
Sentiment Analysis on Starbucks Reviews: Implementation of K-Nearest Neighbors and Support Vector Machine

Jasmine Nabila Ayoedya^{1*}, Ida Ayu Devian Branitasandhini Putra², Heri Wijayanto³

^{1,2,3} Master of Information Technology, Faculty of Engineering, University of Mataram

^{12,3} Jl. Majapahit 62, Mataram, NTB, Indonesia

E-mail: jasmine.nabila03@gmail.com¹, devian.branitasandhini23@gmail.com², heri@unram.ac.id³

Submitted:07/08/2025 , Revision: 09/11/2025, Accepted : 09/15/2025

Abstract

This study analyzes Starbucks customer reviews from the ConsumerAffairs website using K-Nearest Neighbors (KNN) and Support Vector Machine (SVM) for sentiment classification. Reviews were preprocessed with standard NLP techniques and labeled using TextBlob as positive, neutral, or negative. Word cloud analysis shows that most reviews express positive sentiment, with terms like *awesome*, *atmosphere*, and *organic* highlighting appreciation for store environment, product quality, and overall experience, while neutral reviews provide descriptive information and negative reviews indicate health issues or dissatisfaction with service. KNN achieved its best results with a 60%:40% training-to-testing split (87% accuracy, 84% precision, 87% recall, F1-score 84%), whereas SVM performed best with an 80%:20% split (88% accuracy, 86% precision, 88% recall, F1-score 85%), slightly outperforming KNN. The findings demonstrate that both methods effectively classify sentiment and, together with word cloud visualization, provide actionable insights for Starbucks to enhance product quality, service efficiency, and customer satisfaction.

Keywords:

Sentiment Analysis;
K-Nearest Neighbors;
Support Vector Machine;
Starbucks Reviews;

***Corresponding Author: Jasmine Nabila Ayoedya**

E-mail: jasmine.nabila03@gmail.com

1. Introduction

In the ever-evolving digital era, consumer reviews on online platforms play a vital role in shaping public perception of a product or service. Companies increasingly rely on online reviews as a primary source of information to evaluate customer satisfaction and formulate strategies for service quality improvement. Consumer trust in customer reviews is a key factor influencing purchasing decisions, especially when the reviews are numerous and contain in-depth information [1]. The study highlights that honest and informative customer reviews not only enhance customer trust but also provide strategic value for online businesses in expanding their reach and increasing sales conversion [2]. Starbucks, as one of the leading global brands in the food and beverage industry, is frequently the subject of consumer reviews on various digital platforms. The company has embraced digital transformation through its Loyalty Rewards App to increase customer loyalty. Numerous negative reviews on the Google Play Store highlight usability issues such as errors, efficiency, and satisfaction indicating areas in the digital service that require improvement [3]. Meanwhile, reviews from platforms like ConsumerAffairs reflect consumers' direct perceptions of Starbucks' services and products. Monitoring and analyzing sentiment across various review platforms has become a strategic measure to help the company maintain brand reputation while continuously improving service quality.

To gain a deeper understanding of customer sentiment trends toward Starbucks, this study adopts a machine learning approach using the K-Nearest Neighbor (KNN) algorithm and Support Vector Machine (SVM). KNN is one of the machine learning algorithms proven effective in classifying text data based on feature similarity. Its advantages include simplicity and strong performance when dealing with limited datasets, particularly when combined with data representation techniques such as TF-IDF or Word2Vec. KNN achieved up to 99% accuracy in sentiment classification of Bibit investment app reviews, with consistent results on data scraped from both YouTube and the Play Store [4]. The classification process involved text processing, sentiment labeling using the TextBlob library, and analysis based on three sentiment categories: positive, neutral, and negative. These findings further support the relevance of using KNN in this research, especially in classifying Starbucks customer reviews, which are high in both volume and variability providing an accurate distribution of sentiment that is useful for service improvement [5].

Meanwhile, Support Vector Machine (SVM) is a supervised learning algorithm that works by finding the optimal hyperplane to separate data into different classes with the maximum margin. It is particularly effective for high-dimensional data such as text, making it widely used in sentiment analysis. For example, a study on 3,000 Google Play Store reviews of the LinkedIn application applied SVM with preprocessing steps including crawling, cleaning, translation, labeling, tokenization, and stopword

removal[6]. After parameter optimization using grid search, the model achieved 82% accuracy with settings of $C = 1$, $\gamma = 0.1$, and a linear kernel, demonstrating that SVM can classify sentiments in app reviews with fairly high accuracy [6].

The aim of this study is to classify customer reviews of Starbucks into three sentiment categories: positive, negative, and neutral, and to evaluate the performance of the K-Nearest Neighbor (KNN) and Support Vector Machine (SVM) algorithms in this classification process. Through this approach, the research seeks to gain a more comprehensive understanding of customer perceptions regarding Starbucks' services and products. This information will be highly valuable for management in designing responsive strategies to customer feedback and identifying areas in need of improvement.

By conducting sentiment analysis using both the KNN and SVM algorithms on Starbucks customer reviews, this study is expected to make a meaningful contribution to the development of data-driven decision support systems. The sentiment classification results can be used by Starbucks to enhance service quality, address frequently criticized areas, and maintain customer satisfaction and loyalty. By leveraging this data-based approach, the company can become more adaptive to customer needs and retain its competitive advantage in an increasingly dynamic industry.

2. Literature Review

To support the theoretical foundation of this research, a review of several relevant previous studies was conducted, particularly those related to sentiment analysis using machine learning methods such as K-Nearest Neighbors (KNN), Support Vector Machine, Naïve Bayes, transformer-based models, and lexicon-based approaches. A summary of the findings from these prior studies is presented in Table 1 below.

Table 1. Literature Review

No	Journal	Resume
1	Sentiment Analysis Based on Transformer: The Sentiment Analysis about The Boy and the Heron	This study explores sentiment analysis of Weibo comments using a Transformer model to understand public attitudes toward the film The Boy and the Heron. Most comments were positive, indicating strong support for the film. The study highlights the importance of advanced natural language processing (NLP) techniques for sentiment analysis and their implications for social media analytics tools. The research methodology includes dataset preparation, model development, training, evaluation, and result analysis. The study concludes with suggestions for further research and improvements in sentiment analysis techniques [7].

2	Sentiment Analysis of LinkedIn Application Reviews on Google Play Store Using Support Vector Machine	This study investigates public sentiment toward the “Pecat Sri Mulyani” issue on Twitter using Naïve Bayes and Support Vector Machine (SVM) algorithms. A dataset of 1,958 tweets containing the hashtag was analyzed to classify responses into positive and negative sentiments. The results showed that Naïve Bayes achieved an accuracy of 96.14%, correctly classifying 296 negative and 302 positive sentiments, while SVM reached 95.13% accuracy, predicting 254 negative and 342 positive sentiments. These findings indicate that both algorithms are effective for sentiment classification, with Naïve Bayes slightly outperforming SVM [8].
2	Sentiment Analysis of Reksadana on Bibit Applications Using the Naïve Bayes Method and K-Nearest Neighbor (KNN)	This journal discusses sentiment analysis of user opinions on the Bibit mutual fund investment app using the K-Nearest Neighbor (KNN) and Naïve Bayes methods. A total of 30,708 reviews were scraped and processed, with sentiment labeling conducted using the TextBlob library. The classification accuracy using KNN was 99% for both YouTube and app review data, while Naïve Bayes achieved 99% and 98%. The study concludes that neutral sentiments were the most common, followed by positive, and then negative sentiments [4].
3	Moderna's Vaccine Using the K-Nearest Neighbor (KNN) Method: An Analysis of Community Sentiment on Twitter	This study analyzes community sentiment on Twitter about the Moderna vaccine in May 2022 using the K-Nearest Neighbor (KNN) method. KNN proved effective in classifying sentiments into positive, neutral, and negative categories. The study emphasizes KNN's ability to interpret public sentiment, despite challenges in processing the Indonesian language. Performance evaluation using various metrics provides a better understanding of the algorithm's effectiveness in public sentiment analysis [9].
4	Lexicon-Based Sentiment Analysis to Identify Beach Tourism Trends in Yogyakarta Using Twitter Data	This research implements a sentiment analysis system using a lexicon-based approach to identify popular beach tourism destinations in Yogyakarta through hashtags or specific keywords on Twitter. The top five frequently mentioned words were ‘slili’, ‘wisata’, ‘banget’, ‘siung’, and

		‘wajib’. Based on word cloud visualizations, Slili Beach emerged as the most trending destination. Sentiment analysis shows that most users expressed positive opinions about beach tourism in Yogyakarta [10].
--	--	---

Based on the studies summarized in the table above, it can be concluded that sentiment analysis is a valuable approach for understanding public perception across various domains, including entertainment, finance, health, and tourism. Among the methods used, the K-Nearest Neighbors (KNN) algorithm consistently demonstrates high accuracy and effectiveness in processing text-based data, while the Support Vector Machine (SVM) algorithm is widely recognized for its robustness in handling high-dimensional data and strong performance in text classification tasks. These findings reinforce the relevance of applying both KNN and SVM in this research to classify and analyze Starbucks customer reviews, ultimately supporting decision-making processes in service improvement strategies and customer satisfaction enhancement.

3. Method

This research was designed through a series of structured and interconnected methodological stages, starting from data collection to the evaluation of classification results. To provide a more comprehensive understanding, the following flowchart illustrates the entire process carried out in this study.

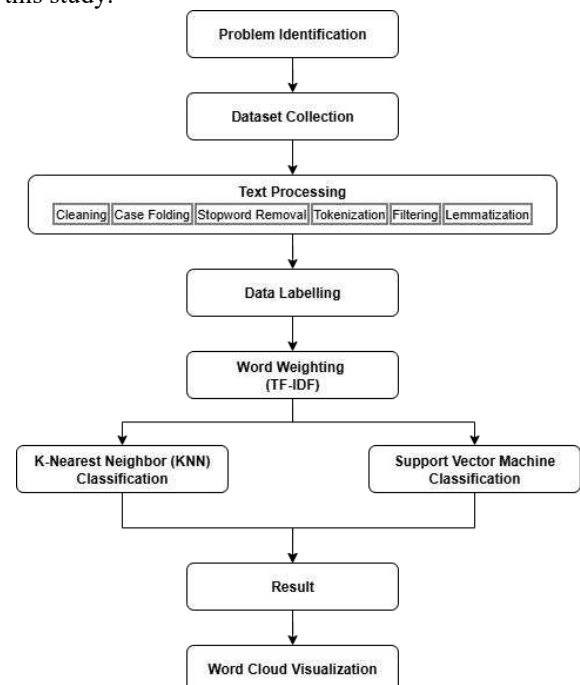


Figure 1. Research Methodology Flowchart

3.1 Dataset Collection

This dataset contains a comprehensive collection of consumer reviews and ratings for Starbucks, a well-known coffeehouse chain. The data was collected through web scraping and includes text reviews, star ratings, location

information, and image links from several pages on the ConsumerAffairs website. It offers valuable insights into customer sentiment and feedback regarding Starbucks locations. The dataset consists of 850 entries and includes the following fields:

- Name: The name of the reviewer, if available.
- Location: The location or city associated with the reviewer, if provided.
- Date: The date when the review was posted.
- Rating: The star rating given by the reviewer, ranging from 1 to 5.
- Review: The textual content of the review, capturing the reviewer's experience and opinions.
- Image Link: A link to an image associated with the review, if available.

The dataset can be accessed through the following link: <https://www.kaggle.com/datasets/harshalhonde/starbucks-reviews-dataset>

4. Results and Discussion

4.1 Text Processing

Text processing is an essential step in preparing raw textual data for effective analysis. It involves several stages, including cleaning to remove irrelevant elements, case folding to standardize text by converting it to lowercase, stopword removal to eliminate common words with little analytical value, tokenization to break text into smaller units, filtering to retain only relevant terms, and stemming to reduce words to their root forms. These steps help make the text more structured, consistent, and suitable for further analysis such as sentiment classification [11].

Table 2. Example Text Processing Results

No	Before	After
1	amber and ladonna at the starbucks on southwest parkway are always so warm and welcoming there is always a smile in their voice when they greet you at the drivethru and their customer service is always spoton they always get my order right and with a smile i would actually give them more than stars if they were available	amber ladonna starbuck southwest parkway always warm welcom always smile voic greet drivethru custom servic alway spoton always get order right smile would actual give star avail
2	at the starbucks by the fire station on in altamonte springs fl made my day and finally helped me figure out the way to make my drink so id love it she took time out to talk to me for minutes to make my	starbuck fire station altamont spring fl made day final help figur way make drink id love took time talk minut make experi better im use much appreci ive bad experi one anoth starbuck that closest work

	experience better than what im used to it was much appreciated ive had bad experiences one after another at the starbucks thats closest to me in my work building with my drinks not being great along with not great customer service from specific baristas niko was refreshing to speak to and pleasant the drink was perfect store	build drink great along great custom servic specif barista niko refresh speak pleasant drink perfect store
3	im on this kick of drinking cups of warm water i work for instacart right now and every location of starbucks i was given free hot water because i asked for it without being charged i really appreciate starbucks for giving me the opportunity to do such thing thats why i give them five stars they fully have my support theyre super nice and professional and the coffee is great go to starbucks	im kick drink cup warm water work instacart right everi locat starbuck given free hot water ask without charg realli appreci starbuck give opportun thing that give five star fulli support theyr super nice profession coffe great go starbuck

The text processing steps significantly transform raw, unstructured reviews into cleaner and more concise forms. As seen in Table 2, the original verbose and informal expressions are converted into simplified, stemmed tokens that retain the essential meaning while eliminating redundant and irrelevant words. This refined version of the text is better suited for computational tasks such as sentiment analysis, enabling more accurate and efficient classification.

4.2 Text Processing

This dataset contains a comprehensive collection of consumer reviews and ratings for Starbucks, a well-known coffeehouse chain. The data was collected through web scraping and includes text reviews, star ratings, location information, and image links from several pages on the ConsumerAffairs website. It offers valuable insights into customer sentiment and feedback regarding Starbucks

4.3 Data Labelling

Data labeling in sentiment analysis is the process of assigning labels to each text based on its emotional content, such as Positive, Negative, or Neutral. In the code provided, this process is carried out using the TextBlob library, which calculates two values: polarity (ranging from -1 for very negative to +1 for very positive) and subjectivity (ranging

from 0 for objective to 1 for subjective). Based on the polarity value, each text is labeled negative values as Negative, zero as Neutral, and positive values as Positive. This results in a dataset with labeled sentiments, which can then be used for further analysis such as visualization, filtering, or training machine learning models. The results of this labeling process can be seen in Table 3.

Table 3. Example Sentiment Analysis Results

No	Text	Polarity	Sentiment
1	The Texas City Starbucks has the best people working. They are always friendly and helpful. But the morning staff goes above and beyond even at their busiest times! Food is excellent.	1.000	Positive
2	There was a long black hair in my blended beverage. Not on my sandwich which I might be able to see before it entered my mouth, but inside my blended beverage which was in turn sucked through my straw into my mouth with who knows how many smaller chopped up pieces... Iâ€™m done with Starbucks as a whole.	-0.004167	Negative
3	You need to do something to help or straighten out the Ukiah, CA store on Perkins St. We went there today, 2/3/12 at 10:00AM. The drive-up line was all the way to the street. I went inside and about a dozen people were in line, but it moved alright. I placed my order and stood with the group waiting for their drinks. My breakfast sandwich was handed to me quickly.	0.0	Neutral

Based on Table 3, the results show the analysis of sentences categorized into positive, negative, and neutral sentiments. The classification is determined by evaluating the polarity score of each sentence, where a polarity greater than 0 indicates positive sentiment, a value less than 0 indicates negative sentiment, and a value equal to 0 indicates neutral sentiment. Positive sentences generally contain praise for the service or product, negative sentences express complaints or dissatisfaction, while neutral sentences tend to be informative without emotional content. This analysis helps systematically categorize customer responses based on their emotional tone.

4.4 Word Weighting

Word weighting using the TF-IDF (Term Frequency–Inverse Document Frequency) method is a text processing technique used to evaluate the importance of a word in a document relative to a collection of documents (corpus). TF (Term Frequency) measures how frequently a term appears in a single document, while IDF (Inverse Document Frequency) reduces the weight of common words that appear across many documents. The combination of these two scores results in the TF-IDF value, which highlights words that are more unique and relevant within a specific context. This makes TF-IDF an effective method for tasks such as text classification, information retrieval, and feature extraction [12]. It can be seen in Table 4.

Table 4. Example Word Weighting Results

No	abl	absolut	accept	..	your	youv	zero
1	0.0	0.0	0.0	..	0.0	0.0	0.0
2	0.0	0.0	0.0	..	0.0	0.0	0.0
3	0.0	0.0	0.0	..	0.0	0.0	0.0
4	0.0	0.0	0.0	..	0.0	0.0	0.0

Based on Table 4, the results of word weighting using the TF-IDF method show the importance level of each word within a document relative to the entire corpus. Low or zero TF-IDF values indicate that the words are either not significant in the specific document or are too common across all documents. Conversely, words with higher TF-IDF values (though not shown in this table) are considered more relevant and unique in context. This demonstrates that the TF-IDF method is effective in filtering out less important terms and highlighting key words for further analysis in tasks such as text classification or information retrieval.

4.5 K-Nearest Neighbor (KNN) Classification

K-Nearest Neighbor (KNN) is a classification algorithm that works by finding the k nearest neighbors of a given data point based on a certain distance metric, such as Euclidean distance. It is a non-parametric and lazy learning algorithm because it does not require an explicit training process instead, it stores all training data and performs classification when new data is input. KNN is widely used in sentiment analysis due to its simplicity and effectiveness in handling labeled data. KNN was successfully applied to predict public sentiment regarding the implementation of government-contracted teachers (PPPK) on Twitter, achieving an accuracy of 73.41% [5]. Therefore, this study also adopts the KNN algorithm to classify public opinions based on the available dataset.

The division of training and testing data in this study was done using three ratios: 60%:40%, 70%:30%, and 80%:20%. This variation was chosen to evaluate how different proportions of training data affect the performance of the K-Nearest Neighbor (KNN) algorithm. A higher percentage of training data may help the model learn better patterns, while a higher percentage of testing data allows for more robust evaluation of the model's generalization

ability. By testing multiple splits, this study aims to find the most optimal balance between training accuracy and testing reliability. It can be seen in Table 5.

Table 5. KNN Results

N o	Training Data (%)	Testing Data (%)	Accur acy (%)	Preci sion (%)	Rec all (%)	F1- Score (%)
1	60	40	87	84	87	84
2	70	30	86	84	86	83
3	80	20	86	84	86	83

Based on Table 5, the best results of the KNN method were obtained with a 60%:40% training-to-testing data split, achieving an accuracy of 87%, precision of 84%, recall of 87%, and an F1-score of 84%. This indicates that in this configuration, the model maintained a good balance between precision, which measures the correctness of positive predictions, and recall, which reflects the ability to identify all relevant data, resulting in the highest F1-score. Meanwhile, at the 70%:30% and 80%:20% splits, accuracy slightly decreased to 86%, with precision remaining at 84%, recall dropping to 86%, and F1-score to 83%, suggesting that increasing training data does not always improve performance since KNN is sensitive to variations and noise in the data.

4.6 Support Vector Machine (SVM) Classification

Support Vector Machine (SVM) is a supervised learning algorithm widely used for classification tasks. It works by finding the optimal hyperplane that separates data into different classes with the maximum margin, making it effective in handling high-dimensional data and well-suited for text classification problems such as sentiment analysis. In this study, SVM was applied to classify “urgent” and “not urgent” messages in WhatsApp group conversations as part of Natural Language Processing (NLP) [13].

The division of training and testing data in this study was done using three ratios: 60%:40%, 70%:30%, and 80%:20%. This variation was chosen to evaluate how different proportions of training data affect the performance of the Support Vector Machine (SVM) algorithm. A higher percentage of training data may help the model learn better patterns, while a higher percentage of testing data allows for more robust evaluation of the model's generalization ability. By testing multiple splits, this study aims to find the most optimal balance between training accuracy and testing reliability. It can be seen in Table 6.

Table 6. SVM Results

N o	Training Data (%)	Testing Data (%)	Accur acy (%)	Preci sion (%)	Rec all (%)	F1- Score (%)
1	60	40	86	83	86	82
2	70	30	87	84	87	84

3	80	20	88	86	88	85
---	----	----	----	----	----	----

Based on Table 6, the best results of the SVM method were obtained with an 80%:20% training-to-testing data split, achieving an accuracy of 88%, precision of 86%, recall of 88%, and an F1-score of 85%. This shows that with more training data, the SVM model is able to construct a more accurate decision boundary, resulting in better overall performance. At the 70%:30% split, the model achieved slightly lower accuracy of 87% with precision of 84%, recall of 87%, and F1-score of 84%, while at the 60%:40% split, the performance further decreased to 86% accuracy, precision of 83%, recall of 86%, and F1-score of 82%. These results indicate that the SVM algorithm benefits from a larger portion of training data, as it improves the model's ability to learn patterns and generalize more effectively.

4.7 Word Cloud Visualization

Word Cloud is a visualization technique used to represent the most frequently occurring words in a collection of text, where the size of each word indicates its frequency or significance. In this study, Word Clouds are generated separately for positive, negative, and neutral sentiments to provide a clearer picture of the dominant words associated with each sentiment category [14]. These visualizations help in understanding customer perceptions and the emotional tone conveyed in Starbucks reviews. The following figures illustrate the Word Clouds for each sentiment class.

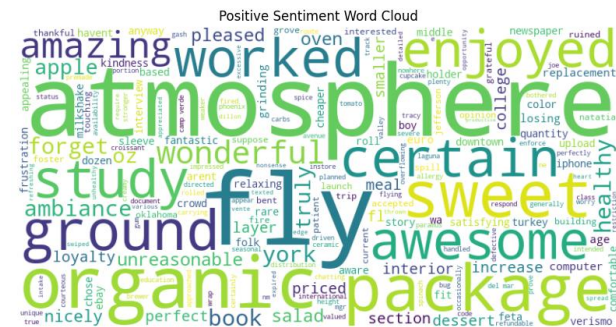


Figure 2. Positive Word Cloud Visualization

Based on the results of the Positive Word Cloud in Figure 2, several prominent words such as *atmosphere*, *organic*, *package*, *fly*, and *awesome* reflect positive aspects that customers value at Starbucks. The word *atmosphere* indicates that many customers appreciate the comfortable and enjoyable environment provided by Starbucks stores. The word *organic* highlights positive perceptions toward product quality, especially for customers who value healthy and natural options. Meanwhile, *package* may represent satisfaction with product packaging or bundled offers that enhance customer experience. The presence of the word *fly* could symbolize convenience, speed, or a pleasant experience during travel-related visits. Additionally,

- [6] M. Makhasinul Maarif and N. Setiyawati, "Analisis Sentimen Review Aplikasi LinkedIn di Google Play Store Menggunakan Support Vector Machine," 2024.
- [7] B. Yang, "Sentiment Analysis Based on Transformer," *International Journal of Computer Science and Information Technology*, vol. 3, no. 1, pp. 144–148, Jun. 2024, doi: 10.62051/ijcsit.v3n1.19.
- [8] S. Berliani and S. Lestari, "Analisis Sentimen Masyarakat Terhadap Isu Pecat Sri Mulyani Pada Twitter Menggunakan Metode Naive Bayes Dan Support Vector Machine," *Jurnal Sains dan Teknologi*, vol. 5, no. 3, pp. 951–960, Apr. 2024, doi: 10.55338/saintek.v5i3.2746.
- [9] M. I. Hutapea and A. P. Silalahi, "Moderna's Vaccine Using the K-Nearest Neighbor (KNN) Method: An Analysis of Community Sentiment on Twitter," *Jurnal Penelitian Pendidikan IPA*, vol. 9, no. 5, pp. 3808–3814, May 2023, doi: 10.29303/jppipa.v9i5.3203.
- [10] A. Rachmadana Ismail, R. Bagus, F. Hakim, and R. Artikel, "Implementasi Lexicon Based Untuk Analisis Sentimen Dalam Mengetahui Trend Wisata Pantai Di DI Yogyakarta Berdasarkan Data Twitter P-ISSN E-ISSN," 2023.
- [11] T. A. Q. Putri, A. Triayudi, and R. T. Aldisa, "Implementasi Algoritma Decision Tree dan Naïve Bayes Untuk Analisis Sentimen Terhadap Kepuasan Pelanggan Starbucks," *Journal of Information System Research (JOSH)*, vol. 4, no. 2, pp. 149, Jan. 2023, doi: 10.47065/josh.v4i2.2949.
- [12] M. D. Rizkiyanto, M. D. Purbolaksono, and W. Astuti, "Sentiment Analysis Classification on PLN Mobile Application Using Random Forest Method and TF-IDF Feature Extraction," *INTEK: Jurnal Penelitian*, vol. 11, no. 1, pp. 37–43, 2024, doi: 10.31963/intek.v11i1.4774.
- [13] D. F. N. Hulu, Jatmika, and Yo'el Pieter Sumihar, "Efikasi Penggunaan Kata Urgent Pada Percakapan Whatsapp Menggunakan Algoritma Support Vector Machine," *JURNAL S DAN KOMPUTER*, vol. 8, no. 02, pp. 43–48, Aug. 2024, doi: 10.179/jurnalinfact.v8i02.579.
- [14] A. I. KABIR, K. AHMED, and R. KARIM, "Word Cloud Sentiment Analysis of Amazon Earphones Reviews with Random Language," *Informatica Economica*, vol. 24, no. 0, pp. 55–71, Dec. 2020, doi: 10.1818/issn14531305/24.4.2020.05.