

SENTIMENT ANALYSIS OF THE EDLINK APPLICATION USING SVM AND NAIVE BAYES

Maya Ramadhani¹, Najmuddin Mubarak MR², Saparudin³, Ahmad Zamsuri⁴

¹ Magister Ilmu Komputer, Universitas Lancang Kuning; maya@unilak.ac.id

² Magister Ilmu Komputer, Universitas Lancang Kuning; mudinnajim@gmail.com

³ Magister Ilmu Komputer, Universitas Lancang Kuning; saparudin0510@gmail.com

⁴ Magister Ilmu Komputer, Universitas Lancang Kuning; ahmadzamsuri@unilak.ac.id

ARTICLE INFO

Keywords:

Sentiment Analysis;
SVM;
Naive Bayes;
NLP;
Edlink

Article history:

Received 2025-07-21

Revised 2025-08-12

Accepted 2025-08-20

ABSTRACT

The growth of Learning Management Systems (LMS) such as Edlink has made sentiment analysis on user feedback crucial for evaluating user satisfaction and improving services. This study aims to implement and compare the performance of Support Vector Machine (SVM) and Naive Bayes algorithms in classifying sentiments in user reviews on the Edlink application. A dataset of 6,000 Indonesian-language comments was collected and preprocessed using standard Natural Language Processing (NLP) techniques such as tokenizing and stemming. Feature extraction was carried out using Term Frequency-Inverse Document Frequency (TF-IDF) and Bag of Words (BoW). The models were evaluated using 10-fold cross-validation. The results show that SVM combined with TF-IDF provides higher accuracy and better classification performance compared to Naive Bayes. These findings support the application of SVM and TF-IDF in Indonesian sentiment analysis tasks and offer insights for Edlink developers to enhance their platform.

This is an open access article under the [CC BY-NC-SA](https://creativecommons.org/licenses/by-nc-sa/4.0/) license.



Corresponding Author:

Maya Ramadhani

Magister Ilmu Komputer, Universitas Lancang Kuning; maya@unilak.ac.id

1. INTRODUCTION

The advancement of digital technology has significantly impacted various aspects of life, including the field of education. One of the fundamental changes brought about by this development is the integration of information technology into the learning process through platforms known as Learning Management Systems (LMS) (Zhafira, Rahayudi, & Indriati, 2021). Edlink is one such LMS platform that has gained popularity in Indonesia and is widely used by educational institutions to facilitate communication, content distribution, and the management of academic activities in an online environment (Nurkumalawati & Rofli, 2023).

As a user-based application, Edlink provides a review or comment feature that allows users—both students and lecturers—to share their experiences while using the platform. These comments reflect

various aspects of usage, such as ease of access, system responsiveness, interface convenience, and application performance stability. This user-generated data serves as a valuable resource for developers to evaluate and improve the quality of Edlink's services. However, user comments on Edlink are generally in the form of free text and unstructured (Susandri et al., 2021). To be effectively utilized in data analysis, they require proper text preprocessing. Two essential techniques in this stage are tokenizing and stemming. Tokenizing breaks the text into units of words (tokens), while stemming reduces words to their root forms. These techniques are foundational to Natural Language Processing (NLP) and must be applied before further analysis, such as sentiment classification, can be performed (Yunitasari & Putera, 2021).

According to Indah and Mardiyanto (2020), the implementation of tokenizing and stemming has been shown to improve model accuracy in sentiment analysis. This is supported by Putri et al. (2021), who emphasized that the quality of classification outcomes heavily depends on the preprocessing stage. Therefore, careful execution of tokenizing and stemming is crucial for obtaining accurate and meaningful data representations (Santoso & Wibowo, 2022). Although many studies have explored these NLP techniques in the context of e-commerce and social media, their application in the education sector—particularly on user comments from LMS platforms like Edlink—remains limited. In fact, sentiment analysis of LMS user feedback holds significant potential for uncovering users' perceptions and experiences of digital learning systems (Suswadi & Erkamim, 2023).

Based on these considerations, this study aims to analyze the sentiment of Edlink user comments by implementing Support Vector Machine (SVM) and Naive Bayes algorithms, and to evaluate the performance of both models (Septian, Fachrudin, & Nugroho, 2019). The study begins with a data preprocessing stage involving tokenizing and stemming to construct textual features, which are subsequently used in classifying sentiment into three main categories: positive, negative, and neutral (Fitri, 2020).

2. METHODS

This study employs a quantitative descriptive approach that integrates Natural Language Processing (NLP) techniques with machine learning-based sentiment classification (Larasati, Ratnawati, & Hanggara, 2022). The primary objective is to develop a system capable of automatically classifying Edlink app reviews into positive and negative sentiment categories based on the textual content provided by users (Hendriyanto, Ridha, & Enri, 2022). To achieve this, a structured methodological framework was implemented, comprising several key stages: data collection, preprocessing, feature extraction, model training, and performance evaluation. User comments were scraped from the Google Play Store using the `google_play_scraper` library, resulting in a dataset of 6,000 Indonesian-language comments.

During the preprocessing phase, standard NLP techniques were applied—these included cleaning (removal of punctuation, numbers, and URLs), lowercasing, tokenizing, stopword removal, and stemming using the `Sastrawi` library. The clean and standardized text was then transformed into numerical representations using two widely adopted feature extraction methods: Term Frequency-Inverse Document Frequency (TF-IDF) and Bag of Words (BoW). The classification task was performed using two machine learning algorithms: Support Vector Machine (SVM) and Naive Bayes. Sentiment labels (positive or negative) were assigned using a lexicon-based approach, supported by manual validation to ensure labeling accuracy. The dataset was split into training and testing subsets using an 80:20 ratio, and performance was evaluated using 10-fold cross-validation. The evaluation metrics used in this study included accuracy, precision, recall, and F1-score, allowing for a comprehensive assessment of model effectiveness. Figure 1 illustrates the overall workflow applied in this study, from raw data acquisition to final sentiment classification.

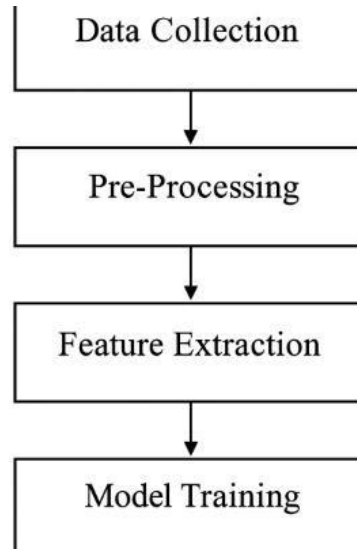


Figure 1. Structured Stages

2.1 Dataset

The dataset used in this study was obtained through an Application Programming Interface (API) call to the Edlink platform, a Learning Management System (LMS) utilized by various educational institutions across Indonesia. The collected data consisted of original user comments—primarily from students and lecturers—providing feedback related to their experience using the application for online learning activities.

Each entry in the dataset includes comment text along with supporting metadata such as date, user role, and comment ID. Since there were no numeric scores or ratings available (as found in e-commerce platforms), sentiment labeling was conducted using a semi-automated approach:

- Lexicon-Based Labeling, by utilizing an Indonesian sentiment lexicon containing positive and negative words.
- Manual validation on a portion of the data to minimize misclassification.

The comments were classified into two main categories: positive sentiment and negative sentiment. Comments with neutral or ambiguous tones were excluded from the dataset to focus the analysis on clearly polarized sentiments. The outcome of this process was a clean dataset, ready for NLP processing and classification tasks.

2.2 Text Preprocessing

Text preprocessing is a critical first step in sentiment analysis using Natural Language Processing (NLP), especially when working with unstructured text data, such as user comments on the Edlink platform. This stage aims to transform raw textual data into a structured, clean format suitable for feature extraction and model training. In this study, preprocessing was carried out in several systematic steps.

First, data cleaning was performed to remove irrelevant elements such as punctuation, numbers, URLs, emoticons, hashtags, mentions, and other special characters that hold no linguistic meaning. This step ensures that only meaningful words are retained.

Next, case folding was applied by converting all text into lowercase to standardize text representation and avoid duplication caused by capitalization differences (e.g., “Edlink” vs. “edlink”). Afterward, the text was segmented into individual tokens using the tokenizing process. For example, the sentence “Edlink sangat membantu pembelajaran daring” would be tokenized into [“edlink”, “sangat”, “membantu”, “pembelajaran”, “daring”]. This process forms the basis for semantic analysis, allowing each word to be analyzed independently.

Stopword removal followed, where common, non-informative words such as “yang”, “dan”, “di”, and “ke” were removed using a predefined stopwords list from NLTK and the Sastrawi library. This was done to focus the model on meaningful words.

The next step was stemming, which converts words to their root forms using the Sastrawi library. Words like “mengakses”, “diakses”, and “pengaksesan” would all be reduced to “akses”. Stemming helps minimize unnecessary word variations and improves consistency in text modeling.

The final stage was label encoding, where comments were labeled based on sentiment polarity. Comments with positive sentiment were assigned the label 1, while those with negative sentiment were labeled 0. Since the data consisted of free text without numeric ratings, sentiment labeling was based on a lexicon approach combined with manual validation to enhance accuracy.

As a result of this preprocessing pipeline, the corpus was cleaned, standardized, and numerically encoded—ready for the feature extraction stage using methods such as TF-IDF and Bag of Words. The success of this process greatly influences the performance of the classification algorithms used in this study, namely Support Vector Machine (SVM) and Naive Bayes.

2.3 Feature Extraction

After preprocessing, user comment texts needed to be converted into numerical representations to be processed by machine learning algorithms. This process, known as feature extraction, employed two popular text representation techniques: Term Frequency-Inverse Document Frequency (TF-IDF) and Bag of Words (BoW).

TF-IDF is a word-weighting technique that evaluates the importance of a word in a document relative to the entire corpus. The Term Frequency (TF) component measures how often a word appears in a comment, while the Inverse Document Frequency (IDF) penalizes words that appear frequently across many documents, as such words tend to carry less unique information. Therefore, TF-IDF is effective in emphasizing truly relevant words in sentiment analysis and reducing the influence of common, non-discriminative words.

In contrast, Bag of Words (BoW) is a simpler method that calculates the frequency of word occurrences without considering their order or context. BoW generates high-dimensional vectors representing the presence of words in a document. Although BoW does not capture word semantics or relationships, it is widely used as a baseline in many text classification tasks due to its ease of implementation and competitive results, especially for small to medium-sized datasets.

In this study, both TF-IDF and BoW were applied in parallel and used as input to two different classifiers—SVM and Naive Bayes—to compare their performance. This comparison aimed to identify which combination of feature extraction technique and classification model was most effective in classifying the sentiment of Edlink user comments.

2.4 Classification Algorithm

This study implemented two machine learning algorithms for sentiment classification of Edlink user comments: Support Vector Machine (SVM) and Naive Bayes. SVM works by constructing an optimal hyperplane that separates two classes of data—positive and negative sentiment—with maximum margin. Its capability to handle high-dimensional data, such as text, makes SVM a popular choice for classification tasks, including sentiment analysis.

Naive Bayes, on the other hand, is an ensemble method based on a collection of decision trees built randomly. By aggregating predictions from multiple trees, Naive Bayes can provide stable and overfitting-resistant predictions. It is also effective in dealing with nonlinear and noisy data. In this study, both algorithms were trained using preprocessed and feature-extracted comment data through two approaches: TF-IDF and BoW. A train-test split method with a ratio of 80:20 was applied—80% for training and 20% for testing. This approach allows for an objective evaluation of each model's performance under different feature combinations.

2.5 Evaluation and Analysis

The classification models in this study were evaluated using four commonly used performance metrics: accuracy, precision, recall, and F1-score.

- Accuracy measures the proportion of correct predictions among all test data.
- Precision indicates the proportion of true positives among all positive predictions made by the model.
- Recall assesses how many actual positive comments were correctly identified by the model.
- F1-score is the harmonic mean of precision and recall and is particularly important when dealing with imbalanced data.

To assess the performance of each feature-model combination, four test scenarios were conducted:

- (1) SVM with TF-IDF,
- (2) SVM with BoW,
- (3) Naive Bayes with TF-IDF, and
- (4) Naive Bayes with BoW.

Each combination was tested to evaluate the models' accuracy and precision in classifying user comments into positive or negative sentiment categories. The evaluation results were presented in tables and graphs to facilitate comparison across models.

In addition to quantitative analysis, manual review of selected predictions was conducted to assess the models' ability to understand informal Indonesian language commonly used by Edlink users and to evaluate the consistency and relevance of their predictions. This comprehensive evaluation offers a well-rounded perspective on each model's effectiveness in performing sentiment analysis on user-generated text data.

his study implemented two machine learning algorithms for sentiment classification of Edlink user comments: Support Vector Machine (SVM) and Naive Bayes.

SVM works by constructing an optimal hyperplane that separates two classes of data—positive and negative sentiment—with maximum margin. Its capability to handle high-dimensional data, such as text, makes SVM a popular choice for classification tasks, including sentiment analysis.

Naive Bayes, on the other hand, is an ensemble method based on a collection of decision trees built randomly. By aggregating predictions from multiple trees, Naive Bayes can provide stable and overfitting-resistant predictions. It is also effective in dealing with nonlinear and noisy data.

In this study, both algorithms were trained using preprocessed and feature-extracted comment data through two approaches: TF-IDF and BoW. A train-test split method with a ratio of 80:20 was applied—80% for training and 20% for testing. This approach allows for an objective evaluation of each model's performance under different feature combinations.

3. FINDINGS AND DISCUSSION

This study aims to detect sentiment in Edlink app user reviews using a Natural Language Processing (NLP) approach with Support Vector Machine (SVM) and Naive Bayes algorithms. In this study, data is classified into two main labels: positive and negative. Positive labels represent reviews indicating user satisfaction, while negative labels represent complaints or dissatisfaction.

3.1 Dataset (Application Review Collection for Edlink)

This study began with the data collection phase, where user reviews of the Edlink application were retrieved from the Google Play Store using a Python library called `google_play_scraper`. A total of 6,000 comments were collected in the Indonesian language and filtered specifically from users located in Indonesia. The filtering process was conducted to ensure that the reviews obtained were relevant to the local user context.

Figure 2. Edlink Application Data Retrieval

The first step in sentiment analysis involved data cleaning and text normalization, which aimed to remove unnecessary elements from the text such as URLs, numbers, punctuation marks, foreign characters, and excessive whitespace. Additionally, all letters were converted to lowercase to avoid differences in meaning between words like “Edlink” and “edlink.”

Figure 3. Edlink Application Review Extraction Results Data (Post-Scraping & Preprocessing)

Following the initial cleaning and normalization, a more advanced filtering process was carried out to eliminate remaining irrelevant data. This step included removing empty rows or rows containing only spaces in the content_token column and eliminating very short words (fewer than three characters) that typically have little semantic value in sentiment analysis. However, exceptions were made for short but relevant terms such as “ktp”, “kk”, “apk”, “app”, and “bug,” as these frequently appear in app reviews and hold specific meanings.

Figure 4. Advanced Screening Data

3.4 Tokenizing

After the comments were cleaned and filtered, the next step was tokenizing, which is the process of splitting review texts into individual word units (tokens) using specific patterns. In this study, RegexpTokenizer from the NLTK library was used with the `\w+` pattern, meaning only word characters (letters and numbers) were extracted while ignoring symbols and punctuation.

```
[90]: # Tokenizing
Regex = RegexpTokenizer('\w+')
data['content_tokenized'] = data['content_token'].apply(Regex.tokenizer)
data
```

C:\Users\62812\AppData\Local\Temp\ipykernel_16812\3474824386.py:3: SettingWithCopyWarning:
A value is trying to be set on a copy of a slice from a DataFrame.
Try using .loc[row_indexer,col_indexer] = value instead

See the caveats in the documentation: https://pandas.pydata.org/pandas-docs/stable/user_guide/indexing.html#returning-a-view-versus-a-copy
data['content_tokenized'] = data['content_token'].apply(Regex.tokenizer)

```
[90]:
```

	reviewId	userName	userImage	content	score	thumbsUpCount	reviewCreatedVersion	at	replyContent	repliedAt	app
0	ce12bb89-d57a-4312-9390-a63472400f32	Jesenia Filomena	https://play-lh.googleusercontent.com/a/ACg8oc...	knp aplikasinya saat di berikan quis oleh dose...	1	0	4.8.10	2025-05-09 15:38:31	Halo Kak, fitur tersebut saat ini sedang dalam...	2025-05-09 15:54:42	
1	6238722f-81f8-4ee2-8d90-7b6576fb7d74	Veronica	https://play-lh.googleusercontent.com/a-/ALV-U...	Pada saat scan barcode absen tidak bisa di zoo...	1	0	4.8.10	2025-05-08 16:21:28	Halo Kak, fitur ini sudah kami kembangkan yah...	2025-05-08 16:23:53	
2	f7af4d52-c5ef-40b3-b957-b4a5f97ca8fd	Muhammad Latiful Ilham	https://play-lh.googleusercontent.com/a/ACg8oc...	saat kuis hampir semua mahasiswa telah selesai...	1	0	NaN	2025-05-05 13:57:06	Halo Kak, terima kasih atas masukannya, akan k...	2025-05-08 13:54:26	

Figure 5. Word-Level Representation-Tokenizing

3.5 Stemming

Stemming is the process of reducing words to their base forms so that machine learning models can more easily recognize word patterns. In this study, stemming was performed using the Sastrawi library, which is specifically designed for the Indonesian language. For example, the words “mengakses,” “diakses,” and “pengaksesan” were all reduced to their root form “akses.”

```
[94]: data
```

```
[94]:
```

	reviewId	userName	userImage	content	score	thumbsUpCount	reviewCreatedVersion	at	replyContent	repliedAt	app
0	ce12bb89-d57a-4312-9390-a63472400f32	Jesenia Filomena	https://play-lh.googleusercontent.com/a/ACg8oc...	knp aplikasinya saat di berikan quis oleh dose...	1	0	4.8.10	2025-05-09 15:38:31	Halo Kak, fitur tersebut saat ini sedang dalam...	2025-05-09 15:54:42	
1	6238722f-81f8-4ee2-8d90-7b6576fb7d74	Veronica	https://play-lh.googleusercontent.com/a-/ALV-U...	Pada saat scan barcode absen tidak bisa di zoo...	1	0	4.8.10	2025-05-08 16:21:28	Halo Kak, fitur ini sudah kami kembangkan yah...	2025-05-08 16:23:53	
2	f7af4d52-c5ef-40b3-b957-b4a5f97ca8fd	Muhammad Latiful Ilham	https://play-lh.googleusercontent.com/a/ACg8oc...	saat kuis hampir semua mahasiswa telah selesai...	1	0	NaN	2025-05-05 13:57:06	Halo Kak, terima kasih atas masukannya, akan k...	2025-05-08 13:54:26	

Figure 6. Stemming Results

3.6 Text Data Visualization (Preprocessing Visualization)

This stage aimed to understand the dominant characteristics of words found in Edlink user reviews. Using WordCloud techniques, a visual representation of the most frequently occurring words in positive and negative sentiment labels was displayed.

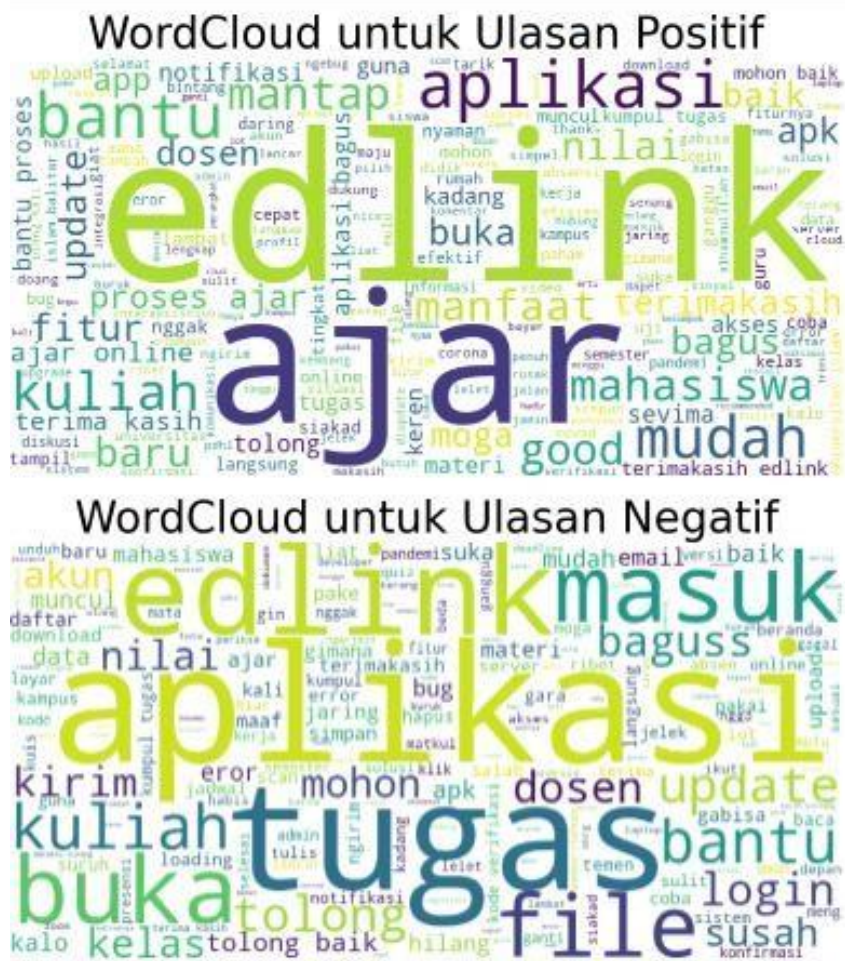


Figure 7. Visualization Results

3.7 Word Frequency Analysis

This phase aimed to identify the most frequently occurring words in positive and negative sentiment comments after stemming. The Counter function from the collections library was used to count the frequency of each token (word).

Top 10 Kata Paling Banyak Muncul - Positif:

ajar: 86
edlink: 85
aplikasi: 68
bantu: 65
mudah: 45
kuliah: 36
bagus: 29
mantap: 27
terimakasih: 26
mahasiswa: 26

Top 10 Kata Paling Banyak Muncul - Negatif:

aplikasi: 79
tugas: 73
edlink: 60
tolong: 44
masuk: 39
buka: 37
file: 33
kuliah: 32
bantu: 30
baik: 30

Figure 8. Top 10 Words Most Frequently Used by Users

3.8 Feature Extraction Using TF-IDF

After cleaning and stemming the user review texts, the next step was to convert the text into a numerical representation that could be processed by machine learning models. This study employed Term Frequency-Inverse Document Frequency (TF-IDF) as the feature extraction method. TF-IDF assigns a weight to each word based on how frequently it appears in a specific document compared to the entire corpus.

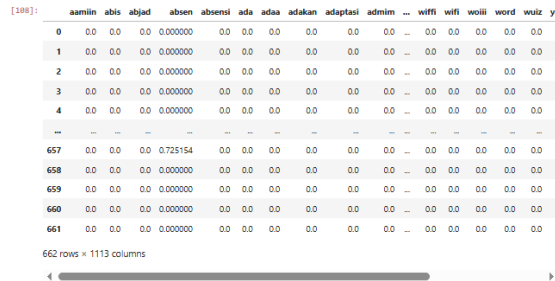


Figure 9. Feature Extraction Output Using TF-IDF

3.9 SVM Model Evaluation Using Cross-Validation and Confusion Matrix

Once feature representations were created using TF-IDF, the Support Vector Machine (SVM) algorithm was applied to classify user sentiment in Edlink reviews. Model performance was evaluated using 10-fold cross-validation to obtain a comprehensive view of the model's accuracy.

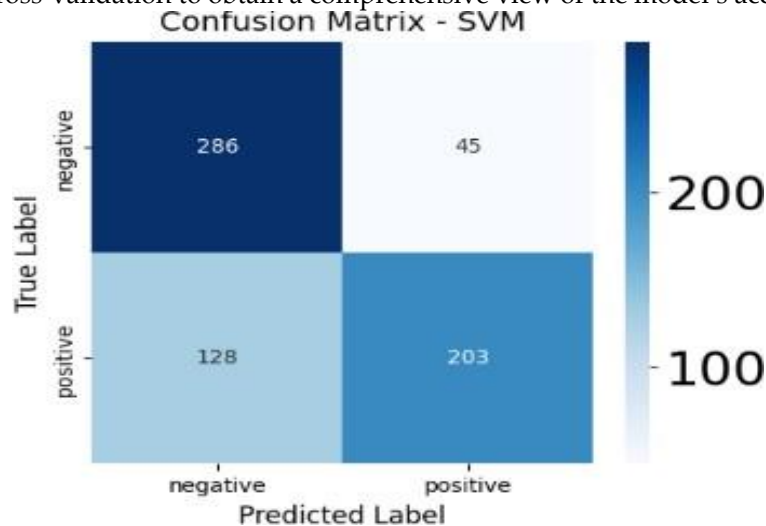


Figure 10. SVM Model Evaluation

3.10 Naive Bayes Model Evaluation Using Cross-Validation and Confusion Matrix

To compare the performance of classification models, this study also implemented the Multinomial Naive Bayes algorithm for sentiment classification of Edlink user comments. Evaluation was conducted using 10-fold cross-validation to determine the overall model accuracy. In addition, a confusion matrix was used to assess the classification precision between positive and negative sentiment labels.

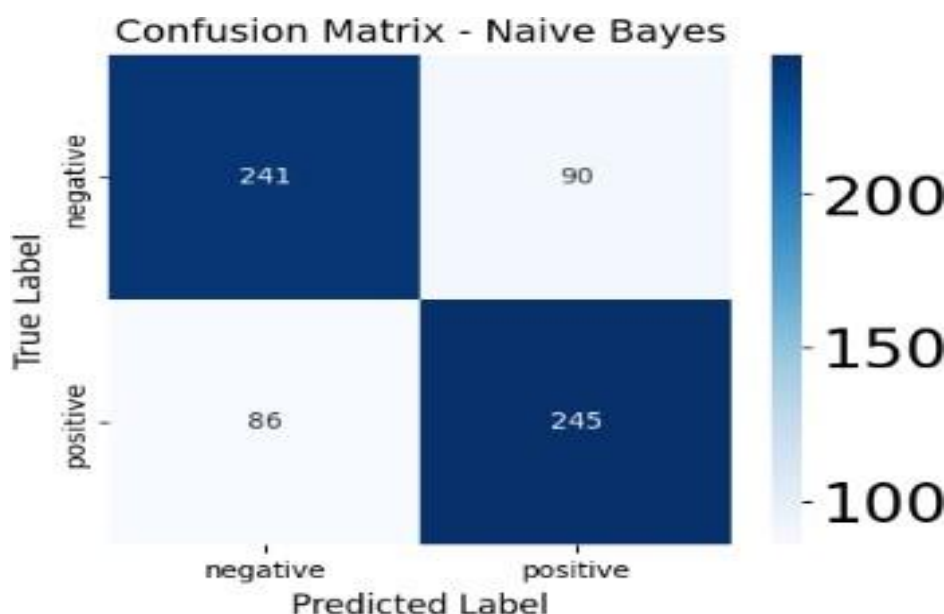


Figure 11. Naive Bayes Model Evaluation

4. CONCLUSION

This study demonstrates that the Natural Language Processing (NLP) approach combined with the Support Vector Machine (SVM) algorithm and TF-IDF features is highly effective for analyzing the sentiment of Edlink app user comments. The evaluation results show that the SVM model with TF-IDF produces higher accuracy than Naive Bayes, especially in accurately distinguishing positive and negative comments. The use of comprehensive text preprocessing techniques, such as cleaning, tokenizing, stopword removal, and stemming, plays a crucial role in improving data quality and model performance. These findings strengthen evidence from previous studies that SVM is a reliable algorithm for classifying Indonesian-language text.

Acknowledgments: The authors thank Universitas Lancang Kuning for support and the lecturers of the NLP course for their guidance throughout this project.

Conflicts of Interest: The authors declare no conflict of interest.

REFERENCES

- Fitri, E. (2020). Sentiment analysis of the Ruangguru application using Naive Bayes, Random Forest and Support Vector Machine algorithms. *Jurnal Transformasi*, 18(1), 71.
- Hendriyanto, M. D., Ridha, A. A., & Enri, U. (2022). Analisis sentimen ulasan aplikasi Mola pada Google Play Store menggunakan algoritma Support Vector Machine. *Jurnal Information Technology and Computer Science*, 5(1), 1–7.
- Larasati, F. A., Ratnawati, D. E., & Hanggara, B. T. (2022). Analisis sentimen ulasan aplikasi Dana dengan metode Random Forest. *Jurnal Teknologi Informasi dan Ilmu Komputer*, 6(9), 4305–4313. Retrieved from <http://j-ptiik.ub.ac.id>
- Nurkumalawati, I., & Rofii, M. S. (2023). Public review of M-Paspor application in Indonesia: Mobile government, digital resilience, cyber security. *Atlantis Press SARL*.
- Santoso, D. P., & Wibowo, W. (2022). Analisis sentimen ulasan aplikasi Buzzbreak menggunakan metode Naïve Bayes classifier pada situs Google Play Store. *Jurnal Sains dan Seni ITS*, 11(2).
- Septian, J. A., Fachrudin, T. M., & Nugroho, A. (2019). Analisis sentimen pengguna Twitter terhadap polemik persepakbolaan Indonesia menggunakan pembobotan TF-IDF dan K-nearest neighbor. *Jurnal Intelligent System and Computing*, 1(1), 43–49.
- Susandri, S., Defit, S., Riandari, F., & Sinaga, B. (2021). Ekplorasi timeline: Waktu respon pesan terbaik WhatsApp group “Gurauan kita STMIK Amik”. *MATRIK Jurnal Manajemen, Teknik Informatika dan Rekayasa Komputer*, 20(2), 317–324.

- Suswadi, S., & Erkamim, M. (2023). Sentiment analysis of Shopee app reviews using random forest and support vector machine. *Ilkom: Jurnal Ilmiah*, 15(3), 427–435.
- Yunitasari, Y., & Putera, A. R. (2021). Analisis sentimen masyarakat di Twitter terkait pandemi Covid-19. *Smatika Jurnal*, 11(1), 22–26.
- Zhafira, D. F., Rahayudi, B., & Indriati, I. (2021). Zhafira, Dhaifa Farah Rahayudi, Bayu Indriati, Indriati. *Jurnal Sistem Informasi, Teknologi Informasi, dan Edukasi Sistem Informasi*, 2(1), 55–63.