264 tweet.

Kata Kunci : kurikulum, Naïve Bayes Classifier, Sentiment, Support Vector Machines

ABSTRACT

Education is a very important aspect in various lives, this cannot be separated from the magnitude of the role and positive impact caused by the advancement of an education system. Education is the main key for a country to excel in global competition. Education is always related to the curriculum. The curriculum is a tool used to achieve educational goals so that it can be said that the curriculum is a reference in the process of organizing education in Indonesia. The Indonesian education curriculum has been changed or revised at least 10 times, namely in 1952, 1964, 1968, 1975, 1984, 1994, 2004, 2006, 2013. The latest curriculum in Indonesia, the independent curriculum is a time when teachers and students can or have freedom in thinking and also free in the burden of thought so that they can develop their educational potential. This research aims to classify the sentiment of Twitter users towards government policies regarding the independent curriculum into positive sentiment and negative sentiment. The method used by the Naïve Bayes Classifier (NBC) and Support Vector Machines (SVM). The Result shows that the percentage of positive sentiment is 47% or 451 tweets while the negative sentiment is 53% or 264 tweets.

Keywords: Curriculum, Naïve Bayes Classifier, Sentiment, Support Vector Machines

PENDAHULUAN

Education is a very important aspect in various lives, this cannot be separated from the magnitude of the role and positive impact caused by the advancement of an education system. The world of education is an effort to improve the quality of human resources in terms of thought and expertise. Education is the main key for a country to excel in global competition (Yaelasari & Yuni Astuti, 2022). Education is always related to the curriculum. The curriculum is a tool used to achieve educational goals so that it can be said that the curriculum is a reference in the process of organizing education in Indonesia (Gyta et al., 2023). The curriculum in education has a very large role in determining the progress of an education, starting from the realm of concepts to applications or practices in the field. The curriculum also has a role as a plan and arrangement regarding the content and teaching materials as well as guidelines for how to organize good education. The education curriculum in Indonesia has implemented several curriculum changes (Iramdan., 2019). This is related to the development of the times starting from the post-independence period to development. After the reformation in 1998, education in Indonesia showed fundamental changes, from a partial to a holistic paradigm. The milestones of curriculum

development in 2004, 2006 and 2013.

The Indonesian education curriculum has been changed or revised at least 10 times, namely in 1952, 1964, 1968, 1975, 1984, 1994, 2004, 2006, 2013 (Hamami, 2021)(Intiana et al., 2023). The latest curriculum in Indonesia, the independent curriculum is a time when teachers and students can or have freedom in thinking and also free in the burden of thought so that they can develop their educational potential (Zulfa Izza & Susilawati, n.d.). The independent learning curriculum is also a refinement of the 2013 curriculum. One of the advantages of the independent curriculum is that teachers can teach according to student achievements and students can develop them. In addition to the advantages, there are also weaknesses, namely the many educational inequalities in socializing, which makes the application of the independent learning curriculum uneven (Voni Nurhidayati, 2022). This curriculum change occurs in line with changes in the political, social, cultural, economic and science and technology systems in the life of the nation and state (Andriani, 2020). Curriculum changes and reforms need to be done because the curriculum has a dynamic nature in order to be able to answer the

development and challenges of the times (Restiana et al., 2022).

Until now, the concept of the Merdeka Curriculum has received diverse responses from various educational institutions that facilitate students' learning, both at the primary, secondary and tertiary education levels (Abidah et al., 2020). Although there are strong reasons regarding Indonesia's independent curriculum policy, there are several things that make this policy become pros and cons. Thus, the implementation of the independent curriculum cannot be separated from pro opinions, counter opinions, and neutral opinions from the community, parents, students and teachers. The existence of various opinions and public comments on the independent curriculum policy of learning an independent campus on social media, especially twitter, is a **sentiment that occurs in society**. Therefore, it is necessary to conduct an in-depth analysis of the sentiment that occurs in the community regarding the satisfaction of the independent curriculum program.

This research aims to classify the sentiment of Twitter users towards government policies regarding the independent curriculum into positive sentiment and negative sentiment. An algorithm for dividing data into training data in order to predict or classify testing data is called supervised machine learning. The

purpose of classification is to categorize data into label classes, where the label class has been determined by the researcher. In supervised machine learning, there are several classification methods used to classify text data. These classification methods include Naïve Bayes Classifier (NBC), K-Nearest Neighbor, Decision Tree and Support Vector Machines (SVM). The NBC method has been widely used in research on Text Mining because it has the advantage of a simple algorithm but has high accuracy (Rachimawan, 2016). While the SVM method is used because this method is very fast and effective in the classification of text data. In addition, SVM also works very well on data with various dimensions and avoids the difficulties of dimensionality problems.

RESEARCH METHODS

Data Source

The data used in this study amounted to 1000 tweets which were a collection of tweets about the Merdeka Curriculum. Data is obtained from the Twitter API (Application Programming Interface) with the keyword "Independent Curriculum". Data retrieval is carried out along with Twitter user comments about the Merdeka Curriculum.

Analysis Steps

The steps in conducting this research are as follows.

1. Retrieve data from tweets related to the independent curriculum.
2. Perform text preprocessing with the following steps.
 - a. Performing cleansing on tweets by removing noise or unimportant words, such as url, username and hashtag.
 - b. Perform case folding by converting all words to lowercase and removing punctuation.
 - c. Delete words in tweets that are on the stopwords list.
 - d. Perform tokenizing to break sentences in tweets into words.
 - e. Perform stemming to remove affixed words
3. Perform cross validation 10 times (10 folds cross validation), namely by dividing the test data into 10 sub samples, for the ratio of test data starting from 10%, increasing 10% each time until 90%. Data classification using the Naïve Bayes Classifier method.
4. Data classification using the Naïve Bayes Classifier method.
 - a. Calculate the probability of y_j in the training data with equation where y_j is the sentiment class (y_1 = positive, y_2 = negative).
 - b. Calculate the probability of word x_i in category y_j
 - c. Calculate the highest probability in the tested sentiment category (YMAP) and then put the data into a positive or negative sentiment class based on the maximum YMAP.
 - d. Calculating the classification performance of the NBC method based on accuracy, recall, precision and F-Measure.
5. Data classification using the Support Vector Machine method.
 - a. Perform text weighting with TF-IDF.
 - b. Determine the parameter weights in SVM on linear and RBF kernels.
 - c. Build SVM models using linear and RBF kernel functions.
 - d. Classify the testing data based on the SVM model.
 - e. Calculate the classification performance of the SVM method based on accuracy, recall, precision and F-Measure.
6. Comparing the performance of NBC and SVM methods based on the average level of accuracy, recall, precision and F-Measure.
7. Visually display tweets using Word Cloud.

Text Mining

Text mining is part of data mining, which is the process of gaining knowledge using a set of analytical tools where users interact with a set of documents over time. The concept of text mining is usually

used in the classification of textual documents where the documents will be classified according to the topic of the document (Hasri & Alita, 2022).

In recent years, text mining methods have been used for analysis in various fields including software development and marketing. Text mining applications are derived from the data mining branch of science which involves document preprocessing such as text categorization, information extraction, and word extraction (Feinerer et al., 2008). Text mining merupakan teknik yang digunakan untuk menangani permasalahan classification, clustering, information extraction dan information retrieval (Firdaus & Firdaus, 2021).

Text classification can be thought of as the process of forming classes of documents based on previously known class groups. The system of retrieving information from a document based on a specific query or keyword is called the Boolean retrieval model, where each document is considered as a collection of words only. Before a text data undergoes data processing, there are several preprocessing steps to get keywords (Saptono et al., 2017) (Shiri, 2004).

Suppose T is the set of all words that describe the content of the document.

$$T = \{t_1, t_2, \dots, t_m\} \quad (1)$$

A document is a subset of all words, where D is the set of all documents.

$$D = \{D_1, D_2, \dots, D_n\} \quad (2)$$

Q is a Boolean model of a query as follows.

$$Q = (W_1 \vee W_2 \vee \dots \vee W_m) \wedge \dots \wedge (W_i \vee W_{i+1} \vee \dots \vee W_m) \quad (3)$$

Where W_i applies to when D_j when $t_i \in D_j$. to find documents that fulfill Q , S_i is given as follows.

$$S_i = \{D_j \mid W_i\} \quad (4)$$

S_i is a document that is relevant to the keyword. So that a collection of documents that fulfill Q is formed as follows.

$$(S_1 \cup S_2 \cup \dots \cup S_m) \cap \dots \cap (S_i \cup S_{i+1} \cup \dots \cup S_m)$$

Text Preprocessing

Text preprocessing is not only an important step to prepare the corpus for modeling but also a key area that directly affects the results (Chai, 2023). Text preprocessing is the first step of sentiment analysis to convert tweets into structured data as needed. The text preprocessing stage in classification aims to improve the accuracy of data classification. Preprocessing in text mining is quite complicated because in Indonesian there are various rules for writing sentences and forming compound words. The rules of word formation in Indonesian are related to text preprocessing because the final result of text preprocessing is expected to get basic words that are

in accordance with the Big Indonesian Dictionary. The stages in text preprocessing are as follows.

1. Cleansing, which is the stage of removing noise or unnecessary words in tweets. Words that are removed are HTML characters, keywords, emotion icons, hashtags (#), username (@username), url (http://website.com), and e-mail (name@website.com) (Saputra et al., 2020).
2. Case Folding, is the process of converting all text characters into lowercase letters and removing punctuation marks and numbers. The way case folding works is to process alphabet letters from "a" to "z" only so that characters other than these letters will be removed (Affandi & Sugiharti, 2023; Saputra et al., 2020).
3. Tokenization is the process of dividing input text into small units called tokens. Tokens or terms commonly referred to can be in the form of words, numbers, or punctuation marks. In this study, punctuation is removed so that it is not considered a token (Roji & Irhamah, 2019). The tokenizing stage starts from separating the parts of the tweet separated by space characters.
4. Stemming, which is the process of obtaining basic words by

removing prefixes, suffixes, inserts, and confixes (combinations of prefixes and suffixes) (Ariadi & Fithriasari, 2015).

5. Stopwords removal, is the stage of removing vocabulary that does not include unique or characteristic words in a document or does not convey any message significantly in the text or sentence (Affandi & Sugiharti, 2023).

Sentiment Analysis

Sentiment analysis or opinion mining can be defined as a field of science that examines how to convey sentiments, opinions or opinions and emotions expressed in text or sentences. There are several topics of discussion in sentiment analysis, one of the most frequently researched is sentiment classification. This topic centers on the activity of grouping sentiments based on opinion texts on the discussion of issues of interest. The sentiment classifier is divided into positive and negative sentiment groups (Hasri & Alita, 2022; Santoso et al., 2017). Sentiment analysis in Bahasa Indonesia is a technique or method used to identify how an opinion is expressed using text and how the sentiment can be categorized as positive sentiment or negative sentiment (Djamaludin et al., 2022).

Naïve Bayes Classifier (NBC)

Naïve Bayes is a classification method using simple probabilities

rooted in Bayes' Theorem and assumes high independence of each condition or event. (Hasri & Alita, 2022). The Naïve Bayes Classification (NBC) method is one method that can classify text. The advantage of NBC is that it is simple but has high accuracy (Rahayu et al., 2022). One of the classification methods that can be used is the naive bayes method method, which is often referred to as the naive bayes classifier (NBC). The advantages of NBC are simple but has high accuracy High. (Saptono et al., 2017) NBC uses probability theory as its theoretical basis. There are two stages to the classification process, the first stage is training the set of examples (training example). While the second stage is the process of classifying data that has no known category.

In the NBC algorithm each document is represented with attribute pairs "x1, x2, x3, ... xi+1, ..., xn" where x1 is the first word, x2 is the second word and so on. While the set of topic categories discussed is symbolized by Y. During classification, the algorithm will look for the highest probability of all tested data categories (YMAP). The YMAP equation is as follows (Astuti, 2016).

$$Y_{MAP} = \arg \max_{y_j \in Y} P(y_j) \prod_i P(x_i | y_j) \quad (6)$$

The value is calculated at the time of training, obtained from the following equation.

$$P(y_j) = \frac{|doc\ j|}{|training|} \quad (7)$$

Where |doc j| is the number of tweets that have category j in training. While |training| is the number of tweets in the example used for training. For each word xi probabilities for each category, are calculated during training.

$$P(x_i | y_j) = \frac{z_i + 1}{|z + kosakata|} \quad (8)$$

Where zi is the number of occurrences of word yi in tweets with category yj, while z is the number of all words in tweets with category yj and |kosakata| is the number of words in the training data.

Support Vector Machine (SVM)

Support Vector Machine (SVM) is a method that studies the area that separates categories within an observation. In SVM terminology, we discuss the distance or margin between categories. Each category has observations where the target variable value is the same. SVM is also known as a learning system that uses hypothesis linear functions in a high-dimensional space and is trained with algorithms based on optimization theory by applying a learning bias derived from statistical theory. The goal of this method is to build an optimum separator called OSH (Optimal Separating Hyperplane) so that it can be used for classification (Ariadi D & Fithriasari K, 2015)

Term Frequency Inverse Document Frequency

Term Frequency Inverse Document Frequency (TF-IDF) is a weighting performed after document extraction. TF-IDF is done so that documents can be analyzed using Support Vector Machine. Term Frequency (TF) summarizes how often a particular word appears in a document, while Inverse Document Frequency (IDF) decreases the size of words that often appear in documents. The process of TF-IDF method is to calculate the weight by integrating TF and IDF. The step in TF-IDF is to find the number of words we know after multiplying by how many news documents where a word appears (Ariadi D & Fithriasari K, 2015). The formula in calculating the weight with TF-IDF is as follows.

$$w_{ij} = tf_{ij} \times idf \quad (9)$$

$$idf = \log \left(\frac{N}{df_j} \right) \quad (10)$$

Where w_{ij} is the weight of word i in document j , N is the total number of documents, tf_{ij} is the number of occurrences of word i in document j , and df_j is the number of documents j that contain word i .

K-fold cross validation is one of the methods used to ensure that each class in the data has been represented by the training and testing data. This method is widely used by researchers because it can reduce the bias that occurs in

sampling. K-fold cross validation repeatedly divides data into partitions where each data gets a chance to become training and testing data. K is the number of data partitions used for training and testing. The following is an illustration of data division using K-fold cross validation.

RESEARCH RESULTS AND DISCUSSION

In this section, it is explained the results of research and at the same time is given the comprehensive discussion. Results can be presented in figures, graphs, tables and others that make the reader understand easily. The discussion can be made in several sub-sections.

Data Characteristics

The first thing to do after tweets about Indonesia's independent curriculum are collected is to apply text preprocessing. This stage includes noise removal, namely removing links, removing hashtags, removing usernames, removing numbers, removing emoticons and removing punctuation. After the tweets are clean from noise, the next stage is to do case folding, stemming and finally do stopwords removal.

Table 1. Application of Text Preprocessing on Tweets

Number	Class	Tweet	Clean
1	Positive	@schfess kurikulum merdeka yah? xixi	kurikulum merdeka iya xixi

		Need joki matkul pendidikan. Yang paham pendidikan dan hal terkait kurikulum merdeka. Ibu Kota Indonesia akan pinda ke Kalimantan dan Need joki tugas matkul	need joki matkul didik paham didik hal kait kurikulum merdeka
2	Negative		
3	Negative	pendidikan. Yang paham pendidikan dan hal terkait kurikulum merdeka. #zon...	need joki tugas matkul didik paham didik hal kait kurikulum merdeka
⋮	...	...	...
4	negative	@convomf Kemaren abis begadang semalaman buat ngerjain tugas kelompok dan kelompokku gak bantuin...	kemarin habis begadang malam buat ngerjain tugas kelompok kelompok bantuin d indah kur

Table 1 shows the comparison of tweets before and after preprocessing. Tweets that have gone through the preprocessing stage are a collection of basic words without stopwords.

The independent curriculum policy certainly reaps pro and con comments from the public. A comparison of positive and negative

sentiments from Twitter users can be seen in Figure 1 below.

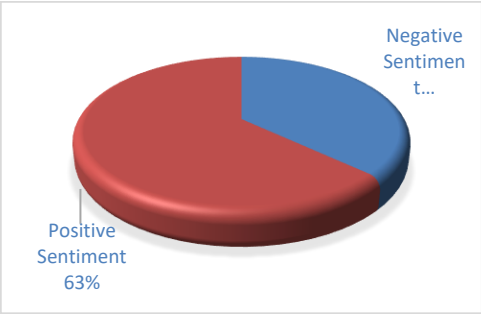


Figure 1. Comparison of Positive and Negative Sentiments

Figure 1 shows that the percentage of positive sentiment is 47% or 451 tweets while the negative sentiment is 53% or 264 tweets. This means that the community's positive response to the independent curriculum policy in Indonesia is more than the negative response.

The percentage of negative sentiment is 37%, meaning that there are still people who disagree with the independent curriculum policy in Indonesia.

The frequency of keywords that appear in positive sentiment and negative sentiment is presented in Figure 2 below.

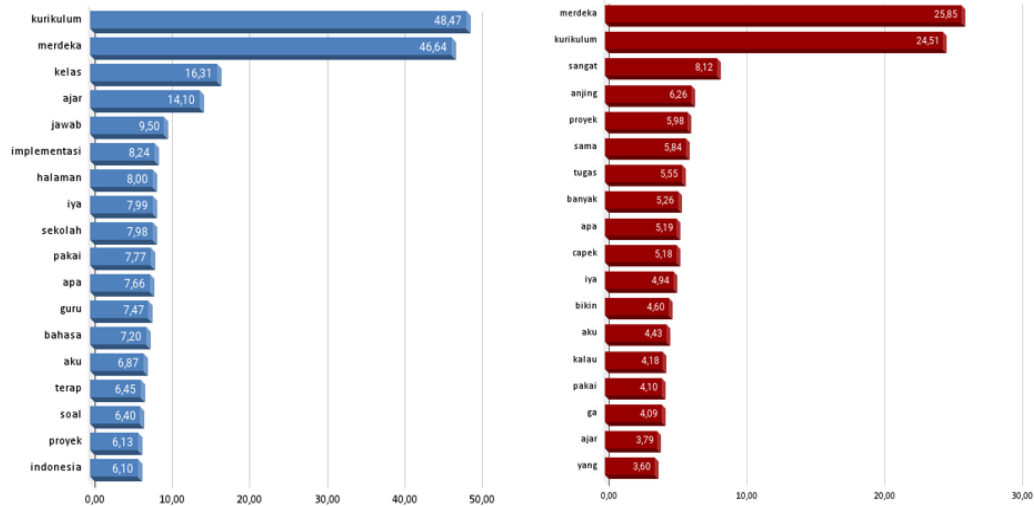


Figure 2. positive word and negative word

the probability of each sentiment class will be sought as follows.

$$P(doc\ negative) = \frac{3}{6} = 0.5$$

$$P(doc\ positive) = \frac{3}{6} = 0.5$$

$P(Y_0)$ is the probability of negative sentiment and $P(Y_1)$ is the probability of positive sentiment. Based on the calculations that have been done, it can be seen that the probability of tweets categorized as positive is 0.5 and the probability of tweets categorized as negative is 0.5. Next, the probability of occurrence of each word in each category will be calculated with the following results.

Table 2. Calculation of Word Probability in Naive Bayes Classifier Method

Number	Word	Class	Probability
1.	Februari	positive	$P(februari positive) = \frac{z_1}{ z + kosakata } = \frac{1 + 1}{ 3 + 48 } = 0,039216$
		negative	$P(februari negative) = \frac{z_1}{ z + kosakata } = \frac{0 + 1}{ 3 + 48 } = 0,019608$
⋮	...	...	...
24	Kurikulum	positive	$P(kurikulum positive) = \frac{z_1}{ z + kosakata } = \frac{2 + 1}{ 3 + 48 } = 0,058824$
		negative	$P(kurikulum negative) = \frac{z_1}{ z + kosakata } = \frac{2 + 1}{ 3 + 48 } = 0,058824$
25	Merdeka	positive	$P(merdeka positive) = \frac{z_1}{ z + kosakata } = \frac{1 + 1}{ 3 + 48 } = 0,039216$

Number	Word	Class	Probability
26	Awal	negative	$P(merdeka negative) = \frac{z_1}{ z + kosakata } = \frac{1 + 1}{ 3 + 48 } = 0,058824$
		positive	$P(awal positive) = \frac{z_1}{ z + kosakata } = \frac{1 + 1}{ 3 + 48 } = v \ 0,039216$
		negative	$P(awal negative) = \frac{z_1}{ z + kosakata } = \frac{0 + 1}{ 3 + 48 } = 0,019608$
⋮	...	...	...
48	Susah	positive	$P(awal positive) = \frac{z_1}{ z + kosakata } = \frac{0 + 1}{ 3 + 48 } = 0,019608$
		negative	$P(awal negative) = \frac{z_1}{ z + kosakata } = \frac{0 + 1}{ 3 + 48 } = 0,019608$

After knowing the probability of occurrence of each word in each category in Table 4.2, we will then look for the highest probability value of the tweets tested. The following is an example of calculation based on equation (2.9) to find the highest probability value on the testing data, namely tweets containing the words many, school, applied, curriculum, and independent.

$$\begin{aligned}
 P(positive) & \prod_i P(x_i|x_y) \\
 & = 0,5 \\
 & \times (0,058824 \\
 & \times 0,039216 \\
 & \times 0,019608 \\
 & \times 0,019608) \\
 & = 6,965 \times 10^{-08} \\
 P(negative) & \prod_i P(x_i|x_y) \\
 & = 0,5 \\
 & \times (0,058824 \\
 & \times 0,019608 \\
 & \times 0,039216 \\
 & \times 0,039216 \\
 & \times 0,039216) \\
 & = 1,3042 \times 10^{-08}
 \end{aligned}$$

Based on these calculations, it can be seen that tweets containing the words many, school, applied, curriculum, and independent are classified in the positive class because they have a greater probability than the calculation results in the negative class. By repeating the equation, the YMAP

value will be obtained which is the maximum output value of the Naïve Bayes classification results from equation (2.9) for each tweet so that the classification accuracy or classification performance using the Naïve Bayes Classifier method will be obtained. In this study, 10-fold cross validation is used, which divides the overall data into 10 parts where each part will be 90% training data and 10% testing data. The following is the classification accuracy in the 4th fold.

Table 3. Classification Accuracy of the 7th Fold with NBC Method

Actual Class	Prediction Class	
	Positive	Negative
Positive	15	12
Negative	4	40

Table 3 shows that the classification accuracy for positive tweets that are correctly classified as positive is 15 tweets, while negative tweets that are correctly classified as negative are 40 tweets. There are 12 positive tweets that are incorrectly classified into negative tweets and there are 4 negative tweets that are classified into positive tweets.

The following is the classification performance obtained

from the Naïve Bayes Classifier method on tweet data about the independent curriculum policy. Classification performance is obtained from the classification results with the calculation of the highest probability value for each tweet in the testing data.

Table 4 .Classification Performance of Naive Bayes Classifier Method

Fold	ccuracy	recision	Recall	F-Measure
1	0,76	0,75	0,71	0,72
2	0,58	0,51	0,50	0,49
3	0,72	0,65	0,65	0,71
4	0,73	0,73	0,67	0,67
5	0,68	0,67	0,60	0,59
6	0,76	0,80	0,70	0,71
7	0,77	0,78	0,73	0,74
8	0,67	0,65	0,60	0,59
9	0,71	0,71	0,65	0,70
10	0,67	0,65	0,60	0,59
Avarage	70	69	64	65

Table 4 shows the results of the accuracy, precision, recall, and F-Measure values with the Naïve Bayes Classifier method for the classification of tweets about the independent curriculum policy on the testing data. Table 4 also shows that the classification performance on each fold is not much different. The highest classification performance on the 7th fold has the best classification performance value, namely with an F-Measure value of 74% while the 2nd fold shows the lowest classification performance with an F-Measure value of 49%.

Classification Using Naive Bayes Classifier

Sentiment classification on tweets about the independent curriculum policy using the SVM method will be carried out in two ways, namely SVM Linear and SVM Kernel RBF. The difference between the two methods lies in the kernel used and will be discussed further in the following subchapters.

Classification Using SVM Linearly Separable Data

Furthermore, sentiment classification is carried out using the Linear SVM method. Data is divided into training and testing using k-fold cross validation. The response variable used is a tweet classification variable that has been categorized into positive and negative based on the perception of the researcher. The independent variable used is the keyword weight of each tweet that has been weighted with TF-IDF. The following is the accuracy of classification in the 2nd fold.

Table 5. Classification Accuracy of 7th Fold for Linear SVM

Actual Class	Prediction Class	
	Positive	Negative
Positive	9	17
Negative	1	44

Table 5 shows that the correctly classified positive tweets are 9 tweets, while the correctly classified negative tweets are 44 tweets. There are 17 positive tweets that are misclassified into negative tweets and there is 1 negative tweet that is classified into positive tweets. Classification performance with the SVM Linear method will be measured

based on accuracy, precision and recall values. The following are the results of the SVM Linear classification performance.

Table 6. Classification Performance of Linear SVM

Fold	Accuracy	Precision	Recall	F-Measure
1	0,66	0,66	0,55	0,52
2	0,75	0,81	0,66	0,67
3	0,75	0,81	0,66	0,67
4	0,73	0,80	0,64	0,64
5	0,70	0,78	0,62	0,60
6	0,69	0,76	0,60	0,57
7	0,65	0,63	0,53	0,54
8	0,70	0,71	0,62	0,61
9	0,66	0,64	0,56	0,53
10	0,73	0,85	0,63	0,62
Average	70	75	61	60

Table 6 shows the results of the accuracy, precision, recall, and F-Measure values with the Linear SVM method on the tweet classification testing data for the independent curriculum policy. Table 4.6 also shows that the classification performance on each fold is almost the same. The highest classification performance in the 2nd fold and 3rd fold has the best classification performance value, namely with an F-Measure value of 67% while the lowest classification performance is shown in the 1st fold with an F-Measure value of 52%.

Classification Using SVM Kernel RBF

Classification analysis using the SVM method with RBF kernel requires cost (C) and gamma (γ) parameters which will later select the optimum parameter values. Before performing classification, first select

the best parameters or tuning parameters on the value of C and gamma. The C parameter used starts from 0.01 to 100 and the gamma parameter used starts from 0.01 to 100. Based on what is done, the best C parameter used is 1 and the best gamma parameter used is 1. Classification with SVM kernel RBF method is tested on training and testing data. Based on the method of dividing training and testing data, the 10 fold cross validation method is used, the data is divided into 90% training and 10% testing. The response variable used is a categorized tweet classification variable. The independent variable used is the keyword of each tweet that has been weighted with TF-IDF. So that the classification accuracy can be arranged as follows.

Table 7. Classification Accuracy of 7th Fold for Kernel RBF SVM

Actual Class	Prediction Class	
	Positive	Negative
Positive	50	3
Negative	11	36

Table 7 shows that there are 50 positive tweets that are correctly classified as positive, while 36 negative tweets are correctly classified as negative. There are 3 positive tweets that are misclassified into negative tweets and there are 11 negative tweets that are classified into positive tweets. Based on the classification accuracy, the classification performance can be measured based on the accuracy, precision, recall and F-Measure

Text	Visualization	Using
Wordcloud		

Figure 3.1 Visualization of Keywords with Wordcloud Positive



Figure 3.1 Visualization of Keywords with Wordcloud Negative

be seen that the words “sangat”, “capek””, “anjing””, “proyek” and “tugas” are words that often appear in negative sentiment tweets. This means that many people complain, disagreeing with the independent curriculum policy.

CONCLUSION

Based on the analysis that has been carried out previously regarding the classification of the sentiments of Twitter users or the public regarding the independent curriculum policy in Indonesia, the conclusions that can be drawn are as follows.

1. Words that often appear as a form of support for the independent curriculum policy in Indonesia besides the independent curriculum are "class", "teach", "language", "jawan" "teacher", "school", "Indonesia" and so on. These words show the enthusiasm of the community in welcoming the independent curriculum policy in Indonesia. Meanwhile, in the negative sentiment, it can be seen that the words "very", "tired", "dog", "project" and "task" are words that often appear.
2. The classification results with the Naive Bayes Classifier method are known to have the best F-Measure value for testing data on the 7th fold, which is 74%. The average accuracy, precision and recall values are 0.77; 0.78 and 0.73, respectively.
3. The classification results with the SVM Linear method are known to have the best F-Measure value for testing data on the 2nd fold and 3rd fold, which is 67%. While the classification results with the SVM

Kernel RBF method are known to have the best F-Measure value for testing data is on the 10th fold, which is 69% with a parameter value of C of 1 and a parameter of 1.

4. In this study, the best classification method is to use Naive Bayes Classifier because it has the highest average classification performance value compared to the other two methods.

REFERENCES

- Abidah, A., Hidaayatullaah, H. N., Simamora, R. M., Fehabutar, D., & Mutakinati, L. (2020). The Impact of Covid-19 to Indonesian Education and Its Relation to the Philosophy of "Merdeka Belajar." *Studies in Philosophy of Science and Education (SiPoSE)*, 1(1), 38–49. <http://scie-journal.com/index.php/SiPoSE>
- Affandi, Y., & Sugiharti, E. (2023). Sentiment Analysis of student on Online Lectured During Covid-19 Pandemic Using K-Means and Naïve Bayes Classifier. *Journal of Advances in Information Systems and Technology*, 5(1), 38–49. <https://doi.org/10.15294/jaist.v5i1.64903>
- Andriani, W. (2020). Pentingnya Perkembangan Pembaharuan Kurikulum dan Permasalahannya. *Universitas Lambung Mangkurat*, 1–12. <https://doi.org/10.35542/osf.io/rkjsg>
- Ariadi, D., & Fithriasari, K. (2015). Klasifikasi Berita Indonesia Menggunakan Metode Naive Bayesian Classification dan Support Vector Machine dengan Confix Stripping Stemmer. In *JURNAL SAINS DAN SENI ITS Vol. 4, No.2* (Vol. 4, Issue 2).
- Ariadi D & Fithriasari K. (2015). *Klasifikasi Berita Indonesia Menggunakan Metode Naive Bayesian Classification dan Support Vector Machine dengan Confix Stripping Stemmer* (Vol. 4, Issue 2).
- Astuti, 2016. (2016). ストレス反応の主成分分析を試みてー 田甫久美子View metadata, citation and similar papers at core.ac.uk. *PENGARUH PENGUNAAN PASTA LABU KUNING (Cucurbita Moschata) UNTUK SUBSTITUSI TEPUNG TERIGU DENGAN PENAMBAHAN TEPUNG ANGKAK DALAM PEMBUATAN MIE KERING*, 15(1), 165–175.
- Chai, C. P. (2023). Comparison of text preprocessing methods. *Natural Language Engineering*, 29(3), 509–553. <https://doi.org/10.1017/S1351324922000213>
- Djamaludin, M. A., Triayudi, A., & Mardiani, E. (2022). Analisis Sentimen Tweet KRI Nanggala 402 di Twitter menggunakan Metode Naïve Bayes Classifier. *Jurnal Teknologi Informasi Dan*

- Komunikasi), 6(2), 2022.
<https://doi.org/10.35870/jti>
- Feinerer, I., Hornik, K., & Meyer, D. (2008). *Journal of Statistical Software Text Mining Infrastructure in R*.
<http://www.jstatsoft.org/>
- Firdaus, A., & Firdaus, W. I. (2021). Text Mining Dan Pola Algoritma Dalam Penyelesaian Masalah Informasi : (Sebuah Ulasan). In *Jurnal JUPITER* (Vol. 13, Issue 1).
- Gyta, D., Harahap, S., Sormin, S. A., Fitrianti, H., & Rafi, M. (2023). *Implementation of Merdeka Curriculum Using Learning Management System (LMS)*.
<https://doi.org/10.55299/ijere.v2i1.439>
- Hamami, T. (2021). National Curriculum Reforms in Indonesia: Moving from Partial to Holistic Curriculum. In *Turkish Online Journal of Qualitative Inquiry (TOJQI)* (Vol. 12, Issue 8).
<https://orcid.org/0000-0002-1490-8883>
- Hasri, C. F., & Alita, D. (2022). PENERAPAN METODE NAÏVE BAYES CLASSIFIER DAN SUPPORT VECTOR MACHINE PADA ANALISIS SENTIMEN TERHADAP DAMPAK VIRUS CORONA DI TWITTER. *Jurnal Informatika Dan Rekayasa Perangkat Lunak (JATIKA)*, 3(2), 145–160.
<http://jim.teknokrat.ac.id/index.php/informatika>
- Intiana, S. R. H., Prihartini, A. A., Handayani, F., Mar'i, M., & Faridi, K. (2023). Independent Curriculum and the Indonesian Language Education throughout the Era of Society 5.0: A Literature Review. *AL-ISHLAH: Jurnal Pendidikan*, 15(1), 911–921.
<https://doi.org/10.35445/alishlah.v15i1.3140>
- Iramdan., L. M. (2019). Sejarah Kurikulum di Indonesia. [https://Jurnal.Peneliti.Net/Index.Php/JIWP/Article/View/98/78,5\(2\)](https://Jurnal.Peneliti.Net/Index.Php/JIWP/Article/View/98/78,5(2)).
- Rachimawan, A. F. (2016). *ADS Filtering Menggunakan Jaringan Syaraf Tiruan Perceptron, Naïve Bayes Classifier Dan Regresi Logistik*. 5(1), 98.
<http://repository.its.ac.id/51346/>
- Rahayu, R., Iskandar, S., & Abidin, Y. (2022). Inovasi Pembelajaran Abad 21 dan Penerapannya di Indonesia. *Jurnal Basicedu*, 6(2), 2099–2104.
<https://doi.org/10.31004/basicedu.v6i2.2082>
- Restiana, S., Agustina, R., Rahman, J., Ananda, R., & Witarsa, R. (2022). Standar Proses Pendidikan Nasional: Implementasi dan Analisis terhadap Komponen Guru Matematika di SD Muhammadiyah 027 Batubelah. *MASALIQ*, 2(4), 489–504.
<https://doi.org/10.58578/masaliq.v2i4.444>

- Roji, M. F. F., & Irhamah, I. (2019). Topic Discovery pada Dokumen Abstrak Jurnal Penelitian di Science Direct Menggunakan Association Rule. *Inferensi*, 2(2), 97.
<https://doi.org/10.12962/j27213862.v2i2.6824>
- Santoso, V. I., Virginia, G., & Lukito, Y. (2017). Penerapan Sentiment Analysis Pada Hasil Evaluasi Dosen Dengan Metode Support Vector Machine. *Jurnal Transformatika*, 14(2), 72.
<https://doi.org/10.26623/transformatika.v14i2.439>
- Saptono, R., Sulisty, M. E., & Trihabsari, N. S. (2017). Text Classification Using Naive Bayes Updateable Algorithm in Sbmpn Test Questions. *Telematika*, 13(2), 123.
<https://doi.org/10.31315/telematika.v13i2.1728>
- Saputra, I., Halomoan, J. A., Raharjo, A. B., & Syavira, C. R. A. (2020). Sentiment Analysis on Twitter of Psbb Effect Using Machine Learning. *Jurnal Techno Nusa Mandiri*, 17(2), 143–150.
<https://doi.org/10.33480/techno.v17i2.1635>
- Shiri, A. (2004). Introduction to Modern Information Retrieval (2nd edition). *Library Review*, 53(9), 462–463.
<https://doi.org/10.1108/00242530410565256>
- Voni Nurhidayati, F. R. M. S. (2022). *VONI NURHIDAYATI, DKK 707*.
- Yaelasari, M., & Yuni Astuti, V. (2022). Implementasi Kurikulum Merdeka Pada Cara Belajar Siswa Untuk Semua Mata Pelajaran (Studi Kasus Pembelajaran Tatap Muka di SMK INFOKOM Bogor). *Jurnal Pendidikan Indonesia*, 3(7), 584–591.
<https://doi.org/10.36418/japendi.v3i7.1041>
- Zulfa Izza, A., & Susilawati, S. (n.d.). *STUDI LITERATUR: PROBLEMATIKA EVALUASI PEMBELAJARAN DALAM MENCAPAI TUJUAN PENDIDIKAN DI ERA MERDEKA BELAJAR*.
<https://proceeding.unikal.ac.id/index.php/kip>