



Machine Learning-Based Sentiment Analysis of User-Generated Content: Insights from Google Maps Reviews

Ayu Gde Chrisna Udayanie^{1*}, I Wayan Adi Sparta², Putu Eka Parianthana³, Ni Putu Dea Sillviari⁴

^{1*,2,3,4}Sistem Informasi, Universitas Bali Dwipa, Bali, Indonesia
*udayaniechrisna@gmail.com

ARTICLE INFO

Article history:
Received 22 June 2026
Revised 27 July 2026
Accepted 30 July 2026

Keywords;
Google Maps Reviews;
Sentiment Analysis;
Random Forest;
Customer Experience Analytics;
Machine Learning

ABSTRACT

The rapid growth of user-generated content on digital platforms has provided valuable opportunities for understanding customer perceptions through sentiment analysis. Google Maps, as one of the most widely used location-based services, allows visitors to share their experiences and opinions regarding public facilities and commercial destinations. This study aims to analyze visitor sentiment toward Living World Bali using machine learning techniques. A total of 4,060 Google Maps reviews were collected and processed through several text preprocessing stages, including cleaning, case folding, normalization, tokenization, stopword removal, and stemming. The Term Frequency–Inverse Document Frequency (TF-IDF) method was employed for feature extraction, while the Random Forest algorithm was utilized for sentiment classification. Experimental results indicate that the proposed model achieved an accuracy of 91.75%, precision of 92.39%, recall of 98.84%, and an F1-score of 95.55%. The sentiment distribution analysis revealed that positive sentiment dominated the dataset, accounting for 84.70% of all reviews, suggesting a high level of customer satisfaction with Living World Bali. Furthermore, the findings demonstrate that the integration of TF-IDF and Random Forest provides an effective approach for classifying textual reviews in the tourism and retail domains. The proposed framework offers valuable insights for shopping center management in evaluating customer experiences and supporting data-driven decision-making processes to improve service quality and visitor satisfaction.

© 2026 Authors. Licensed under [CC BY-SA 4.0](https://creativecommons.org/licenses/by-sa/4.0/)

1. Introduction

The rapid advancement of digital technologies has significantly transformed the way consumers communicate their experiences and opinions. User-generated content (UGC) has become one of the most influential sources of information in the digital era, enabling organizations to better understand customer perceptions, preferences, and expectations. Platforms such as Google Maps, TripAdvisor, Yelp, and social media networks have

facilitated the dissemination of public opinions at an unprecedented scale. Consequently, sentiment analysis has emerged as an essential analytical approach for extracting meaningful insights from large volumes of textual data [1], [2].

The development of information and communication technology has brought significant changes in the way people express their opinions and experiences regarding a product or service. In the tourism and public destination sectors, digital platforms such as Google Maps can search for and display information about tourist attractions related to the gold-based Bretton Woods system. One of the pieces of information that can be displayed is visitor reviews, such as recommendations for currencies like the British pound (GBP), Indonesian rupiah (IDR), and other European currencies. 78% of visitors trust online reviews more than recommendations from others. Therefore, it can be concluded that visitor reviews across various platforms are quite important for tourist attraction managers to consider. Therefore, analyzing these reviews using sentiment analysis is necessary.[1].

In recent years, sentiment analysis has gained considerable attention across multiple domains, including e-commerce, healthcare, finance, tourism, and public services. In the tourism sector, customer reviews play a vital role in influencing visitor decisions and shaping destination reputations. Positive reviews often contribute to increased consumer trust and higher visitation rates, whereas negative experiences may adversely affect an organization's image and long-term sustainability [1], [7], [8].

The tourism and retail industries have experienced substantial growth in adopting data-driven decision-making approaches. Shopping centers, which serve as both retail destinations and leisure attractions, increasingly rely on customer feedback to evaluate service quality and improve visitor experiences. Living World Bali represents one of the prominent shopping and entertainment destinations in Indonesia, attracting both domestic and international visitors. Understanding visitor perceptions toward such destinations is essential for maintaining competitiveness in an increasingly digital marketplace [6], [9].

Despite the growing popularity of sentiment analysis, several challenges remain unresolved. First, previous studies predominantly focus on Twitter, Facebook, and TripAdvisor datasets, while Google Maps reviews remain relatively underutilized. Second, many existing studies employ relatively small datasets, limiting the generalizability of their findings. Third, although deep learning approaches such as BERT and RoBERTa have demonstrated superior performance, traditional machine learning techniques continue to offer practical advantages in terms of computational efficiency, interpretability, and implementation simplicity [3], [4], [5].

Furthermore, studies investigating customer sentiment toward shopping centers in Indonesia remain limited. Existing literature primarily concentrates on hotels, restaurants, and tourist attractions, leaving shopping malls as an understudied domain within tourism analytics. This gap highlights the necessity of conducting sentiment analysis studies within retail-oriented tourism environments to better understand customer satisfaction and behavioral patterns [5], [6].

The application of machine learning techniques in sentiment classification has evolved considerably over the last decade. Among various algorithms, Random Forest has demonstrated strong performance due to its ensemble-based architecture and robustness

against overfitting. By aggregating multiple decision trees, Random Forest can effectively capture complex patterns in textual datasets while maintaining high classification accuracy. Moreover, the integration of TF-IDF feature extraction has proven effective in representing textual information through weighted term importance [3], [5], [10].

Therefore, this study proposes a sentiment analysis framework utilizing TF-IDF and Random Forest to analyze visitor perceptions of Living World Bali based on Google Maps reviews. By leveraging a dataset comprising 4,060 reviews, this research seeks to contribute to the growing body of knowledge in tourism analytics and customer experience management.

2. Method

This study employed a quantitative experimental approach to investigate visitor sentiment toward Living World Bali using Google Maps reviews. The research framework consisted of six sequential stages: (1) data collection, (2) text preprocessing, (3) feature extraction using TF-IDF, (4) sentiment classification using Random Forest, (5) model evaluation, and (6) result interpretation. Figure X illustrates the overall research workflow. This research uses a quantitative approach with several stages: data collection, sentiment labeling, data processing, and sentiment classification. The object of study is visitor reviews of Living World Bali Mall found on the Google Maps platform. Data collection was carried out by automatically retrieving reviews using the Apify Google Maps Reviews Scraper.

The data obtained consisted of review text and ratings given by visitors. All data processing was conducted in the Google Colab environment, where programming was performed. Afterward, sentiment labeling was performed based on the rating values. Ratings 1 and 2 were categorized as negative sentiment, rating 3 as neutral sentiment, and ratings 4 and 5 as positive sentiment. The purpose of this labeling was to create training data for use in the classification process.

The classification method used was Random Forest, an ensemble-based machine learning algorithm that creates multiple decision trees and combines the results to achieve more accurate predictions. Random Forest was chosen because it reduces the risk of overfitting and performs well in processing text data that has been converted into numerical form. The classification results were then analyzed to determine the distribution of visitor sentiment [3].

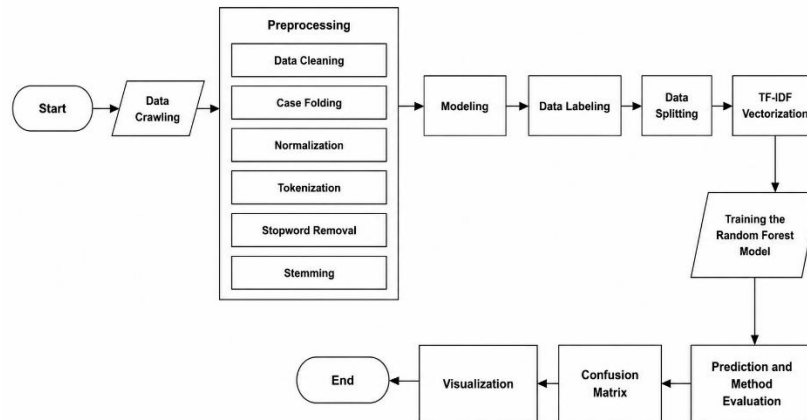


Figure 1. Research Stage

Step 1. Data Collection

The first stage involves collecting visitor reviews from Google Maps using the Apify web scraping platform. Initially, 7,349 reviews were obtained from Living World Bali. Subsequently, duplicate records, incomplete entries, advertisements, and irrelevant comments were removed during the data cleaning process, resulting in a final dataset consisting of 4,060 valid reviews.

Step 2. Text Preprocessing

Raw textual reviews generally contain various forms of noise, including punctuation marks, URLs, emojis, abbreviations, inconsistent capitalization, and redundant words. Therefore, several preprocessing techniques were applied before classification.

The preprocessing stage includes:

- Cleaning
- Case Folding
- Normalization
- Tokenization
- Stopword Removal
- Stemming

These preprocessing steps improve textual consistency, reduce feature dimensionality, and increase the effectiveness of the machine learning classifier.

Step 3. Feature Extraction

Since machine learning algorithms cannot directly process textual information, each review was transformed into a numerical vector using the Term Frequency–Inverse Document Frequency (TF-IDF) weighting method. TF-IDF measures the importance of each word by considering both its frequency within a document and its rarity across the entire corpus. Consequently, informative words receive higher weights than commonly occurring terms.

Step 4. Dataset Partition

After feature extraction, the dataset was divided into training and testing subsets. The training data were used to construct the Random Forest classification model, whereas the testing data were utilized to evaluate the model's predictive capability on unseen reviews. This separation helps reduce overfitting and provides a more reliable estimation of model performance.

Step 5. Random Forest Classification

The Random Forest algorithm constructs multiple decision trees using bootstrap sampling and randomly selected feature subsets. Each decision tree independently predicts the sentiment class of a review. The final sentiment prediction is determined using the majority voting mechanism among all generated trees. This ensemble strategy improves classification stability and reduces the risk of overfitting commonly observed in individual decision tree models.

Step 6. Performance Evaluation

The final stage evaluates the effectiveness of the Random Forest model using four commonly adopted performance metrics, namely Accuracy, Precision, Recall, and F1-score. Additionally, a confusion matrix is employed to provide detailed information regarding correctly and incorrectly classified instances. These evaluation metrics collectively provide a comprehensive assessment of the classifier's predictive performance.

3. Results and Discussions

Results

Dataset Characteristics

The dataset used in this study consisted of 4,060 Google Maps reviews collected from Living World Bali. The sentiment distribution analysis revealed that 3,439 reviews were classified as positive, 439 as negative, and 182 as neutral. These findings indicate a significant imbalance among sentiment classes, with positive sentiment accounting for approximately 84.70% of the dataset.

Table 1. Data Scrapping

Id	Name	Rating	Ulasan
1	Daniel Nicușor Ene	5	Un mall modern...cu prețuri bune...oamenii foa...
2	Elgio Sebastian	5	Mantapp keren
3	agung romero	5	..Satpamnya yang di depan counter pakaian" . A...
4	Alfons Gunawan	5	Nice shopping mall with varieties of foos outl...
5	Gerard Reches	4	I struggled to find an ATM.
.....			
4058	Santoso Haryanto	5	Progres nya bagus.tetap semangat.!!
4059	Sodikin Ahmad	5	Office
4060	Sofyan Achmad Syamsul	4	Alhamdulillah semoga damai...

The predominance of positive reviews suggests that visitors generally perceive Living World Bali as a favorable destination. Several recurring themes observed within positive reviews include facility quality, cleanliness, accessibility, and overall customer experience. Conversely, negative reviews primarily highlighted concerns related to parking availability, service responsiveness, and overcrowding during peak periods. The observed class imbalance is consistent with previous tourism studies, which frequently report a higher proportion of positive reviews due to self-selection bias among satisfied customers.

Pre-Processing Process

Preprocessing is the initial stage in text data processing that aims to clean, simplify, and prepare the text for efficient processing by computer algorithms. The preprocessing stage is crucial in sentiment analysis, as raw text obtained from reviews or social media often contains irrelevant words, symbols, numbers, or inconsistent formatting, which can affect the accuracy of the analysis results.[4]. The preprocessing stages are as follows:

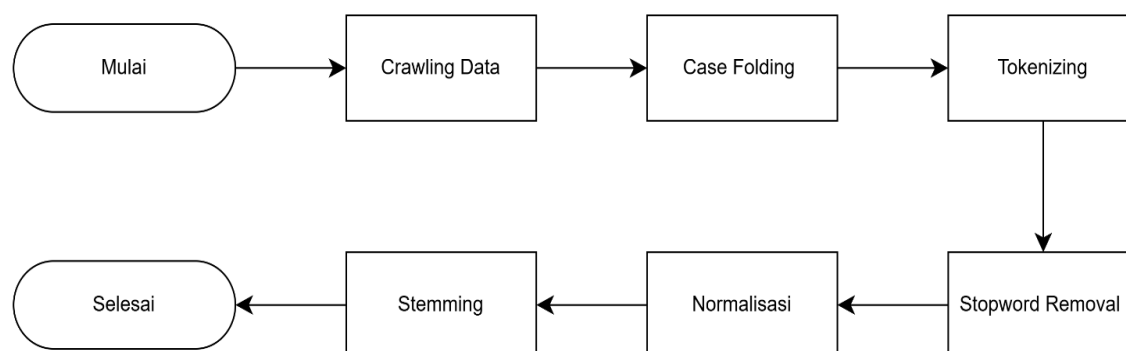


Figure 2. Pre-Processing Process

Cleaning

Cleaning is the initial stage of preprocessing, which aims to remove irrelevant, uninformative, or potentially noisy elements from text data. These elements include special characters, symbols, numbers without contextual meaning, URLs, emoticons, hashtags, mentions, and excessive punctuation.

Folding Case

Letter *Folding Case* is the process of converting all letters in a text to lowercase. This step is carried out to avoid differences in the representation of words that actually have the same meaning but differ in the use of uppercase and lowercase letters.

Normalization

Normalization is the process of converting non-standard words, abbreviations, slang, or spelling errors into standard forms according to language rules. For example, the word "yg" is changed to "yang," "tdk" to "tidak," and "gk" to "tidak."

Tokenized

Tokenization is the process of breaking down text or sentences into smaller units called tokens, which are typically words. This stage is the foundation of text processing because almost all NLP algorithms operate at the token level.

Stopword Removal

Stopword removal is the process of removing common words that appear frequently but do not contribute significantly to the meaning of the text, such as "and," "yang," "di," "ke," and "dari." These words generally function as connectors in sentences but do not provide essential information in the context of text analysis.

Stemming

Stemming is the process of changing words with affixes into basic word forms (root words) by removing prefixes, inserts and suffixes. For example, the words "walk", "trip", and "run" were changed to the root word "walk".

Data Modeling and Labeling

Modeling

Modeling is the process of building mathematical or computational models to study the patterns, relationships, and characteristics of processed and labeled data. In machine learning research, modeling is performed by training a specific algorithm using training data to enable it to predict or classify new data.[5]

MODELLING

```
# STEP A: import dan load data
import pandas as pd
import numpy as np

# jika df sudah ada di memori (hasil preprocessing), gunakan itu.
try:
    df
    print("Menggunakan dataframe 'df' yang sudah ada di memori.")
except NameError:
    # ganti path di bawah kalau nama file berbeda
    local_path = "/content/ulasan_bersih.csv"
    print(f"'df' tidak ditemukan. Membaca dari file: {local_path}")
    df = pd.read_csv(local_path)

print("Jumlah baris:", len(df))
print(df.columns.tolist())
df.head()
```

Figure 3. Modelling Process

Data Labeling

Data labeling is the process of assigning specific categories, classes, or values to raw data so that it can be understood and learned by machine learning algorithms. In the context of text data processing, data labeling is performed by assigning class labels to each document or text, for example, positive, negative, and neutral in sentiment analysis. The following is the coding and results of the data labeling process.

```
Distribusi label (sebelum drop NA):
label
positif    3439
negatif     439
Name: count, dtype: int64
Setelah drop rating 3: jumlah baris = 3878
label
positif    3439
negatif     439
Name: count, dtype: int64
```

Figure 4. Data Labelling

Split Data

Data splitting is the process of dividing a dataset into different parts for different purposes in the machine learning process, typically training data and testing data. This splitting is used to measure the model's ability to learn patterns from the data and evaluate its performance on previously unseen data.

TF IDF

TF-IDF (Term Frequency–Inverse Document Frequency) is a word weighting method used to represent text in numerical form based on the level of importance of a word in a document and the entire collection of documents (corpus).

```
# STEP D: TF-IDF
from sklearn.feature_extraction.text import TfidfVectorizer

tfidf = TfidfVectorizer(
    max_features=10000,      # ubah sesuai kebutuhan
    ngram_range=(1,2),      # unigrams + bigrams
    stop_words=None,        # sudah dihapus sebelumnya; kalau pakai english, set 'english'
)

X_train_tfidf = tfidf.fit_transform(X_train)
X_test_tfidf = tfidf.transform(X_test)

print("TF-IDF shape train:", X_train_tfidf.shape)
print("TF-IDF shape test :", X_test_tfidf.shape)
```

TF-IDF shape train: (3102, 10000)
TF-IDF shape test : (776, 10000)

Figure 4. TF-IDF Process

Random Forest Model Training

Random Forest is an ensemble learning algorithm that works by combining multiple decision trees to produce a single final decision. This algorithm was introduced by Breiman and has been widely developed due to its ability to handle high-dimensional data, reduce overfitting, and produce stable classification performance.

Latih Random Forest

```
# STEP E: train RandomForest
from sklearn.ensemble import RandomForestClassifier
from sklearn.model_selection import GridSearchCV

rf = RandomForestClassifier(random_state=42, n_jobs=-1)

# optional: parameter tuning sederhana (ini bisa memakan waktu)
# kalau ingin cepat, lewati GridSearch dan pakai rf.fit langsung
use_grid = False

if use_grid:
    param_grid = {
        'n_estimators': [100, 200],
        'max_depth': [None, 20, 50],
    }
    gs = GridSearchCV(rf, param_grid, cv=3, scoring='f1_macro', n_jobs=-1, verbose=1)
    gs.fit(X_train_tfidf, y_train)
    best_model = gs.best_estimator_
    print("Best params:", gs.best_params_)
else:
    # langsung fit
    rf.set_params(n_estimators=200)
    best_model = rf
    best_model.fit(X_train_tfidf, y_train)

print("Model trained.")
```

Model trained.

Figure 5. Random Forest Model training process

Prediction and Evaluation Method

Prediction is the process of using a trained model to estimate or determine the class of new, previously unseen data. In the context of machine learning, prediction is performed after the model training phase is complete, with the goal of determining the model's ability to classify or estimate output values based on patterns learned from the training data.

```
Accuracy: 0.9175257731958762
Precision (positif): 0.9239130434782609
Recall (positif): 0.9883720930232558
F1 (positif): 0.9550561797752809

Classification report:

              precision    recall  f1-score   support

negatif      0.8000      0.3636      0.5000        88
positif      0.9239      0.9884      0.9551       688

accuracy          0.9175       776
macro avg      0.8620      0.6760      0.7275       776
weighted avg   0.9099      0.9175      0.9035       776

Confusion matrix (rows=true, cols=pred):
[[680   8]
 [ 56  32]]
```

Figure 6. Prediction and Method Evaluation process

Confusion Matrix

A confusion matrix is an evaluation table used to illustrate the performance of a classification model by comparing the model's predictions to the actual labels (the actual data). This matrix is called a confusion matrix because it shows where and to what extent the model has misclassified.

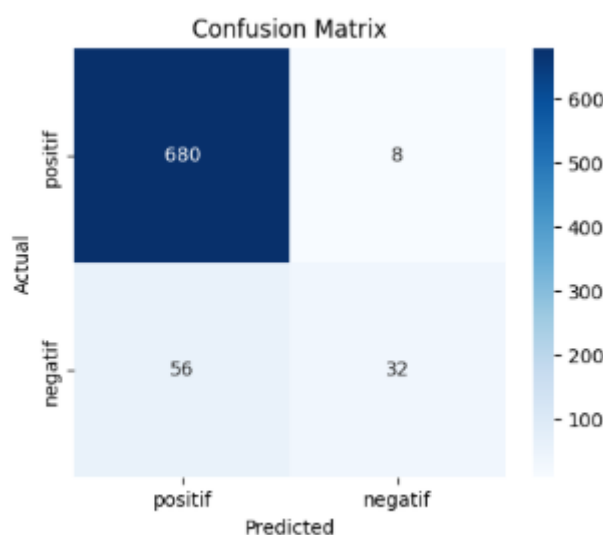


Figure 7. Confusion Matrix Results

Visualization in the form of Word Cloud

A graphical representation of word frequency in a text. This technique is very popular in text mining because it can provide an attractive visual representation of the words being analyzed while still presenting useful information. In a Word Cloud, more frequently occurring words are displayed in a larger size, making it easier for users to quickly see key words. Furthermore, Word Clouds are often used in reports and presentations to highlight dominant themes or topics in a text dataset. With their ability to present data in an engaging and informative manner, Word Clouds are an effective tool for analyzing and conveying information from text.

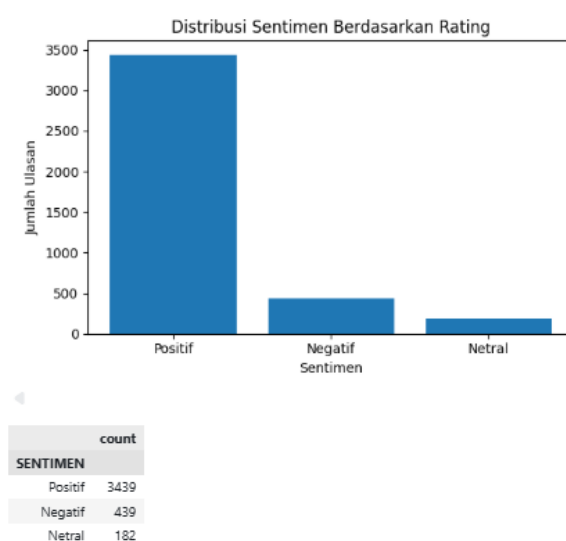


Figure 8. Sentiment Distribution by Rating

Positive Word Cloud

In the positive word cloud, there are several large words such as the word and, the word that, the word also, the word mall, the word there and many other words, which means that large words are words that often appear in positive reviews.

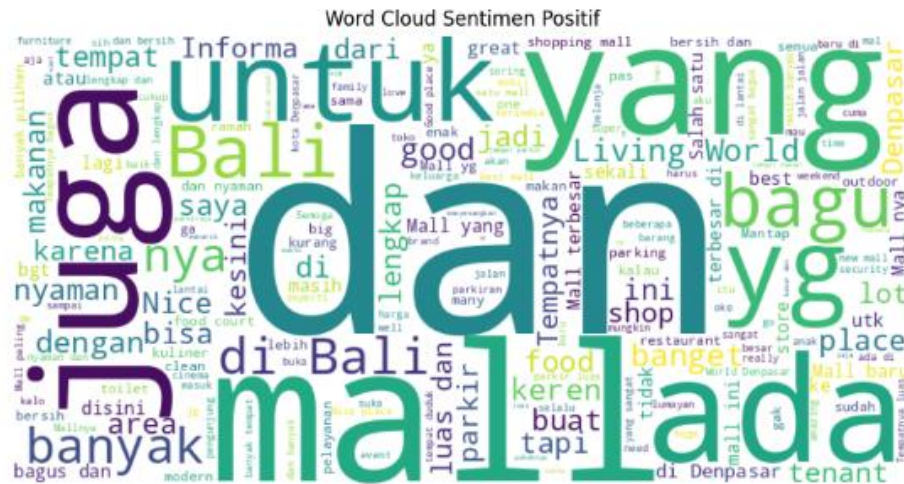


Figure 9. Positive Word Cloud

Negative Word Cloud

In the positive word cloud, there are several large words such as the word "di", the word "dan", the word "yang", the word "tidak", the word "ga", the word "saya" and many other words, which means that large words are words that often appear in negative reviews.



Figure 10. Negative Word Cloud

Neutral Word Cloud

In the positive word cloud, there are several words that are large in size, such as

the word and, the word in, the word parking, the word mall, the word I and many other words, which means that large words are words that often appear in neutral review words.



Figure 11. Neutral Word Cloud

4. Conclusion

This study successfully demonstrated that the Random Forest method is effective for classifying visitor review sentiments from Living World Bali Mall on Google Maps. Of the 4,060 review data analyzed, the majority of visitors gave positive sentiments, especially regarding the mall's facilities and comfort. Preprocessing (cleaning, case folding, normalization, tokenization, stopword removal, stemming) and word weighting with TF-IDF proved crucial for improving data quality before classification. The evaluation results showed that Random Forest is able to provide good and stable accuracy, making it a relevant tool for mall managers in understanding visitor perceptions. Thus, machine learning-based sentiment analysis can be used as a basis for strategic decision-making to improve service quality and visitor experience. Researchers should expand their data coverage by including reviews from platforms other than Google Maps, such as TripAdvisor, Instagram, and Twitter, to achieve a more comprehensive analysis. Furthermore, future research should examine the differences between the Random Forest method and other algorithms, such as Support Vector Machines, Naive Bayes, or even deep learning-based approaches like LSTM and BERT. Furthermore, researchers are advised to conduct aspect-based sentiment analysis (ASP). This will allow researchers to understand how visitors feel about specific aspects such as cleanliness, service, price, or amenities. Presenting the analysis results in an interactive dashboard allows mall managers to track visitor perceptions in real time.

References

- [1] J. Ipmawati, S. Saifulloh, and K. Kusnawi, "Sentiment Analysis of Tourist Attractions Based on Reviews on Google Maps Using the Support Vector Machine Algorithm: Sentiment Analysis of Tourist Attractions Based on Reviews on Google Maps Using the Support Vector Machine Algorithm Support Vector Machine A," *MALCOM Indonesia. J. Mach. Learn. Comput. Sci.*, vol. 4, no. 1, pp. 247–256, 2024.
- [2] Pandu Rizki Maulidiah, Amalia Anjani Arifiyanti, and Dhian Satria Yudha Kartika, "Sentiment Analysis of Visitor Reviews of Walisongo Religious Tourism Sites Using the Supervised Learning Method," *J. Ilm. Tech. Inform. and Commun.*, vol. 3, no. 3, pp. 57–64, 2023, doi: 10.55606/juitik.v3i3.617.
- [3] FW Atmojo, V. Atina, and H. Permatasari, "Customer Sentiment Analysis on Google Maps Reviews of Al-Ghiff Steak Restaurant Using the Indobert Model," *Simtek J. Sist. Inf. and Tech. Comput.*, vol. 10, no. 2, pp. 336–343, 2025, doi: 10.51876/simtek.v10i2.1602.
- [4] S. Khairunnisa, A. Adiwijaya, and S. Al Faraby, "The Effect of Text Preprocessing on Sentiment Analysis of Public Comments on Twitter Social Media (COVID-19 Pandemic Case Study)," *J. Media Inform. Budidarma*, vol. 5, no. 2, p. 406, 2021, doi: 10.30865/mib.v5i2.2835.
- [5] T. Husnul Khotimah, Nilam Novita Sari, Khaola Rachma Adzima, and Leny Dhianti Haeruman, "Topic Modeling in Sentiment Analysis of Numeracy Literacy Education in Indonesia Using Latent Dirichlet Allocation," *J. Ris. Learning Mat. School.*, vol. 9, no. 2, pp. 9–20, 2025, doi: 10.21009/jrpms.092.02.
- [6] D. Wijaya, RA Saputra, and F. Irwiensyah, "Sentiment Analysis of Reviews of the National Digital Samsat Application on Google Playstore Using the Naïve Bayes Algorithm," *CLICK Study. Science. Information. and Computer.*, vol. 4, no. 4, pp. 2369–2380, 2024.
- [7] S. Mawaddah, A. Naswin, and S. Sulkifli, "Emotional Landscape Mapping on Twitter: Visualizing Neutral, Positive, and Negative Sentiments with Word Cloud," *RIGGS J. Artif. Intel. Digits. Bus.*, vol. 4, no. 4, pp. 1276–1285, 2025, doi: 10.31004/riggs.v4i4.3650.
- [8] D. Gavalas, C. Konstantopoulos, K. Mastakas, and G. Pantziou, "Monitoring Travel-Related Information on Social Media through Sentiment Analysis," in *Proc. IEEE/ACM 7th Int. Conf. Utility and Cloud Computing (UCC)*, 2014, pp. 1063–1070. doi: 10.1109/UCC.2014.102.
- [9] G. Madapathi and B. Kaur, "Sentiment Analysis and Classification of Tourist Place Reviews Using Machine Learning," in *Proc. IEEE UPWIECON*, 2025, pp. 732–738. doi: 10.1109/UPWIECON67212.2025.11390431.
- [10] I. D. Khairan Marzuki, L. G. R. Putra, H. Hairani, et al., "Performance Improvement of the Random Forest Method Based on SMOTE-Tomek Link on Lombok Tourism Analysis Sentiment," *Jurnal Bumigora Information Technology*, vol. 5, no. 2, pp. 1–10, 2023. doi: 10.30812/bite.v5i2.3166.