

Sentiment Analysis of Indonesian Tourism Social Media Using Naïve Bayes for Decision Support

Nathaniel Felix Fraderic¹, Aditiya Hermawan², Junaedi³

^{1,2}Departement of Informatics Engineering, Universitas Buddhi Dharma, Tangerang, Indonesia

^{1,2}Departement of Information System, Universitas Buddhi Dharma, Tangerang, Indonesia

¹nathaniel.felix@ubd.ac.id

²aditiya.hermawan@ubd.ac.id

³junaedi@ubd.ac.id

Accepted on April 07th, 2025
Approved on January 07th, 2026

Abstract— The tourism sector is still the leading alternative sector to boost the Indonesian economy. However, the development of Indonesia's tourism sector still faces challenges, particularly in communication and publication, which remain inadequate. The purpose of this research is to improve the quality of tourist attractions by analyzing responses from visitors to assess the community's feedback. This research aims to employ a data analysis technique that automatically processes these responses, enabling effective sentiment classification and providing actionable insights for enhancing tourism experiences. The method of this study is sentiment analysis using text mining, specifically utilizing the Naïve Bayes algorithm to classify sentiment efficiently. The data for this process is collected using the API provided by the X platform. X was selected as the research object because it facilitates rapid data retrieval through keywords or hashtags, offering a rich source of public opinion specifically on Indonesian tourist destinations. Clearly, the results of this study demonstrate the successful implementation of a website created using the Python programming language. This website visualizes the analysis of people's responses to tourist attractions in graphical form, presenting the percentage of each sentiment class. The text mining process achieved an accuracy of 87%. The conclusion of this study highlights its novelty in applying the Naïve Bayes algorithm for sentiment analysis of Indonesian-language tweets related to tourism, a context that has not been widely explored before. Future research can enhance the sentiment analysis process by incorporating a comprehensive list of positive and negative words in Indonesian, further improving classification precision.

Index Terms— Analyzing sentiment; Naïve Bayes; Text mining; Tourism.

I. INTRODUCTION

The tourism sector serves as a crucial alternative driver for Indonesia's economic growth. In 2022, Indonesia's tourism sector contributed 3.6% to the national GDP, reflecting a decline from 4.2% in the previous year [1]. This upward trend demonstrates the sector's gradual recovery and growing significance to

the national economy after a period of challenges. However, the development of Indonesia's tourism sector still faces various challenges that must be addressed to support national economic growth [2]. Key among these challenges are deficiencies in communication and publication, which are critical for promoting tourist destinations and attracting visitors. Effective communication and publication are essential to ensure positive feedback and sentiments from visiting tourists, which play a crucial role in shaping perceptions of tourist destinations.

Understanding consumer perceptions is particularly vital in the context of tourism destination marketing, as it influences decisions to visit these destinations [3]. The growing presence of digital tourists highlights the importance of media in shaping interests and preferences, contributing to the popularity of certain themes and destinations [4]. This has led to the rise of social media as a powerful platform for gauging public opinion. Analyzing traveler sentiments, therefore, is essential for enhancing the visibility and appeal of tourist destinations. Leveraging the rapid development of information technology, social media platforms like X have emerged as valuable data sources. With Indonesia ranking fifth globally in terms of X users (18.45 million users as of 2022) [5], X offers rich data for sentiment analysis, facilitated by its public API, which enables efficient text data retrieval. This makes X an ideal medium to explore real-time public opinions on tourist attractions, contributing to a dynamic approach to tourism promotion.

Sentiment analysis, or opinion mining, involves analyzing opinions, evaluations, attitudes, judgments, and feelings about products, services, organizations, events, or topics. Although the field is well-established, its application in tourism, especially in Indonesia, has been insufficiently explored. In recent years, there has been an increase in sentiment analysis studies applied to tourism; however, these have primarily focused on broader or international contexts without addressing the unique linguistic and cultural

nuances of Indonesian tourism discourse [6]. For example, Mehraliyev et al, reviewed sentiment analysis methodologies in tourism, emphasizing the need for localized studies in emerging markets [7]. Similarly, Puh and Bagić Babac (2023) employed machine learning for sentiment analysis in tourism but focused on European destinations, leaving a gap in Indonesian-specific research [8].

This study aims to address these gaps by applying the Naïve Bayes algorithm to analyze sentiments in Indonesian-language tweets about tourist destinations. Previous research using the Naïve Bayes method has demonstrated promising results, such as an accuracy of 81.77% in sentiment analysis [9]. However, despite these advancements, there is still limited research specifically targeting Indonesian-language tweets in the tourism sector. Al-Bakri et al. (2022) compared Naïve Bayes with other algorithms, such as KNN, highlighting the superior performance of Naïve Bayes in specific contexts, though not in Indonesian tourism [10]. Similarly, Singgalen (2024) demonstrated the effectiveness of the Naïve Bayes Classifier in classifying sentiments related to over-tourism issues, highlighting its ability to handle nuanced public perceptions and facilitate sustainable tourism practices by addressing imbalances in data distribution through techniques like SMOTE [11]. Kalaivani et al. (2023) further highlighted the reliability of machine learning techniques like Naïve Bayes in hotel review sentiment analysis [12]. Additionally, Bi et al. (2024) provided an extensive review of text analysis methods in tourism, emphasizing the use of text mining techniques like sentiment analysis and topic modeling. Their study highlighted the utility of these approaches for understanding complex tourist behaviors, predicting industry trends, and managing dynamic challenges in tourism and hospitality [13].

The study differs from prior works in several key aspects. First, it focuses exclusively on Indonesian-language tweets, addressing the linguistic and cultural context that has been largely overlooked in previous studies. Second, it integrates recent advancements in machine learning, specifically in the area of data preprocessing. Techniques like case folding, tokenization, stopword removal, stemming, and TF-IDF weighting were employed to improve data quality and feature extraction. These preprocessing steps ensure that the text data is cleaned and structured effectively, enhancing the quality of the sentiment classification. Third, this study emphasizes practical applications, such as visualizing sentiment analysis results to aid decision-making for potential tourists and industry stakeholders. By leveraging social media data from X, this research not only contributes to the growing body of sentiment analysis literature but also offers practical tools for enhancing tourism marketing strategies and improving tourism experiences in Indonesia.

II. METHODOLOGIES

The research object used in this study is the response from the public obtained from X. This method follows a similar approach described by Hasanati [14], where sentiment analysis was performed using data from X to assess public opinions, and a series of preprocessing steps, such as text normalization and classification, were employed to structure the data for analysis. The research method is depicted in the following diagram in Fig. 1 below.

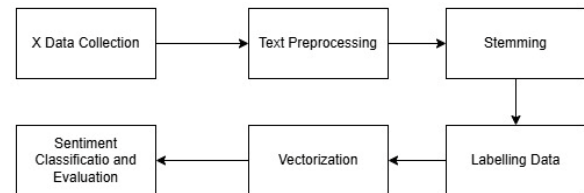


Fig. 1. Research methods diagram

A. X Data Collection

The data for this study was retrieved using the X platform's public API, requiring an API Key and an Access Token to access the data. Data collection was conducted using a predefined set of keywords related to popular Indonesian tourist destinations, such as "Candi Prambanan," "Borobudur," and "Malioboro," which were selected based on their relevance to tourism discussions on X. The data was gathered over a period from March 2023 to June 2023, during which a total of 2,717 tweets were collected. The following is the raw data crawling results shown in Table I.

TABLE I. RAW DATA

Timestamp	Raw_Text	Label	Hashtag
2023-03-17 19:33:05	berapa harga tiket candi prambanan	NaN	Candi Prambanan
2023-03-17 19:33:05	RT @media_twc: PT TWC melakukan penutupan dest...	NaN	Candi Prambanan
2023-03-17 19:33:05	RT @media_twc: PT TWC melakukan penutupan dest...	NaN	Candi Prambanan
2023-03-17 19:33:05	RT @cineyar: Terima kasih telah meliburkan are...	NaN	Candi Prambanan
2023-03-17 19:33:05	RT @Fredy_siTom: @VeelaaRhie Candi Borobudur d...	NaN	Candi Prambanan

Table I presents the initial results of the data collection process, consisting of raw textual data retrieved from the X social media platform through its public API using the keyword "Candi Prambanan." The raw text attribute contains the original, unprocessed content of each social media post as published by users, including informal language, abbreviations, and contextual expressions that reflect authentic public discourse. As such, this attribute may include noise such as reposts, user mentions, hashtags, URLs, and other non-informative elements, and therefore has not yet undergone text cleaning or sentiment classification.

The hashtag attribute indicates the tourism-related keyword or destination associated with each post and is used to link individual texts to specific tourist attractions during the analysis. The Label column is intentionally left empty at this stage, as sentiment categories (positive, neutral, or negative) are assigned only after the preprocessing and lexicon-based labeling procedures are completed. Additionally, the timestamp attribute records when each post was collected, enabling temporal traceability of public responses. Overall, this table illustrates the role of raw social media data as the foundational input for subsequent sentiment analysis using the Naïve Bayes algorithm to computationally assess public perceptions of Indonesian tourist destinations.

B. Text Preprocessing Process

After the data crawling process from X, text processing is conducted as a crucial step in sentiment analysis. This stage aims to clean and transform unstructured data into a structured format that can be efficiently processed by the system, enabling algorithms to calculate and classify sentiments with greater accuracy [15]. The preprocessing process involves several stages to prepare the text for analysis. Firstly, text cleaning is performed to eliminate non-textual characters like numbers, punctuation, and unnecessary spaces. Subsequently, case folding is applied to ensure uniformity by converting all text to lowercase, preventing distinctions based on capitalization. Tokenization follows, breaking down sentences into individual words or tokens, facilitating further processing. This step enables thorough cleaning of each component within the text, such as tweets. Finally, stopword removal occurs, where common connecting words like "in," "to," "which," and "and" are filtered out using a predefined stoplist, refining the text for subsequent analysis. Each stage contributes to enhancing the quality and relevance of the text data for subsequent tasks.

The results of this process will be displayed in Table II below:

TABLE II. PREPROCESSING RESULTS

Timestamp	Text_clean
2023-03-17 19:33:05	berapa harga tiket candi prambanan
2023-03-17 19:33:05	PT TWC melakukan penutupan destinasi Taman Wi...
2023-03-17 19:33:05	PT TWC melakukan penutupan destinasi Taman Wi...
2023-03-17 19:33:05	Terima kasih telah meliburkan area Candi Pram...
2023-03-17 19:33:05	Candi Borobudur dibangun oleh Raja dari Dinas...

C. Stemming

After the text preprocessing in Table II is performed, the stemming process continues, converting

the words into their root words. The stemming results are displayed in Table III below:

TABLE III. STEMMING RESULTS

Text_clean	Text_stemming
berapa harga tiket candi prambanan	berapa harga tiket candi prambanan
PT TWC melakukan penutupan destinasi Taman Wi...	pt twc laku tutup destinasi taman wisata candi...
PT TWC melakukan penutupan destinasi Taman Wi...	pt twc laku tutup destinasi taman wisata candi...
Terima kasih telah meliburkan area Candi Pram...	terima kasih telah libur area candi prambanan...
Candi Borobudur dibangun oleh Raja dari Dinas...	candi borobudur bangun oleh raja dari dinasti...

D. Labelling Data

During the data labeling process, words are assigned scores based on a predefined list of positive and negative words. These scores determine the polarity of the text, thereby classifying each sentence accordingly: sentences with polarity scores > 0 are classified as positive, those with scores $= 0$ are considered neutral, and those with scores < 0 are labeled as negative [16].

E. Vectorization

After the words resulting from Preprocessing and Stemming are collected, then the words will be weighted using the TF-IDF (Term Frequency-Inverse Document Frequency) algorithm. The purpose of this process is to assign a weight value to the words in the document so that the classification process can be carried out. The following is an example of the top value of the TF-IDF weighting after weighting as shown in Table IV below:

TABLE IV. TOP TERM TF-IDF

Term	TF-IDF
Borobudur	1.000000
Malioboro	1.000000
Makutmu	0.968713
Enter	0.968713
Puji	0.963584

F. Sentiment Classification and Evaluation

The data collection phase involved utilizing the X API to gather information, focusing on five specific keywords: Malioboro Street, Kraton Yogyakarta, Borobudur Temple, Prambanan Temple, and Tugu Yogyakarta. This endeavor yielded a dataset comprising 2717 entries. Following data collection, the text preprocessing stage ensued. Here, collected text underwent several procedures including text cleaning,

case folding, noise removal, tokenization, filtering (stopword removal), and stemming. Next, in the data labeling step, sentiment analysis was conducted based on lexicon-based methodologies. This involved associating sentiments with their semantic values and assessing the intensity of words within sentences. Utilizing a predetermined lexicon consisting of positive and negative words, the data labeling process was automated. Subsequently, each sentence was assigned a score based on the presence of positive and negative words, ultimately categorizing them into positive, negative, or neutral classes.

Moving forward, the model classification phase employed the Multinomial Naïve Bayes algorithm for sentiment analysis [17]. The general equation 1 for Naïve Bayes classification is given by:

$$P(c|d) = P(c) * \prod_{(w \in d)} P(w|c) \quad (1)$$

Here, $P(c|d)$ represents the posterior probability of class c given document d . To handle the issue of zero probabilities, Laplace smoothing is applied, which adds a +1 to the numerator and $|V|$ (the size of the vocabulary) to the denominator:

$$P(w_i|c) = (\text{count}(w_i, c) + 1) / (\text{count}(c) + |V|) \quad (2)$$

In this formula, $P(w_i|c)$ is the probability of the word w_i occurring in class c . $\text{count}(w_i, c)$ is the number of times the word w_i occurs in class c , while $\text{count}(c)$ is the total number of occurrences of all words in class c , and $|V|$ is the total number of unique terms (or features) in the vocabulary.

The Naïve Bayes algorithm uses these formulas to classify sentiments in tourism-related content by constructing a predictive model, assigning probability scores to individual words, and categorizing them into distinct sentiment classes: positive, negative, or neutral. In this study, the Multinomial Naïve Bayes variant was employed with Laplace smoothing to address zero-probability issues and enhance model robustness when handling sparse textual features. The classification process was preceded by a parameterized feature extraction stage using TF-IDF weighting with unigram representations, allowing the model to emphasize discriminative terms relevant to tourism discourse.

Utilizing a supervised machine learning approach, this study categorized sentiments within tourism-related content by constructing a predictive model. This model assigned probability scores to individual words, enabling the accurate classification of text into distinct sentiment categories, thereby offering insights into feedback and perceptions [18]. The training labels were generated through an Indonesian sentiment lexicon-based labeling strategy, which serves as an improvement over manual labeling by ensuring consistency and scalability across large datasets. The dataset was divided into training (75%) and testing (25%) data subsets for model training and evaluation.

Finally, evaluation of the sentiment analysis model was conducted. This involved assessing the performance metrics including accuracy, precision, recall, and F1-Score through the computation of the confusion matrix. Additionally, the evaluation results underwent validation using the 10 K-Fold Cross Validation technique. This validation scheme was selected to ensure performance stability across different data partitions and to demonstrate the generalizability of the proposed parameterized sentiment analysis framework. The reason for using 10-fold cross-validation is to provide a reliable estimate of model performance by splitting the data into 10 subsets, ensuring that each subset is used for both training and testing. This approach helps reduce the risk of overfitting and ensures the model's generalizability. The choice of 10 folds is based on standard practice in machine learning, where it offers a good balance between computation time and performance reliability [19].

III. RESULT AND DISCUSSION

System implementation and analysis of the results is carried out using the Python programming language. The implementation results in each stage are as follows:

A. Classification Process

The Multinomial Naïve Bayes algorithm was utilized to classify sentiments in tourism-related content. By leveraging Scikit-learn, an open-source library, the model was trained to calculate word probability scores, enabling efficient categorization of documents into their respective sentiment classes [20]. Scikit-learn offers a wide range of algorithms including KNN, SVM, Random Forest, and Naïve Bayes. Beyond classification, it supports preprocessing, regression, clustering, and model selection. The training data comprises 75%, while 25% is reserved for testing purposes. Multinomial naive Bayes is a method that works by calculating the frequency of each term in the document. In Multinomial Naive Bayes, the order in which words appear in the document is ignored so that the document is treated as a "bag of words", so that each word is processed using a multinomial distribution.

Classification results using the Multinomial Naïve Bayes algorithm are shown by the following Table V confusion matrix:

TABLE V. NAÏVE BAYES CONFUSION MATRIX RESULTS

	Precision	Recall	F1-score	Support
Negative	0.94	0.58	0.71	417
Netral	0.85	0.97	0.91	1526
Positive	0.90	0.84	0.87	774
Accuracy	-	-	0.87	2717
Macro average	0.90	0.79	0.83	2717

Weighted average	0.88	0.87	0.87	2717
------------------	------	------	------	------

B. Evaluation

The first evaluation, performed with matrix confusion to measure the results of the experiment. To find true positive (TP), true negative (TN), false positive (FP), and false negative values (FN). Below is a matrix confusion by comparing columns of labelling data with lexicon-based and classification results with Naïve bayes algorithms.

From the results of Table VI above there are three classes, namely, Negative, Neutral, and Positive. In this case, True positive = 647, True negative = 240, False positive = 127, False negative = 177, True neutral = 1481 False neutral = 45 and Total = 2717. To find an accuracy value, you can use the following Formula (3), Formula (4) to find precision, Formula (5) to find recall and last formula (6) to get f1 score :

$$Accuracy = \frac{TP+TN}{TP+TN+FP+FN} * 100\% \quad (3)$$

$$Accuracy = \frac{647+240+1481}{647+240+1481+127+177+45} * 100\% \\ = 87,15\%$$

$$Precision = \frac{TP}{TP+FP} * 100\% \quad (4)$$

$$Precision = \frac{647}{647+127} * 100\% \\ = 83,59\%$$

$$Recall = \frac{TP}{TP+FN} * 100\% \quad (5)$$

$$Recall = \frac{647}{647+177} * 100\% \\ = 78,51\%$$

$$F1 - Score = \frac{2*precision*recall}{precision+recall} \quad (6)$$

$$F1 - Score = \frac{2*0,8359*0,7851}{0,8359+0,7851} \\ = 80,97\%$$

TABLE VI. CONFUSION MATRIX NAÏVE BAYES EVALUATION RESULTS

	Prediction
--	------------

		Negative	Netral	Positive	Total
Actual	Negative	240	146	31	417
	Netral	4	1481	41	1526
	Positive	12	115	647	774
	Total	269	1742	719	2717

This evaluation was carried out using the 10 K-fold cross validation method to obtain average values of accuracy, recision, recall and F1-Score. With the 10 K-fold cross validation method, the data is divided into 10 parts. With the 10-K-fold Cross Validation method, it is calculated by dividing the data into K subsets, In this case K = 10, the dataset will be divided equally into random positions, where each subset is alternately used as test data while the other subset is used as training data. This process is carried out repeatedly K times or 10 times [21], so that each subset is used as test data once.

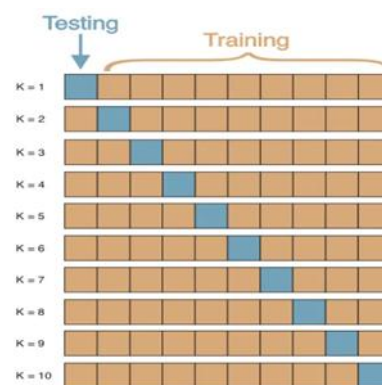


Fig. 2. 10-k-fold cross validation process [21]

So by using the 10 K-fold cross validation evaluation method like in **Error! Reference source not found.**, it can detect the occurrence of overfitting or underfitting. Overfitting occurs when a model can have excellent performance on a specific subset in several iterations, instead Underfitting happens if a model has poor performance in a particular subset.

Then a random test is performed on testing data and training data that yields the following Table VII:

TABLE VII. EVALUATION 10 K-FOLD VALIDATION RESULTS

Fold	Accuracy	F1-Score	Recall	Precision
1	0.875000	0.835944	0.80248	0.890807
2	0.845588	0.812332	0.782112	0.875484
3	0.900735	0.87692	0.85291	0.919553
4	0.841912	0.787672	0.75646	0.873181
5	0.889706	0.85527	0.824056	0.911976
6	0.830882	0.769654	0.730701	0.86676
7	0.900735	0.853745	0.816694	0.922976
8	0.878229	0.839525	0.809758	0.902984
9	0.878229	0.825891	0.793823	0.896928

10	0.874539	0.833785	0.79608	0.89552
Average	0.871556	0.829074	0.796508	0.895617

After testing 10 times, the accuracy, F1-score, precision and recall values will be calculated as the average of each iteration carried out using 10 K-Fold Cross Validation. The following are the results of K-fold Cross Validation, as shown in Table VIII below:

TABLE VIII. K-FOLD CROSS VALIDATION RESULTS

Evaluation Method	Accuracy	F1-Score	Recall	Precision
Manual	0.871549	0.809703	0.785194	0.835917
10-fold	0.871556	0.829074	0.796508	0.895617

From Table VIII, it is obtained and then evaluated using K-Fold Cross Validation, with a total of 10 folds. The validation results show an average accuracy of 0.871556, which is considered to be a reliable and good performance based on K-Fold Cross Validation techniques, as reported by several studies. For instance, Majbur et al. [22], found that Naïve Bayes achieved an average accuracy of 82.33% using similar methods. Based on these findings, several suggestions are proposed to improve the performance of the Naive Bayes classification model. First, to improve the accuracy, F1-score, precision, and recall produced by the model, it is recommended to enhance the application of this method by involving the addition of a list of positive and negative words. This would enrich the labeling process of the training data used, allowing the model to become more sensitive in recognizing positive and negative nuances within the text. Second, regarding the size of the training data, it is advisable to ensure a balance between the amount of data with positive and negative labels. By utilizing a large amount of training data that maintains a balanced label distribution, the model would become more effective in detecting sentences that reflect either positive or negative sentiment. Thus, this approach would enable an improvement in accuracy when classifying these texts more effectively.

To better understand the implications of the model's performance and the proposed improvements, Figure 3 presents a visualization of sentiment classification results for the hashtag Candi Prambanan. This visualization offers a clear representation of the sentiment distribution and complements the previous evaluation results.

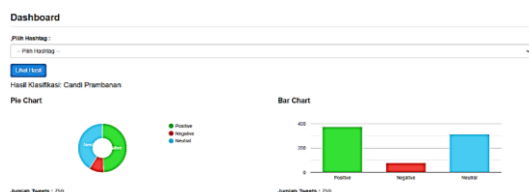


Fig. 3 Visualization of sentiment classification results

This study presents significant contributions to the field of sentiment analysis in tourism, particularly through its contextual and methodological advancements. The primary novelty lies in the development of an Indonesian-language corpus derived from the X platform, collected systematically using targeted tourism-related keywords such as Borobudur, Malioboro, and Candi Prambanan. This corpus captures authentic public discourse about Indonesian tourist destinations, thereby addressing a notable gap identified in previous international studies [7], [8] that predominantly focused on Western tourism contexts and English-language data. By localizing sentiment analysis to the Indonesian linguistic and cultural framework, this research enhances the applicability of machine learning models in understanding domestic tourism sentiment and provides a foundation for regional digital analytics. Furthermore, the study operationalizes a complete text-mining pipeline spanning data cleaning, lexicon-based sentiment labeling, TF-IDF vectorization, and classification using the Multinomial Naïve Bayes algorithm to achieve high computational efficiency and reproducibility. Rather than solely emphasizing model performance, this study positions the sentiment classification results as a means to interpret public perceptions and communication effectiveness within Indonesian tourism. The model's performance, validated through 10-fold cross-validation with an average accuracy of 87.15%, surpasses comparable studies utilizing SVM [14] or Doc2Vec approaches [16], thus reinforcing the robustness of Naïve Bayes for sentiment classification in morphologically rich languages like Indonesian.

In terms of substantive findings, the sentiment distribution extracted from 2,717 social media posts reveals that public perception toward Indonesian tourist destinations is predominantly positive and neutral, with approximately 84% of discussions reflecting favorable impressions or neutral informational exchanges, and only about 15% expressing dissatisfaction or critical opinions. The positive sentiments commonly highlight cultural heritage, hospitality, and natural beauty, indicating that Indonesia's core tourism assets are generally well perceived by the public. Neutral sentiments are largely informational, reflecting inquiries about ticket prices, access, and operational details, which underscores the role of social media as an important communication channel for tourism-related information. Meanwhile, negative sentiments are mostly associated with infrastructure limitations, crowding, and inconsistent pricing, directly mirroring the communication and service management challenges identified in the

introduction. These findings explicitly answer the research problem stated in the title and introduction by demonstrating how public sentiment reflects both the strengths and persistent challenges of Indonesian tourism. Accordingly, sentiment analysis is shown to function not merely as a classification task, but as an empirical diagnostic tool for identifying priority areas in tourism communication, service improvement, and destination management. Moreover, the integration of a Python-based web visualization system transforms these findings into actionable intelligence, enabling tourism authorities and stakeholders to monitor public sentiment trends and adapt promotional and managerial strategies accordingly. Thus, beyond methodological innovation, this research provides substantive, data-driven evidence of how sentiment analysis can contribute to enhancing Indonesia's tourism competitiveness through a deeper understanding of visitor perceptions and service gaps.

IV. CONCLUSION

This research successfully developed a system for sentiment analysis of tweets, achieving an accuracy rate of 87.15% with the Naïve Bayes algorithm. The use of TF-IDF for text vectorization enhanced the model's ability to identify sentiment as positive, negative, or neutral, offering valuable insights for decision-making in tourism. The system demonstrated reliability through rigorous 10 K-Fold cross-validation, yielding an F1-score of 82%, recall of 79%, and precision of 89%. These results highlight the system's robustness in handling varied tweet content and ensuring consistent sentiment classification.

The system's ability to process diverse social media data efficiently makes it a promising tool for sentiment-based decision support, particularly in the tourism industry. While the integration of Naïve Bayes is effective, future research should explore other machine learning techniques such as Random Forest or Deep Learning for comparison. Additionally, expanding to multilingual sentiment analysis and incorporating advanced features like emotion detection could further enhance the system's performance and applicability. This work lays a solid foundation for developing more comprehensive sentiment analysis systems in various domains.

REFERENCES

- [1] I. M. Hasibuan, S. Mutthaqin, R. Erianto, and I. Harahap, "Kontribusi Sektor Pariwisata Terhadap Perekonomian Nasional," *urnal Masharif al-Syariah J. Ekon. dan Perbank. Syariah*, vol. 8, no. 2, pp. 1200–1217, 2023.
- [2] R. Kurniawan, "Strategi Pemasaran Pariwisata Untuk Meningkatkan Pariwisata Lokal," vol. 7, pp. 9867–9877, 2024, doi: <https://doi.org/10.31004/jrpp.v7i3.31431>.
- [3] D. Setiadi and Y. I. Mukti, "Electronic Tourism Using Decision Support Systems to Optimize the Trips," *MATRIK J. Manajemen, Tek. Inform. dan Rekayasa Komput.*, vol. 23, no. 1, pp. 183–200, 2023, doi: [10.30812/matrik.v23i1.3331](https://doi.org/10.30812/matrik.v23i1.3331).
- [4] Y. A. Singgalen, "Analisis Sentimen dan Pemodelan Topik dalam Optimalisasi Pemasaran Destinasi Pariwisata Prioritas di Indonesia," *J. Inf. Syst. Informatics*, vol. 3, no. 3, pp. 459–470, 2021, doi: [10.51519/journalisi.v3i3.171](https://doi.org/10.51519/journalisi.v3i3.171).
- [5] R. T. B. Ardianto and Y. Nataliani, "Analisis Jaringan Sosial Pengguna Twitter Dalam Skema Penyediaan Laptop Bagi Toko Komputer," vol. 02, no. 3, pp. 208–218, 2023, doi: <https://doi.org/10.24246/itexplore.v2i3.2023.pp208-218>.
- [6] Y. Mao, Q. Liu, and Y. Zhang, "Sentiment analysis methods, applications, and challenges: A systematic literature review," *J. King Saud Univ. - Comput. Inf. Sci.*, vol. 36, no. 4, p. 102048, 2024, doi: <https://doi.org/10.1016/j.jksuci.2024.102048>.
- [7] F. Mehraliyev, I. C. C. Chan, and A. Kirilenko, "Sentiment analysis in hospitality and tourism: a thematic and methodological review," *Int. J. Contemp. Hosp. Manag.*, vol. ahead-of-p, 2021, doi: [10.1108/IJCHM-02-2021-0132](https://doi.org/10.1108/IJCHM-02-2021-0132).
- [8] K. Puh and M. Bagić Babac, "Predicting sentiment and rating of tourist reviews using machine learning," *J. Hosp. Tour. Insights*, vol. 6, no. 3, pp. 1188–1204, 2023, doi: [10.1108/JHTI-02-2022-0078](https://doi.org/10.1108/JHTI-02-2022-0078).
- [9] C. Villavicencio, J. J. Macrohon, X. A. Inbaraj, J. H. Jeng, and J. G. Hsieh, "Twitter sentiment analysis towards covid-19 vaccines in the Philippines using naïve bayes," *Inf.*, vol. 12, no. 5, 2021, doi: [10.3390/info12050204](https://doi.org/10.3390/info12050204).
- [10] N. F. Al-Bakri, J. F. Yonan, A. T. Sadiq, and A. S. Abid, "Tourism companies assessment via social media using sentiment analysis," *Baghdad Sci. J.*, vol. 19, no. 2, pp. 422–429, 2022, doi: [10.21123/BSJ.2022.19.2.0422](https://doi.org/10.21123/BSJ.2022.19.2.0422).
- [11] Y. Afrianto Singgalen, "Sentiment Classification of Over-Tourism Issues in Responsible Tourism Content using Naïve Bayes Classifier," *J. Comput. Syst. Informatics*, vol. 5, no. 2, pp. 275–285, 2024, doi: [10.47065/josyc.v5i2.4904](https://doi.org/10.47065/josyc.v5i2.4904).
- [12] M. V. Kalaivani, Shamini B, S. R., and T. S. D., "Hotel Reviews Sentiment Analysis Using Machine Learning," *Int. Res. J. Mod. Eng. Technol. Sci.*, no. 03, pp. 3047–3056, 2023, doi: [10.56726/irjmets34902](https://doi.org/10.56726/irjmets34902).
- [13] J. W. Bi, X. E. Zhu, and T. Y. Han, "Text Analysis in Tourism and Hospitality: A Comprehensive Review," *J. Travel Res.*, no. May, 2024, doi: [10.1177/00472875241247318](https://doi.org/10.1177/00472875241247318).
- [14] N. Hasanati, Q. Aini, and A. Nuri, "Implementation of Support Vector Machine with Lexicon Based for Sentiment Analysis on Twitter," in *2022 10th International Conference on Cyber and IT Service Management (CITSM)*, 2022, pp. 1–4, doi: [10.1109/CITSM56380.2022.9935887](https://doi.org/10.1109/CITSM56380.2022.9935887).
- [15] A. M. Pravina, "Sentiment Analysis of Delivery Service Opinions on Twitter Documents using K-Nearest Neighbor," *JATISI (Jurnal Tek. Inform. dan Sist. Informasi)*, vol. 9, no. 2, pp. 996–1012, 2022, doi: [10.35957/jatisi.v9i2.1899](https://doi.org/10.35957/jatisi.v9i2.1899).
- [16] T. H. Jaya Hidayat, Y. Ruldeviyani, A. R. Aditama, G. R. Madya, A. W. Nugraha, and M. W. Adisaputra, "Sentiment analysis of twitter data related to Rinca Island development using Doc2Vec and SVM and logistic regression as classifier," *Procedia Comput. Sci.*, vol. 197, pp. 660–667, 2022, doi: <https://doi.org/10.1016/j.procs.2021.12.187>.
- [17] A. Sabrani, I. G. W. Wedashwara W., and F. Bimantoro, "Multinomial Naïve Bayes untuk Klasifikasi Artikel Online tentang Gempa di Indonesia," *J. Teknol. Informasi, Komputer, dan Apl. (JTika)*, vol. 2, no. 1, pp. 89–100, 2020, doi: [10.29303/jtika.v2i1.87](https://doi.org/10.29303/jtika.v2i1.87).
- [18] A. R. Alaei, S. Becken, and B. Stantic, "Sentiment Analysis in Tourism: Capitalizing on Big Data," *J. Travel Res.*, vol. 58, no. 2, pp. 175–191, 2021, doi: [10.1177/0047287517747753](https://doi.org/10.1177/0047287517747753).
- [19] A. Rifa'i, H. Sujaini, and D. Prawira, "Sentiment Analysis Objek Wisata Kalimantan Barat Pada Google Maps Menggunakan Metode Naive Bayes," *J. Edukasi dan Penelit. Inform.*, vol. 7, no. 3, p. 400, 2021, doi: [10.26418/jp.v7i3.48132](https://doi.org/10.26418/jp.v7i3.48132).

- [20] C. Muhamad, S. Ramdani, A. N. Rachman, and R. Setiawan, "Comparison of the Multinomial Naive Bayes Algorithm and Decision Tree with the Application of AdaBoost in Sentiment Analysis Reviews PeduliLindungi Application," *Int. J. Inf. Syst. Technol. Akreditasi*, vol. 6, no. 158, pp. 419–430, 2022, doi: <https://doi.org/10.30645/ijistech.v6i4.257>.
- [21] H. Belyadi and A. Haghighat, "Machine Learning Guide for Oil and Gas Using Python: A Step-by-Step Breakdown with Data, Algorithms, Codes, and Applications.," H. Belyadi and A. B. T.-M. L. G. for O. and G. U. P. Haghighat, Eds., Gulf Professional Publishing, 2021, p. iii. doi: <https://doi.org/10.1016/B978-0-12-821929-4.01001-5>.
- [22] R. F. Majbur, F. Yanuar, D. Devianto, D. Matematika, and U. Andalas, "Penerapan Metode Naïve Bayes Classifier Dalam Menganalisis Sentimen Pada Media Sosial X Terhadap Pilpres 2024 Di Indonesia," *J. Ilm. Pendidik. Mat. Mat. dan Stat.*, vol. 5, no. 3, pp. 2012–2025, 2024, doi: 10.46306/lb.v5i3.

