

Sentiment Analysis of Netizens on Constitutional Court Rulings in the 2024 Presidential Election

^{*}¹Wahyudi Ariannor, ²Sami M A B Alshalwi, ¹Budi Susarianto

¹Informatics Engineering, STMIK Banjarbaru, Indonesia

²Science, Mathematics, and Computing Faculty, Bard College at Simon' Rock, United States

Corresponding Email: wahyu.arian@stmik-banjarbaru.ac.id

Abstract:

Online conversations among netizens play an important role in forming collective opinions and views about important events, including judicial decisions such as those taken by the Constitutional Court (MK). This research explores sentiment analysis of the Constitutional Court's decisions, especially in the context of the presidential election, using the Support Vector Machine (SVM), Logistic Regression, and Naive Bayes algorithms. Previous studies on public sentiment toward the Constitutional Court's decision provide a basis. Still, this research focuses on a different context, analysing sentiment toward the Constitutional Court's decision in the 2024 presidential election dispute. This study adopts an experimental methodology, involving several key stages such as data collection through Twitter web scraping, labelling, pre-processing, TF-IDF weighting, and algorithm testing. Evaluation using a confusion matrix shows comparable accuracy among SVM, Logistic Regression, and Naive Bayes, with SVM and Logistic Regression demonstrating superior precision and F1 scores. Negative sentiment carries greater weight than neutral and positive sentiment, highlighting potential social tensions and the need for effective communication and deeper analysis to understand the root causes of negativity. The SVM and logistic regression algorithms have proven effective in understanding public sentiment towards the Constitutional Court's decisions in a political context, providing valuable insights for understanding the dynamics of public opinion.

Keywords: *analysis, machine learning, netizens, precision, sentiment*

IJIE (Indonesian Journal of Informatics Education)

Vol 8 Issue 2 2024

DOI: <http://dx.doi.org/10.20961/ijie.v8i2.94614>

Received: 30/10/24 Revised: 27/12/24 Accepted: 28/12/24 Online: 31/12/24

© 2024 The Authors. Published by Universitas Sebelas Maret. This is an open-access article under the CC BY license (<http://creativecommons.org/licenses/by/4.0/>).

Introduction

In the ever-growing digital era, netizen activities and online conversations play a pivotal role in shaping collective opinions and perspectives on critical events, including decisions made by judicial institutions such as the Constitutional Court (MK). The significance of these decisions cannot be understated, as they have far-reaching implications for political stability, legal justice, and public trust in judicial institutions. Trust in judges and the judiciary influences the relationship between public confidence and the perception of justice (Widłak, 2022). However, there remains an informational gap regarding public sentiment toward the MK's decisions, particularly in the context of presidential elections.

Previous research has demonstrated the significant role of social media in capturing public sentiment on policies and events. For instance, a study on the Jakarta Online Zoning System utilized Twitter data to evaluate public reactions to zoning policies and highlighted social media's effectiveness as a medium for virtual community engagement and policy evaluation (Ratnawati & Iljas, 2021). Similarly, analysing netizen sentiment toward judicial decisions provides valuable insights into public opinion dynamics and societal reactions, particularly in politically sensitive contexts like presidential elections.

Sentiment is central to various communication studies, where negativity and polarization in political communication, product reviews, and public comments significantly influence economic conditions, social life, politics, and development issues integrated with information technology (Atteveldt et al., 2021). Sentiment analysis techniques determine a user's positive or negative impression of a topic from online communication platforms such as Twitter, Facebook, or YouTube (Vig et al., 2022). The application of sentiment analysis spans business and social domains, as opinions often drive individual behaviour and organizational decision-making (Liu, 2022). As social media users adopt more holistic communication habits, sentiment-based contextual analysis becomes increasingly crucial (Singgalen, 2021).

The inclusion of emojis has further enhanced the accuracy of sentiment analysis. Early efforts in this domain focused on single emojis Novak et al. (2015), but recent studies, such as "The Sentiment of Emoji Sets" (Othman et al., 2022), have explored sequences of multiple emojis, known as Emoji Sets. This more comprehensive approach to analysing emoji combinations has significantly improved sentiment score accuracy. However, our study diverges by not currently incorporating emojis into the model, leaving this as a consideration for future research.

As demonstrated in Sujana (2023), the selection of appropriate algorithms significantly influences the performance of sentiment analysis, especially in capturing contextual information within reviews. Sentiment analysis techniques also rely heavily on algorithmic advancements. The choice of algorithms significantly affects performance, especially in capturing contextual nuances in data. For example, Support Vector Machine (SVM), Logistic Regression, and Naive Bayes are widely used machine learning techniques for sentiment classification due to their ability to handle complex tasks effectively (Yuan et al., 2020). SVM, based on the structural theory of statistical learning, excels in pattern recognition and natural language processing (Oktavia, 2023; Wang, 2022). Logistic Regression is often employed to analyse relationships between predictor variables and a binary outcome (Husen et al., 2023). The logistic regression response variable has the values 1 (yes) and 0 (no), which produces a Bernoulli response variable (Husen et al., 2023). Meanwhile, Naive Bayes is particularly effective in classifying sentiments within texts that mix positive and negative statements (Ariannor et al., 2024).

Hariyanti et al. (2024) conducted research examining public reactions to the decision of the Indonesian Constitutional Court (MK), which maintains a minimum age limit of 35 years for presidential and vice-presidential candidates. The dataset used comprises 1,090 entries categorized into three groups: positive, neutral, and negative. The dataset was then divided into training and test sets with ratios of 60%:40%, 70%:30%, and 80%:20%. The Naïve Bayes method was used to analyse sentiment from Twitter data, achieving an accuracy of 67.98%.

Ningsih et al. (2024) discusses Twitter Sentiment Analysis on the Use of Electric Cars in Indonesia by comparing the SVM and Naïve Bayes Algorithms. This sentiment data was obtained from the Twitter social network, totalling 1,517 pieces of data. The data is labelled as positive or negative. The results show that the Support Vector Machine algorithm is more accurate than the Naïve Bayes algorithm, with an SVM accuracy of 70.83% and a Naïve Bayes accuracy of 63.02%.

Afandi and Isnaini (2024) applied the Support Vector Machine (SVM) algorithm and logistic regression to examine public trust in the context of a presidential election survey. The research analysed 1,778 Instagram comments and 985 Twitter tweets, each categorized as either positive or negative. The results show that SVM, with a scenario of 80% training data and 20% test data, provides high accuracy: 93.19% from Instagram and 91.19% from Twitter. Logistic regression also showed high accuracy, with 89.79% from Instagram and 88.01% from Twitter in the same scenario. Thus, it can be concluded that the SVM algorithm has better accuracy than the logistic regression algorithm.

Muzaki and Witanti (2021) conducted research titled "Sentiment Analysis of the Community on Twitter Regarding the 2020 Election During the COVID-19 Pandemic Using the Naive Bayes Classifier Method." Twitter was used as a medium to represent people's responses to public issues, particularly regarding the 2020 Pilkada amid the COVID-19 pandemic. Sentiment analysis involved the classification of tweets containing public sentiment about the issue using the Naive Bayes Classifier (NBC) and Support Vector Machine (SVM) methods. The test results show that Naive Bayes achieved better accuracy than Support Vector Machine, with an accuracy rate of 92.2%. This research also included pre-processing, tokenizing, filtering, stemming, and analysing processes to obtain accurate sentiment analysis results. The findings reveal that more than 50% of the data indicated negative sentiment towards the implementation of the 2020 Pilkada amid the COVID-19 pandemic.

The topics discussed in analysing sentiment in this study are similar to those examined by Hariyanti et al. (2024), which focused on the Constitutional Court's decision regarding the age limitation for presidential candidates, but with a different case—the Constitutional Court's decision on the 2024 presidential election dispute. This study compares the performance of the Support Vector Machine (SVM), Logistic Regression, and Naive Bayes algorithms in classifying sentiments, which are labelled as positive, neutral, or negative. This approach differs from those used by Ningsih et al. (2024) and Afandi and Isnaini (2024), which focused on specific algorithmic comparisons without

including the broader sentiment dynamics related to judicial decisions in a politically sensitive context. The primary objectives are to analyse netizen sentiment towards the Constitutional Court's decision and to identify the most effective algorithm for sentiment classification.

Thus, this sentiment analysis study aims to expand knowledge about how netizens respond to the Constitutional Court's decision, particularly those related to the presidential election, using the Support Vector Machine (SVM), Logistic Regression, and Naive Bayes algorithms.

Research Method

This study adopts an experimental methodology, which is well-suited for analysing structured and unstructured data through machine learning models (Johal & Mohana, 2020). It involves several key stages: data collection, labelling, pre-processing, weighting, and classification. The workflow is illustrated in Figure 1.

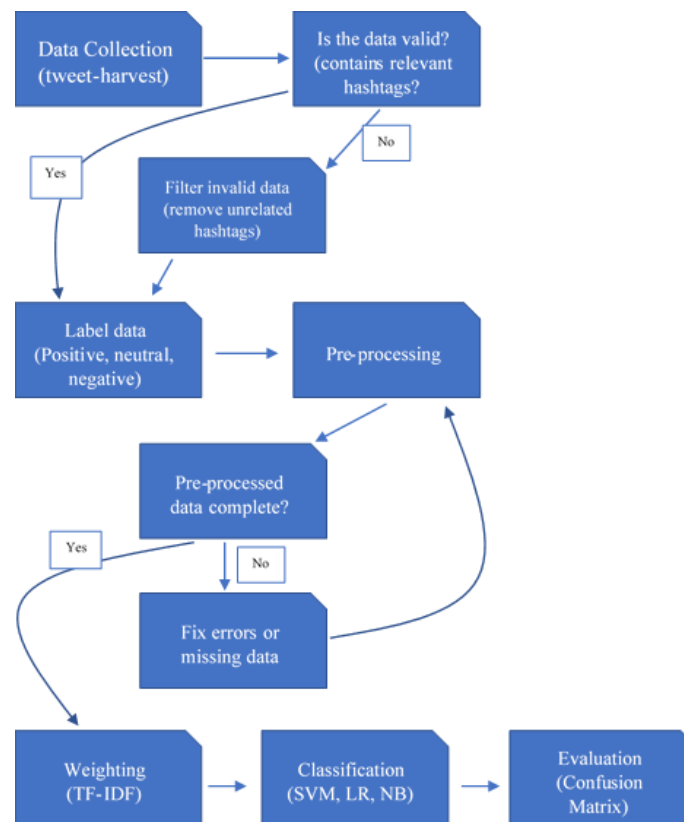


Figure 1. Sentiment Analysis Workflow

Data Collection

Data was collected from Twitter using the "tweet-harvest" Python library. A total of 397 valid tweets discussing MK rulings were obtained through web scraping. Invalid entries, such as promotional content or unrelated hashtags, were excluded (Bale et al., 2022).

According to the information available on the Twitter X page, there were initially 1,004 tweets from netizens. However, after scraping, only 397 valid data points were obtained from the tweet data. Valid data in this context refers to tweets from netizens that provide genuinely relevant responses to the Constitutional Court's decision. Unrelated tweets, such as those containing endorsements, promotions, or simply using hashtags like #putusanmk but not discussing the MK decision, are considered invalid for this research.

Labelling

The obtained data is then labelled based on comments from X users. We then labelled the obtained X users' data points—tweets' comments. Labelling in sentiment analysis is the process of determining the opinion or feeling of a

text and assigning it a label (positive, negative, or neutral) (Alkabkabi & Taileb, 2019). In this study, positive sentiment was assigned a value of 1, neutral sentiment was assigned a value of 0, and negative sentiment was assigned a value of -1. The labelling process involved human judgment, ensuring that each comment was carefully evaluated and categorized based on the nuanced understanding of sentiment by human annotators.

An experienced annotator classified all datapoints as Negative, Neutral, or Positive. The process involved examining words and sentences: contradictory ones were labelled negative; those neither conflicting nor strongly supportive were labelled neutral; and fully supportive ones were labelled positive. 100% of datapoints were accurately classified.

Pre-processing

Pre-processing steps included case folding, text cleaning, tokenization, stop word removal, and stemming. Emoticon symbols are converted into meaningful text to optimise the stemming process (Ariannor et al., 2024). These steps optimised data quality and improved classification accuracy (Johal & Mohana, 2020). The pre-processing steps involve case folding, text cleaning, tokenization, removal of stop words, and stemming. These stages of pre-processing are outlined in Table 1.

Table 1. Pre-processing

No.	Operation	Description
1	Case folding	Change text to lowercase (lowercase)
2	Cleaning text	Clean text from unnecessary characters and sentences, such as punctuation, symbols, URLs, emoticons, and so on. Characters or sentences containing this will be deleted, except emoticons. Emoticons will be converted into text that represents the meaning of the emoticon. Example: 😊 → smile, 👍 → thumbs up, 😞 → sad
3	Tokenization	Separate each word in a sentence using commas. Example: i like the constitutional court's decision → [i, like, the, constitutional, court's, decision]
4	Stop words	Delete conjunctions and ineffective words in sentences. Example: [i, like, the, constitutional, court's, decision] → [like, constitutional, court's, decision]
5	Stemming	Convert words to base words. Example: the constitutional court's decision suggests that improvements be made to make it better → the constitution court decis suggest that improv be made to make it better

Weighting

Each word is then assigned a weight using Term Frequency (TF) and Inverse Document Frequency (IDF) techniques. TF-IDF in sentiment analysis aids in reducing the number of irrelevant features by assigning weight values to features based on their strong correlation (Ririanti & Purwinarko, 2021).

TF represents the frequency of a word's occurrence in a document, and the weight of the word is proportional to this frequency. Therefore, the more often a word appears, the greater its weight. In contrast, in the IDF concept, a word that appears frequently across documents is assigned a smaller weight (Yutika et al., 2021). The TF-IDF equation formula can be seen in Formula (1) (Tanggraeni & Sitokdana, 2022):

$$TF.IDF_{std}(t) = tf_d^t \times \log \frac{n}{df_t} \quad (1)$$

Testing and Evaluation

In the testing stage, there are several steps, including classification using the Support Vector Machine (SVM) algorithm, Logistic Regression, and Naive Bayes. Subsequently, testing is conducted using a confusion matrix based on these three algorithms. Performance evaluation testing employs a confusion matrix to measure accuracy, precision, and recall. The confusion matrix is a table utilized to assess the performance of classification models in binary classification problems, comparing actual and predicted samples (Fahmy, 2022). The calculation of the confusion matrix follows Formula (2).

$$Accuracy = \frac{TP+TN}{TP+FP+TN+FN} \quad (2)$$

$$Precision = \frac{TP}{TP+FP} \quad (3)$$

$$Recall = \frac{TP}{TP+FN} \quad (4)$$

Description:

TP: *True Positive*

TN: *True Negative*

FP: *False Positive*

FN: *False Negative*

Result and Discussion

Result

Data Scraping

At this stage, the researcher employed web scraping techniques with the assistance of a library in the Python programming language to collect data from the Twitter X page. The collected data was saved in a .csv file format. A total of 397 reviews were obtained, with an average comment length of 166 characters. The maximum comment length was 348 characters, while the minimum was 29 characters. Additionally, the average number of words per tweet was 24, with a maximum of 56 words and a minimum of 4 words. The data was collected between April 22nd and April 23rd, 2024. Data sample can be seen in Table 2.

Table 2. Data Sample

No.	Text	English translation
1	Pasca Putusan MK Bakal Banyak Peristiwa Politi...	After the Constitutional Court Decision, There Will Be Many Political Events...
2	@ZulkifliLubis69 @officialMKRI MK sebagaimana ...	@ZulkifliLubis69 @officialMKRI MK as ...
3	Capres pemenang Pilpres 2024 Prabowo Subianto ...	The winning presidential candidate for the 2024 presidential election, Prabowo Subianto ...
4	@Boediantar4 Taubat Hakim MK? Kesalahan putusa...	@Boediantar4 MK Judge's Repentance? Decision Error...
5	Usai Putusan MK Anies Akan Kunjungi Kantor PKS...	After the Constitutional Court Decision, Anies Will Visit the PKS Office...
...	...	...
397	Me setelah membaca 3 paragraf muqodimah putusa...	After reading the 3 paragraphs of the introductory verdict...

Labelling

Out of the total dataset, there are 397 entries. Each review in the dataset is assigned a label, with negative labels defined as -1, neutral labels as 0, and positive labels as 1. The sample data labelling results are displayed in Table 3.

Table 3. Labelling

No.	Text	Label
1	Pasca Putusan MK Bakal Banyak Peristiwa Politi...	0
2	@ZulkifliLubis69 @officialMKRI MK sebagaimana ...	-1
3	Capres pemenang Pilpres 2024 Prabowo Subianto ...	0
4	@Boediantar4 Taubat Hakim MK? Kesalahan putusa...	-1
5	Usai Putusan MK Anies Akan Kunjungi Kantor PKS...	0
...	...	...
397	Me setelah membaca 3 paragraf muqodimah putusa...	-1

Pre-processing

Case Folding and Cleaning Text

Sample results of case folding and cleaning text can be seen in Table 4.

Table 4. Case Folding and Cleaning Text

No.	Text	Label
1	pasca putusan mk bakal banyak peristiwa politi...	0
2	mk sebagaimana yang telah terjadi sebelumnya s...	-1
3	capres pemenang pilpres prabowo subianto angka...	0
4	taubat hakim mk kesalahan putusan yg perbuat t...	-1
5	usai putusan mk anies akan kunjungi kantor pks...	0
...	...	...
397	saya setelah membaca paragraf muqodimah putusa...	-1

Tokenization

The next process is tokenization, which involves separating each word in the sentence using commas. Sample tokenization results are shown in Table 5.

Table 5. Tokenization

No.	Text	Label
1	[pasca, putusan, mk, bakal, banyak, peristiwa,...	0
2	[mk, sebagaimana, yang, telah, terjadi, sebelu...	-1
3	[capres, pemenang, pilpres, prabowo, subianto,...	0
4	[taubat, hakim, mk, kesalahan, putusan, yg, pe...	-1
5	[usai, putusan, mk, anies, akan, kunjungi, kan...	0
...	...	...
397	[saya, setelah, membaca, paragraf, muqodimah, ...	-1

Stop words

The goal of removing stop words is to eliminate conjunctions and ineffective words. Sample stop word removal results can be seen in Table 6.

Table 6. Stop words

Komentar	Stop words
[pasca, putusan, mk, bakal, banyak, peristiwa,...	[pasca, putusan, mk, peristiwa, politik, prabowo]
[mk, sebagaimana, yang, telah, terjadi, sebelu...	[mk, sarangnya, pengkhianat, bangsa, memutuska...
[capres, pemenang, pilpres, prabowo, subianto,...	[capres, pemenang, pilpres, prabowo, subianto,...
[taubat, hakim, mk, kesalahan, putusan, yg, pe...	[taubat, hakim, mk, kesalahan, putusan, yg, pe...
[usai, putusan, mk, anies, akan, kunjungi, kan...	[putusan, mk, anies, kunjungi, kantor, pks]
....	

Stemming

The stemming process is carried out to convert words into their base forms. Sample stemming results can be seen in Table 7.

Table 7. Stemming

No.	Stemming
1	pasca putus mk bakal banyak peristiwa politik ...
2	mk bagaimana yang telah jadi belum sarang khia...
3	capres menang pilpres prabowo subianto angkat ...
4	taubat hakim mk salah putus yg buat tidak hany...
5	usai putus mk anies akan kunjung kantor pks ha...
...	...
397	saya telah baca paragraf muqodimah putus mk ok...

Testing and Evaluation

Classification

The classification process utilizes three data splitting models to partition the dataset into training and testing sets: 90:10, 80:20, 70:30, and 60:40. In the 90:10 split, 90% of the data is allocated for training, and 10% for testing. Similarly, the 80:20 split designates 80% for training and 20% for testing. This pattern is consistently applied to the remaining ratios.

Confusion Matrix Evaluation

Confusion matrices are used to evaluate a classification model's ability to distinguish between positive and negative classes. Various performance metrics, including accuracy, precision, recall (sensitivity), and F1-score, can be derived from this matrix. The evaluation results based on the confusion matrix can be seen in Table 8.

Table 8. Confusion Matrix Measurement Results

Algorithm	Result	Split Model				Average accuracy
		90:10	80:20	70:30	60:40	
Support Vector Machine (SVM)	Accuracy	68%	68%	71%	69%	69,00%
	Precision	56%	47%	66%	68%	
	Recall	55%	54%	62%	61%	
	F1-score	53%	50%	60%	60%	
Logistic Regression	Accuracy	68%	69%	71%	65%	68,25%
	Precision	43%	46%	81%	66%	
	Recall	53%	55%	60%	55%	
	F1-score	47%	50%	57%	59%	
Naïve Bayes	Accuracy	68%	69%	69%	67%	68,25%
	Precision	43%	46%	79%	63%	
	Recall	53%	55%	59%	56%	
	F1-score	47%	50%	53%	51%	

Based on the measurement results in the Table 8, the following conclusions can be drawn:

- Accuracy: All three algorithms have similar accuracy performance, with results around 68–69%. There is no significant difference in accuracy between the three algorithms.
- Precision: SVM tends to have higher precision in the 70:30 and 60:40 data split schemes, while logistic regression has better precision in the 80:20 data split scheme. However, Naive Bayes shows consistently lower precision compared to the other two algorithms.
- Recall: Recall results for SVM and logistic regression tend to be stable, while Naive Bayes shows lower recall, especially in the 60:40 data split scheme.
- F1-score: The F1-score, which is the harmonic average of precision and recall, indicates the balance between the two. SVM and logistic regression have consistently better F1 scores compared to Naive Bayes.

Based on each data split, Figure 2, Figure 3, Figure 4, and Figure 5 illustrates the accuracy of the three algorithms.

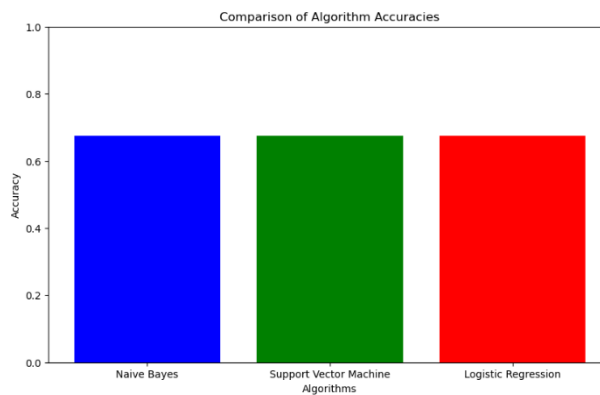


Figure 2. Accuracy Of Data Split: 90:10

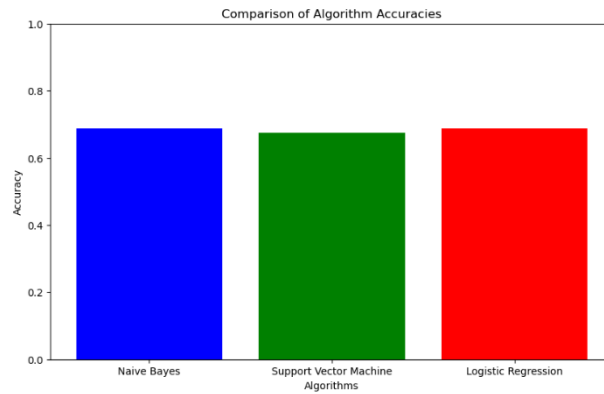


Figure 3. Accuracy Of Data Split: 80:20

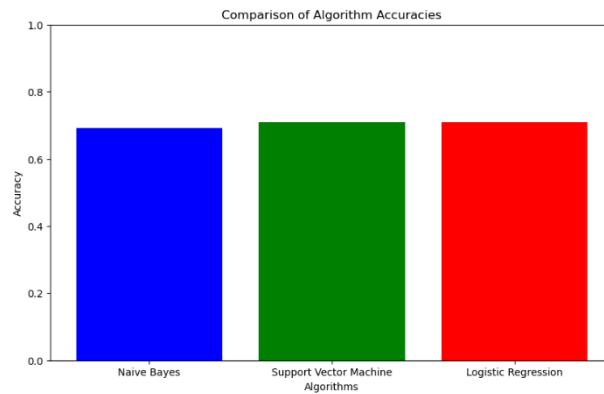


Figure 4. Accuracy Of Data Split: 70:30

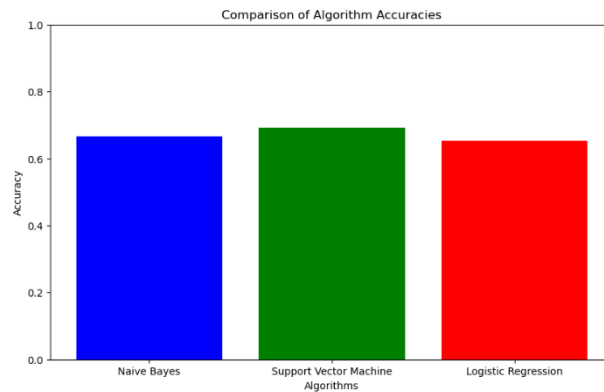


Figure 5. Accuracy Of Data Split: 60:40

Based on the results shown in Figure 2, Figure 3, Figure 4, and Figure 5, it can be concluded that both SVM and logistic regression perform better than Naive Bayes in this classification task, especially in terms of precision and F1 score. Therefore, the SVM and logistic regression algorithms can be effectively applied to analyse netizen sentiment towards the Constitutional Court (MK) decision on 397 netizen sentiment data points, which include 121 (30.48%) positive sentiments, 93 (23.43%) neutral sentiments, and 183 (46.10%) negative sentiments.

The percentage of negative sentiment is relatively greater compared to neutral and positive sentiment. This reflects possible social tensions, questions regarding public trust in the legal process, and the need for more effective communication and in-depth analysis to understand the root causes of negative attitudes.

Visualization

A word cloud visualization was created to capture netizens' opinions on the Constitutional Court's decision regarding the disputed 2024 presidential election. The visualization can be seen in Figure 6, Figure 7, and Figure 8.



Figure 6. Positive Sentiment

It can be seen in Figure 6 which is a visualization of positive netizen sentiment, where the larger word size represents words that appear more frequently in the neutral sentiment category.



Figure 7. Neutral Sentiment

It can be seen in Figure 7 which is a visualization of neutral netizen sentiment, where the larger word size represents words that appear more frequently in the neutral sentiment category.



Figure 8. Negative Sentiment

It can be seen in Figure 8 which is a visualization of negative netizen sentiment, where larger word sizes indicate words that appear more frequently in the negative sentiment category.

Discussion

This study enhances public understanding of the judiciary's role in the current political context, particularly concerning the Constitutional Court's (MK) decision on the 2024 presidential election dispute. The findings address two key research questions: (1) the sentiment of netizens toward the MK decision and (2) the most effective algorithm for

sentiment classification. Negative sentiment dominates, representing 46.10% of total sentiment, compared to 23.43% neutral and 30.48% positive sentiment. This indicates potential public distrust toward the MK decision, aligning with theories that public trust in judicial institutions is crucial for social stability, as discussed by Widłak (2022). Among the algorithms tested, Support Vector Machine (SVM) and Logistic Regression outperform Naive Bayes in terms of precision and F1-score, highlighting their effectiveness in analysing public sentiment in politically sensitive contexts.

While previous research has explored public sentiment toward MK rulings, focusing on this case sheds new light on the dynamics between the judiciary and society in contemporary democracy. These findings are significant as they underscore the urgency for better communication strategies by judicial institutions to address public concerns and mitigate potential social tensions. Moreover, the findings provide valuable insights for policymakers, particularly in understanding the public's reaction to critical judicial decisions.

A comprehensive approach to data collection, processing, and analysis from social media—including web scraping and data visualization techniques—provides new insights into understanding public opinion online, supporting the effectiveness of pre-processing and TF-IDF in improving sentiment classification accuracy, as noted by Johal and Mohana (2020). However, limited data from Twitter affects the breadth of sentiment analysis coverage, and future research should include multi-platform data, consistent with Yutika et al.'s (2021) recommendations. Additionally, ensemble learning approaches, as supported by Alkabkabi and Taileb (2019), should be considered to further improve classification performance on unevenly labelled or large text datasets.

Conclusion

This study has revealed that negative sentiment dominates netizen reactions to the Constitutional Court's decision on the 2024 presidential election dispute, representing 46.10% of the total sentiment, compared to 23.43% neutral and 30.48% positive sentiment. This finding highlights potential public distrust toward the judiciary, emphasizing the need for improved communication strategies to foster public confidence in judicial processes. Among the algorithms tested, Support Vector Machine (SVM) and Logistic Regression outperform Naive Bayes in terms of precision and F1-score, demonstrating their effectiveness in analysing public sentiment in politically sensitive contexts. These results provide valuable insights into the dynamics of public opinion and can guide policymakers and communication practitioners in addressing social tensions related to judicial decisions. Future research should expand data collection beyond Twitter to include other social media platforms, enabling a more comprehensive analysis of public sentiment. Additionally, incorporating ensemble learning approaches or advanced sentiment analysis techniques could further improve the accuracy and robustness of sentiment classification.

Acknowledgement

The authors express gratitude to colleagues at STMIK Banjarbaru and Bard College for their support and collaboration during this research.

References

- Afandi, M., & Isnaini, K. N. (2024). Analyzing Public Trust in Presidential Election Surveys: A Study Using SVM and Logistic Regression on Social Media Comments. *Journal of Computer Science and Engineering (JCSE)*, 5(1), 1-11. <https://doi.org/10.36596/jcse.v5i1.791>
- Alkabkabi, A., & Taileb, M. (2019). Ensemble Learning Sentiment Classification for Un-labeled Arabic Text. *ICC 2019: Advances in Data Science, Cyber Security and IT Applications*, (pp. 203–210). Cham. https://doi.org/10.1007/978-3-030-36365-9_17
- Ariannor, W., Kusuma, E. A., Fadilah, F., & Arsyad, M. (2024). Analyzing User Sentiments in Motor Vehicle Tax Applications Using the Naïve Bayes Algorithm. *Progresif: Jurnal Ilmiah Komputer*, 20(1), 91-101. <https://doi.org/10.35889/progresif.v20i1.1694>
- Atteveldt, W. v., Velden, M. A., & Boukes, M. (2021). The Validity of Sentiment Analysis: Comparing Manual Annotation, Crowd-Coding, Dictionary Approaches, and Machine Learning Algorithms. *Communication Methods and Measures*, 15(2), 121-140. <https://doi.org/10.1080/19312458.2020.1869198>
- Bale, A. S., Ghorpade, G., Naveen, N., S, R., Kamalesh, S., R, R., & S, R. B. (2022). Web Scraping Approaches and their Performance on Modern Websites. *2022 3rd International Conference on Electronics and Sustainable Communication Systems (ICESC)*, (pp. 956-959). Coimbatore, India. <https://doi.org/10.1109/ICESC54411.2022.9885689>
- Fahmy, M. M. (2022). Confusion Matrix in Binary Classification Problems: A Step-by-Step Tutorial. *Journal of Engineering Research (ERJ)*, 6(5), T0-T12. <https://doi.org/10.21608/erjeng.2022.274526>

- Hariyanti, Y., Kacung, S., & Santoso, B. (2024). Analisis Sentimen Terhadap Putusan Mahkamah Konstitusi Tentang Batasan Umur Capres Dan Cawapres Menggunakan Metode Naïve Bayes. *Multidisciplinary Indonesian Center Journal (MICJO)*, 1(1), 517-525. <https://doi.org/10.62567/micjo.v1i1.61>
- Husen, R. A., Astuti, R., Marlia, L., Rahmaddeni, R., & Efrizoni, L. (2023). Sentiment Analysis of Public Opinion on Twitter Toward BSI Bank Using Machine Learning Algorithms. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 3(2), 211-218. <https://doi.org/10.57152/malcom.v3i2.901>
- Johal, S. K., & Mohana, R. (2020). Effectiveness of Normalization Over Processing of Textual Data Using Hybrid Approach Sentiment Analysis. *International Journal of Grid and High Performance Computing (IJGHPC)*, 12(3), 43-56. <https://doi.org/10.4018/IJGHPC.2020070103>
- Liu, B. (2022). *Sentiment Analysis and Opinion Mining*. Toronto: Springer Nature.
- Muzaki, A., & Witanti, A. (2021). Sentiment Analysis Of The Community In The Twitter To The 2020 Election In Pandemic Covid-19 By Method Naive Bayes Classifier. *Jurnal Teknik Informatika (JUTIF)*, 2(2), 101-107. <https://doi.org/10.20884/1.jutif.2021.2.2.51>
- Ningsih, W., Alfianda, B., Rahmaddeni, R., & Wulandari, D. (2024). Perbandingan Algoritma SVM dan Naïve Bayes dalam Analisis Sentimen Twitter pada Penggunaan Mobil Listrik di Indonesia. *MALCOM: Indonesian Journal of Machine Learning and Computer Science*, 4(2), 556-562. <https://doi.org/10.57152/malcom.v4i2.1253>
- Novak, P. K., Smailović, J., Sluban, B., & Mozetič, I. (2015). Sentiment of Emojis. *PLoS ONE*, 10(12), 1-22. <https://doi.org/10.1371/journal.pone.0144296>
- Oktavia, D. (2023). Analisis Sentimen Terhadap Penerapan Sistem E-Tilang Pada Media Sosial Twitter Menggunakan Algoritma Support Vector Machine (SVM). *KLIK: Kajian Ilmiah Informatika dan Komputer*, 4(1), 407-417. <https://doi.org/10.30865/klik.v4i1.1040>
- Othman, S., Alshalwi, S., & Caushi, E. (2022). Sentiment of EmojiSets: How Emoji Sequences Improve Sentiment Cognition. *2022 IEEE 21st International Conference on Cognitive Informatics & Cognitive Computing (ICCI*CC)*, (hal. 72-79). Toronto. https://doi.org/10.1109/ICCI*CC57084.2022.10101644
- Ratnawati, T., & Iljas, N. (2021). Twitting the Public Sentiment on the Jakarta Online Zoning System. *IJIE (Indonesian Journal of Informatics Education)*, 5(1), 28-36. <https://doi.org/10.20961/ijie.v5i1.49718>
- Ririanti, N. P., & Purwinarko, A. (2021). Implementation of Support Vector Machine Algorithm with Correlation-Based Feature Selection and Term Frequency Inverse Document Frequency for Sentiment Analysis Review Hotel. *Scientific Journal of Informatics*, 8(2), 297-303. <https://doi.org/10.15294/sji.v8i2.29992>
- Singgale, Y. A. (2021). Pemilihan Metode dan Algoritma dalam Analisis Sentimen di Media Sosial : Sistematis Literature Review. *Journal of Information Systems and Informatics*, 3(2), 278-302. <https://doi.org/10.33557/journalisi.v3i2.125>
- Sujana, Y. (2023). A Comparative Study of Machine Learning Models for Sentiment Analysis of Dana App Reviews. *IJIE (Indonesian Journal of Informatics Education)*, 7(2), 161-166. <https://doi.org/10.20961/ijie.v7i2.93132>
- Tanggraeni, A. I., & Sitokdana, M. N. (2022). Analisis Sentimen Aplikasi E-Government Pada Google Play Menggunakan Algoritma Naïve Bayes. *Jurnal Teknik Informatika dan Sistem Informasi*, 9(2), 785-795. <https://doi.org/10.35957/jatisi.v9i2.1835>
- Vig, E. V., Kumar, V., Trehan, E. M., & Sharma, E. R. (2022). An Evaluation of Sentiment Analysis Techniques, Processes and Challenges. *2022 International Conference on Fourth Industrial Revolution Based Technology and Practices (ICFIRTP)*, (pp. 145-149). Uttarakhand, India. <https://doi.org/10.1109/ICFIRTP56122.2022.10059420>
- Wang, Q. (2022). Support Vector Machine Algorithm in Machine Learning. *2022 IEEE International Conference on Artificial Intelligence and Computer Applications (ICAICA)*, (pp. 750-756). Dalian, China. <https://doi.org/10.1109/ICAICA54878.2022.9844516>
- Widlak, T. (2022). Trust and Trustworthiness as Judicial Virtues. *Krytyka Prawa. Niezależne Studia nad Prawem*, 14(4), 151-166. <https://doi.org/10.7206/kp.2080-1084.562>
- Yuan, J., Wu, Y., Lu, X., Zhao, Y., Qin, B., & Liu, T. (2020). Recent advances in deep learning based sentiment analysis. *Science China Technological Sciences*, 63, 1947-1970. <https://doi.org/10.1007/s11431-020-1634-3>
- Yutika, C. H., Adiwijaya, A., & Faraby, S. A. (2021). Analisis Sentimen Berbasis Aspek pada Review Female Daily Menggunakan TF-IDF dan Naïve Bayes. *Jurnal Media Informatika Budidarma*, 5(2), 422-430. <https://doi.org/10.30865/mib.v5i2.2845>