

Sentiment Analysis Of Indosat's Mobile Operator Services On Twitter Using The Naïve Bayes Algorithm

Sufajar Butsianto^{1*}, Sifa Fauziah², Candra Naya³, Futuh Maulana⁴

^{1,2,3,4}Informatics Engineering Study Program, Faculty of Engineering, Pelita Bangsa University, Indonesian

¹sufajar@pelitabangsa.ac.id, ²sifa_fauziah@pelitabangsa.ac.id, ³candranaya@pelitabangsa.ac.id,

⁴futuh.maulana@gmail.com



*Corresponding Author

Article History:

Submitted: 12-06-2024

Accepted: 13-06-2024

Published: 28-06-2024

Keywords:

Twitter, Data Mining, Naïve

Bayes, Rstudio

Brilliance: Research of

Artificial Intelligence is licensed
under a Creative Commons

Attribution-NonCommercial 4.0

International (CC BY-NC 4.0).

ABSTRACT

Twitter is a social media that allows users to share information with others in real time. Information that is shared on Twitter is usually referred to as a tweet. Sentiment analysis is a branch of research in the text mining domain where the process of identifying and extracting sentiment data will usually be categorized based on its polarity, whether it is positive, negative or neutral. We can process data from opinions on Twitter using data mining techniques, namely classification. The algorithm that will be used in this research is the Naïve Bayes Algorithm. This research will also use the RStudio application. It is a computer programming language that allows users to program algorithms and use tools that have been developed through R by other users. R is a high-level programming language and is also an environment for data and graph analysis. Based on the experimental results, using a comparison of training data and test data of 20%: 80%, 40%: 60%, 60%: 40%, 80%: 20% and 90%:10%, the results of sentiment classification using the Naïve Bayes method are obtained. and using 10-fold cross validation obtained an average value of 85.00% accuracy and The decrease in machine learning performance occurs in the ratio of 80:20 or 1440 training data: 360 data testing, while the ratio of 20%:80% and 90%:10% has the same accuracy value, namely 85.41%.

INTRODUCTION

The development of the world of technology and information is increasingly moving in a digital direction day by day. The digital era has made humans enter a new lifestyle that cannot be separated from electronic devices. Technology is a tool that helps human needs, with technology everything can be done more easily. Challenges in the digital era have also entered various fields such as politics, economics, social culture, defense, security and information technology itself. The digital era was born with the emergence of digital, internet networks, especially computer information technology. The new media of the digital era has the characteristics of being able to be manipulated, being network or internet in nature. . The media capabilities of this digital era make it easier for people to receive information more quickly.(Megawati, 2021)

The role of technology is so important that it is starting to bring civilization into the digital era and increasingly rapidly and the current development of communication technology has changed people's habitual tendencies in expressing their opinions on social media. One of the social media that is popular among internet users today is Twitter. The most popular social media for expressing opinions is Twitter (Syarifuddin, 2020). Twitter is a social media that prioritizes socializing using text, although the new version supports video and photo formats to support tweets. In this way, Twitter is the right tool for collecting public sentiment data on the internet.(Asro'i & Februariyanti, 2022)

Twitter is a social media that allows users to share information with other people in real time. Information shared on Twitter is usually called a tweet (Tweett).(Elsa Annisa Batu Bara et al., 22 C.E.). Indonesia is in fifth position as the country with the most Twitter users after England and other large countries (Kominfo.go.id, 2020). With so many social media users, social media marketing has become a strategic key in marketing activities in the world. This is proven by 94% of companies in the world using social media for marketing purposes. Likewise, cellular operator companies use Twitter as a promotional medium, one of which is Indosat. Indosat currently has 1.5 million followers with official usernames. Indosat is one of the largest cellular operators in Indonesia. However, Indosat also does not escape the comments of its users with various kinds of positive and negative comments. The large number of Indosat cellular operator users who submit comments can be used to search for information to analyze comments on Twitter using sentiment analysis.(Syailendra Reza Irwansyah Rezeki et al., 2020)

Sentiment analysis is a branch of research in the text mining domain where the process of identifying and extracting sentiment data is usually categorized based on its polarity, whether positive, negative or neutral. We can process opinion data on Twitter using data mining techniques, namely classification. To group sentiment, the author divides 3 indicators, namely positive, negative and neutral sentiment with indicators based on tweets. Currently Twitter

is a good indicator for influencing research. This sentiment analysis is carried out to determine public sentiment about something using a machine learning approach (Putranti & Winarko, 2014)

LITERATURE REVIEW

Sentiment Analysis on Twitter Regarding Post-Disaster Using the Naïve Bayes Method with the N-Gram Feature (Rozi et al., 2023), explained that the Naïve Bayes Classifier Algorithm can be used to classify tweets into positive or negative, especially tweets regarding post-disaster. And testing the accuracy of the algorithm which was carried out by manually labeling 15 respondents, it was found that the results from unigrams and bigrams had quite significant differences. From these four tests, the highest accuracy results were obtained for unigrams, namely 93.33% and bigrams, 86.67%. Classification of public opinion towards the MyRepublicId ISP service using a dataset on Twitter and carried out by applying the naïve Bayes method produces accuracy in the categories of positive 0.976%, neutral 0.833%, and negative 0.82895%, with an average value of 0.87949%, the theory is presented by Hafiz Irsyad, Ahmad Farisi, Muhammad Rizky Pribadi in his research entitled "Classification of Public Opinion on MyRepublic ISP Services using Naive Bayes" (Hafiz Irsyad et al., 2019)

Dedi Darwis (Darwis et al., 2021) in their research "Application of the Naive Bayes Algorithm for Sentiment Analysis Review of National BMKG Data" explains that the process of extracting data from National BMKG Twitter uses the Python 3.74 programming language with preprocessing stages which include casefolding, filtering, tokenization, slang replacement and stopword removal. Then the sentiment analysis carried out was to create three classes, namely positive, negative and neutral using the Naive Bayes classification algorithm with an accuracy level based on the tests carried out which was 68.97%.

"Sentiment Analysis of INDOSAT Services on Twitter Social Media During the Pandemic" (Aulawi et al., 2022) is research that aims to determine the sentiment of Indosat users on Twitter social media regarding the services provided during the Covid-19 pandemic. This research was carried out at the Indosat company by conducting a review on Twitter social media to find out the problems that were the topic of public discussion. The method used in collecting data is by scraping and crawling. Scraping is used to collect information data. Scraping is a technique that creates new forms of data in data collection, visualization and analysis. Based on the results of data processing and discussion in this research, the general picture of the sentiment of Telkomsel users in Indonesia during the COVID-19 pandemic on social media regarding the services provided is varied. From the results of data collection and processing, information was also obtained that there were 12,278 positive opinions and 6,952 negative opinions. From the general picture that has been obtained, information was obtained regarding problems and complaints that often appear on social media using the naïve Bayes classifier method and negative text association. Several complaints were obtained that often arise, namely regarding poor and lost signals, then regarding the high price of the package, internet from Indosat, internet connection which is sometimes slow and broken and activation of packages or starter cards which is difficult.

Ronny Julianto and Evi Dianti Bintari, Indrianti (Ronny Julianto et al., 2017) in the research "Sentiment Analysis of Cellular Telephone Provider Services on Twitter using the Naïve Bayesian Classification Method" using a collection of Indonesian language tweets taken from the official Twitter accounts of Telkomsel, Indosat, and XL Axiata. This tweets data was obtained manually from the Twitter API without using a crawler program. The tweets data taken contain the words "telkomsel", "kartu_as", "simpati", "indosatmania", "indosat", "indosatcare", "xl123", "xlcare", "xlandme" as data in this study. Based on the test results that have been carried out on the sentiment analysis system using the naïve Bayes classification method, this application is able to classify sentiments randomly or manually, the classification process is more accurate by using the confusion matrix method as an accuracy calculation process, and the naïve Bayes classification method is able to classify tweets on a sentiment analysis system. (Maharani Nirmala Dewi & Ricky Eka Putra, 2023)

Sentiment analysis regarding #Telkomsel or @Telkomsel tweets on Twitter shows that many people have "negative sentiments" towards Telkomsel services. (Cut Medika Zellatifanny, 2019) The results of the processing of 151 training data concluded that the sentiment classification results obtained were negative sentiment of 51 tweets, neutral sentiment of 51 tweets, and positive sentiment of 49 tweets where neutral is an opinion that has no sentiment. The level of accuracy in determining categories is 70.21% and the micro average is 70.20% in determining sentiment, with a precision level of 70.11% and recall of 70.33%. The results of this research are research from Nizam Haqqizar and Tika Nur Larasyanti in their research entitled "Sentiment Analysis of Telkomsel Telecommunication Provider Services on Twitter Using the Naïve Bayes Method". (Nizam Haqqizar & Tika Nur Larasyanti, 2019) Research on sentiment analysis towards provider The method used in this research is the KDD (Knowledge Discovery in Databases) method which consists of the Data Selection, Preprocessing, Transformation, Data mining and Evaluation stages. The research results show that the performance of the Naïve Bayes Classifier algorithm in classifying user sentiment of provider XYZ has an accuracy value of 92%. Based on manual labeling, sentiment in the opinions of service provider XYZ users contains more negative sentiment. It can be seen that there are many users of service provider XYZ who complain about the price, internet speed, and internet network availability. In this case, XYZ can use this research as a reference for making improvements to the services provided. (Muliani et al., 2024)

METHOD

Problem Analysis

The types of comments on Twitter are in the form of large amounts of unstructured text. This condition can result in missing useful information from a set of text documents. Knowing the sentiments of Twitter users manually can cost you time and energy. Therefore, it is necessary to carry out research using text mining analysis.

In this research, the research object used is Indonesian language tweet data found on the social media Twitter. The tweets used were tweets containing positive, negative and neutral sentiments obtained from October 5 2021 to October 8 2021 on the official Indosat account and data was obtained for 9252 tweets. From each of these sentiments, 1600 data were taken for each positive and negative sentiment so that the total tweets used as data were 3200. Neutral sentiments were not included in this research because they did not contain important meaning.

The data search was carried out using the keyword @indosat on the official Indosat account, then labeling was carried out automatically, namely using an Indonesian lexicon dictionary, the selected tweets were then saved to a text file in CSV format. Then the text file is used as input for further processing.

Method Used

The following is a flowchart that shows the stages of the research method, starting from data collection to reaching conclusions. In designing this system the author used the Knowledge Discovery Database (KDD) method.(Wahidah et al., 2022)

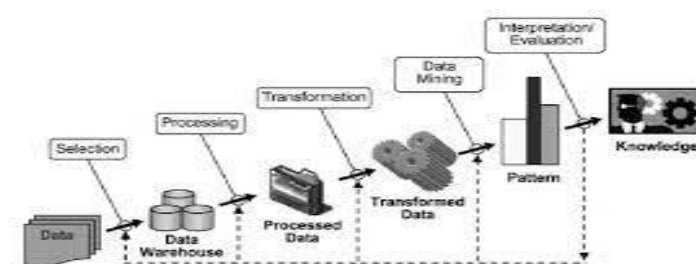


Figure 1. Knowledge Discovery Database Method

a. Data Selection (Data Used)

In retrieving data from Twitter using the scrapping method using the Twitter API in the period from October 5 2021 to October 8 2021 using the keyword "@indosat". In this research, a method using web scrapping techniques was used and the tool used for data crawling was R/RStudio. To crawl data on Twitter, you need a code obtained from the Twitter API to access the Twitter data. Twitter API is an application created by Twitter with the aim of making it easier for developers to access Twitter web information. API registration is used to confirm to Twitter to grant permission to explore more widely regarding data related to Twitter.

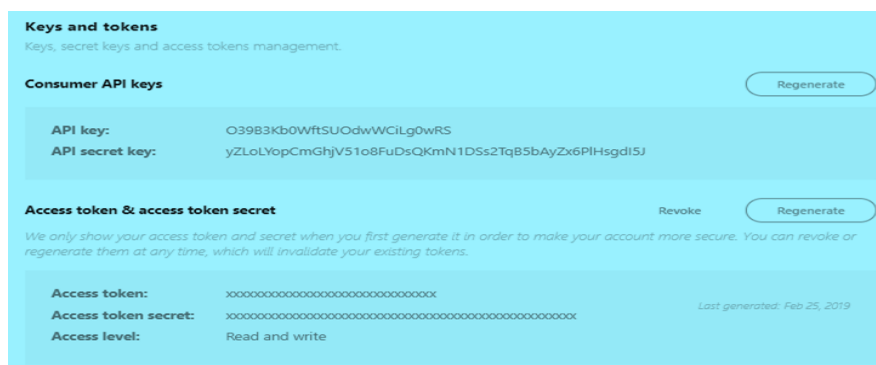


Figure 2. Twitter API Configuration Registration

After registering and joining the Twitter API, you will get several codes in the form of consumer key, consumer secret, access token and access key from Twitter. The API code is a bridge between Twitter and other applications. In this research, the code can be used for the integration process between the Twitter API and R.

b. Data Preprocessing

For example, there are four documents that have gone through the preprocessing stage, two documents are taken from the positive class and two documents are taken from the negative class. The document is as follows.

Table 1. Preprocessing

Document 1: Jokowi was cool when he had Indosat buyback
Document 2: easy, fast, no hassle Indosat credit via SMS
Document 3: Indosat net is slow
Document 4: Indosatcare Indosat's bad network is really slow

c. Data Transformation

Data Transformation is the process of changing the initial data format into a standard data format for the data reading process using the method in the program or tool used. The following are the results of initial data processing after going through the stages above to be used as a dataset in the next stage, shown in Table 2. Data Transformation.

Table 2. Data Transformation

Query	Dokument			
	1	2	3	4
jokowi	1	0	0	0
Cool	1	0	0	0
O'clock	1	0	0	0
he	1	0	0	0
indosat	1	1	1	1
buyback	1	0	0	0
credit	0	1	0	0
easy	0	1	0	0
fast	0	1	0	0
No	0	1	0	0
complicated	0	1	0	0
via	0	1	0	0
text	0	1	0	0
net	0	0	1	1
slow	0	0	1	1
indosatcare	0	0	0	1
Bad	0	0	0	1
mulu	0	0	0	1
very	0	0	0	1

The data set in this research is response data regarding the Indosat account with the keyword @indosat on Twitter social media from October 5 2021 to October 8 2021. The response data obtained was 9252 responses. Determining the amount of training and testing data by comparing the ratio for each positive and negative data is 1800 positive response data and 1800 negative response data, so that the total response data used is 3600 response data while neutral data is not used because it does not provide important information. . In this research, all data is divided into 5 combinations of ratios of training data and testing data as in the table below.

Table 3. Ratio of training data and testing data

No	Amount of data	Ratio (%) Training: Testing	Data Training	Data Testing
1	3600	20:80	720	2880
2		40:60	1440	2160
3		60:40	2160	1440
4		80:20	2880	720
5		90:10	3240	360

d. Data Mining (Testing)

The results of data set classification using K-fold cross validation and the k value 10 times are presented by a 2x2 matrix and is called a confusion matrix. There are several values that can be used to measure performance in classification in this research, including accuracy, precision and recall values which are usually displayed in percentages.

Table 4. Confusion Table

		Actual Value	
		TRUE	FALSE
Predicted Value	TRUE	TP (True Positive) Correct Result	FP (False Negative) Unexpeted Result
	FALSE	FN (False negative) Missing Result	TN (True Negative) Correct Absence of Result

The simplest performance measure is accuracy. The overall effectiveness of the algorithm is calculated by dividing the correct labels over all classifications. Accuracy is a measure of recognition that states the percentage of the amount of data in the test data that is correctly classified by the classifier.

$$\text{Accuracy Formula} = \frac{TP+TN}{TP+TN+FP+FN} * 100\%$$

Recall, also called True Positive Rate (TPR), is the ratio of the number of positive examples classified over all positive examples.

$$\text{Recall Formula} = \frac{TP}{TP+FN} * 100\%$$

Precision is the number of positive examples classified above all classified examples.

$$\text{Precision Formula} = \frac{TP}{TP+FP} * 100\%$$

Receiver Operating Characteristics (ROC) is used to evaluate accuracy results in graphical form. ROC is a curve that will produce an Area Under Cover (AUC) value. AUC is the accuracy value of the area under the curve produced by ROC.

e. Interpretation / Evaluation

To illustrate how well machine learning performs in classifying data, one of the methods used to measure classification performance is the confusion matrix. The confusion matrix contains true positive, true negative, false positive and false negative values. The values of true positive and true negative provide information that when the classifier is classifying the data the value is correct, and while false negative and false positive provide information that when the classifier is wrong in classifying the data. Effective measurements can be made by calculating accuracy values, precision values and recall values.

RESULT

Simulation of Manual Algorithm Calculations

For example, there are four documents that have gone through the preprocessing stage, two documents are taken from the positive class and two documents are taken from the negative class. The document is as follows.

Document 1: Jokowi was cool when he had Indosat buyback

Document 2: easy, fast, no hassle Indosat credit via SMS

Document 3: Indosat net is slow

Document 4: Indosatcare Indosat's bad network is really slow

Table 5. Frequency of Word Appearance

Query	Document			
	1	2	3	4
jokowi	1	0	0	0
Cool	1	0	0	0
O'clock	1	0	0	0
he	1	0	0	0
indosat	1	1	1	1
buyback	1	0	0	0
credit	0	1	0	0
easy	0	1	0	0
fast	0	1	0	0
No	0	1	0	0

complicated	0	1	0	0
via	0	1	0	0
text	0	1	0	0
net	0	0	1	1
slow	0	0	1	1
indosatcare	0	0	0	1
Bad	0	0	0	1
mulu	0	0	0	1
very	0	0	0	1

Based on the table above, it is known that the positive class consists of 2 documents with a total of 14 words from 19 existing vocabularies, while the 2 negative class documents consist of 10 words from 19 existing vocabularies. Based on the number of words, the probability value for each class can be calculated using the following equation.

a. Probability of positive class words

Example of probability calculation for the word "indosat" which is in the positive class.

$$P(w_i|c) = \frac{\text{count}(w_i, c) + 1}{(\sum_{w \in V} \text{count}(w, c) + |V|)}$$

$$P(\text{indosat positif}) = \frac{2 + 1}{14 + 19} = 0,09091$$

Where:

w_i : number of words "indosat" in positive class documents

w : the number of all words in the positive class document

$|V|$: number of words in the training phase

The probability values for other words, using the same calculation method, are shown in the following table.

Table 6. Probability of Positive Word Classes

Query	Probabilitas	Query	Probabilitas
jokowi	0,06061	complicated	0,06061
Cool	0,06061	via	0,06061
O'clock	0,06061	text	0,06061
he	0,06061	net	0,03030
indosat	0,09091	slow	0,03030
buyback	0,06061	indosatcare	0,03030
credit	0,06061	Bad	0,03030
easy	0,06061	mulu	0,03030
fast	0,06061	very	0,03030
no	0,06061		

b. Negative class word probability

Example of calculating the probability of the word "indosat" being negative.

$$P(w_i|c) = \frac{\text{count}(w_i, c) + 1}{(\sum_{w \in V} \text{count}(w, c) + |V|)}$$

$$P(\text{indosat negatif}) = \frac{2 + 1}{10 + 19} = 0,10345$$

Where,

w_i : number of words "indosat" in negative class documents

w : the number of all words in the negative class document

$|V|$: number of words in the training phase

The probability values for other words in the negative class are shown in the table below.

Table 7. Probability of Negative Word Classes

Query	Probabilitas	Query	Probabilitas
Jokowi	0,03448	ribet	0,03448
Keren	0,03448	via	0,03448
Jam	0,03448	sms	0,03448
Beliau	0,03448	jaring	0,10345
Indosat	0,10345	lemot	0,10345
Buyback	0,03448	indosatcare	0,06897
Pulsa	0,03448	jelek	0,06897
Mudah	0,03448	mulu	0,06897
Cepat	0,03448	banget	0,06897
Nggak	0,03448		

Test Results

The results of the response data analysis are used to find out a general picture of the response data about Indosat after going through text preprocessing and automatic labeling by the lexicon dictionary. The response data obtained from 5 October 2021 to 8 October 2021 was categorized into three categories, namely positive responses, negative responses and neutral responses. Overall, an overview of the responses obtained based on each category is shown in the following graph.

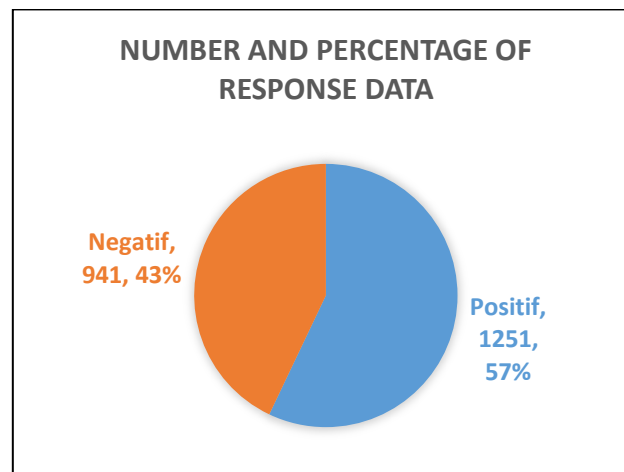


Figure 3 Graph of Number and Percentage of Response Data

Based on the graph above, it is known that of the total responses obtained from Twitter, namely 2192 responses, as many as 1251 (57%) responses were positive responses, namely responses that supported or liked Indosat, while 941 (43%) responses were negative responses, namely responses that did not like something matters related to Indosat. Judging from the percentage of response data regarding Indosat, positive responses are greater than negative responses with a difference of 310 (14%) responses.

To test the performance of machine learning in classifying, cross validation was carried out using 10-fold cross validation so that an average accuracy result of 85.00% was obtained, in more detail shown in the table and graph below.

Table 8. Accuracy of Indosat using 10-fold cross validation

No	Comparison Ratio (%)	Accuracy Naive Bayes
1	20:80	0,8541667
2	40:60	0,8518519
3	60:40	0,8486111
4	80:20	0,8416667
5	90:10	0,8541667
Average		0,8500926

Based on Table 8, it can be presented in the following graphical form.

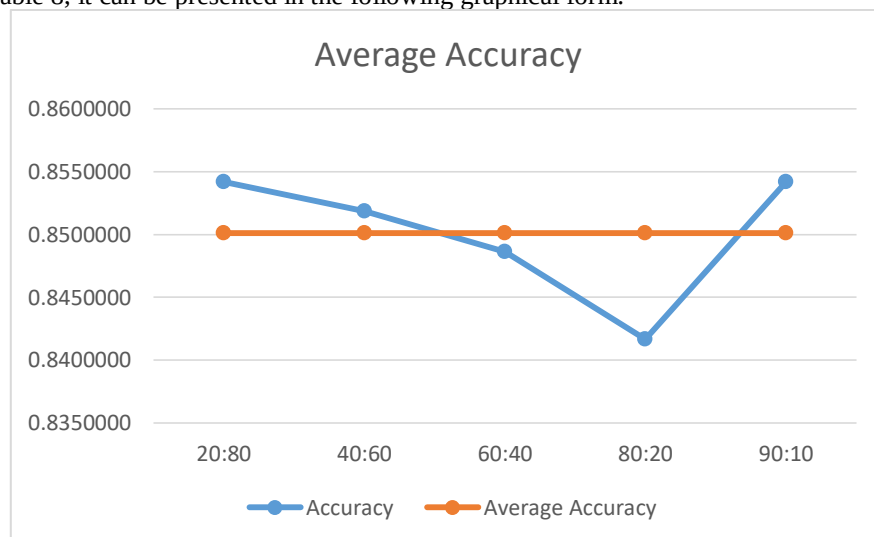


Figure 4. Accuracy graph Indosat uses 10-fold cross validation

Based on Figure 4.2. It is known that the accuracy level of the Naïve Bayes algorithm ranges from 84.16% - 85.41% or an average accuracy level of 85.00%. The Naive Bayes algorithm experienced the highest increase in accuracy when using a ratio of 20:80 and 90:10, then experienced a decrease in the lowest level of accuracy at a ratio of 80:20.

In the results of the classification of positive responses regarding Indosat, from a total of 1251 positive reviews, the 10 words that appeared the most were the words "gb" with a frequency of 3615 times, "bonus" 725 times, "quota" 618 times, "Hockey" 514 times, "cheap" 274 times, "jam" 273 times, "wa" 266 times, "full" 260 times, "check" 245 times, and "credit" 214 times. The words that appear in Figure 4.3. are 10 words that have a positive sentiment and are the topics of conversation most discussed by Twitter users regarding Indosat.

DISCUSSION

Sentiment analysis of cellular operator services, such as Indosat, using data from Twitter can provide important insights into customer perceptions. The Naïve Bayes algorithm is a method that is often used for sentiment analysis because of its simplicity and effectiveness. The following are general steps that can be taken to carry out sentiment analysis of Indosat services on Twitter using the Naïve Bayes algorithm:

- Data Collection**
Twitter API: Use the Twitter API to collect tweets that mention "Indosat" or other related keywords.
Date and Location: Make sure to include filters based on date range and location if necessary.
- Data Preprocessing**
Data Cleaning: Remove noise such as irrelevant URLs, hashtags and mentions.
Tokenization: Split text into individual words or tokens.
Lowercasing: Convert all text to lowercase for consistency.
Stop Words: Remove common words that don't carry much meaning (for example, "and", "or").
Stemming/Lemmatization: Reducing words to their basic forms.
- Data Labeling**
Manual Labeling: Manually label a number of tweets as positive, negative, or neutral to build a training dataset.
Crowdsourcing: Use a crowdsourcing platform if the data is too large to label manually.
- Data Sharing**
Training and Test Set: Split the dataset into a training set and a test set (for example, 80% for training and 20% for testing).
- Feature Extraction**
Bag of Words (BoW): Represent text as a vector based on word frequency.
TF-IDF: Term Frequency-Inverse Document Frequency to give more weight to more significant words.
- Model Training**
Naïve Bayes: Train a Naïve Bayes model using training data. This algorithm relies on Bayes' Theorem with the assumption of independence between features.
Multinomial Naïve Bayes: Suitable for text data where features are word frequencies.
- Model Evaluation**

Accuracy: Percentage of correct predictions out of total predictions.

Precision, Recall, F1-Score: Other measures to evaluate model performance especially on imbalanced data.

Confusion Matrix: A matrix that describes the performance of a classification model.

8. Analysis of Results

Visualization: Use graphs to display sentiment distribution.

Insight: Analyze what sentiment is showing up the most (positive, negative, neutral), specific times when sentiment changes, etc.

9. Implementation and Monitoring

Real-Time Monitoring: Deploy models to monitor sentiment in real-time.

Feedback Loop: Use feedback from monitoring results to periodically improve the model.

CONCLUSION

Sentiment analysis using Naïve Bayes on Twitter data can provide a clear picture of customer perceptions of Indosat services. These steps can be adapted and further adjusted according to specific project needs and data availability. Based on the results of the analysis and discussion in the previous chapter, the author can draw several conclusions that by using a comparison of training data and test data of 20% : 80%, 40% : 60%, 60% : 40%, 80% : 20% and 90% : 10% obtained from sentiment classification results using the Naïve Bayes method and using 10-fold cross validation obtained an average accuracy value of 85.00%. The decrease in machine learning performance occurs in a ratio of 80:20 or 1440 training data: 360 testing data, while the ratio ratio of 20%:80% and 90%:10% has the same accuracy value, namely 85.41%.

REFERENCES

- Asro'i, A., & Februariyanti, H. (2022). ANALISIS SENTIMEN PENGGUNA TWITTER TERHADAP PERPANJANGAN PPKM MENGGUNAKAN METODE K-NEAREST NEIGHBOR. *Jurnal Khatulistiwa Informatika*, 10(1), 17–24. <https://doi.org/10.31294/jki.v10i1.12624>
- Aulawi, H., Kurniawan, W. A., & Syaeful Azhar, A. (2022). Analisis Sentimen Terhadap Layanan INDOSAT pada Media Sosial Twitter Selama Pandemi. *Jurnal Kalibrasi*, 20(1), 53–59. <https://doi.org/10.33364/kalibrasi/v.20-1.1114>
- Cut Medika Zellatifanny. (2019). Respon Pengguna Twitter terhadap Regulasi Pengendalian Akses Ponsel Ilegal melalui Validasi IMEI (Twitter User's Response to Regulation of Contraband Cell Phone Access Control through IMEI Validation). *Jurnal Ilmu Pengetahuan Dan Teknologi Komunikasi*, 21(2).
- Darwis, D., Siskawati, N., & Abidin, Z. (2021). PENERAPAN ALGORITMA NAIVE BAYES UNTUK ANALISIS SENTIMEN REVIEW DATA TWITTER BMKG NASIONAL. *Jurnal Tekno Kompak*, 15(1), 131. <https://doi.org/10.33365/jtk.v15i1.744>
- Elsa Annisa Batu Bara, Kartika Amelia Nasution, Rafika Zahara Ginting, & Kartini Kartini. (22 C.E.). Penelitian tentang Twitter. *Jurnal Edukasi Non Formal*, 3(2), 167–172.
- Hafiz Irsyad, Ahmad Farisi, & Muhammad Rizky Pribadi Mail. (2019). Klasifikasi Opini Masyarakat Terhadap Jasa ISP MyRepublic dengan Naïve Bayes. *Jurnal Nasional Teknik Elektro Dan Teknologi Informasi*, 8(1), 30–34.
- Maharani Nirmala Dewi, & Ricky Eka Putra. (2023). Pengembangan Sistem Analisis Sentimen Masyarakat Terhadap Chatgpt Pada Twitter Dengan Perbandingan Metode Naive Bayes Classifier Dan K-Nearest Neighbors. *Journal of Informatics and Computer Science*, 5(3), 344–357.
- Megawati, S. (2021). Pengembangan Sistem Teknologi Internet of Things Yang Perlu Dikembangkan Negara Indonesia. *Journal of Information Engineering and Educational Technology*, 5(1), 19–26. <https://doi.org/10.26740/jieet.v5n1.p19-26>
- Muliani, R., Solehudin, A., & Jamaludin, A. (2024). ANALISIS SENTIMEN TERHADAP PROVIDER XYZ DI TWITTER MENGGUNAKAN ALGORITMA NAÏVE BAYES CLASSIFIER. *JATI (Jurnal Mahasiswa Teknik Informatika)*, 7(4), 2841–2848. <https://doi.org/10.36040/jati.v7i4.7191>
- Nizam Haqqizar, & Tika Nur Larasyanti. (2019). Analisis Sentimen Terhadap Layanan Provider Telekomunikasi Telkomsel Di Twitter Dengan Metode Naïve Bayes. *Prosiding TAU SNAR-TEK 2019 Seminar Nasional Rekayasa Dan Teknologi*, 30–33.
- Putranti, N. D., & Winarko, E. (2014). Analisis Sentimen Twitter untuk Teks Berbahasa Indonesia dengan Maximum Entropy dan Support Vector Machine. *IJCCS (Indonesian Journal of Computing and Cybernetics Systems)*, 8(1), 91. <https://doi.org/10.22146/ijccs.3499>
- Ronny Julianto, Evi Dianti Bintari, & Indrianti. (2017). Analisis Sentimen Layanan Provider Telepon Seluler pada Twitter menggunakan Metode Naïve Bayesian Classification. *Big Data Analytic and Artificial Intelligence*, 3(1), 23–30.
- Rozi, I. F., Firdausi, A. T., & Khalimatul Islamiyah. (2023). Analisis Sentimen Pada Twitter Mengenai Pasca Bencana Menggunakan Metode Naïve Bayes Dengan Fitur N-Gram. *Jurnal Informatika Polinema*, 6(2), 33–39. <https://doi.org/10.33795/jip.v6i2.316>

-
- Syailendra Reza Irwansyah Rezeki, Yuliana Restiviani, & Rita Zahara. (2020). Penggunaan Sosial Media Twitter dalam Komunikasi Organisasi (Studi Kasus Pemerintah Provinsi Dki Jakarta Dalam Penanganan Covid-19). *Journal Of Islamic And Law Studies*, 4(2), 63–78.
- Wahidah, A. R., Bachtar, Y., & Wulan, R. (2022). Sistem Pendukung Analisa Key Performance Indicator (KPI) Menggunakan Metode Data Mining Berbasis Web Python Programming. *JRKT (Jurnal Rekayasa Komputasi Terapan)*, 2(03). <https://doi.org/10.30998/jrkt.v2i03.7971>