

Analysis of Customer Perception on Google Maps Reviews of Klinik Kopi Samarinda Using Extreme Gradient Boosting (Xgboost)

M.Ariya Parengrengi

Information System, STMIK Widya
Cipta Dharma
Samarinda, 75123, Indonesia
2241029@wicida.ac.id

Heny Pratiwi*

Information System, STMIK Widya
Cipta Dharma,
Samarinda, 75123, Indonesia
henypratiwi@wicida.ac.id
**Corresponding author*

Muhammad Fahmi

Information System, STMIK Widya
Cipta Dharma,
Samarinda, 75123, Indonesia
mfahmi@wicida.ac.id

 Submitted: 2026-02-11; Accepted: 2026-04-02; Published: 2026-06-01

Abstract— This study addresses the challenge of analyzing large volumes of unstructured customer reviews on Google Maps to understand customer perception of Klinik Kopi Samarinda. The objective is to apply the XGBoost algorithm for sentiment classification and evaluate its effectiveness in identifying positive and negative customer sentiments to support data-driven service improvement. This study uses a quantitative experimental approach with customer review data collected from Google Maps. Text preprocessing techniques, including case folding, tokenization, stopword removal, and stemming, were applied. TF-IDF was used for feature extraction, and sentiment classification was performed using the XGBoost algorithm. Model performance was evaluated using accuracy, precision, recall, and F1-score metrics. The results show that the XGBoost model successfully classified customer sentiment with good performance. Most reviews were categorized as positive, indicating high customer satisfaction with Klinik Kopi Samarinda. The confusion matrix evaluation confirms that the model can effectively identify sentiment patterns and provide reliable insights into customer perception based on Google Maps reviews. This study confirms that the XGBoost algorithm is effective for classifying customer sentiment from Google Maps reviews and provides valuable insights into customer perception. These findings can support service evaluation and improvement. Future research is recommended to use larger datasets and compare multiple algorithms to enhance sentiment classification accuracy and reliability.

Keywords— Sentiment Analysis, Xgboost, Machine Learning, Customer Perception, Google Maps

I. INTRODUCTION

The rapid advancement of digital technology has significantly changed how customers share their experiences and evaluate products and services. Online review platforms such as Google Maps allow customers to provide feedback in the form of ratings and textual

comments, which can influence business reputation and consumer decision-making. These reviews represent a valuable source of information that reflects customer satisfaction, service quality, and overall experience. Customer feedback shared through digital platforms plays an important role in shaping consumer perceptions and supporting business evaluation processes. Recent studies indicate that online reviews have a strong impact on customer trust and purchasing decisions, especially in service-based industries such as food and beverage businesses

The increasing number of customer reviews generated on digital platforms has created challenges in analyzing the data effectively. Customer reviews are typically presented in unstructured textual format, which makes manual analysis inefficient, time-consuming, and prone to subjective interpretation. Traditional analysis methods are not capable of processing large volumes of textual data efficiently, resulting in limited utilization of valuable customer feedback. As a result, businesses may fail to identify important patterns related to customer satisfaction, service weaknesses, and areas that require improvement. Therefore, automated approaches using computational techniques are necessary to analyze textual data efficiently and accurately

Sentiment assessment emerges as a common technique in text mining that identifies and groups feelings expressed within written material. This method permits the automatic arrangement of customer feedback into feeling groups such as favorable, unbiased, and unfavorable. Sentiment assessment assists organizations in understanding user perspectives, judging program performance, and promoting choices supported by evidence. Current research has suggested that sentiment assessment is highly capable in extracting useful understandings from internet reviews and supplying details that may boost service quality and client satisfaction. Through employing sentiment assessment, businesses can discover trends in user comments and handle user needs in a better manner.

Machine learning methods are widely employed for improving how well sentiment analysis tools work. These methods allow for the automatic gaining of understanding from data and can sort text quite accurately. Several machine learning ways, such as Naive Bayes, Support Vector Machine, and Random Forest, have been successfully used in jobs involving sentiment categorization. More recent studies indicate that ensemble learning approaches produce better outcomes compared to standard machine learning methods. These techniques combine multiple models to improve prediction accuracy and reduce classification errors.

Among the premier ensemble learning methods is Extreme Gradient Boosting (XGBoost), which earned substantial acclaim for its outstanding results and great impact (Wardana et al., 2025). XGBoost represents an advanced form of gradient boosting that utilizes decision trees as its base estimators and boosts predictive precision through repeated improvement steps. This approach incorporates regularization methods to lessen the possibility of overfitting and boost the model's capacity to generalize. Furthermore, XGBoost manages complex data dimensions and very large data collections with notable speed. Current studies suggest that XGBoost outperforms other machine learning procedures in jobs like sentiment categorization, particularly when assessing feedback from online patrons

Its ability to maximize categorization outcomes makes it suitable for scrutinizing sentiment data derived from text. Most previous studies have focused on large-scale platforms such as e-commerce websites, social media platforms, or global service providers. Local businesses, such as coffee shops, have unique characteristics and customer interaction patterns that require specific analysis. Understanding customer perception in these businesses is important to support service improvement and maintain competitiveness. (Zhang et al., 2022).

Klinik Kopi Samarinda is one of the coffee shops that has received numerous customer reviews on Google Maps. These reviews reflect customer experiences related to service quality, product quality, and overall satisfaction. The large volume of reviews makes it difficult to manually identify overall sentiment patterns and extract meaningful insights from the data. These reviews have not been analyzed systematically using machine learning techniques to extract meaningful insights. Without proper analysis, valuable information contained in customer reviews cannot be utilized effectively to support business evaluation and improvement. This situation highlights the need for an automated sentiment analysis approach to analyze customer feedback efficiently. (Tribuana et al., 2025).

This research addresses the existing gap by implementing the XGBoost algorithm to classify customer sentiment based on Google Maps reviews of Klinik Kopi Samarinda. The novelty of this study lies in the application of XGBoost for sentiment analysis using real-world review data from a local coffee shop. Unlike previous research that mainly focused on large commercial platforms, this study emphasizes sentiment analysis in a

local business context using Google Maps reviews. This study evaluates the performance of the XGBoost algorithm in sentiment classification tasks involving textual customer feedback.

The results of this study include developing a sentiment analysis system employing the XGBoost algorithm, assessing shopper feelings found in Google Maps reviews, and presenting evidence based on facts to support choices for commerce. It is hoped that the discoveries from this project will let company owners secure a clearer understanding of client perspectives and identify areas needing improvement. Furthermore, this examination demonstrates the capability of machine learning techniques for scrutinizing text-based feedback data and assists in the progress of sentiment analysis tools for neighborhood enterprise environments (Rahman et al., 2025).

II. LITERATURE REVIEW

A. Sentiment analysis

1. Definition of sentiment analysis

Sentiment analysis is a text mining technique used to identify and classify opinions expressed in textual data. This method allows researchers and organizations to analyze subjective information such as customer opinions, emotions, and experiences toward products or services. According to (Rahman et al., 2025), sentiment analysis enables automated processing of textual feedback and provides insights into customer satisfaction and service performance. This technique is widely applied to online reviews, social media data, and customer feedback platforms to evaluate public perception.

Sentiment analysis typically classifies textual data into three main categories: positive, neutral, and negative sentiment. Positive sentiment indicates customer satisfaction, negative sentiment reflects dissatisfaction, and neutral sentiment represents opinions that do not express strong emotions. It is stated that sentiment classification helps organizations identify customer perception trends and evaluate service quality effectively.

2. Importance of sentiment analysis in customer review evaluation

Customer reviews provide valuable information about customer experiences and satisfaction. Analyzing these reviews manually is inefficient due to the large volume of data generated on digital platforms. Sentiment analysis provides an automated solution that enables efficient analysis of large-scale textual data. It is explained that sentiment analysis helps organizations extract meaningful insights from customer reviews and supports data-driven decision-making.

Sentiment analysis allows businesses to identify strengths and weaknesses in their services. By understanding customer sentiment, businesses can improve service quality and enhance customer satisfaction. It has become an essential tool for evaluating customer feedback in digital environments.

B. Text preprocessing and feature extraction

1. Text preprocessing

Text preprocessing is an important step in sentiment analysis that aims to improve data quality before classification. Raw textual data often contains irrelevant information, inconsistencies, and noise that can affect classification accuracy. Therefore, preprocessing techniques are applied to clean and standardize the data. explained that preprocessing improves data consistency and enhances machine learning model performance.

Common preprocessing techniques include case folding, tokenization, stopwords removal, and stemming. Case folding converts all characters into lowercase to ensure consistency. Tokenization divides text into individual words or tokens. Stopword removal eliminates common words that do not contribute meaningful information. Stemming reduces words to their root form to simplify analysis. These techniques improve the quality of textual data and enhance classification performance

2. Feature extraction using TF-IDF

Feature extraction is used to convert textual data into numerical form so that it can be processed by machine learning algorithms. One of the most commonly used feature extraction techniques is Term Frequency–Inverse Document Frequency (TF-IDF). This method measures the importance of words based on their frequency and distribution in the dataset. (Maharani et al., 2026).

TF-IDF assigns higher weights to words that appear frequently in specific documents but rarely in other documents. This approach helps highlight important words and improve classification accuracy. (Donita Maharani et al., 2026) TF-IDF is an effective method for representing textual data in sentiment analysis tasks.

C. Machine learning and XGBoost for sentiment classification

1. Machine learning and XGBoost for sentiment classification

Machine learning algorithms enable automatic classification of textual data by learning patterns from training data. These algorithms can analyze large datasets efficiently and provide accurate classification results (Gifari et al., 2022). stated that machine learning improves sentiment classification performance compared to manual analysis methods.

Several machine learning algorithms have been widely used in sentiment analysis, including Naive Bayes, Support Vector Machine, and Random Forest. These algorithms provide good classification performance, but recent studies have shown that ensemble learning methods provide better accuracy and reliability

2. Extreme Gradient Boosting (XGBoost)

Extreme Gradient Boosting (XGBoost) is an advanced ensemble learning algorithm based on gradient boosting techniques. This algorithm builds multiple decision tree models iteratively to improve classification performance. explained that XGBoost improves prediction accuracy by minimizing classification errors through optimization techniques. (Merpatika & Chandra, 2025).

XGBoost includes regularization mechanisms that prevent overfitting and improve model generalization. The algorithm is also capable of handling large datasets and high-dimensional data efficiently. Previous studies have shown that XGBoost provides high performance in classification tasks involving textual data.

Recent studies have shown that XGBoost achieves better performance compared to traditional machine learning algorithms in sentiment classification tasks. Its efficiency, accuracy, and scalability make it suitable for analyzing customer sentiment based on online review data (Noorunnahar et al., 2023).

III. METHODS

This study uses a quantitative experimental approach to analyze customer sentiment from Google Maps reviews of Klinik Kopi Samarinda using the XGBoost algorithm. The process involves transforming textual data into numerical features, training a classification model, and evaluating its performance using standard metrics. The research consists of several stages, including data preparation, text preprocessing, feature extraction, model training, and evaluation.

The research methodology used in this study is shown in Figure 1.

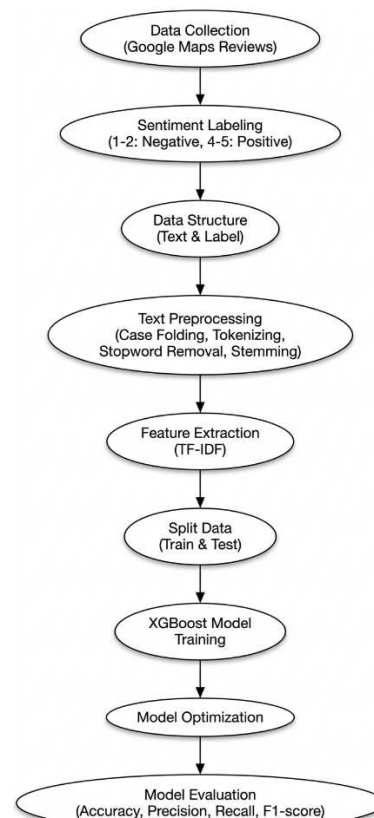


Figure 1. Sentiment classification workflow

1. Data Source

The dataset used in this study consists of customer reviews collected from Google Maps related to Klinik Kopi Samarinda. The data include textual review content and rating scores provided by customers. These reviews

reflect customer experiences and perceptions of service quality, product quality, and overall satisfaction. The rating values range from 1 to 5, which are used as indicators for sentiment labeling.

2. Sentiment Labeling

The sentiment classification process is performed using rating-based labeling to ensure objective categorization. Reviews with ratings of 4 and 5 are classified as positive sentiment, rating 3 is classified as neutral sentiment, and ratings of 1 and 2 are classified as negative sentiment. This labeling approach ensures consistency and enables the classification model to learn sentiment patterns based on customer evaluations. (Latif & Wildah, 2025).

3. Data Structure

Each record in the dataset consists of two main attributes, namely review text and sentiment label. The review text serves as the input feature, while the sentiment label becomes the target variable for classification. The dataset is then prepared for preprocessing and feature transformation.

4. Text Preprocessing

1) Case Folding

Case folding is performed to convert all characters in the review text into lowercase format. (Rabbani et al., 2023) This process ensures uniformity and prevents duplication of words that differ only in letter capitalization.

2) Tokenization

Tokenization is the process of dividing text into smaller units called tokens, typically in the form of individual words. This step allows the model to examine each element of the sentence independently and detect meaningful patterns.

3) Stopword Removal

Stopword removal eliminates common words that do not carry significant meaning, such as conjunctions, prepositions, and frequently occurring words. Removing stopwords helps improve model performance by reducing irrelevant features.

The preprocessing stage plays an important role in ensuring the quality of the data before classification. Raw text data often contain inconsistencies, unnecessary symbols, and irrelevant words that may affect classification accuracy. By applying preprocessing techniques, the text becomes cleaner, more structured, and easier for the model to analyze.

Preprocessing helps reduce the number of irrelevant features, allowing the model to focus on words that contribute to sentiment classification. This improves the efficiency and accuracy of the classification process. Proper preprocessing is essential to ensure that the model can learn meaningful patterns from the data and produce reliable results.

4) Stemming

Stemming is applied to convert words into their root form. This process reduces variations of words with similar meanings into a single representation, improving feature consistency and reducing dimensionality. (Gifari et al., 2022).

5. Feature Extraction Using TF-IDF

After preprocessing, textual data are converted into numerical form using Term Frequency-Inverse Document Frequency (TF-IDF). This method measures the importance of words in a document relative to the entire dataset.

6. Term Frequency (TF)

The Term Frequency (TF) measures how frequently a term appears in a document. It is defined as (1):

$$TF(t, d) = \sum_{w \in df(w, d)} f(t, d) \quad (1)$$

$TF(t, d)$ represents the frequency of term t in document d , $f(t, d)$ represents the number of occurrences of term t in document d , and $\sum f(w, d)$ represents the total number of terms in document d .

The Inverse Document Frequency (IDF) is calculated using Equation (2):

$$IDF(t) = \log(N / df(t)) \quad (2)$$

N represents the total number of documents and $df(t)$ represents the number of documents containing term t .

The $TF - IDF$ value is obtained by combining TF and IDF , as shown in Equation (3):

$$TF - IDF(t, d) = TF(t, d) \times IDF(t) \quad (3)$$

This transformation converts textual data into numerical vectors that can be processed by the classification algorithm.

7. Classification Using XGBoost Algorithm

1) XGBoost Model Implementation

Extreme Gradient Boosting (XGBoost) is a machine learning algorithm that operates within the gradient boosting framework. It constructs classification models by sequentially combining multiple decision trees to improve predictive accuracy and minimize errors. Each newly generated tree focuses on correcting the errors of the preceding trees, resulting in a more reliable and effective prediction model. (Wati et al., 2025).

The prediction of the XGBoost model can be expressed using Equation (4):

$$\hat{y}_i = \sum f_k(x_i) \quad (4)$$

$\hat{y}_i$ represents the predicted value, f_k represents the decision tree function, and x_i represents the input feature vector.

2) Model Training

The dataset is divided into training and testing sets. The training set is used to build the classification model, while the testing set is used to evaluate the model's performance. Using the training data, the XGBoost algorithm learns patterns and generates optimal decision trees to classify sentiment accurately.

The dataset is partitioned into training and testing data. The training data are used to construct the classification model, whereas the testing data are used to

are treated as the same feature. Tokenization breaks text into individual words, enabling more detailed analysis, while stopword removal eliminates common words with limited meaning. Stemming further improves consistency by reducing words to their root forms, minimizing variations of similar terms. (Mao et al., 2023).

Preprocessing helps improve the efficiency of feature extraction. By removing irrelevant and redundant information, the dataset becomes more focused on meaningful words that reflect customer sentiment. This allows TF-IDF to assign more appropriate weights to important terms, which supports better classification performance.

Clean and structured data also help the XGBoost model identify patterns between words and sentiment categories more accurately. (Yang et al., 2025). As a result, preprocessing directly affects both the quality of feature representation and the performance of the classification model.

This stage ensures that the dataset is properly prepared for analysis by improving data quality, reducing noise, and supporting more accurate sentiment classification. Proper preprocessing is essential to obtain reliable results in sentiment analysis using machine learning methods. (Surala et al., 2025).

C. Term Weighting Results Using TF-IDF

During the feature extraction stage, the Term Frequency-Inverse Document Frequency (TF-IDF) method was applied to convert textual data into numerical representations suitable for machine learning algorithms. The process resulted in 808 unique features derived from the corpus after undergoing preprocessing steps, including case folding, filtering, and stemming. These features define the dimensional space used as input for the classification model.

TF-IDF assigns weights to each term by considering both its frequency within a document and its distribution across the entire dataset. Terms that frequently appear in specific documents but occur less often in others are given higher weights, as they carry stronger discriminative information. In contrast, terms that appear consistently across many documents receive lower weights due to their limited contribution to distinguishing between sentiment categories.

Based The resulting feature set enables the model to capture meaningful textual patterns and differences between sentiment classes. By utilizing these weighted features, the model is expected to improve classification performance and maintain computational efficiency. This finding is consistent with previous studies, which indicate that TF-IDF is effective in representing textual data for sentiment classification tasks. (Jelita & Sa'ad, 2025).

The model's performance was evaluated using test data represented in a Confusion Matrix to examine its ability to classify sentiments accurately. Based on the testing results involving 64 data points, the model shows strong performance in identifying positive sentiment, with 53 instances correctly classified as True Positive.

The model correctly classified 5 instances of negative sentiment as True Negatives. Several misclassifications

were still observed, with 5 negative sentiment instances incorrectly predicted as positive (False Positives) and 1 positive sentiment instance misclassified as negative (False Negative).

These results indicate that the model performs well in detecting positive sentiment, yet it is less reliable in identifying negative sentiment. The presence of misclassified negative instances suggests that the model tends to assign positive labels more frequently. This tendency may be influenced by an imbalance in the distribution of training data, where positive sentiment is more dominant. As a result, the model learns patterns associated with positive sentiment more effectively than those related to negative sentiment, as illustrated in Figure 3.

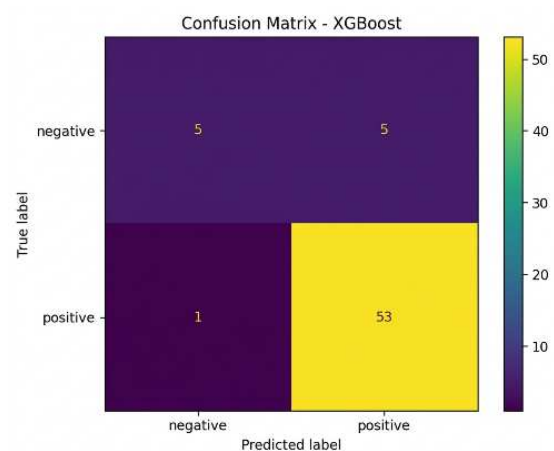


Figure 3. Confusion Matrix – XGBoost

The results of this study show that sentiment analysis using the XGBoost algorithm is able to provide meaningful insights into customer perception based on Google Maps reviews. Through systematic preprocessing and feature extraction, textual data were transformed into structured numerical representations that can be effectively processed by the classification model. This approach enables the identification of sentiment patterns that reflect customer satisfaction and dissatisfaction, while also allowing large volumes of unstructured data to be analyzed more efficiently than manual methods.

The classification results indicate that the XGBoost algorithm is capable of learning relationships between textual features and sentiment categories. The model is able to distinguish between positive and negative sentiment based on patterns identified in the review data. This demonstrates that XGBoost is a reliable method for sentiment classification, particularly for analyzing customer reviews from online platforms. The combination of preprocessing techniques, TF-IDF feature extraction, and ensemble learning contributes to consistent classification performance (Rabbani et al., 2023).

An analysis of TF-IDF feature weights shows that positive reviews are frequently associated with terms such as "taste," "comfortable atmosphere," and "friendly service." These terms indicate that product quality and service experience are key factors influencing positive

customer perception. In contrast, negative reviews are often related to terms such as “waiting time” and “crowded,” suggesting that service efficiency and space management during peak hours may contribute to customer dissatisfaction. These findings provide practical implications for Klinik Kopi Samarinda, particularly in maintaining product quality while improving service responsiveness and operational efficiency.

Customer reviews also serve as an important source of information for business evaluation. Feedback provided by customers reflects their experiences, expectations, and satisfaction levels. By analyzing this information systematically, businesses can identify strengths that need to be maintained and recognize aspects that require improvement. This supports more informed decision-making and contributes to improving service quality.

The application of machine learning-based sentiment analysis also supports the development of data-driven evaluation approaches. This method allows businesses to monitor customer perception continuously and respond to feedback more effectively. Sentiment analysis can therefore be utilized as a practical tool to understand customer perception and support service improvement strategies.

The integration of text preprocessing, TF-IDF feature extraction, and XGBoost classification is effective in transforming unstructured Google Maps reviews into quantifiable insights. The dominance of positive sentiment indicates that customers generally have a favorable perception of Klinik Kopi Samarinda, while the presence of negative sentiment highlights aspects that still require attention and improvement.

Table 2 presents the evaluation metrics of the XGBoost classifier. The model achieves an accuracy of 0.91, indicating strong classification performance. The recall for positive sentiment reaches 0.98, while the recall for negative sentiment is 0.50, showing that the model is more effective in identifying positive reviews. The macro-averaged F1-score of 0.79 reflects the overall balance of model performance across both sentiment classes.

Table 2. Evaluation metric

Metric	Negative (0)	Positive (1)	Macro Avg	Weighted Avg
Precision	0.83	0.91	0.87	0.90
Recall	0.50	0.98	0.74	0.91
F1-Score	0.62	0.95	0.79	0.90
Accuracy	-	-	-	0.91

D. Discussion and Implications

The results of this study indicate that the XGBoost algorithm shows satisfactory performance in classifying customer sentiment based on Google Maps reviews. Several methodological aspects need to be considered. The dataset consists of 341 reviews after excluding neutral ratings. Although larger than the previous dataset, the data size remains relatively limited, which may affect generalization and increase the risk of overfitting. While XGBoost includes regularization

mechanisms, limited data variability can still influence model robustness.

The exclusion of neutral (3-star) reviews simplifies the task into binary classification. This approach helps the model distinguish between classes more clearly, but it reduces the representation of real customer perception, which is often more complex. Future studies may consider multi-class classification to provide a more complete understanding of sentiment patterns.

Rating-based labeling assumes that star ratings are consistent with the sentiment expressed in the review text. In practice, this assumption does not always hold, as some reviews may contain mixed or contradictory opinions. Applying manual validation or measuring agreement between annotators can improve the reliability of the labeling process.

Considering The distribution of sentiment shows an imbalance between positive (80.65%) and negative (19.35%) reviews. This imbalance can influence model performance, especially in detecting negative sentiment. Therefore, evaluation metrics such as precision, recall, and F1-score, particularly in macro and weighted averages, are important to provide a more balanced assessment. The use of cross-validation in future work can also help improve model stability and reduce bias. (Markus & Fenriana, 2025).

From a managerial perspective, the analysis of TF-IDF features shows that positive reviews are often associated with taste quality, comfortable ambiance, and friendly service. On the other hand, negative reviews tend to highlight waiting time and service speed. These findings suggest that Klinik Kopi Samarinda should maintain product quality and service friendliness while improving operational efficiency, especially during busy periods.

This study demonstrates the use of ensemble learning methods for sentiment analysis in a local business context. Further improvements can be made by using larger datasets, applying cross-validation, and comparing different classification algorithms to strengthen the reliability and generalization of the results.

V. CONCLUSION

This study examined customer perceptions of Klinik Kopi Samarinda by applying sentiment analysis to Google Maps reviews using the XGBoost algorithm. The dataset consisted of textual reviews and star ratings provided by customers. The data were processed through several preprocessing steps, including case folding, tokenization, stopword removal, and stemming, to improve data quality and consistency. The processed text was then transformed into numerical features using the TF-IDF method, allowing the model to analyze the data effectively.

The results show that most customer reviews fall into the positive sentiment category. This indicates that Klinik Kopi Samarinda is generally well perceived by its customers. Positive sentiment is associated with satisfaction related to service quality, product quality, and the overall atmosphere of the coffee shop. A smaller

portion of reviews is classified as negative, suggesting that certain aspects still need improvement.

Based on the confusion matrix, the XGBoost model demonstrates satisfactory performance in classifying sentiment. Most reviews are correctly classified, indicating that the model is able to capture relevant patterns in the data. The use of TF-IDF contributes to this performance by assigning appropriate importance to key terms within the dataset.

The findings show that sentiment analysis can be used to extract meaningful information from online review data. The results can support business evaluation by providing insights into customer perceptions and identifying factors that influence satisfaction. These insights can be used to maintain aspects that are already well received and improve areas that receive negative feedback.

The implementation of the XGBoost algorithm also shows that it can handle text classification effectively with stable performance. The model is able to learn patterns from customer reviews and produce consistent predictions. This makes XGBoost a suitable approach for sentiment analysis, especially for analyzing customer feedback from platforms such as Google Maps. The use of machine learning helps reduce the limitations of manual analysis, which is often time-consuming and prone to bias.

Customer review data also play an important role as a source of information for business development. Online reviews reflect real customer experiences that can be used to evaluate service quality and understand customer expectations. Through systematic analysis, businesses can make more informed decisions and develop strategies that align with customer needs.

Future studies can improve this work by using a larger dataset to enhance model generalization, including additional sentiment categories such as neutral sentiment, and comparing different classification algorithms to obtain more comprehensive results.

REFERENCES

- Gifari, O. I., Adha, M., Hendrawan, I. R., & Durrand, F. F. S. (2022). Analisis Sentimen Review Film Menggunakan TF-IDF dan Support Vector Machine. *Journal of Information Technology*, 2(1), 36-40. <https://doi.org/10.46229/jifotech.v2i1.330>
- Hendrawan, I. R. (2022). Perbandingan Algoritma Naïve Bayes, Svm Dan Xgboost Dalam Klasifikasi Teks Sentimen Masyarakat Terhadap Produk Lokal Di Indonesia. *Transformasi*, 18(1). <https://doi.org/10.56357/jt.v18i1.295>
- Jelita, H. P., & Sa'ad, M. I. (2025). Penerapan Algoritma Naïve Bayes Dalam Analisis Sentiment Masyarakat Terhadap STMIK Widya Cipta Dharma. *Bulletin Of Information Technology (BIT)*, 6(2), 148-160. <https://doi.org/10.47065/bit.v6i2.2029>
- Latif, A., & Wildah, S. K. (2025). Analisis Kinerja Algoritma Ensemble Dalam Prediksi Perilaku Pembelian Pelanggan. *Jati (Jurnal Mahasiswa Teknik Informatika)*, 9(1), 557-563. <https://doi.org/10.36040/jati.v9i1.12413>
- Maharani, M. D., Indradewi, I. G. A. A. D., & Wijaya, I. N. S. W. (2026). Perbandingan Performa Tf-Idf Dan Bow Pada Analisis Sentimen Bpjs Kesehatan Menggunakan Xgboost. *Jurnal Pendidikan Teknologi Dan Kejuruan*, 23(1), 13-24. <https://doi.org/10.23887/jptk-undiksha.v23i1.109604>
- Mao, C., Xu, W., Huang, Y., Zhang, X., Zheng, N., & Zhang, X. (2023). Investigation Of Passengers' Perceived Transfer Distance In Urban Rail Transit Stations Using Xgboost And SHAP. *Sustainability*, 15(10), 7744. <https://doi.org/10.3390/Su15107744>
- Marcellina, M., & Mukhlason, A. (2024). Analisis Prediktif Churn Untuk Meningkatkan Tingkat Retensi Pelanggan Pada Perusahaan Saas Menggunakan Machine Learning. *ILKOMNIKA*, 6(1), 21-32. <https://doi.org/10.28926/ilkomnika.v6i1.598>
- Markus, P., & Fenriana, I. (2025). Prediksi Dan Interpretasi Tingkat Churn Pelanggan Di Pt. Bpr Magga Jaya Utama Menggunakan Extreme Gradient Boosting Dan Shap Untuk Mendukung Pengambilan Keputusan. *Poters (Proceedings Of Technology, Engineering And Computers)*, 1(2), 301-307.
- Merpatika, D. F., & Chandra, A. Y. (2025). Analisis Sentimen Publik Di Media Sosial Terhadap Kasus Dugaan Korupsi Impor Minyak Pertamina Menggunakan Xgboost. *Rabit: Jurnal Teknologi Dan Sistem Informasi Univrab*, 10(2), 576-593. <https://doi.org/10.36341/rabit.v10i2.6307>
- Noorunnahar, M., Chowdhury, A. H., & Mila, F. A. (2023). A Tree Based Extreme Gradient Boosting (Xgboost) Machine Learning Model To Forecast The Annual Rice Production In Bangladesh. *Plos One*, 18(3), E0283452. <https://doi.org/10.1371/journal.pone.0283452>
- Rabbani, S., Safitri, D., Rahmadhani, N., Sani, A. A. F., & Anam, M. K. (2023). Perbandingan Evaluasi Kernel SVM Untuk Klasifikasi Sentimen Dalam Analisis Kenaikan Harga BBM: Comparative Evaluation Of SVM Kernels For Sentiment Classification In Fuel Price Increase Analysis. *MALCOM: Indonesian Journal Of Machine Learning And Computer Science*, 3(2), 153-160. <https://doi.org/10.57152/malcom.v3i2.897>
- Rahman, A., Yusnita, A., & Ekawati, H. (2025). Penerapan Algoritma Xgboost Dalam Prediksi Harga Sewa Kos Di Kota Samarinda. *Bulletin Of Information Technology (BIT)*, 6(4), 379-390. <https://doi.org/10.47065/bit.v7i1.2304>
- Surala, L., Hananto, A. L., Novalia, E., & Nurapriani, F. (2025). Analisis Status Pembayaran Group Order NooBlue Menggunakan Algoritma XGBoost. *Jurnal Sistem Informasi Triguna Dharma (JURSI TGD)*, 4(4), 838-847. <https://doi.org/10.53513/jursi.v4i4.11279>

- Tribuana, D., Baharuddin, B., & Resky, A. M. (2025). Penerapan Algoritma XGBoost Untuk Prediksi Kepuasan Pelanggan Pada Layanan E-Commerce: Studi Pada Dataset Transaksi Nyata. *Jurnal Teknologi dan Bisnis Cerdas*, 1(1), 50-59. <https://doi.org/10.64476/jtbc.v1i1.5>
- Vidiantara, I. R. A., Yanuarti, R., & Rintyarna, B. S. (2025). Analisis Sentimen Pasien terhadap Layanan Antarmedika Dentalcare Menggunakan XGBoost. *Jurnal Teknologi Informasi dan Ilmu Komputer*, 12(6), 1337-1384. <https://doi.org/10.25126/jtiik.2025126>
- Wardana, D. A. B., Muflikhah, L., & Perdana, R. S. (2025). Analisis Sentimen Ulasan Pengguna Aplikasi MyPertamina Pada Google Play Store Menggunakan Metode Xgboost dan Word2Vec Embedding. *Jurnal Pengembangan Teknologi Informasi dan Ilmu Komputer*, 9(8). <https://j-ptiik.ub.ac.id/index.php/j-ptiik/article/view/15247>
- Wati, H. L., Wiradinata, M. T. I. R., Anggraeni, N., Kolbiah, S., Hendar, U., & Agustina, N. (2025). Perbandingan Algoritma Random Forest Dan XGBoost Untuk Analisis Sentimen Publik Terhadap Rumah Makan Payakumbuh. *Naratif: Jurnal Nasional Riset, Aplikasi dan Teknik Informatika*, 7(1), 64-71. <https://doi.org/10.53580/naratif.v7i1.307>
- Wibowo, J. A., Mawardi, V. C., & Sutrisno, T. (2024). Penerapan Support Vector Machine untuk Analisis Sentimen Fitur Layanan pada Ulasan Gojek. *Jurnal Ilmu Komputer dan Sistem Informasi*, 12(1). <https://doi.org/10.24912/jiksi.v12i1.28211>
- Yang, H., Lee, W., & Kim, J. (2025). Identification of key factors influencing teachers' self-perceived AI literacy: an XGBoost and SHAP-based approach. *Applied Sciences*, 15(8), 4433. <https://doi.org/10.3390/app15084433>
- Zhang, P., Jia, Y., & Shang, Y. (2022). Research and application of XGBoost in imbalanced data. *International Journal of Distributed Sensor Networks*, 18(6), 15501329221106935. <https://doi.org/10.1177/1550132922110693>