

Classification of Bank Syariah Indonesia (BSI) Customer Sentiments on Twitter Using Naive Bayes Algorithm

Muhammad Alfarisi¹, Munawir², Samsuddin³

^{1,2,3}Teknik Komputer Fakultas Teknik, Universitas Serambi Mekkah

¹alfarisimuhammad2002@gmail.com, ²munawir@serambimekkah.ac.id, ³samsuddin@serambimekkah.ac.id

Article Info

Article history:

Received May 27, 2025

Revised May 28, 2025

Accepted May 29, 2025

Keywords:

Bank Syariah Indonesia
Naïve Bayes Classifier
Sentiment Analysis
Social Media Analytics
Tweets

ABSTRACT

Bank Syariah Indonesia (BSI) is an active topic of conversation on Twitter, but customer sentiment patterns towards the bank's services have not been quantitatively analysed. This study performs positive and negative sentiment classification on 24,401 Indonesian tweets collected on May 17, 2023. The preprocessing stage includes text cleaning, nonstandard word normalization, stopword removal, and stemming with the Sastrawi library. The data was labeled based on the affection dictionary and verified manually. Text representation is done with word frequency-based unigram-bigram method using CountVectorizer, then trained using Multinomial Naive Bayes algorithm. Evaluation of the model against test data resulted in an accuracy of 94%, with precision, recall, and F1-score of 93% each. Words that commonly appear in positive sentiments include easy and fast service, while negative sentiments are dominated by the words error and maintenance. These results show that the Naive Bayes-based approach and word frequency representation are effective for rapid analysis of public opinion towards BSI on social media.



Corresponding Author:

Munawir

Universitas Serambi Mekkah

Email: munawir@serambimekkah.ac.id

1. INTRODUCTION

Twitter is the main space for customers to express satisfaction or complaints about banking services. For Bank Syariah Indonesia (BSI), monitoring sentiment on this platform is crucial as digital reputation directly affects public trust. The high volume of conversations makes manual analysis impractical; machine learning-based automated techniques are required. Multinomial Naive Bayes is known to be lightweight and effective for Indonesian text classification [1][2][3].

This mini-study classifies positive, negative, and neutral sentiments from 24 187 tweets about BSI dated May 17, 2023. The process includes cleaning, slang normalization, stop-word removal, literary stemming, and unigram-bigram TF-IDF representation [4][5]. The model was evaluated with 10-fold cross-validation to assess accuracy as well as sentiment-triggering keywords. The research aims to provide a quantitative picture of public opinion and a foundation for BSI to proactively manage reputation and improve digital services [6].

2. METHOD

The method to be applied is Knowledge Discovery in Databases (KDD), which is considered an effective and suitable approach [7]. This method is able to process data in the right format, making it easier to analyze and make decisions. By using KDD, raw data will be transformed into useful and relevant information, thus supporting the decision-making process more accurately and systematically [8]. This approach also helps to find important patterns in the data that can provide valuable insights for policy making.

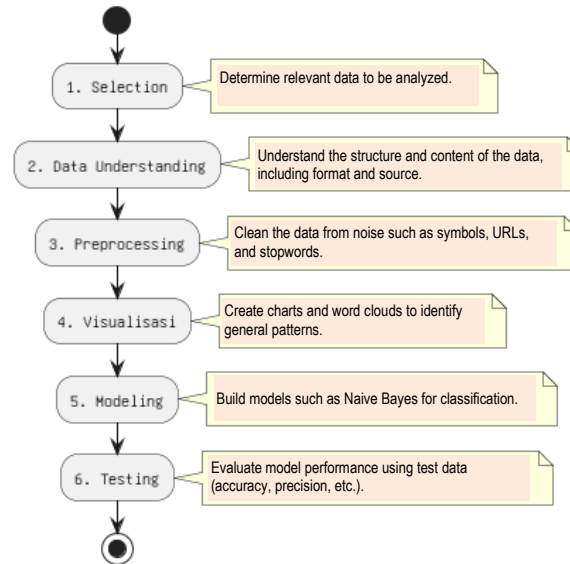


Figure 1. Stages of Methodology

The research stages that use the Knowledge Discovery in Databases (KDD) methodology include the following processes:

2.1 Data Selection

Data Selection is the first stage in the KDD process is data selection, which includes collecting, filtering, and labeling data. The data used is obtained from Twitter social media. After that, labeling is done to classify sentences based on their positive or negative values [1][5].

2.2 Data Understanding

Data understanding is the initial stage in the data analysis process, or Knowledge Discovery in Databases (KDD), which aims to recognise and understand the data to be used. At this stage, researchers collect information about data characteristics, such as data type, data amount, data quality, and relationships between attributes in the dataset. In addition, this stage also involves exploring the data to find initial patterns, detect problems such as missing data or anomalies, and ensure the data is suitable for the research objectives. With a good understanding of the data, subsequent processes such as preprocessing and modelling can be done more effectively and on target.

2.3 Pre-Processing

The preprocessing stage serves to process the raw data that has been collected by removing irrelevant or problematic data such as noisy and missing values so that the data is ready to be used in the next process. The preprocessing steps applied include:

a) Cleaning

The cleaning process aims to clean the text by removing various unnecessary elements, such as punctuation marks (e.g., periods and commas), numbers in tweets, and other elements such as HTML code, URLs, and hashtags [9]. This is done so that the text data becomes cleaner and ready for analysis.

b) Normalization

The text normalization process starts by preparing a conversion dictionary to replace abbreviations or slang (e.g. blm → belum), then homogenizing various spellings of the same word (such as (gk, gak, ga, to tidak), followed by removing excess letter repetition (baaaagus) → (bagus) and correcting common typos (mantp → mantap) after which the writing of numbers and symbols is normalized for example changing 50rb to 50000

and then double-checking to ensure that there are no non-standard forms left, so that the text is finally clean, consistent, and ready to be processed in the next stage of analysis.

c) Labeling

Labeling is the process of assigning a mark or category to data, such as positive or negative to text, in order to facilitate analysis and model training in recognizing patterns and making appropriate predictions.

d) Stopwords

Filtering is done to filter out words in the text by removing terms that appear too frequently and have no significant information value. Such words are known as stopwords and are often ignored in text analysis.

e) Tokenized

Tokenized is the process of separating a sentence or text into small parts in the form of single words. The purpose of this stage is to obtain word units that will be used as basic elements in forming a document matrix that is the basis for further analysis.

f) Stemming

Stemming is the process of obtaining the basic form of a compound word by removing affixes such as prefixes, inserts, suffixes, or a combination of prefixes and suffixes. In text analysis, this stage has an important role because it greatly affects the accuracy and quality of data analysis results.

2.4 Data Mining

Data mining is the process of analyzing large-scale data sets to find useful patterns, relationships, and hidden information. The main goal is to extract new knowledge from data through the application of statistical, mathematical, artificial intelligence (AI), and machine learning techniques. In this mini journal, the data mining process is done by implementing the Naïve Bayes algorithm using the Google Colab platform, which allows writing cloud-based Python code interactively. Google Colab was chosen because it supports various data science libraries and provides an efficient and flexible computing environment without the need for additional software installation.

2.5 Visualization

Data visualization is a technique to convey information or analysis results in graphical form, such as diagrams, charts, or data maps. The goal is to facilitate understanding of complex data, reveal hidden patterns or trends, and help the decision-making process based on information that is presented visually and intuitively.

2.6 Modeling

The modeling stage is done by applying the Naive Bayes algorithm, specifically the MultinomialNB variant, which is suitable for text data such as documents or tweets. This algorithm works by calculating the probability of each class based on word frequency, so it can be used to perform classification, such as determining positive, negative, or neutral sentiment efficiently and quickly.

2.7 Evaluation

In this study, an evaluation was conducted on the sentiment classification model of Bank Syariah Indonesia (BSI) customers on Twitter using the Naive Bayes algorithm. The data used is divided into two sentiment categories, namely positive and negative. Prior to model training, the text data was converted into a numerical representation using the CountVectorizer method, which converts each tweet into a feature vector based on word frequency. The data was then separated into training and test data, and the model was evaluated using accuracy metrics as well as classification reports. The evaluation results showed that the model achieved an accuracy of 94.01%, with high precision, recall, and f1-score values in both classes. Based on the AUC classification standard, the model's performance falls into the Excellent Classification category (0.90-1.00) [1][2][3], indicating that the model is highly effective in classifying customer sentiment towards BSI on Twitter.

3. RESULTS AND DISCUSSION

This research uses a dataset consisting of a total of 3,175 tweets, with an unbalanced sentiment distribution between the two classes. Negative sentiment data (0) dominated as many as 2,191 tweets (69%), while positive sentiment (1) amounted to 984 tweets (31%). This data imbalance indicates that complaints or negative feedback from BSI customers are more common than positive responses. Nonetheless, the Naïve Bayes model still performed very well with an accuracy of 94.02%, proving the robustness of this algorithm in handling data imbalance when the extraction and preprocessing features are done properly. This result also indicates that despite more customer complaints, the model can still recognize positive sentiment patterns well

(precision 92% and recall 90%). The following is a discussion of the research stages that have been carried out:

3.1 Data Processing

This research uses a public dataset available on the Kaggle platform, containing 24,401 tweets related to Bank Syariah Indonesia (BSI) customers. The data is collected and uploaded by the dataset provider through a previous tweet crawling process, so researchers can directly utilize the data that is already available without the need to do web scraping again. The dataset in CSV format was then manually grouped into positive and negative then continued the labeling process. The results of data crawling and labeling are as in table 1.

Table 1. Data Labeling

No	Text	Label
1	sudah bisa blm ya ,kok jadi repot begini ðŸ	Negative
2	Jangan lupa perbaiki mbanking untuk topup!!!!	Positive
3	Rekrutmen berdasarkan identitas bukan kualitas bisa merusak tatanan sistem operasi perbankan di Indonesia.	Negative
4	gila sampe tetep nih top up aja sulit banget	Negative
5	error aplikasi bsi mobil muncul komentar pedas	Negative

3.2 Data Understanding

This stage is carried out to understand the characteristics of the data, including the type and structure of the dataset, making it easier to determine the most appropriate analysis method and model to use, as can be seen in the following figure:

Variable	Data Type
URL	object
Date	object
Tweet	object
ID	int64
Username	object
Replies	int64
Retweets	int64
Likes	int64
Quotes	int64
conversationId	int64
Language	object
Links	object
Media	object
Retweeted Tweet	float64
Bookmarks	int64

dtype: object

Figure 2. Specifying the Data Type

3.3 Preprocessing

This stage aims to prepare the data to be suitable and ready for processing in the next step. An example of the results of the preprocessing process that has been carried out can be seen in Table 2.

Table 2. Preprocessing Results

Process	Result
Data Review	Sudah bisa blm ya ,kok jadi repot begini ðŸ
Cleaning	Layan bayar blm ya repot
Normalization	Layan bayar belum ya repot
Stopwords	Layan bayar ya repot
Tokenizing	[layan, bayar, ya, repot]
Stemming	Layan bayar repot

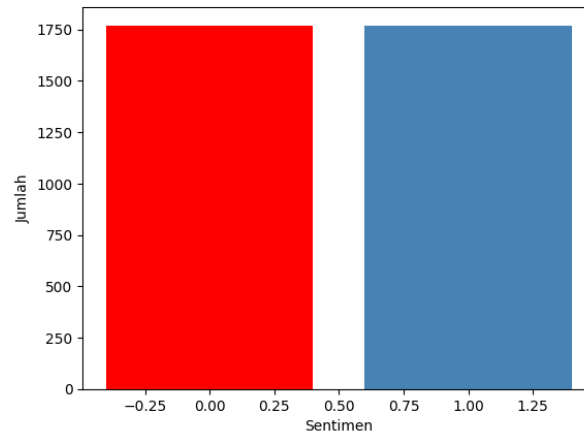


Figure 8. Balanced Training Data Results before testing

3.6 Evaluation

In the evaluation stage, the performance of the classification model is measured using metrics such as accuracy, precision, recall, and f1-score. Before the model training and testing process, the data is analyzed based on the number of each sentiment class. The results show that there are 2,191 data with negative sentiment and 984 data with positive sentiment, resulting in data imbalance. After training the Naive Bayes model, the model was evaluated using test data and resulted in an accuracy of 94%. The model also showed good classification performance with high precision and recall in both classes. The full evaluation results are shown in the following figure.

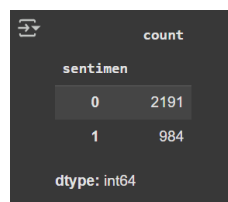


Figure 9. Number of Negative and Positive Sentiments

Furthermore, testing with new data that has first balanced the training data can be seen in the following figure.

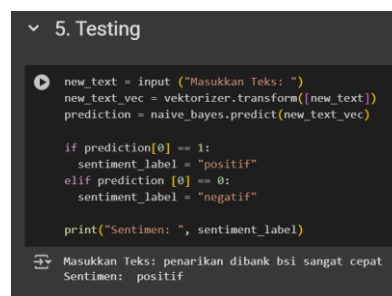


Figure 10. Evaluation results using new data with Naive Bayes algorithm

4. CONCLUSION

Based on the research that has been conducted, it is found that sentiment analysis of public opinion about Bank Syariah Indonesia (BSI) on Twitter social media using the Naïve Bayes method produces the following conclusions:

1. The research was conducted using 3,175 tweets classified into two sentiment categories, namely positive and negative. The data was then processed using the Multinomial Naïve Bayes algorithm with preprocessing, vectorization using CountVectorizer, model training, and model performance evaluation. The training results show an overall accuracy of 94.01% with high precision, recall, and f1-score values for each sentiment class, indicating that the model performs very well in classification.

2. Based on the results of sentiment analysis, it was found that opinions with negative sentiments dominated, which amounted to 2,191 tweets, while opinions with positive sentiments amounted to 984 tweets. This shows that the majority of public opinions expressed through Twitter social media towards BSI are negative, which can be an evaluation material for related parties in improving services and public perception of the institution.

REFERENSI

- [1] I. Amelia, A. M. Adinda, and I. Santoso, "Analisis Sentimen Opini Publik Terhadap Pengambil Alihan Tmii Oleh Pemerintah Dengan Algoritma Naïve Bayes," *J. IKRAITH-INFORMATIKA*, vol. 7, no. 2, pp. 142–148, 2023. [Online]. Available: <https://journals.upi-yai.ac.id/index.php/ikraith-informatika/issue/archive>
- [2] D. Puspitasari, S. S. Al Khautsar, and W. P. Mustika, "Algoritma Naïve Bayes Untuk Memprediksi Kredit Macet Pada Koperasi Simpan Pinjam," *J. Informatika Upgris*, vol. 4, no. 2, 2019. [Online]. Available: <https://doi.org/10.26877/jiu.v4i2.2919>
- [3] T. Arifin and S. Syalwah, "Prediksi Keberhasilan Immunotherapy Pada Penyakit Kutil Dengan Menggunakan Algoritma Naïve Bayes," *J. Responsif: Riset Sains dan Informatika*, vol. 2, no. 1, pp. 38–43, 2020. [Online]. Available: <https://doi.org/10.51977/jti.v2i1.177>
- [4] E. Karyadiputra, "Analisis Algoritma Naive Bayes Untuk Klasifikasi Status Kesejahteraan Rumah Tangga Keluarga Binaan Sosial," *Technologia: J. Ilm.*, vol. 7, no. 4, pp. 199–208, 2016. [Online]. Available: <https://doi.org/10.31602/tji.v7i4.653>
- [5] I. Novitasari, T. B. Kurniawan, D. A. Dewi, and Misinem, "Analisis Sentimen Masyarakat Terhadap Tweet Ruang Guru Menggunakan Algoritma Naive Bayes Classifier (NBC)," *J. Mantik*, vol. 6, no. 3, pp. 2685–4236, 2022.
- [6] C. Supriyanto and I. N. Dewi, "Klasifikasi Teks Pesan Spam Menggunakan Algoritma Naïve Bayes," *Simantik* 2013, pp. 156–160, 2013.
- [7] M. Y. Haffandi, E. Haerani, F. Syafria, and L. Oktavia, "Klasifikasi Penyakit Paru-Paru Dengan Menggunakan Metode Naïve Bayes Classifier," *J. Tek. Inf. dan Komputer (Tekinkom)*, vol. 5, no. 2, p. 176, 2022. [Online]. Available: <https://doi.org/10.37600/tekinkom.v5i2.649>
- [8] A. M. Rahat, A. Kahir, and A. K. M. Masum, "Comparison of Naive Bayes and SVM Algorithm Based on Sentiment Analysis Using Review Dataset," in *Proc. 8th Int. Conf. on System Modeling and Advancement in Research Trends (SMART)*, 2019, pp. 266–270. [Online]. Available: <https://doi.org/10.1109/SMART46866.2019.911751>
- [9] P. P. M. Surya and B. Subbulakshmi, "Sentimental Analysis Using Naive Bayes Classifier," in *Proc. Int. Conf. on Vision Towards Emerging Trends in Communication and Networking (ViTECoN)*, 2019, pp. 1–5. [Online]. Available: <https://doi.org/10.1109/ViTECoN.2019.8899618>
- [10] S. D. Prasetyo, S. S. Hilabi, and F. Nurapriani, "Analisis Sentimen Relokasi Ibukota Nusantara Menggunakan Algoritma Naïve Bayes Dan KNN," *J. KomtekInfo*, vol. 10, pp. 1–7, 2023. [Online]. Available: <https://doi.org/10.35134/komtekinfo.v10i1.330>
- [11] A. V. Sudiantoro and E. Zuliarso, "Analisis Sentimen Twitter Menggunakan Text Mining Dengan Algoritma Naïve Bayes Classifier," *Pros. SINTAK*, pp. 398–401, 2018.
- [12] A. R. Isnain, N. S. Marga, and D. Alita, "Sentiment Analysis Of Government Policy On Corona Case Using Naive Bayes Algorithm," *IJCCS (Indonesian J. Comput. and Cybern. Syst.)*, vol. 15, no. 1, p. 55, 2021. [Online]. Available: <https://doi.org/10.22146/ijccs.60718>
- [13] F. Ratnawati, "Implementasi Algoritma Naive Bayes Terhadap Analisis Sentimen Opini Film Pada Twitter," *INOVTEK Polbeng - Seri Informatika*, vol. 3, no. 1, p. 50, 2018. [Online]. Available: <https://doi.org/10.35314/isi.v3i1.335>
- [14] W. A. Prabowo and C. Wiguna, "Sistem Informasi UMKM Bengkel Berbasis Web Menggunakan Metode SCRUM," *J. Media Inform. Budidarma*, vol. 5, no. 1, p. 149, 2021. [Online]. Available: <https://doi.org/10.30865/mib.v5i1.2604>
- [15] S. Suryono and E. T. Luthfi, "Analisis Sentimen Pada Twitter Dengan Menggunakan Metode Naïve Bayes Classifier," *Jnanaloka*, pp. 81–86, 2021. [Online]. Available: <https://doi.org/10.36802/jnanaloka.2020.v1-no2-81-86>