

K-Means clustering analysis of public satisfaction with 50% electricity tariff reduction

Muhammad Fitrah Affandi Harahap¹, Abdul Halim Hasugian²

^{1,2}Computer Science Study Program, Faculty of Science and Technology, Universitas Islam Negeri Sumatera Utara, Medan, Indonesia

ARTICLE INFO	ABSTRACT
Article history: Received Jul 1, 2025 Revised Jul 9, 2025 Accepted Jul 15, 2025 Keywords: Electricity Tariff; K-Means Clustering; Public Satisfaction; Sentiment Analyze; Social Media X.	At the beginning of 2025, the Indonesian government implemented a policy to reduce electricity tariffs by 50% for household customers with power capacities of up to 2,200 VA. This policy aims to boost public purchasing power and stimulate economic growth, particularly among lower-middle-income groups. However, public responses to the policy have been varied and widely expressed on social media, especially on platform X (formerly known as Twitter). This study aims to evaluate public satisfaction with the electricity tariff reduction policy through sentiment analysis on social media X using the K-Means Clustering method. Data were collected through a crawling process using specific relevant keywords, followed by preprocessing steps such as cleansing, case folding, tokenizing, stemming, and conversion into numerical form using TF-IDF. The clustering results show that Cluster 1 dominates with 662 tweets (68.74%), predominantly reflecting positive sentiment, indicating that the majority of the public responded favorably to the 50% electricity tariff reduction policy. Cluster 2 consists of 165 tweets (17.13%) expressing negative sentiment, suggesting that some members of the public voiced concerns or dissatisfaction with the policy. Meanwhile, Cluster 0 includes 136 tweets (14.12%) containing neutral sentiment, representing moderate responses without a strong stance. These findings indicate that, overall, the policy received a generally positive reception from the public, although there are still critical and neutral perspectives.

This is an open access article under the CC BY-NC license.



Corresponding Author:

Muhammad Fitrah Affandi Harahap,
Computer Science Study Program,
Faculty of Science and Technology,
Universitas Islam Negeri Sumatera Utara,
Jl. William Iskandar Ps. V, Medan Estate, Kec. Percut Sei Tuan, Kabupaten Deli Serdang, Sumatera Utara
20371, Indonesia
Email: fitrah.affandiharahap@gmail.com

1. INTRODUCTION

Electricity is one of the basic necessities in daily life. The stability of electricity tariffs significantly affects household purchasing power and overall economic activity (Mutumba et al., 2024). Acknowledging this, at the beginning of 2025, the Indonesian government introduced a strategic policy to reduce electricity tariffs by 50% for residential customers with a power capacity of up to 2,200 VA. This policy aims to support economic recovery, increase household consumption, and provide an economic stimulus to lower-middle-income groups (Amiri et al., 2022). However, public responses to this policy have been varied and widely circulated on social media, particularly on platform X (formerly known as Twitter), which has now become a public space for expressing opinions in real time (Gilardi et al., 2021; Hutagalung et al., 2023; Weng et al., 2021).

However, public responses to this policy have not been homogeneous. Some groups responded positively, feeling economically supported, while others questioned the sustainability

and equitable distribution of the policy's benefits (Gilardi et al., 2021; Supriadi et al., 2020). Uncertainty in information and the lack of effective public communication have also led to misunderstandings or resistance from certain segments of society (Satria & Nurmandi, 2024). In today's digital era, public opinion is no longer shaped solely through conventional media but is strongly influenced by discussions on social media platforms (Alwaan et al., 2025; Sjoraida et al., 2024). One of the most actively used platforms by Indonesian society is social media platform X (formerly known as Twitter), which serves as an open space for the public to express opinions, criticisms, and expectations regarding government policies (Gumelar & Girsang, 2024).

Alongside the advancement of big data analytics and artificial intelligence, more sophisticated approaches are now available to capture and understand public perceptions through social media data analysis (Qi et al., 2024). One effective method for grouping opinions or perceptions based on pattern similarities is the K-Means Clustering algorithm (Xiao, 2024). This method allows researchers to identify key clusters within textual data, which represent various segments of public opinion. The interpretation of public sentiment in this study is not limited to surface-level reactions but is contextualized through the lens of socio-economic realities, particularly the direct implications of the policy on household-level conditions. Specifically, reductions in monthly electricity bills are assumed to influence household financial planning, discretionary spending, and overall economic resilience among lower-income segments. By acknowledging these tangible effects, the clustering of sentiment is expected to reflect not only subjective opinion but also concrete socio-economic experiences that shape how people perceive and respond to the policy. In the context of public policy, this approach serves as a powerful and responsive evaluation tool, capable of capturing public sentiment in real time and based on large-scale data. Previous studies have emphasized the importance of social media as an alternative data source for evaluating government policies. For instance, (Rochman et al., 2020) applied the K-Means Clustering method to analyze the customer satisfaction index of PT PLN (Persero) Jember Area. The data were classified into four distinct clusters, with the highest satisfaction index recorded at 86.04% and the lowest at 40.34%. This demonstrated the ability of the K-Means method to systematically cluster public service data based on customer satisfaction levels. Similarly, (Miranti et al., 2025) employed the K-Means Clustering method to analyze public sentiment toward K-Pop fans on Twitter. The dataset used consisted of 1,000 tweets categorized into three sentiment clusters: positive (33.15%), neutral (51.75%), and negative (15.09%).

The main focus of this study is to evaluate public satisfaction with the government's policy to reduce electricity tariffs by 50%, implemented at the beginning of 2025. This policy is a highly relevant and impactful issue, particularly for households with low electricity consumption capacities (Ilyas et al., 2022; Wang et al., 2020). Although the policy has direct implications for energy consumption and household purchasing power, there is still a lack of in-depth research evaluating it especially using social media-based approaches (Enrich et al., 2024).

Therefore, this study utilizes real-time data from platform X to more accurately and responsively capture public opinion. The rationale for selecting platform X as the sole source of public opinion data lies in its popularity, openness, and real-time interaction features that enable researchers to capture emerging narratives and sentiments as they unfold. In Indonesia, platform X has become a widely used medium for expressing social and political views, especially among younger, tech-savvy demographics. However, it is important to acknowledge its inherent limitations. Social media users particularly those active on platform X do not represent the entire population. The data may be biased toward individuals with internet access, specific age groups, or those who are more vocal online. Consequently, while platform X offers a rich and timely source of sentiment data, its representativeness is partial and should be interpreted with caution when generalizing findings to the broader public. The approach adopts the principles of computational social science by combining big data analytics with social insights based on digital public sentiment (Desiderio & Pablo, 2022). For the analysis process, the study employs the Term Frequency-Inverse Document Frequency (TF-IDF) method to extract features from textual data and convert them into numerical representations ready for processing. All procedures were carried out using Google Colab, which supports the integration of analytical libraries such as Scikit-learn and Pandas and allows for interactive and efficient data visualization (Setiawan & Adnyana, 2023). The results of this study are expected to contribute significantly to the evaluation of policy implementation effectiveness and provide strategic input for the government in designing more targeted, adaptive, and participatory public communication strategies in the future.

2. RESEARCH METHOD

This study employs a quantitative methodology, aiming to reveal phenomena holistically and contextually by collecting data from natural occurrences on social media platforms. The use of a quantitative approach enables the acquisition of more measurable and objective data, allowing for systematic analysis (Itua & Monday, 2025). To carry out this research, several stages were undertaken, as illustrated in Figure 1.

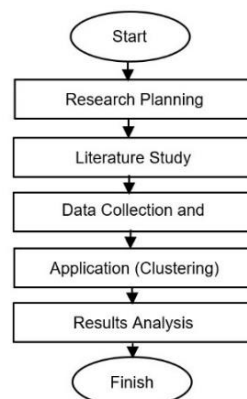


Figure 1. Research stage

The initial stage of this study was the research planning phase, which included formulating objectives, identifying problems, and determining the appropriate methodological approach. The purpose of this study is to cluster public opinions on a specific topic based on data obtained from platform X (formerly known as Twitter), using the K-Means Clustering method. This method was chosen for its capability to group large amounts of data into several clusters based on the similarity of characteristics among data points (Xiao, 2024)

The next step involved a literature review, conducted to gain a comprehensive theoretical and conceptual understanding of clustering methods, unstructured data processing, and the application of the K-Means algorithm in the context of public opinion analysis. This stage included reviewing scientific journals, articles, theses, and dissertations relevant to the research topic, as well as recent developments in the fields of text mining and unsupervised learning

In the data collection and processing stage, the dataset was obtained using a web crawling technique from platform X, employing specific keywords relevant to the research issue. This crawling process successfully gathered over 1.000 tweets containing public opinions in unstructured textual form. To ensure the reliability of the clustering results, the collected tweets were assessed quantitatively in terms of volume, temporal distribution, and topical relevance. Duplicate tweets, bot-generated content, and irrelevant entries were removed during the cleaning phase. Additionally, the dataset was reviewed for time-based diversity to minimize temporal or topic bias, ensuring that the data reflect a broad and dynamic range of public sentiment. The Elbow Method was also applied to determine the optimal number of clusters before the K-Means algorithm was executed.

The collected tweets then underwent a pre-processing phase, which included cleaning, case folding, normalization, tokenization, stopword removal, and stemming. After pre-processing, word weighting was performed using the Term Frequency Inverse Document Frequency (TF-IDF) method to convert textual data into a numerical representation for further analysis.

The weighted data was subsequently analyzed using the K-Means Clustering algorithm, which groups the data into a number of clusters based on feature similarity measured using the Euclidean Distance formula:

$$Distance(a,b) = \sqrt{\sum_{i=1}^n (a_i - b_i)^2}$$

Each document is measured by calculating its distance to all centroids, and it is then assigned to the cluster with the smallest distance. Subsequently, the centroid values are updated

based on the average position of all data points within each cluster. This process is repeated iteratively until the centroids converge (i.e., no longer change) (Kusuma & Nugroho, 2021).

The evaluation of clustering results is carried out using the Davies-Bouldin Index (DBI) method to assess the quality of the clustering. DBI calculates the ratio between the inter-cluster distance and the intra-cluster distance (Rahmawati et al., 2024). A lower DBI value indicates a better cluster structure, as it reflects that the elements within a cluster are compact and well-separated from other clusters. The DBI formula is as follows:

$$DBI = \frac{1}{k} \sum_{i=1}^k \max_{j \neq i} \left(\frac{S_i + S_j}{M_{ij}} \right)$$

The final stage of this research is the result analysis, which involves interpreting the clustering outcomes to identify the dominant sentiment patterns or public opinions within each group. This analysis provides insights into public perceptions of the examined issue and serves as the basis for drawing objective, data-driven conclusions.

3. RESULTS AND DISCUSSIONS

List of User Sentiments Regarding the 50% Electricity Tariff Reduction Policy

The comments used in this study were collected from the social media platform X through a web scraping process conducted in January 2025. Data collection was carried out using specific keywords such as: #TarifListrikTurun, #DiskonListrik50, #PenurunanTarifListrik, #ListrikLebihMurah, #PLNDiskonListrik, and others. The data scraping process was facilitated using the Python programming language via the Google Colab platform. Examples of the comments used in the sentiment analysis process are presented in Table 1.

Table 1. Examples of user comments regarding the 50% electricity tariff reduction

Data Set
@Suka2akula @irwndfrry santai aja bro gak perlu komen masalah nama. ini diskusi point of view. Poin saya jika liat inflasi liat ada 3 komponennya. Inflasi januari ini unik karena meski low banget namun disumbang faktor diskon tarif listrik. sementara inflasi inti naik. Ya gini ini kalau baca berita tapi dari judulnya aja. Inflasi turun ke level yang rendah itu salah satunya karena diskon tarif listrik di Januari. Makanya buka dulu itu rilis BPS sebelum bacot panjang lebar. Inflasi inti masih sekitar 2.4%-an. flashback awal tahun ini pemrinth memberikan diskon tarif listrik 50% diskon yang diberikan bkn berupa potongan hrg tp penambahan isi token but its ok (positif thinking) dan skrn gas langkah apa kt hrs masak menggunakan kompor listrik? Kejauhan ga sih? #PeringatanDarurat

Preprocessing Data

The pre-processing stage is a crucial initial step in sentiment analysis aimed at improving data quality before further analysis. This process begins with cleaning, which involves removing irrelevant elements such as URLs, emojis, punctuation marks, and numbers—ensuring that only informative text is retained. Next is case folding, which converts all letters to lowercase so that variations like “Listrik” and “listrik” are treated as identical by the system.

Normalization is then applied to standardize informal or non-standard words into formal language according to the Kamus Besar Bahasa Indonesia (KBBI). For instance, “gak” is converted to “tidak” and “bgt” to “banget,” using a normalisasi.csv reference file obtained from GitHub.

Following this, the text is broken down into individual word units using tokenization—for example, the sentence “tarif listrik turun” becomes [“tarif”, “listrik”, “turun”]. Common words that do not carry sentiment value such as “yang,” “dan,” “di,” and “untuk” are removed through a stopword removal process to ensure the analysis focuses only on opinion-laden terms.

Finally, stemming is performed to reduce words to their root form—for instance, “penurunan,” “menurunkan,” and “turunnya” are all reduced to “turun.” This step minimizes word variation and enhances sentiment classification accuracy. Through this comprehensive pre-processing, unstructured text data becomes well-prepared for in-depth analysis.

Table 2. Text pre-processing

Stage	Result
Initial Data	@Suka2akula @irwndfrry santai aja bro gak perlu komen masalah nama. ini diskusi point of view. Poin saya jika liat inflasi liat ada 3 komponennya. Inflasi januari ini unik karena meski low banget namun disumbang faktor diskon tarif listrik. sementara inflasi inti naik.
cleaning	santai aja bro gak perlu komen masalah nama ini diskusi point of view Poin saya jika liat inflasi liat ada komponennya Inflasi januari ini unik karena meski low banget namun disumbang faktor diskon tarif listrik sementara inflasi inti naik

Stage	Result
case folding	santai aja bro gak perlu komen masalah nama ini diskusi point of view poin saya jika liat inflasi liat ada komponennya inflasi januari ini unik karena meski low banget namun disumbang faktor diskon tarif listrik sementara inflasi inti naik
normalization	santai saja bro tidak perlu komen masalah nama ini diskusi point of view poin saya jika lihat inflasi lihat ada komponennya inflasi januari ini unik karena meski low banget namun disumbang faktor diskon tarif listrik sementara inflasi inti naik
tokenizing	['santai', 'saja', 'bro', 'tidak', 'perlu', 'komen', 'masalah', 'nama', 'ini', 'diskusi', 'point', 'of', 'view', 'poin', 'saya', 'jika', 'lihat', 'inflasi', 'lihat', 'ada', 'komponennya', 'inflasi', 'januari', 'ini', 'unik', 'karena', 'meski', 'low', 'banget', 'namun', 'disumbang', 'faktor', 'diskon', 'tarif', 'listrik', 'sementara', 'inflasi', 'inti', 'naik']
stopword removal	['santai', 'bro', 'komen', 'nama', 'diskusi', 'point', 'of', 'view', 'poin', 'lihat', 'inflasi', 'lihat', 'komponennya', 'inflasi', 'januari', 'unik', 'low', 'banget', 'disumbang', 'faktor', 'diskon', 'tarif', 'listrik', 'inflasi', 'inti']
stemming	santai bro komen nama diskusi point of view poin lihat inflasi lihat komponen inflasi januari unik low banget sumbang faktor diskon tarif listrik inflasi inti

Sentiment Labeling Using VADER Lexicon

In this study, sentiment labeling was performed using the VADER (Valence Aware Dictionary and Sentiment Reasoner) method, a lexicon-based approach specifically developed for analyzing sentiment in informal text such as that found on social media. VADER calculates a compound score ranging from -1 (extremely negative) to +1 (extremely positive), which reflects both the direction and intensity of sentiment expressed in a text (Hoti & Ajdari, 2023). To classify sentiments, a threshold-based approach was applied: scores equal to or greater than 0.05 were labeled as positive, scores equal to or less than -0.05 as negative, and scores falling between -0.05 and 0.05 were classified as neutral.

For the sentiment labeling process, this study employs the VADER (Valence Aware Dictionary and sEntiment Reasoner) algorithm, a lexicon-based sentiment analysis tool. Although VADER was initially developed for English texts, it has been adapted for this research by incorporating translated lexicons and contextual adjustments to accommodate Indonesian linguistic structures and emotive expressions. The choice of VADER is based on its practicality in handling social media texts that are often short, informal, and unstructured, without requiring extensive labeled training data. Compared to supervised learning methods such as Support Vector Machine (SVM) or Naïve Bayes, VADER offers faster deployment and is more suitable when manually labeled datasets are unavailable or limited. Nevertheless, future research may consider hybrid or supervised models to enhance classification precision and handle complex sentiment nuances in the Indonesian language.

Table 3. Labeling result

Stemming	Sentimen Score	Label
santai bro komen nama diskusi point of view poin lihat inflasi lihat komponen inflasi januari unik low banget sumbang faktor diskon tarif listrik inflasi inti	4	Positive
ya baca berita judul inflasi turun level rendah salah satu diskon tarif listrik januari buka rilis bps bacot lebar inflasi inti an	1	Positive
flashback pemrinth diskon tarif listrik diskon potong hrg tambah isi token but its ok positif thinking skrn gas langkah masak kompor listrik jauh sih peringatandarurat	2	Positive
akutu degan diskon tarif listrik kejut ya sulit suudzon perintah	1	Positive
rasai diskon tarif listrik	2	Positive

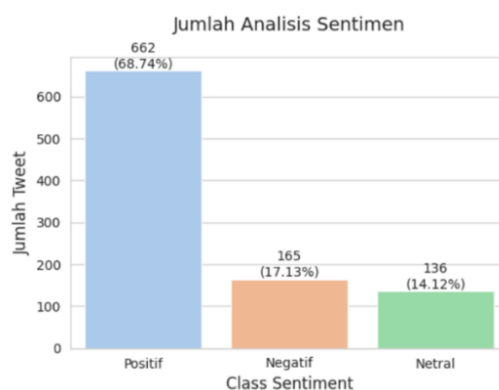


Figure 2. Visualization of the Amount of sentiment analysis

The sentiment analysis results indicate that positive sentiment dominates, with 662 tweets or 68.74% of the total data. Negative sentiment accounts for 165 tweets (17.13%), while neutral sentiment comprises 136 tweets (14.12%). This distribution shows that the majority of public responses are positive. Nevertheless, further analysis is needed to understand the context behind each sentiment category, particularly in identifying common patterns within positive tweets and the potential factors contributing to the emergence of negative and neutral sentiments.

TF-IDF Weighting

The initial step in TF-IDF weighting is calculating the TF (Term Frequency) value for each word in the document, which represents how frequently a word appears in that specific document. After calculating TF, the IDF (Inverse Document Frequency) value is then computed for each word (Setiawan & Adnyana, 2023). The final TF-IDF value is obtained by multiplying the TF and IDF scores. In this example, four sentences are used as follows: a) santai bro komen nama diskusi point of view poin lihat inflasi lihat komponen inflasi januari unik low banget sumbang faktor diskon tarif listrik inflasi inti; b) ya baca berita judul inflasi turun level rendah salah satu diskon tarif listrik januari buka rilis bps bacot lebar inflasi inti an; c) flashback pemrinth diskon tarif listrik diskon potong hrg tambah isi token but its ok positif thinking skrn gas langkah masak kompor listrik jauh sih peringatan darurat; d) akutu degan diskon tarif listrik kejut ya sulit suudzon perintah. For the term "diskusi", which appears in only one document (DF = 1), the IDF value is:

$$IDF_{\text{diskusi}} = \log_{10} = \left(\frac{4}{1}\right) = 0,602$$

Meanwhile, terms such as "listrik", "diskon", and "tarif" appear in all four documents (DF = 4), so their IDF value is:

$$IDF = \log_{10} = \left(\frac{4}{4}\right) = \log_{10} (1) = 0$$

For a term that appears in three documents (DF = 3), such as the term "inflasi", the IDF value is calculated as:

$$IDF_{\text{inflasi}} = \log_{10} = \left(\frac{4}{3}\right) = \log_{10} (1,33) = 0,125$$

This indicates that terms frequently appearing across all documents have a low discriminative power, resulting in small or even zero IDF values. The detailed calculation process for each word's IDF value is not discussed in full; instead, the final IDF results are presented in Table 4.

Table 4. IDF value calculation

Term	DF	IDF
listrik	4	0.000
diskon	4	0.000
tarif	4	0.000
santai	1	0.602
inflasi	3	0.125
diskusi	1	0.602
berita	1	0.602
subsidi	0	0.000
baca	1	0.602
level	1	0.602
rilis	1	0.602
lebar	1	0.602
flashback	1	0.602
positif	1	0.602
ada	0	0.000
sulit	1	0.602
tahun	1	0.602
gas	1	0.602
lipat	0	0.000
masak	1	0.602
kompor	1	0.602
melihat	1	0.602
pemerintah	0	0.000

Application of K-means Clustering

This stage was conducted after the sentiment data had been processed and represented in a 6-dimensional TF-IDF vector form (Nurchayawati et al., 2025). Four short documents were used to represent the data. Centroid initialization was performed randomly by selecting three initial

documents as the cluster centers. The clustering process was then carried out by calculating the Euclidean distance between each document and each centroid. Each data point was assigned to the cluster with the shortest distance to its centroid. The following are the distance measurements of each data point to the respective centroids:

1) Calculate the distance of D1 to the centroid

➤ **D1 to C1 (Same)**

Jarak = 0

➤ **D1 to C2**

$$\sqrt{(0,602 - 0)^2 + (0,250 - 0,125)^2 + (0,602 - 0)^2 + 0^2 + 0^2 + 0^2}$$

$$\sqrt{0,740} = 0,860$$

➤ **D1 to C3**

$$\sqrt{(0,602 - 0)^2 + (0,250 - 0)^2 + (0,602 - 0)^2 + 0^2 + (0 - 0,602)^2}$$

$$\sqrt{+ (0 - 0,602)^2}$$

$$\sqrt{1,511} = 1,229$$

D1 enters Cluster 1 (distance = 0)

2) Calculate the distance of D2 to the centroid

➤ **D2 to C1 (Same)**

$$\sqrt{(0 - 0,602)^2 + (0,125 - 0,250)^2 + (0 - 0,602)^2 + (0 - 0,602)^2}$$

$$\sqrt{1,102} = 1,049$$

➤ **D2 to C2 (to)**

Distance = 0

➤ **D2 to C3**

$$\sqrt{0^2 + 0,125^2 + 0^2 + (0,602 - 0)^2 + 0^2 + 0^2}$$

$$\sqrt{0,378} = 0,615$$

D2 is closest to C2 (distance = 0) → Cluster 2

3) Calculate the distance of D3 to the centroid

➤ **D3 to C1 (Same)**

$$\sqrt{(0 - 0,602)^2 + (0 - 0,250)^2 + (0 - 0,602)^2}$$

$$\sqrt{0,787} = 0,887$$

➤ **D3 to C2 (same)**

$$\sqrt{(0 - 0)^2 + (0 - 0,125)^2 + (0 - 0)^2 + (0 - 0,602)^2}$$

$$\sqrt{0,378} = 0,615$$

➤ **D3 to C3**

$$\sqrt{(0 - 0)^2 + (0 - 0)^2 + (0 - 0)^2 + 0^2 + (0 - 0,602)^2 + (0 - 0,602)^2}$$

$$\sqrt{0,724} = 0,851$$

D3 is closest to C2 (distance = 0.615) → Cluster 2

4) Calculate the distance of D4 to the centroid

➤ **D4 to C1 (Same)**

$$\sqrt{(0 - 0,602)^2 + (0 - 0,250)^2 + (0 - 0,602)^2 + (0 - 0)^2 + (0,602 - 0)^2}$$

$$\sqrt{+ (0,602 - 0)^2}$$

$$\sqrt{1,511} = 1,229$$

➤ **D4 to C2 (same)**

$$\sqrt{(0 - 0)^2 + (0 - 0,125)^2 + 0^2 + (0 - 0,602)^2 + (0,602)^2}$$

$$\sqrt{1,102} = 1,049$$

➤ **D4 to C3**

Distance = 0

D4 goes to Cluster 3 (distance = 0)

After calculating the Euclidean distance between each document and the three initial centroids, each document was then assigned to the cluster with the nearest centroid. This initial

clustering result is presented in Table 4.13. This grouping represents the preliminary outcome of the clustering process using the K-Means algorithm, prior to the evaluation of its alignment with sentiment labels.

Table 5. Data on clusters

Document	Nearest distance	Cluster
D1	C1	Cluster 1
D2	C2	Cluster 2
D3	C2	Cluster 2
D4	C3	Cluster 3

Based on Table 4.14, of the four test documents, only one document (D4) was successfully classified according to the initial label (positive). The other three documents experienced classification mismatches. This shows that the accuracy of the K-Means model is only 25%. The low accuracy can be caused by the limited amount of data, simple feature representation, and the nature of K-Means that does not use labels in the training process.

Table 6 Test results on applying k-means clustering

Dokumen	Tweet	Label Awal	Label Akhir
D1	listrik tarif santai kepala	Netral	Negatif
D2	mantap baca inflasi indonesia	Positif	Netral
D3	diskusi sulit gas kompor	Negatif	Netral
D4	subsidi masak melihat rilis positif pemerintah	Positif	Positif

Result

After the entire training and testing process is carried out, the system produces predictions in the form of sentiment labels for the test data. The labels predicted by the model are then compared with the actual labels in the dataset to calculate model evaluation metrics including accuracy, precision, recall, and f1-score. These metrics are used to assess how well the system recognizes and distinguishes negative, neutral, and positive sentiments overall.

Due to the incomparable data between positive, neutral and negative data, the calculation of precision, recall, and f1-score values is focused on positive sentiments, where sentiments with positive values are much more than neutral and negative sentiments. The following shows the confusion matrix from the results of sentiment analysis conducted using the K-means Clustering algorithm.

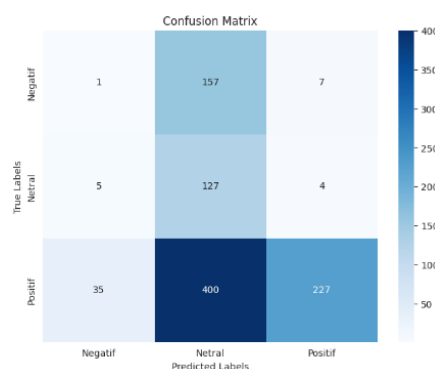


Figure 3. Confusion matrix

Table 7 shows that the model successfully classified 227 data points as true positives (TP), meaning these were actual positive cases correctly predicted as positive. There were 11 false positives (FP), where non-positive data were incorrectly predicted as positive. Additionally, the model failed to detect 435 actual positive data points, resulting in a high number of false negatives (FN). Meanwhile, 290 data points were true negatives (TN), indicating that non-positive cases were accurately identified. The high FN value reveals a significant weakness of the model in accurately detecting positive sentiment.

Table 7. Determination of TP, FP, TN and FP values

Label	Value
TP (True Positive)	227
FP (False Positive)	11
FN (False Negative)	435
TN (True Negative)	290

Based on the results, the model achieved a precision of 95.4%, indicating that most of the data predicted as Positive were indeed correct. However, the recall was only 34.3%, which suggests that the model failed to identify a large portion of actual Positive data. The overall accuracy of the model was 53.7%, while the F1-score stood at 50.3%, reflecting an imbalanced trade-off between precision and recall. Therefore, although the model demonstrates high precision in predicting the Positive class, its sensitivity to actual Positive data remains low. This implies that improvements such as class balancing or model parameter tuning are needed to enhance the overall classification performance.

Discussion

This study evaluates public satisfaction with the 50% electricity tariff reduction policy by analyzing public opinion gathered from the social media platform X (formerly Twitter). By applying the K-Means Clustering algorithm, the data was successfully grouped into three main clusters based on textual similarity. The clustering results show that Cluster 1 dominates with 662 tweets (68.74%), which, based on manual analysis, were classified as positive sentiment. This indicates that the majority of the public responded enthusiastically to the tariff reduction policy, likely because it is perceived to ease household financial burdens.

Cluster 2 includes 165 tweets (17.13%) representing negative sentiment. Although smaller in number, this cluster indicates the presence of individuals who are dissatisfied or concerned about the effectiveness of the policy's implementation. Meanwhile, Cluster 0 consists of 136 tweets (14.12%) expressing neutral sentiment, reflecting moderate responses without strong opinions, possibly due to the public still waiting to experience the policy's real impact or due to limited information dissemination. The unequal distribution across clusters reflects a dominance of positive opinion, with far fewer negative and neutral sentiments. This suggests that the electricity tariff reduction policy is generally well received by the public on social media, although there is still room for improvement, particularly in terms of communication and transparency. From a technical perspective, the application of the K-Means Clustering algorithm is considered effective in grouping data based on similar text patterns, as evidenced by evaluations using the confusion matrix and performance metrics such as True Positive, False Positive, and True Negative ((Ikotun et al., 2023) Despite the imbalance in sentiment label distribution, the model was still able to identify the main characteristics of each group. Overall, the findings of this study indicate that the electricity tariff reduction policy received predominantly positive public responses, although some critical and neutral voices remain. These insights are valuable for the government to strengthen public communication strategies, build public trust, and ensure that the policy is implemented effectively and equitably. Furthermore, this social media-based analytical approach proves to be a relevant and rapid evaluation tool for assessing public perception of policy in real time.

4. CONCLUSIONS

This study successfully clustered public opinion regarding the 50% electricity tariff reduction policy using the K-Means Clustering algorithm applied to data obtained from the social media platform X (formerly Twitter). After undergoing preprocessing and data representation using the TF-IDF method, the data was grouped into three main clusters. The clustering results showed that the majority of data fell into Cluster 1 (68.74%), which was dominated by positive sentiment, followed by Cluster 2 (17.13%) with negative sentiment, and Cluster 0 (14.12%) with neutral sentiment. These findings indicate that, in general, the public responded positively to the policy, although there were still a smaller portion of neutral and negative opinions.

The evaluation of model performance shows that despite the imbalance in sentiment label distribution, the K-Means algorithm was able to cluster the data effectively and representatively, especially in identifying the dominant sentiment. However, it is important to note that K-Means has certain limitations when applied to social media data, particularly due to the high level of noise, informal language, and contextual ambiguity inherent in user-generated content. These factors can

affect clustering accuracy, and therefore, alternative approaches such as DBSCAN, hierarchical clustering, or supervised deep learning models (e.g., LSTM, BERT) may offer improved performance for future research. Moreover, to strengthen the analytical depth in future studies, researchers may explore hybrid models that combine lexicon-based and machine learning approaches, as well as integrate temporal sentiment analysis to track how public opinion evolves over time. These advancements can help policymakers not only understand static public perception, but also monitor dynamic shifts in sentiment in response to policy implementation stages or external socio-political events.

This study highlights that social media can serve as an effective alternative data source for quickly and in real-time assessing public perception of public policies. Therefore, the findings can serve as a valuable reference for policymakers to better understand public opinion and to design more targeted and responsive communication and policy implementation strategies.

ACKNOWLEDGEMENTS

The author would like to express sincere gratitude to all parties who have provided support in the completion of this research. Special thanks are extended to colleagues and academic supervisors for their guidance, scholarly input, and unwavering encouragement throughout the research process. Every form of assistance and collaboration has played a crucial role in ensuring the smooth progress and success of this study. It is the author's hope that the results of this research may offer meaningful benefits and tangible contributions to the advancement of knowledge, particularly in the field of data analysis and public opinion mapping.

REFERENCES

- Alwaan, A. Z., Muhammad, B., Wahyuni, S., & Hidayah, H. (2025). Analisis Persepsi Masyarakat Terhadap Aktivitas Media Sosial dalam Pembentukan Opini Publik di Era Digital. *Analisis Persepsi Masyarakat Terhadap Aktivitas Media Sosial Dalam Pembentukan Opini Publik Di Era Digital*, 6(1), 459–462.
- Amiri, F., Abbasi, S., & Babaie Mohamadeh, M. (2022). Clustering Methods to Analyze Social Media Posts during Coronavirus Pandemic in Iran. *Technology Journal of Artificial Intelligence and Data Mining*, 10(2), 159–169. <https://doi.org/10.22044/JADM.2022.11270.228>
- Desiderio, R. J., & Pablo, del R. (2022). Analysing the drivers of the efficiency of households in electricity consumption. *Energy Policy*, 164, 112828. <https://doi.org/10.1016/j.enpol.2022.112828>
- Enrich, J., Li, R., Mizrahi, A., & Reguant, M. (2024). Measuring the impact of time-of-use pricing on electricity consumption: Evidence from Spain. *Journal of Environmental Economics and Management*, 123, 102901. <https://doi.org/10.1016/j.jeem.2023.102901>
- Gilardi, F., Gessler, T., Kubli, M., & Müller, S. (2021). Social Media and Policy Responses to the COVID-19 Pandemic in Switzerland. *Swiss Political Science Review*, 27(2), 243–256. <https://doi.org/10.1111/spsr.12458>
- Gumelar, G., & Girsang, A. S. (2024). Understanding Public Opinions of Government Measures Against COVID-19 Through Twitter Sentiment Analysis. *Journal of Logistics, Informatics and Service Science*, 11(2). <https://doi.org/10.33168/JLISS.2024.0201>
- Hoti, M. H., & Ajdari, J. (2023). Sentiment Analysis Using the Vader Model for Assessing Company Services Based on Posts on Social Media. *SEEU Review*, 18(2), 19–33. <https://doi.org/10.2478/seeur-2023-0043>
- Hutagalung, W. M. S. N., Tony, T., & Jaya Perdana, N. (2023). Analisis Sentimen Pada Opini Kenaikan Harga Bahan Bakar Minyak Pada Media Sosial Twitter. *Simtek: Jurnal Sistem Informasi Dan Teknik Komputer*, 8(2), 280–284. <https://doi.org/10.51876/simtek.v8i2.207>
- Ikotun, A. M., Ezugwu, A. E., Abualigah, L., Abuhaija, B., & Heming, J. (2023). K-means clustering algorithms: A comprehensive review, variants analysis, and advances in the era of big data. *Information Sciences*, 622, 178–210. <https://doi.org/10.1016/j.ins.2022.11.139>
- Ilyas, R., Hussain, K., Ullah, M. Z., & Xue, J. (2022). Distributional impact of phasing out residential electricity subsidies on household welfare. *Energy Policy*, 163, 112825. <https://doi.org/10.1016/j.enpol.2022.112825>
- Itua, E., & Monday, O. S. (2025). Utility of Quantitative Research Methods: Implication for Public Administration. *International Journal of Social Science and Human Research*, 08(05). <https://doi.org/10.47191/ijsshr/v8-i5-33>
- Kusuma, A., & Nugroho, A. (2021). Analisa Sentimen Pada Twitter Terhadap Kenaikan Tarif Dasar Listrik Dengan Metode Naïve Bayes. *Jurnal Ilmiah Teknologi Informasi Asia*, 15(2), 137. <https://doi.org/10.32815/jitika.v15i2.557>

- Miranti, A. Z. L., Nurul, A. S. W., Muhammad, S. R., & Galuh, W. S. (2025). Penerapan Metode K-Means Clustering dalam Menganalisis Sentimen Masyarakat terhadap K-Popers pada Twitter. *Progresif: Jurnal Ilmiah Komputer*, 18(2), 201–210. <https://databoks.katadata.co.id/>
- Mutumba, G. S., Mubiinzi, G., & Amwonya, D. (2024). Electricity consumption and economic growth: Evidence from the East African community. *Energy Strategy Reviews*, 54, 101431. <https://doi.org/10.1016/j.esr.2024.101431>
- Nurcahyawati, V., Mustaffa, Z., & Khalaf, M. (2025). Exceeding Manual Labeling: VADER Lexicon as an Accurate Alternative to Automatic Sentiment Classification. *The International Arab Journal of Information Technology*, 22(2). <https://doi.org/10.34028/iajit/22/2/2>
- Qi, M., Zhao, J., & Feng, Y. (2024). An optimized public opinion communication system in social media networks based on K-means cluster analysis. *Heliyon*, 10(24), e40033. <https://doi.org/10.1016/j.heliyon.2024.e40033>
- Rahmawati, N., Buaton, R & Ambarita. I. (2024). Klasifikasi Tingkat Kepuasan Masyarakat Penggunaan BPJS Kesehatan di Kota Binjai Menggunakan K-Means Clustering. *Switch: Jurnal Sains Dan Teknologi Informasi*, 2(4), 86–98. <https://doi.org/10.62951/switch.v2i4.188>
- Rochman, R. S., Retnani, W. E. Y., & Juwita, O. (2020). Rancang Bangun Aplikasi Analisis Indeks Kepuasan Pelanggan pada PT. PLN (PERSERO) Area Jember dengan Menggunakan Pendekatan Metode Servqual dan K-Means Clustering. *Berkala Saintek*, 8(3), 96. <https://doi.org/10.19184/bst.v8i3.12571>
- Satria, T., & Nurmandi, A. (2024). Behavioral Patterns of Social Media Users toward Policy: A Scientometric Analysis. *Policy & Governance Review*, 8(2), 192. <https://doi.org/10.30589/pgr.v8i2.902>
- Setiawan, G. H., & Adnyana, I. M. B. (2023). Improving Helpdesk Chatbot Performance with Term Frequency-Inverse Document Frequency (TF-IDF) and Cosine Similarity Models. *Journal of Applied Informatics and Computing*, 7(2), 252–257. <https://doi.org/10.30871/jaic.v7i2.6527>
- Sjoraida, D. F., Guna, B. W. K., Nungraha, A. R., Pasaribu, D., & Djafri, N. (2024). Public Opinion Formation in the Digital Age : A Review of Literature. *Indonesia Journal of Engineering and Education Technology (IJEET)*, 2(2), 290–297. <https://doi.org/10.61991/ijeet.v2i2.52>
- Supriadi, A. Y., Ikhsan, M., Mahi, B. R., & Girianna, M. (2020). Impact of Changes to Subsidised Electricity Tariffs on The Welfare Of Households. *ETIKONOMI*, 18(2), 169–184. <https://doi.org/10.15408/etk.v18i2.12041>
- Wang, Z., Li, H., Deng, N., Cheng, K., Lu, B., Zhang, B., & Wang, B. (2020). How to effectively implement an incentive-based residential electricity demand response policy? Experience from large-scale trials and matching questionnaires. *Energy Policy*, 141, 111450. <https://doi.org/10.1016/j.enpol.2020.111450>
- Weng, S., Schwarz, G., Schwarz, S., & Hardy, B. (2021). A Framework for Government Response to Social Media Participation in Public Policy Making: Evidence from China. *International Journal of Public Administration*, 44(16), 1424–1434. <https://doi.org/10.1080/01900692.2020.1852569>
- Xiao, E. (2024). Comprehensive K-Means Clustering. *Journal of Computer and Communications*, 12(03), 146–159. <https://doi.org/10.4236/jcc.2024.123009>