



Text Mining-Based Sentiment Analysis of ChatGPT Users on X Platform Using Naïve Bayes Algorithm

Chaerulsyah Alkamal^a, Dede Kurniadi^{a,*}

^a Department of Computer Science, Institut Teknologi Garut, Garut, Indonesia

Corresponding author: dede.kurniadi@itg.ac.id

Abstract— ChatGPT (Generative Pre-trained Transformer) is a natural language processing model based on Artificial Intelligence (AI) that is currently trending and rapidly transforming human–computer interactions. ChatGPT is widely used by the public because it is considered very helpful in completing various tasks such as writing, information retrieval, programming assistance, and problem-solving faced by society. Its growing popularity demonstrates the increasing reliance on AI-driven tools in everyday digital activities. However, as the use of ChatGPT expands, questions have arisen about how people perceive, evaluate, and respond to interactions with this AI system. The use of ChatGPT not only creates new opportunities for innovation, efficiency, and automation but also introduces challenges in understanding user trust, perceptions, and sentiments toward emerging AI technologies. Various controversies have also emerged regarding its reliability, ethical implications, and potential misuse. Therefore, this research aims to determine the public sentiment, particularly among users of social media platform X (formerly Twitter), toward ChatGPT, and to analyze whether most users perceive it positively, negatively, or neutrally. By conducting sentiment analysis and implementing text mining, the study identifies the dominant sentiment trends within public discourse. The research applies the SEMMA (Sample, Explore, Modify, Model, Assess) methodology with Naïve Bayes as the classification algorithm. A Confusion Matrix is used to evaluate the model's performance. The results show that out of a total of 1,314 data points, 39.4% of sentiments were positive, 37.7% were neutral, and 22.9% were negative. The classification model achieved an accuracy rate of 72.78%, which is considered quite good. These findings indicate that the Indonesian public generally has a positive attitude toward ChatGPT, while also highlighting areas for improvement in understanding public perceptions of AI-based conversational systems.

Keywords— ChatGPT, Naive Bayes, SEMMA, Sentiment Analysis, Text Mining.

Manuscript received 25 Oct. 2025; revised 25 Oct. 2025; accepted 13 Nov. 2025. Date of publication 13 Nov. 2025.

International Journal of Applied Information systems and Informatics is licensed under a Creative Commons Attribution-Share Alike 4.0 International License.



I. INTRODUCTION

Artificial Intelligence (AI) technology is rapidly advancing, particularly in the field of Natural Language Processing (NLP), where the ability to process and understand human language is rapidly improving. One significant development in NLP is the emergence of ChatGPT (Generative Pre-trained Transformer) [1]. ChatGPT is a robot or chatbot that utilizes artificial intelligence, can interact, and perform various tasks [2]. ChatGPT can be used for various purposes, including translating languages, generating original text, supporting programmers in solving coding problems, simplifying concept explanations, drafting and outlining articles, and various other functions that can simplify the user's work [3].

However, despite its advantages, the use of AI technology also carries several risks. One of these is the inability of AI to replace human capabilities in tasks that require creativity and empathy. Furthermore, data security and privacy are also major issues that need to be considered in the use of AI technology [4]. Another example is the controversy surrounding ChatGPT in education or academia. Many

believe it is morally wrong to use ChatGPT to obtain answers to tests. Many also consider this an innovative advancement that deserves appreciation, because students also have more variety in obtaining sources and tools that can be used to facilitate their learning activities [5]. In this regard, it is necessary to conduct sentiment analysis to gain a deeper understanding of how the Indonesian public, especially users of social network X, respond to ChatGPT.

X is a social media platform primarily used for sending messages called tweets. Research conducted by We Are Social and Meltwater indicates that X had 24.69 million users in Indonesia in 2024, making it one of the top 10 most widely used social media platforms in Indonesia [6]. This research data can serve as a reference for utilizing comments or tweets from X users as a source of research data.

Sentiment analysis refers to the automated process of interpreting, extracting, and analyzing textual data to gain insights about the orientation of opinions toward a particular object, determining whether those opinions are positive, negative, or neutral. [7]. The significant contribution and relevance of sentiment analysis in various fields have driven significant improvements in the development of research and

its implementation [8]. Sentiment analysis is closely related to the field of text mining. Text mining plays a crucial role in sentiment analysis, particularly as a means of identifying the emotional content of a statement. Therefore, various studies related to sentiment analysis have been conducted [9].

The main objective of text mining is to extract useful information from a collection of documents, while data mining focuses on supporting the process of discovering knowledge from large datasets [10]. The text mining process generally consists of four stages. The first stage is folding, which converts all letters in the document to lowercase and accepts only characters from “a” to “z.” The second stage is tokenizing, which breaks the text into individual words. The third stage, called filtering, involves selecting important words after the tokenization process. The fourth or final stage is analysis, which aims to determine the extent of the relationships between words within the documents [11]. As a branch of data mining, text mining is considered to have a higher commercial value since approximately 80% of company information is stored in text form [12].

Several previous studies have been used as references for this research. The first study, conducted [13], focused on sentiment analysis related to cryptocurrency trends—specifically Bitcoin—by employing data crawling techniques from Twitter, resulting in 1,998 collected data points. The findings revealed that 46.69% of the sentiments were neutral, 43.54% positive, and 9.75% negative. Another study by [14] compared various classification algorithms and feature selection methods applied in sentiment analysis, including Deep Learning, Decision Tree, Naïve Bayes, and K-Nearest Neighbour (KNN), as well as feature selection techniques such as Information Gain and Chi-Square. The results demonstrated that Deep Learning achieved the highest performance with an accuracy rate of 78.43% and a kappa value of 0.625. Meanwhile, Information Gain produced the best feature selection results, achieving an average accuracy of 63.79% and a kappa value of 0.382. The combination of the most effective classification algorithm and optimal feature selection technique resulted in an accuracy of 78.63% and a kappa value of 0.626, which falls under the fair classification category. Additionally, a third study by Hasri and Alita [15] examined public sentiment toward the effects of the coronavirus and compared the performance of the Support Vector Machine (SVM) algorithm with the Naïve Bayes Classifier. Using a dataset of 1,104 tweets, the study found that the Naïve Bayes algorithm achieved an accuracy of 81.07%, while the Support Vector Machine reached 79.96%.

The fourth study, conducted by [16], explored sentiment analysis on the National BMKG using the Naïve Bayes Classifier algorithm. The research utilized 1,179 data points gathered from posts or tweets made by users on social media platform X. The developed sentiment analysis model classified the data into three categories—positive, negative, and neutral—with an achieved accuracy rate of 68.97% based on the evaluation results. Meanwhile, the final study by [17] examined Indonesian public sentiment toward the performance of the House of Representatives (DPR) using the Naïve Bayes Classifier approach. The findings revealed 95 positive tweets with a polarity value of 0.75 (equivalent to 75% positive sentiment), 693 neutral tweets with a polarity of 0.79 (79% neutral sentiment), and 758 negative tweets with a

polarity of 0.82 (82% negative sentiment). The sentiment analysis achieved an accuracy score of 0.8 (80%), obtained from testing that used 20% of the total dataset, which comprised 1,546 data entries.

This study aims to analyze sentiment related to ChatGPT based on the opinions of Indonesian people on social media X by implementing text mining and creating a classification model with the Naïve Bayes algorithm from the results of the sentiment analysis. The aim of this study is to determine the sentiment of the Indonesian people, particularly social media X users, regarding ChatGPT, specifically whether it is perceived as positive, negative, or neutral.

II. MATERIALS AND METHODS

The research stages in this study follow the SEMMA methodology. The SEMMA data mining technique is divided into five stages: Sample, Explore, Modify, Model, and Assess [18]. The SEMMA method is suitable for sentiment analysis because it provides a systematic workflow that supports the entire text data processing pipeline, from sampling review data, exploring word patterns and frequencies, modifying data through text cleaning and transformation (such as tokenization and stopword removal), to modeling using classification algorithms like Naive Bayes or SVM, and finally evaluating model performance to ensure sentiment accuracy. The step-by-step structure of SEMMA helps researchers maintain consistency, efficiency, and validity in sentiment analysis results [19]. The following illustration shows the stages of the SEMMA methodology.

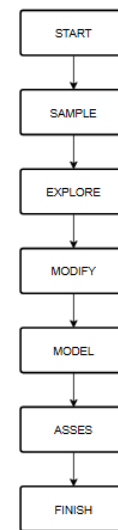


Fig. 1 Stages of the SEMMA Methodology

A. Sample

At this stage, a literature study and data collection will be carried out as a basis for compiling the research:

1) Literature Study

At this stage, a literature review was conducted on topics relevant to the research, including concepts such as text mining, sentiment analysis, and the Naïve Bayes algorithm. Furthermore, a literature review of several relevant previous studies was conducted, examining both the application of sentiment analysis in social media and

the use of the Naïve Bayes algorithm in classifying short texts, such as tweets.

2) Data Collection

In this study, data collection was conducted using a data crawling method from social media platform X using Google Colab. To obtain data in the form of tweets relevant to the keyword ChatGPT, publicly available third-party source code was utilized.

B. Explore

At this stage, data identification and selection are performed from the previously obtained crawling results. Although search parameters were defined during the data collection process with specific keywords and language restrictions (lang:id), the crawling results still contained some inconsistencies. Therefore, a manual data selection process was performed during the exploration stage.

C. Modify

In the Modify stage, the obtained data is prepared through a series of processes to ensure it is ready for use in modeling. This process involves data cleaning, case folding, tokenization, stopword removal, and filtering to ensure the text is clean and relevant to the research context.

1) Data Cleaning

In the Data Cleaning stage, data is cleaned from noise, symbols, characters, and numbers so that the data is ready to be used[20].

2) Case Folding

Then, the case folding activity is carried out, which means converting all the text into lowercase letters[21].

3) Tokenizing

Tokenizing is the process of dividing text into smaller units called tokens, such as words or phrases, which are used as the basic elements for further text analysis[22].

4) Stopword Removal

Stopword removal is the process of eliminating common words that do not carry significant meaning, such as "and," "the," or "is," to improve the focus and efficiency of text analysis[23].

5) Filtering

Filtering is the process of refining the tokenized text by applying certain criteria to remove unnecessary or irrelevant tokens. Specifically, filtering by length involves eliminating tokens that are too short or too long, so that only words with meaningful lengths are kept for further analysis[23].

Then after the data goes through the Preprocessing stage, data labeling is performed to assign sentiment categories to each data item. Data labeling is the process of assigning predefined categories or labels to each piece of data, such as marking text as positive, negative, or neutral, to prepare it for model training.

D. Model

Then, at this model stage, the model creation process will be carried out using the naive Bayes algorithm. Naïve Bayes is a probability- and statistics-based classification method introduced by the English scientist Thomas Bayes to predict

future probabilities based on past experiences[24]. The steps in the Naïve Bayes process are as follows [25]:

1) Calculating the Prior Value

Calculate the number of classes by finding the average using equation (1):

$$P = \frac{X}{A} \quad (1)$$

Description:

P = Prior value

X = Number of data in each class

A = Total number of data across all classes

2) Calculating the Likelihood Value

Calculate the number of cases that match within the same class using equation (2):

$$L = \frac{F}{B} \quad (2)$$

Description:

L = Likelihood

F = Number of feature data in each class

B = Total number of features in each class

3) Calculating the Posterior Value

Multiply all variable results from each existing class using equation (3):

$$P(C|A) = P(C) \times P(A|C) \quad (3)$$

Description:

P(C|A) = Posterior value

P(C) = Prior value for each class

P(A|C) = Likelihood value

E. Asses

At this asses stage, a model evaluation will be carried out using a confusion matrix. Confusion Matrix is a table that states the classification of the number of correct test data and the number of incorrect test data [26]. Based on the confusion matrix process, the values of precision, recall, F1-score for the three sentiment classes and the overall accuracy are obtained [27].

- 1) Precision is the ratio of correctly predicted positive observations to the total predicted positive observations, as expressed in the following equation:

$$\text{Precision} = \frac{TP}{TP + FP} \times 100\% \quad (4)$$

- 2) Recall is the ratio of correctly predicted positive observations to all actual positive observations, formulated as follows:

$$\text{Recall} = \frac{TP}{TP + FN} \times 100\% \quad (5)$$

- 3) F1-Score represents the weighted average between precision and recall, expressed by the formula:

$$F1\ Score = 2 \times \frac{Recall \times Precision}{Recall + Precision} \quad (6)$$

- 4) Accuracy is the ratio of correctly predicted observations (positive, negative, and neutral) to the total number of observations, defined as:

$$Accuracy = \frac{TP+TN+TNt}{TP+TN+TNt+FP+FN+FNt} \times 100\% \quad (7)$$

III. RESULT AND DISCUSSION

A. Sample

At this stage, a literature study and data collection will be carried out as a basis for compiling the research:

1) Literature Study

The stages of literature study carried out produced several theories that can be used as references in research.

2) Data Collection

The dataset to be used is the result of crawling from the limited social media site X from January 1, 2025, to August 30, 2025, comprising a total of 2,000 data points. The data frame contains three attributes, the attributes are:

1. Created_at: contains the time when the user made the tweets.
2. Full_text: contains tweets that contain sentiment towards ChatGPT
3. Tweet_url: contains links to tweets created by users.

created_at	full_text	tweet_url
Sat Mar 29 23:56:38 +0000 2025	@Tchaikovsky Same same. Bahkan aku gak update ada AI apa aja selain chatgpt sm cai	https://x.com/undefined/status/1906133429814730951
Sat Mar 29 23:42:46 +0000 2025	memang betul pun One of the concerns GenAI ni adalah fitnah. But tech bros mana ada comprehension skill. Kita cakup lain dia cari gaduh pasal is	https://x.com/undefined/status/1906129939784720721
Sat Mar 29 23:29:26 +0000 2025	@lincheFaris masalahnya buat jem trafik geng2 ni penuh dkt tl https://x.com/undefined/status/1906126583087399253	https://x.com/undefined/status/1906126583087399253
Sat Mar 29 23:27:03 +0000 2025	ChatGPT nuk pot qojka ne finale. #bbvipa4	https://x.com/undefined/status/190612598434632308
Sat Mar 29 23:25:05 +0000 2025	ChatGPT ni nga def sama kanam bi	https://x.com/undefined/status/1906125489925864058
Sat Mar 29 23:23:58 +0000 2025	gugel translate sm chatgpt aja ga sejelek ini	https://x.com/undefined/status/1906123209545139164
Sat Mar 29 23:23:10 +0000 2025	Ku bingung sm org yg buat fanfic otp pake chatgpt kek lu yakin https://x.com/undefined/status/1906125007820071258	https://x.com/undefined/status/1906125007820071258
Sat Mar 29 23:20:21 +0000 2025	@yonimiran Pelo ChatGpt bb	https://x.com/undefined/status/190612429707899722
Sat Mar 29 23:16:09 +0000 2025	@Frederett. Tidak hari tanpa curhat dengan chatgpt. Dia resp https://x.com/undefined/status/1906123240206160205	https://x.com/undefined/status/1906123240206160205

Fig. 2 Crawl Result Data

B. Explore

During the exploration stage, an attribute selection process is conducted to determine which features will be used for model building. At this stage, only the full text attribute is selected, as it contains the main textual data required for sentiment analysis. Other attributes are excluded because they are not relevant to the classification process and may introduce noise into the model.

full_text
@Tchaikovsky Same same. Bahkan aku gak update ada AI apa aja selain chatgpt sm cai
memang betul pun One of the concerns GenAI ni adalah fitnah. But tech bros mana ada comprehension skill. Kita cakup lain dia cari gaduh pasal is
@lincheFaris masalahnya buat jem trafik geng2 ni penuh dkt threads amik kmpatan buat diit.... chatgpt pun jd slow org nk convert suka2 buat kni
ChatGPT nuk pot qojka ne finale. #bbvipa4
ChatGPT ni nga def sama kanam bi
gugel translate sm chatgpt aja ga sejelek ini
Ku bingung sm org yg buat fanfic otp pake chatgpt kek lu yakin itu ga otp lu g jd ooc?? #6,
@yonimiran Pelo ChatGpt bb
@Frederett. Tidak hari tanpa curhat dengan chatgpt. Dia responsi sih. Nggak sabar ya sabar ya aja.

Fig. 3 Explore Result Data

C. Modify

In the Modify stage, the obtained data is prepared through a series of processes to ensure it is ready for use in modeling. This process includes data cleaning, case folding, tokenizing, stopword removal, and filtering to ensure the text is clean.

1) Data Cleaning

In the data cleaning stage, several preprocessing steps are performed to ensure that the dataset is free from noise and inconsistencies. This process includes removing unnecessary characters, punctuation marks, numbers, and duplicate entries that may affect the accuracy of the analysis.

Table 1 below is the result of the data cleaning stage:

TABLE I
DATA CLEANING RESULT

Before	After
@lucllabie Jujur gw kurang tau coba tanya chatgpt aja	Jujur gw kurang tau coba tanya chatgpt aja

2) Case Folding

In the case folding stage, all letters in the text are converted to lowercase to ensure uniformity in word representation. This process prevents the model from treating words with different letter cases as distinct terms, such as “ChatGPT” and “chatgpt.” By applying case folding, the text becomes more consistent and easier to process in subsequent steps such as tokenization.

Table 2 below is the result of the case folding stage:

TABLE 2
RESULTS OF CASE FOLDING

Before	After
Jujur gw kurang tau coba tanya chatgpt aja	jujur gw kurang tau coba tanya chatgpt aja

3) Tokenizing

In the tokenization stage, the cleaned text data is divided into smaller units called tokens, which usually represent individual words. This process helps transform unstructured text into a structured format that can be analyzed computationally. By splitting sentences into tokens, the model can better understand the frequency and distribution of words within the dataset.

TABLE 3
DATA TOKENIZATION RESULTS

Before	After
jujur gw kurang tau coba tanya chatgpt aja	['jujur' 'gw' 'kurang' 'tau' 'coba' 'tanya' 'chatgpt' 'aja']

4) Stopword Removal

In the stopword removal stage, commonly used words that do not contribute significant meaning to the analysis are removed from the text. These words are typically filtered out because they appear frequently but provide little value in determining sentiment or context. By eliminating stopwords, the dataset becomes more focused on meaningful words that carry emotional or contextual weight, thereby improving the

efficiency and accuracy of the sentiment classification process.

TABLE 4
RESULTS OF STOPWORD REMOVAL

Before	After
jujur gw kurang tau coba tanya chatgpt aja	jujur gw kurang tau coba tanya chatgpt

5) Filtering

In the filtering stage, specifically the filter by length process, words are filtered based on the number of characters they contain. This stage is intended to eliminate very short words, typically those consisting of one or two letters, as they generally provide little to no meaningful contribution to the sentiment analysis. By setting a minimum word length threshold, only words with sufficient semantic value are retained in the dataset. This process helps reduce noise, improve text quality, and ensures that only relevant tokens are used in the subsequent stages of model building.

TABLE 5
RESULTS OF FILTER BY LENGTH

Before	After
jujur gw kurang tau coba tanya chatgpt	jujur kurang tau coba tanya chatgpt

6) Data Labeling

The data labeling process in this study was conducted with the assistance of a chatbot. The chatbot used to assist the labeling process here is the Gemini chatbot. The Gemini chatbot was chosen because it has demonstrated good accuracy in comparisons with sentiment validated by Ms. Sri Utami, S.Pd., an Indonesian language expert and graduate of the Indonesian University of Education (UPI).

The following is a comparison of the responses of several chatbots with sentiments that have been validated by linguists.

1. full_text	sentimen ChatGPT	Gemini	Claude	Copilot	Chatgpt	Gemini	Claude	Copilot
2. gila betul chatgpt ni	positif	positif	positif	positif	1	1	1	1
3. chatgpt nglag banget	negatif	negatif	negatif	negatif	1	1	1	1
4. Pengetahuan chatgpt limitnya di data terakhir yg terkum	negatif	netral	netral	netral	0	0	0	0
5. Chatgpt gw kok gk bisa dibuka ya	netral	negatif	negatif	negatif	0	0	1	0
6. Mencoba OpenAI ChatGPT yang akan Menghapus Pekerja	negatif	negatif	negatif	netral	1	1	0	0
7. cobain chatgpt kak sangat membantu tugas kuliah	positif	positif	positif	positif	1	1	1	1
8. ChatGPT bisa dipake buat curhat dan responnya lumayan	positif	positif	positif	positif	1	1	1	1
9. ChatGPT Bakal Menggelamkan Google	positif	negatif	positif	netral	0	1	0	0
10. Chat GPT akan menggantikan konsultan	negatif	negatif	negatif	netral	1	1	0	0

Fig. 4 Labeling Comparison Data

Based on Figure 4 above, we can see that the comparison data between the chatbot's responses and the linguist's responses contains full_text, which is crawled data including sentiment toward ChatGPT and sentiment validation from the linguist. Meanwhile, ChatGPT, Gemini, Claude, and Copilot represent some of the responses from each chatbot regarding the existing sentiment. A column containing the number 1 indicates that the chatbot's response has been validated by the interview, while a column containing 0 indicates the opposite.

The following is a comparison of the responses from the linguist and several chatbots.

1. full_text	sentimen ChatGPT	Gemini	Claude	Copilot	Chatgpt	Gemini	Claude	Copilot
92. Metode ini tidak terbatas untuk SEO saja tapi juga bisa un	netral	positif	positif	positif	0	0	0	0
93. ChatGpt ini emang bener bener gokil	positif	positif	positif	positif	1	1	1	1
94. Buat kelas manfaatn chatgpt dong tum	positif	positif	positif	positif	1	1	1	1
95. gue udah coba chatgpt dan ngaco banget sih wajar google	positif	positif	positif	positif	1	1	1	1
96. Cara Pakai ChatGPT yang Bisa Menjawab Semua Pertanya	positif	netral	netral	netral	0	0	1	0
97. Aplikasi ChatGPT yang Viral Fitur Chatbot yang Digadang	positif	netral	netral	netral	1	0	0	1
98. Dibalik ChatGPT ternyata ada python	netral	netral	netral	netral	1	1	1	0
99. busseeeet chatgpt lebih tau apa yang kita mau harusnya	positif	positif	positif	positif	1	1	1	1
100. Bill Gates memprediksi AI seperti ChatGPT yang semakin	negatif	negatif	negatif	negatif	1	1	1	1
Accuracy					70	74	66	74

Fig. 5 Labeling Comparison Data

Based on Figure 5 above, the comparison of responses from several chatbots shows that the Gemini and Copilot chatbots achieved the highest accuracy, with 74% and 70%, respectively, followed by the ChatGPT chatbot with 70%, and finally, the Claude chatbot with 66%.

Based on these results, two chatbots can be used as tools for the labeling process. However, in this study, Gemini was chosen because it can process more data simultaneously.

D. Model

After going through the pre-processing and sentiment labeling stages, a total of 1,314 data entries were obtained. Based on the initial classification results, the data were divided into three sentiment categories: positive sentiment (39.4%), neutral sentiment (37.7%), and negative sentiment (22.9%). If each percentage is converted to the actual number of data points, then positive sentiment amounts to approximately 518 data points, neutral sentiment at 496 data points, and negative sentiment at 301 data points. These results indicate that the positive sentiment category dominates, followed by neutral sentiment, while the negative sentiment class has the fewest data points. To maintain a balance in the number of data points in each sentiment category, 300 data points were taken proportionally from each class. Thus, a total of 900 data points were obtained, which were then used as the main dataset in the training and testing processes of the sentiment classification model.

Next, the RapidMiner Studio application was used as the data analysis platform for creating the sentiment classification model. This process was carried out by following the workflow steps shown in the following image.

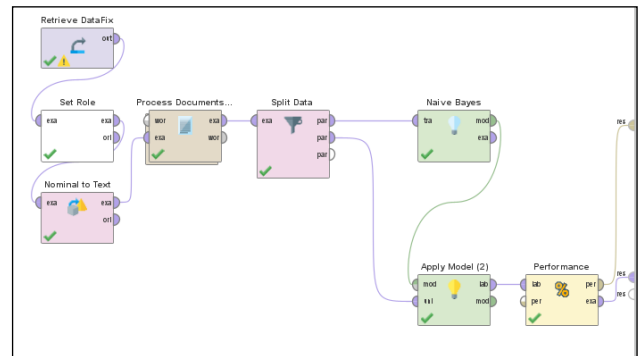


Fig. 6 Model Creation Stage

At this stage, the pre-processed dataset, consisting of 900 data points, 300 each for positive, neutral, and negative sentiment, was imported into Altair AI Studio for training and testing. The flowchart in the figure systematically illustrates the sequence of processes, from data loading and text processing to data splitting and the application of the classification algorithm used in this study. The data was divided using the split validation method, with a proportion of 80% for the training data and 20% for the testing data.

E. Asses

During the assessment stage, model evaluation is conducted using a confusion matrix to assess the performance

of the sentiment classification model. The confusion matrix provides detailed information on the number of correct and incorrect predictions for each sentiment class positive, neutral, and negative. From this matrix, several evaluation metrics are calculated, including precision, recall, F1-score, and overall accuracy. These metrics help determine how well the model distinguishes between different sentiment categories. By analyzing the confusion matrix, it becomes possible to identify potential weaknesses in the model and assess its overall reliability in accurately predicting sentiment.

The detailed results of the confusion matrix are shown in Figure 7.

accuracy: 72.78%				
	true negatif	true netral	true positif	class precision
pred. negatif	44	5	3	84.62%
pred. netral	8	37	7	71.15%
pred. positif	8	18	50	65.79%
class recall	73.33%	61.67%	83.33%	

Based on the model testing results shown in the figure, an accurate value of 72.78% was obtained. This value indicates that the model can classify data with an accurate level of approximately 73 out of every 100 tested data. Referring to the general model accuracy assessment category, an accuracy value of 72.78% is considered quite good, as it falls within the 70% to 80% range. This means that the model can recognize data patterns with an adequate level of reliability, although there is still room for improvement in its performance.

IV. CONCLUSION

Based on sentiment analysis of Indonesian public opinion data regarding ChatGPT, this study shows that most users have a positive view of ChatGPT. Of the total data analyzed, the proportions of positive, neutral, and negative sentiment were 39.4%, 37.7%, and 22.9%, respectively. These results indicate that the Indonesian public generally welcomes ChatGPT. The conclusion should summarize the discussion and outline the implications for future research.

Based on the results of the sentiment classification model testing, an accuracy value of 72.78% was obtained. This value indicates that the model has a fairly good level of accuracy in classifying data into positive, neutral, and negative sentiment categories. Therefore, it can be concluded that the model built using the applied naive Bayes algorithm performs quite reliably, although there is still room for improvement in accuracy through refinements in the data preprocessing stage or adjustments to model parameters.

For future research, it is recommended to expand the dataset to include more diverse social media platforms and time periods to capture a broader range of public sentiment. Researchers may also consider comparing multiple machine learning or deep learning algorithms, such as Support Vector Machines (SVM), Random Forest, or LSTM, to identify models with higher precision and recall. Additionally, incorporating linguistic features such as sarcasm detection, context analysis, and emotion classification could enhance sentiment interpretation accuracy. Further studies could also explore the correlation between sentiment trends and external factors such as media coverage or policy changes related to

AI, providing deeper insights into public perception dynamics toward ChatGPT and similar technologies.

REFERENCES

- [1] I. T. Julianto, D. Kurniadi, Y. Septiana, and A. Sutedi, "Alternative Text Pre-Processing using Chat GPT Open AI," *Janapati*, vol. 12, no. 1, pp. 67–77, 2023, [Online]. Available: <https://wjaets.com/content/artificial-intelligence-ai-based-chatbot-study-chatgpt-google-ai-bard-and-baidu-ai>
- [2] D. Setiawan, E. Ayu Dewi Karuniawati, S. Imelda Janty, and P. Bintan Cakrawala, "Peran Chat Gpt (Generative Pre-Training Transformer) Dalam Implementasi Ditinjau Dari Dataset," *Innovative: Journal Of Social Science Research*, vol. 3, no. 3, pp. 9527–9539, 2023.
- [3] W. Suharmawan, "Pemanfaatan Chat GPT dalam dunia pendidikan," *Education Journal: Journal Educational Research and Development*, vol. 7, no. 2, 2023, doi: 10.31537/ej.v7i2.1248.
- [4] Misnawati Misnawati, "ChatGPT: Keuntungan, Risiko, Dan Penggunaan Bijak Dalam Era Kecerdasan Buatan," *Prosiding Seminar Nasional Pendidikan, Bahasa, Sastra, Seni, Dan Budaya*, vol. 2, no. 1, pp. 54–67, Apr. 2023, doi: 10.55606/mateandrau.v2i1.221.
- [5] T. U. S. H. Wibowo, F. Akbar, S. R. Ilham, and M. S. Fauzan, "Tantangan dan Peluang Penggunaan Aplikasi Chat GPT Dalam Pelaksanaan Pembelajaran Sejarah Berbasis Dimensi 5.0," *JURNAL PETISI (Pendidikan Teknologi Informasi)*, vol. 4, no. 2, pp. 69–76, Jul. 2023, doi: 10.36232/jurnalpetisi.v4i2.4226.
- [6] S. Kemp, "DIGITAL 2024: INDONESIA," *DATAREPORTAL*. Accessed: Mar. 17, 2024. [Online]. Available: <https://datareportal.com/reports/digital-2024-indonesia>
- [7] R. Parlita, S. Ilham Pradika, A. M. Hakim, and R. N. M. Kholilul, "Analisis Sentimen Twitter Terhadap Bitcoin dan Cryptocurrency Berbasis Python TextBlob," *Jurnal Ilmiah Teknologi Informasi dan Robotika*, vol. 2, pp. 33–37, Dec. 2020, Accessed: Sep. 03, 2023. [Online]. Available: <https://scholar.archive.org/work/4klvx7lt5ndbhovnaol4exyym/acces/wayback/http://jifti.upnjatim.ac.id/index.php/jifti/article/download/22/24>
- [8] D. Alita and A. Isnain, "Pendeteksian Sarkasme pada Proses Analisis Sentimen Menggunakan Random Forest Classifier," *Jurnal Komputasi*, vol. 8, no. 2, Oct. 2020, doi: 10.23960/komputasi.v8i2.2615.
- [9] A. Z. Amrullah, A. Sofyan Anas, M. Adrian, and J. Hidayat, "Analisis Sentimen Movie Review Menggunakan Naive Bayes Classifier Dengan Seleksi Fitur Chi Square," *Jurnal Bumigora Information Technology (BITE)*, vol. 2, no. 1, pp. 40–44, Jul. 2020, doi: 10.30812/bite.v2i1.804.
- [10] S. Sukriadi, I. Ismail, and A. M. Andzar, "Penerapan Text Mining Dalam Klasifikasi Judul Skripsi Yang Diusulkan Mahasiswa Menggunakan Metode Naïve Bayes," *Jurnal Ilmiah Sistem Informasi Dan Teknik Informatika (JISTI)*, vol. 6, no. 2, pp. 184–196, 2023, doi: 10.57093/jisti.v6i2.174.
- [11] T. Darwansyah and I. Mansyur, "Implementasi Algoritma Text Mining dan Cosine Similarity untuk Desain Sistem Aspirasi Publik Berbasis Mobile," *Komputika Jurnal Sistem Komputer*, vol. 9, no. 2, pp. 169–176, 2022, doi: 10.34010/komputika.v11i2.6501.
- [12] A. Firdaus and W. Firdaus, "Text Mining Dan Pola Algoritma Dalam Penyelesaian Masalah Informasi: (Sebuah Ulasan)," *Jurnal Penelitian Ilmu dan Teknologi Komputer (JUPITER)*, vol. 13, no. 1, p. 66, Apr. 2021, Accessed: Sep. 04, 2023. [Online]. Available: <https://jurnal.polsri.ac.id/index.php/jupiter/article/view/3249>
- [13] I. T. Julianto, D. Kurniadi, M. R. Nashrulloh, A. Mulyani, and J. I. Komputer, "Twitter Social Media Sentiment Analysis Against Bitcoin Cryptocurrency Trends Using Rapidminer," *Jurnal Teknik Informatika (JUTIF)*, vol. 3, no. 3, 2022, doi: 10.20884/1.jutif.2022.3.3.343.
- [14] I. T. Julianto, D. Kurniadi, M. R. Nashrulloh, and A. Mulyani, "Comparison of Classification Algorithm and Feature Selection in Bitcoin Sentiment Analysis," *Jurnal Teknik Informatika (Jutif)*, vol. 3, no. 3, pp. 739–744, Jun. 2022, doi: 10.20884/1.jutif.2022.3.3.343.
- [15] C. F. Hasri and D. Alita, "Penerapan Metode Naïve Bayes Classifier Dan Support Vector Machine Pada Analisis Sentimen Terhadap Dampak Virus Corona Di Twitter," *Jurnal Informatika dan Rekayasa Perangkat Lunak*, vol. 3, no. 2, pp. 145–160, 2022.

- [16] D. Darwis, N. Siskawati, and Z. Abidin, "Penerapan Algoritma Naive Bayes Untuk Analisis Sentimen Review Data Twitter Bmkg Nasional," *Jurnal Tekno Kompak*, vol. 15, no. 1, pp. 131–145, 2021.
- [17] D. D. Putri, G. F. Nama, and W. E. Sulistiono, "Analisis Sentimen Kinerja Dewan Perwakilan Rakyat (DPR) Pada Twitter Menggunakan Metode Naive Bayes Classifier," *Jurnal Informatika dan Teknik Elektro Terapan*, vol. 10, no. 1, pp. 34–40, 2022.
- [18] Medika Risnasari, *Konsep Dasar Data Mining Teori dan Praktik dengan Python*. Malang: CV. Literasi Nusantara Abadi, 2022.
- [19] I. T. Julianto and M. A. Al Sidqi, "Automatic Sentiment Annotation Using Grok AI for Opinion Mining in a University Learning Management System," *Journal of Intelligent Systems Technology and Informatics*, vol. 1, no. 3, pp. 72–77, Nov. 2025, doi: 10.64878/jistics.v1i3.42.
- [20] Junadhi, Agustin, M. Rifqi, and M. K. Anam, "Sentiment Analysis of Online Lectures using K-Nearest Neighbors based on Feature Selection," *Jurnal Nasional Pendidikan Teknik Informatika (JANAPATI)*, vol. 11, no. 3, pp. 216–225, Dec. 2022, doi: 10.23887/janapati.v11i3.51531.
- [21] D. S. Utami and A. Erfina, "ANALISIS SENTIMEN PINJAMAN ONLINE DI TWITTER MENGGUNAKAN ALGORITMA SUPPORT VECTOR MACHINE (SVM)," in *SISMATIK (Seminar Nasional Sistem Informasi dan Manajemen Informatika)*, Aug. 2021, pp. 299–305.
- [22] J. Sangeetha and U. Kumaran, "A hybrid optimization algorithm using BiLSTM structure for sentiment analysis," *Measurement: Sensors*, vol. 25, p. 100619, Feb. 2023, doi: 10.1016/j.measen.2022.100619.
- [23] M. L. K. Harsono, Y. Alkhalifi, N. Nurajijah, and W. Gata, "Analisis Sentimen Stakeholder Atas Layanan HAIDJPB pada Media Sosial Twitter dengan Menggunakan Metode Support Vector Machine dan Naive Bayes," *Infoman's*, vol. 14, no. 1, 2020, doi: 10.33481/infomans.v14i1.126.
- [24] Y. B. Widodo, S. A. Anggraeni, and T. Sutabri, "Perancangan Sistem Pakar Diagnosis Penyakit Diabetes Berbasis Web Menggunakan Algoritma Naive Bayes," *Jurnal Teknologi Informatika dan Komputer*, vol. 7, no. 1, pp. 112–123, Mar. 2021, doi: 10.37012/jtik.v7i1.507.
- [25] L. Affandi, A. N. Pramudhita, and M. P. Sasmita, "Sistem Pakar Klasifikasi Kecanduan Gadget Menggunakan Teori Arthurt T. Hovart Dengan Metode Naive Bayes Classifier Untuk Anak Sekolah Dasar," in *Seminar Informatika Aplikatif Polinema (SIAP)*, Malang, Aug. 2020, pp. 102–106.
- [26] D. Normawati and S. A. Prayogi, "Implementasi Naive Bayes Classifier Dan Confusion Matrix Pada Analisis Sentimen Berbasis Teks Pada Twitter," *Jurnal Sains Komputer & Informatika*, vol. 5, no. 2, pp. 697–711, 2021, Accessed: Apr. 20, 2024. [Online]. Available: <http://ejurnal.tunasbangsa.ac.id/index.php/jsakti/article/view/369>
- [27] N. G. A. Dasriani, L. A. G. Pariandi, and I. M. Y. Dharma, "Analisis Sentimen Program Jaminan Kesehatan Nasional Menggunakan Multiclass Support Vector Machine," *BIOS: Jurnal Teknologi Informasi dan Rekayasa Komputer*, vol. 6, no. 1, pp. 20–30, Nov. 2024, doi: 10.37148/bios.v6i1.136.