



Sentiment Analysis of Online Lending Services Using Support Vector Machine and Logistic Regression

Ardita Isnanda Rahayu*, Slamet Kacung

Faculty of Engineering, Informatics, Dr. Soetomo University, Surabaya, Indonesia

Email: ^{1,*}ardita.rahayu583@gmail.com, ²slamet@unitomo.ac.id

Correspondence Author Email: ardita.rahayu583@gmail.com

Abstract—This research examines public sentiment toward online lending services in Indonesia by analyzing opinions from social media platforms, specifically YouTube and Twitter, collected from January 2021 to January 2024. The objective of this study is to develop an accurate sentiment classification system that can effectively categorize public opinions into positive, negative, and neutral sentiments, thereby providing valuable insights for regulatory bodies and service providers to understand consumer concerns and improve service quality. The collected data underwent thorough preprocessing, semi-automatic labeling, and Term Frequency-Inverse Document Frequency (TF-IDF) weighting. Four classification models were evaluated: Support Vector Machine (SVM) with Linear, Polynomial, and Radial Basis Function (RBF) kernels, and Logistic Regression. Results demonstrate that Linear SVM achieves the best performance with an accuracy of 90.17% and an F1-score of 0.902, effectively categorizing sentiments across all classes while excelling particularly in negative and neutral categories. The expected impact of this analysis is to provide evidence-based recommendations for policymakers in financial technology regulation and help online lending service providers understand consumer satisfaction levels to improve their service delivery. This study offers valuable insights for service providers and regulatory bodies seeking to better understand and address public concerns in this domain.

Keywords: Sentiment Analysis; Online Lending; Social Media; Support Vector Machine; Logistic Regression

1. INTRODUCTION

The rapid growth of financial technology (fintech) in Indonesia has led to the proliferation of online lending services, creating both opportunities and concerns among consumers. Understanding public sentiment toward these services is crucial for regulatory bodies, service providers, and consumers themselves. With the increasing volume of user-generated content on social media platforms, automated sentiment analysis has become essential for processing large-scale public opinions efficiently. However, analyzing sentiment in Indonesian social media presents unique challenges due to informal language, slang usage, and cultural context that require specialized preprocessing and classification approaches [1].

Previous studies have shown significant efforts to improve sentiment analysis approaches for Indonesian text. Ahmad (2023) analyzed sentiments for Akulaku and Kredivo online loans using SVM and found that most online loan service ratings were negative, with classification accuracy reaching 88.2% for Kredivo [2]. A recent study in the fintech domain revealed that SVM achieved an average accuracy of 87.32% when analyzing sentiment toward emerging versus legacy financial technologies, reinforcing its effectiveness for text classification in financial contexts [3]. Ilham et al. (2024) created an online classification system for online loan applications, where the soft voting classification approach optimized with particle swarm optimization achieved the highest accuracy compared to individual models such as SVM and XGBoost [4]. Ranti et al. (2023) compared logistic regression, SVM, and gradient boosting algorithms for student comment analysis and found that all models were effective, with gradient boosting generally performing best [5]. Budianita et al. (2022) compared dictionary-based and machine-learning approaches for sentiment analysis of social media data, finding that SVM produces clearer and more stable classification than lexical approaches [6].

Additionally, recent studies by Islam et al. (2024) demonstrated the effectiveness of deep learning methods in sentiment analysis, achieving comprehensive results in financial sentiment analysis across different approaches [7]. Jelodar et al. (2023) explored comprehensive survey on sentiment analysis and opinion mining, covering various tasks, approaches, and applications in the field [8]. Taboada et al. (2012) investigated lexicon-based methods for sentiment analysis, providing foundational understanding of dictionary-based approaches that remain relevant for current research [9]. Du et al. (2024) analyzed financial sentiment analysis techniques and applications, concluding that specialized approaches are needed for financial domain text processing [10].

Despite these advances, several research gaps remain unaddressed in the current literature. First, most existing studies focus on individual classification algorithms without a comprehensive comparison of multiple kernel functions within SVM models, specifically for Indonesian online lending sentiment analysis. Second, limited research has been conducted on the integration of multi-platform social media data (Twitter and YouTube) for financial sentiment analysis in the Indonesian context. Third, existing studies lack a detailed analysis of sentiment distribution patterns and their implications for regulatory policy making. Fourth, there is insufficient exploration of semi-automatic labeling approaches that balance efficiency and accuracy for large-scale Indonesian social media data processing.

This study aims to analyze public sentiment toward online lending services in Indonesia by comparing the performance of several machine learning algorithms. The research focuses on using data from multiple social media platforms, applying a semi-automatic labeling method, and evaluating sentiment patterns. The results are expected to support policymakers and service providers in better understanding public perceptions and improving digital financial services.



2. RESEARCH METHODOLOGY

2.1 Literature Review

Sentiment analysis has evolved significantly in recent years, particularly in the financial services domain. Support Vector Machine (SVM) remains one of the most effective algorithms for text classification due to its ability to handle high-dimensional data and provide robust performance across different datasets [2], [6]. The choice of kernel function in SVM significantly impacts classification performance, with linear kernels being effective for linearly separable data, while polynomial and RBF kernels handle non-linear relationships [5], [11]. Logistic regression, despite its simplicity, has shown consistent performance in sentiment analysis tasks, particularly when dealing with binary and multi-class classification problems [4], [5]. In a recent comparative study, both SVM and logistic regression achieved classification accuracy exceeding 86% when applied to social media sentiment datasets, confirming their reliability across varied informal textual domains [12].

Term Frequency-Inverse Document Frequency (TF-IDF) has been widely adopted as a feature extraction method for text classification, as it effectively captures the importance of terms within documents while reducing the impact of common words [13], [14]. The preprocessing stage is crucial for social media sentiment analysis, as informal language, abbreviations, and cultural expressions significantly impact classification accuracy [15], [16]. Semi-automatic labeling approaches have gained attention as they balance manual annotation accuracy with computational efficiency, which is particularly important for large-scale social media datasets [17], [18].

Online lending sentiment analysis presents unique challenges due to domain-specific terminology, emotional expressions related to financial stress, and regulatory language usage [2], [19]. Previous studies have shown that financial sentiment analysis requires specialized preprocessing and feature extraction techniques to handle the complexity of financial discourse [10], [20]. The integration of multi-platform social media data provides comprehensive insights but requires careful consideration of platform-specific language patterns and user behavior differences [21], [22].

2.2 Research Stages

The research methodology follows a systematic approach to implementing sentiment analysis, comprising several sequential stages to ensure comprehensive data processing and accurate classification results. The research stages are illustrated in the flowchart below, which describes the sequence and stages in conducting research, method application stages, and testing methods to obtain research results according to expectations and research descriptions. Figure 1 illustrates the sentiment distribution of user comments across all platforms.

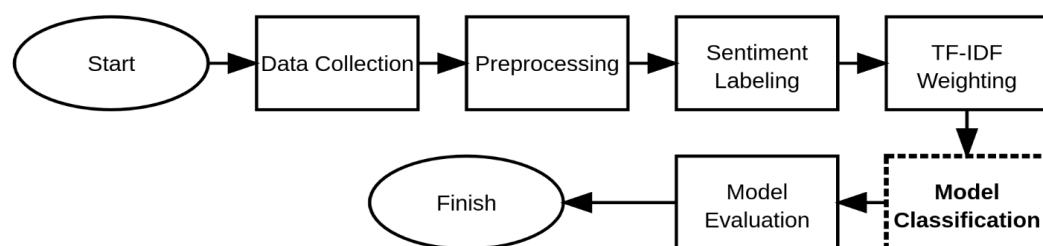


Figure 1. Research stages flowchart

The research begins with data collection from social media platforms, followed by comprehensive preprocessing stages, sentiment labeling, feature extraction using TF-IDF, model training and evaluation, and finally results analysis and discussion.

2.3 Data Collection

The first research phase collects data from two popular social media platforms: Twitter and YouTube. Data is retrieved using Python-based crawling tools and scripts. The SNSCRAPE and BeautifulSoup libraries access public comments that discuss online credit services such as Akulaku, Kredivo, EasyCash, and Credit Pintar [2], [19]. The crawling process uses a multikeyword approach to obtain relevant representative data, with specific keywords acting as search parameters.

Keywords include variations such as "illegal pinjol" (illegal online loan), "OJK" (financial services authority), "online loan", "online credit", "fast loan", and "legal pinjol". Keyword selection was based on a preliminary analysis of general terms commonly used by the public to discuss online credit services on social media [2], [19]. Twitter data was accessed using extended search parameters to ensure data relevance, including language (Indonesian), date range, and geographical location (Indonesia). The YouTube comment collection focuses on financial topics, especially channels and videos that discuss online lending.

This strategy follows the method used by Khairunnisa et al. in their research on Indonesian Twitter message preprocessing [15]. Data collected between January 2021 and January 2024 provided 25,068 comments for further processing. The platform distribution includes 18,457 YouTube comments and 6,611 Twitter comments. These comments represent Indonesian public opinions on various aspects of online credit services, including user experience, interest rates, payment processes, and collection guidelines [2], [10].



2.4 Preprocessing

Preprocessing is critical in text data processing, ensuring data quality and consistency before weighting and classification. This study conducted six preprocessing steps: cleansing, case folding, tokenization, normalization, stopword removal, and stemming. These phases address the unstructured, informal nature of social media comments [15], [16], [22].

Cleansing eliminates unimportant text elements such as symbols, numbers, emojis, punctuation marks, HTML tags, and URLs. The goal is to create clean data that can be processed further. According to previous research [15] Effective cleaning can significantly reduce data complexity and help accelerate the text analysis phase.

Case Folding converts all characters in the text to lowercase to prevent redundancy in word recognition due to capitalization differences. For example, "Pinjol" and "pinjol" are considered identical. This process combines semantically identical word entries that differ only in written form [15], [22].

Tokenization breaks down text word-by-word, allowing independent analysis of each word. Tokenization facilitates word frequency calculation and the application of methods such as TF-IDF. Khairunnisa et al. [15] demonstrated that tokenization significantly influences the success of text vector representation for classification.

Normalization standardizes non-standard word forms. For example, "gk", "ga", or "gak" are normalized to "tidak". This technique is essential for handling informal language variations on social media, according to Aufar et al. [16] Word normalization can improve sentiment classification accuracy by reducing vocabulary diversity and text noise.

Stopword Removal eliminates words with no significant information value, such as "yang" (which), "dan" (and), or "di" (in), using the Indonesian Stopwords dictionary. This process reduces irrelevant features and accelerates machine learning processing [16]. The stopwords dictionary combines Indonesian stopwords lists with functional words unique to online finance services discussions. This approach fits the research context. Certain words usually take stop words in general contexts (such as "money," "interest," and "tenor") [23], [24].

Stemming converts words to their base form. For example, "pembayaran", "bayar", "membayar" become "bayar". This study uses the Nazief-Adriani algorithm, which is popular in Indonesian texts. Khairunnisa et al. [15] showed that the Nazief-Adriani algorithm significantly improved Indonesian text classification performance [14], [23].

2.5 Sentiment Labeling and TF-IDF

The labeling process uses a semi-automatic, weighted dictionary-based method. The Indonesian Sentiment Dictionary (InSet - Indonesian Sentiment Lexicon) identifies word polarity based on specific weights. Each comment receives scores based on positive and negative word contributions from -5 to +5 per word. Comments with positive overall values are classified as positive sentiment, negative scores as negative sentiment, and zero as neutral sentiment. This approach efficiently identifies essential data volumes while maintaining accuracy and consistency [2], [10].

The Inset Dictionary has been enriched with terms specific to financial and online credit issues to improve label accuracy in the context of this study. The concentration of dictionary enrichment included manual extraction from representative comment samples with domain experts (financial experts) to determine the polarity and sentiment intensity of typical words that occur in the discussion of online loans. In addition to the lexicon-based approach, this study uses rule-based methods for exceptional cases such as negation ("not good," "not bad") and intensifiers ("very good," "very slow"). These linguistic rules change the assessment of word emotion based on the use of sentence contexts [17], [18].

To verify the accuracy of automatic labeling, data subgroups (10% of the total data records) were manually checked by three independent annotators. The agreement between the annotators reached 0.78 and was measured using Cohen's Kappa, which showed good consistency. Labeling results showed 45% negative, 32% positive, and 23% neutral comments, reflecting a range of public opinion compared to online credit services. This semi-automated labeling method uses the approach of Khotimah and Sarno [17] in analyzing hotel aspects that coordinate the context of financial services. This strategy allows for efficient processing of large data records and accuracy of labels required reliable model training.

For feature extraction, the TF-IDF method (Term Frequency-Inverse Document Frequency) measures word importance within a document relative to the entire document corpus. Words that appear frequently in a particular document but rarely across all documents receive high TF-IDF weights. The mathematical formula is:

$$TF-IDF(t, d) = f_{t,d} \times \log\left(\frac{N}{df_t}\right) \quad (1)$$

Where $f_{t,d}$ is the frequency of the word t in the document, d N is the total number of documents in the corpus, and df_t is the number of documents containing the word t . The vector representations of documents generated from this method are used as inputs for classification algorithms such as support vector machine and logistic regression [6], [15].

The TF-IDF application in this study was tailored to take into account special features of social media. The modified TF-IDF version handles very low-frequency words (Hapax Legomata) and weights in N-Gram, which contain 1-3 words. This approach has improved the ability of the model to capture broader semantic contexts in informal social media languages recommended by Wahid and SN [14] and further supported by research on the importance of a tailored model for short-form user-generated content with ambiguity, irony, and abbreviations.

2.6 Classification Models

Classification represents the vital process of sentiment analysis in which data is analyzed and categorized into specific groups based on characteristics or patterns. This study categorizes public comments characterized by emotions (positive,



neutral, negative) to determine public opinion regarding online credit services. The two classification methods used are the Support Vector Machine (SVM) and the Logistic Regression. Each offers clear approaches and distinct benefits for treating complex, unstructured textual data [19]. The following is a detailed explanation of each method.

Support Vector Machine (SVM) is an effective supervised learning algorithm, especially for classifying labeled text data. The SVM identifies the optimal hyperplane that separates two or more classes at the maximum margin. This study uses three kernel types: linear, polynomial, and radial basis functions (RBF), to allow SVM to process data at different nonlinearity levels.

The Linear kernel separates the data in the original feature space without transformation. The decision function is:

$$K(x_i, x_j) = x_i^T x_j \quad (2)$$

Polynomial kernels deal with nonlinear properties and map data to higher-dimensional spaces. Polynomial kernel functions are:

$$K(x_i, x_j) = (\gamma x_i^T x_j + r)^d \quad (3)$$

Radial basis functions (RBF) or Gaussian kernels are one of the most popular nonlinear kernels due to their ability to manage complex data distributions. The RBF kernel functions are:

$$K(x_i, x_j) = \exp(-\gamma \|x_i - x_j\|^2) \quad (4)$$

Where γ is a parameter that controls how much influence one data point has on another. The larger the value of γ , the less influence one data point has on its surroundings. These three kernels are implemented and tested to determine the best classification performance for determining sentiment from public comments about online loans. Adequate kernel selection affects the accuracy and stability of classification, depending on the distribution and linearity of the data [5], [11].

Logistic Regression predicts class membership probability based on logistic (sigmoid) functions. The primary logistic regression function uses:

$$P(y = 1 | x) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_n x_n)}} \quad (5)$$

Where β_0 Is the intercept and β_1 Is the regression coefficient for the feature x_i . The model maps the input values to the range [0,1], resulting in class membership probabilities, and uses a threshold to determine the class of observations [4]. Logistic regression was chosen because it provides a probabilistic interpretation of classification results. This allows for further analysis of model confidence in sentiment classification. Furthermore, logistic regression shows greater resilience than some complex models compared to overfitting, especially when dealing with significant standard features, particularly in text classification. A study by Ranti et al [5] demonstrated the effectiveness of logistic regression in sentiment analysis in the Indonesian context, especially in domains with specific vocabulary.

2.7 Model Evaluation

Following model development, classification performance is evaluated using standard metrics in text classification: confusion matrix, accuracy, precision, recall, F1 score, and ROC curve. The confusion matrix forms an evaluation base and shows each class's correct predictions (positive, neutral, negative). From this matrix, additional quantitative metrics evaluate the overall model performance.

Accuracy measures the percentage of correct predictions:

$$\text{Accuracy} = \frac{TP + TN}{TP + TN + FP + FN} \quad (6)$$

Precision indicates model accuracy in classifying positive data:

$$\text{Precision} = \frac{TP}{TP + FP} \quad (7)$$

Recall indicates how effectively the model captures all actual positive data:

$$\text{Recall} = \frac{TP}{TP + FN} \quad (8)$$

F1-score represents the harmonic mean of precision and recall:

$$\text{F1 - Score} = \frac{2 \cdot \text{Precision} \cdot \text{Recall}}{\text{Precision} + \text{Recall}} \quad (9)$$

For social media data that shows imbalanced sentiment classes (e.g., negative comments are dominant compared to neutral or positive), F1 scores are essential to avoid bias toward the majority class. Finally, the ROC curve visually evaluates the classification performance of the model, especially in multi-class classification. These metrics provide a comprehensive overview of model efficiency and effectiveness, providing a more accurate and fair classification of public opinion [4], [5].



3. RESULT AND DISCUSSION

3.1 Data Collection

The study collected data from Twitter and YouTube social media platforms using Python-based crawling techniques from January 2021 to January 2024. Keywords were selected based on topics frequently discussed in public conversations about online credit services, including mentions of regulatory bodies such as OJK. Table 1 presents the sample results after applying the text preprocessing steps.

Table 1. Online Loan Comment Crawling Results

Platform	Comment
Twitter	OJK Reminds People to Beware of Illegal Online Loans #terkiniid #terkinidotid #makassarterkiniid #makassarterkinidotid #ojk #pinjamanonline https://t.co/36QbnDfTWR
Youtube	Why doesn't the government prohibit/disallow these online loans completely? Isn't it clear that these online loans cause so many problems in society? Please, President Prabowo, eliminate these online loans immediately. Because both society and government are harmed, right? So why are they still allowed? Please help, Sir. Thank you ❤️❤️❤️

These comments indicate that public opinion includes complaints about interest rates, collection practices, and demands for regulatory interventions in online credit monitoring.

3.2 Preprocessing

Comments obtained through crawling typically contain many unimportant elements and noise. Therefore, several preprocessing stages prepare data for analysis, including cleaning, case folding, tokenization, normalization, stopword removal, and stemming. Table 2 shows the accuracy performance of the four classification models applied in this study. Table 2 provides an example of how one raw comment is transformed across each preprocessing step.

Table 2. Example of Preprocessing Stages

Stages	Transformation result
Raw text	Why doesn't the government prohibit/disallow these online loans completely? Isn't it clear that there are so many problems in society because of these online loans? Please, President Prabowo, eliminate these online loans immediately. Because both society and government are both harmed, right? So why are they still allowed??? Please help, Sir. Thank you ❤️❤️❤️
Cleansing	Why doesn't the government prohibit disallow these online loans completely Isn't it clear that there are so many problems in society because of these online loans Please President Prabowo eliminate these online loans immediately Because both society and government are both harmed right So why are they still allowed Please help Sir Thank you
Case folding	why doesn't the government prohibit disallow these online loans completely isn't it clear that there are so many problems in society because of these online loans please president prabowo eliminate these online loans immediately because both society and government are both harmed right so why are they still allowed please help sir thank you
Tokenizing	['why', 'doesn't', 'the', 'government', 'prohibit', 'disallow', 'these', 'online', 'loans', 'completely', 'isn't', 'it', 'clear', 'that', 'there', 'are', 'so', 'many', 'problems', 'in', 'society', 'because', 'of', 'these', 'online', 'loans', 'please', 'president', 'prabowo', 'eliminate', 'these', 'online', 'loans', 'immediately', 'because', 'both', 'society', 'and', 'government', 'are', 'both', 'harmed', 'right', 'so', 'why', 'are', 'they', 'still', 'allowed', 'please', 'help', 'sir', 'thank', 'you']
Normalization	['why', 'doesn't', 'the', 'government', 'prohibit', 'disallow', 'these', 'online', 'loans', 'completely', 'isn't', 'it', 'clear', 'that', 'there', 'are', 'so', 'many', 'problems', 'in', 'society', 'because', 'of', 'these', 'online', 'loans', 'please', 'president', 'prabowo', 'eliminate', 'these', 'online', 'loans', 'immediately', 'because', 'both', 'society', 'and', 'government', 'are', 'both', 'harmed', 'right', 'so', 'why', 'are', 'they', 'still', 'allowed', 'please', 'help', 'sir', 'thank', 'you']
Stopword removal	['why', 'government', 'prohibit', 'disallow', 'online', 'loans', 'completely', 'clear', 'problems', 'society', 'because', 'online', 'loans', 'please', 'president', 'prabowo', 'eliminate', 'online', 'loans', 'immediately', 'because', 'society', 'government', 'harmed', 'right', 'why', 'still', 'allowed', 'please', 'help', 'sir', 'thank']

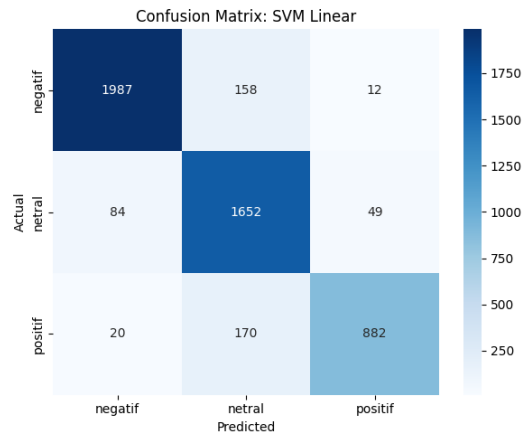


Figure 3. Visualization of Linear SVM Confusion Matrix

The Linear SVM model shows excellent predictive performance for all three sentiment classes. Of 2,157 negative sentiment test data, 1,987 (92.12%) were correctly predicted. In the neutral class, 1,652 (92.55%) of 1,785 data points were correctly predicted. The positive class showed moderate accuracy, correctly predicting 882 (82.28%) out of 1,072 data points. The most common misclassification occurred when positive data was predicted as neutral (170 data points). Figure 4 presents the confusion matrix for the SVM with RBF kernel.

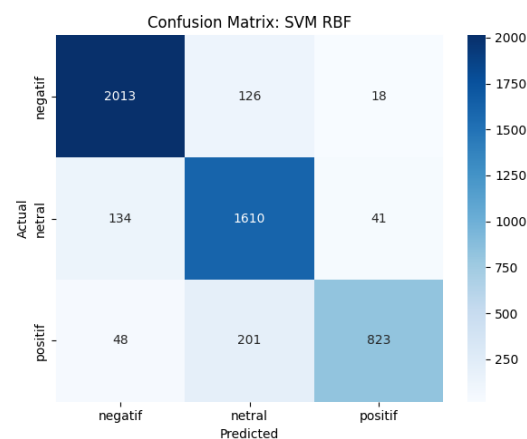


Figure 4. Visualization of SVM RBF Confusion Matrix

The SVM-RBF model performs very well for negative classes, with 2,013 (93.32%) out of 2,157 data points correctly predicted. In the neutral class, prediction accuracy reached 90.20% (1,610 out of 1,785 data points). However, positive class accuracy decreased to 76.77% (823 of 1,072 data points). Error patterns indicate the model's tendency to classify positive data as neutral. This suggests that the RBF kernel is not sensitive enough to capture positive nuances in data with low sentiment intensity.

Figure 5 presents the confusion matrix for the SVM with a Polynomial kernel

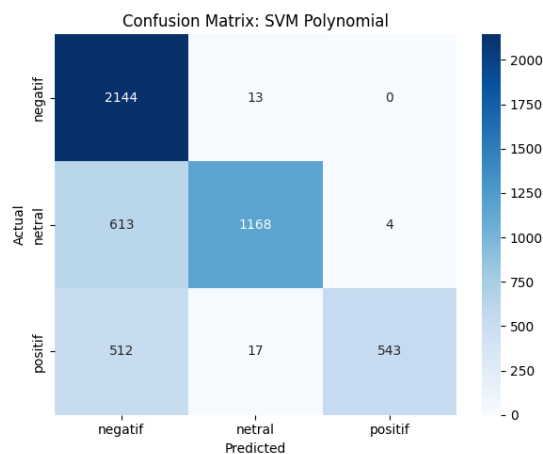


Figure 5. Visualization of SVM Polynomial Confusion Matrix



The Polynomial SVM model shows a strong bias toward the negative class. For the negative class, prediction accuracy reaches 99.40% (2,144 out of 2,157 data points), which is very high. However, this comes at the cost of performance in other classes. For the neutral class, accuracy decreased to 65.43% (1,168 out of 1,785 data points), while for the positive class, accuracy was only 50.65% (543 out of 1,072 data points). This model tends to classify neutral and positive data as negative, misclassifying 613 neutral and 512 positive data points. This pattern suggests that polynomial transformations tend not to capture sentiment nuances in data and tend to oversimplify them as negative class predictions. Figure 6 presents the confusion matrix for the Logistic Regression model.

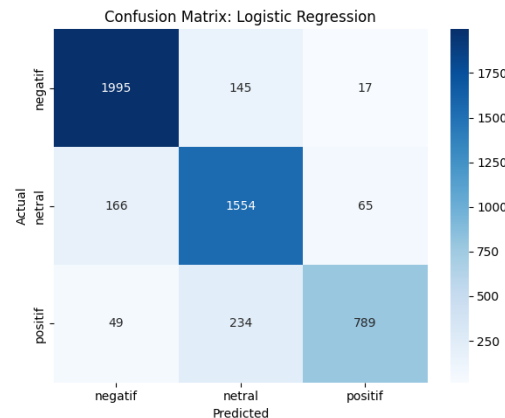


Figure 6. Visualization of Logistic Regression Confusion Matrix

The Logistic Regression model performs well in classifying negative data, reaching 92.49% (1,995 out of 2,157 data points). For the neutral class, accuracy reaches 87.06% (1,554 out of 1,785 data points). However, this model shows a significant weakness in classifying positive data at 73.60% (789 out of 1,072 data points). Serious misclassification occurred when positive data were predicted as neutral (234 data points). This pattern reflects logistic regression properties that tend to be more sensitive to classes with larger sample sizes (in this case, negative and neutral classes).

3.6 Discussion

Experimental results demonstrate that the Support Vector Machine model with a linear kernel provides the most reliable approach for classifying public opinion into three sentiment classes. This model's main advantage lies in its ability to maintain accuracy and stability when handling unstructured and informal social media data. This finding aligns with research by Ahmad [2] and other studies [10] highlighting SVM effectiveness in handling high-dimensional and complex data, particularly in Indonesian textual contexts. Logistic Regression provides a good probabilistic interpretation but still shows limitations in detecting minority classes, such as positive opinions. The sentiment distribution indicates dominance of negative comments regarding online credit services, suggesting that people express considerable concern about service provider practices, high interest rates, and psychological stress. The dominant words such as "Pinjol" (online lending), "person", "payment", "bill", "debt", "fear", and "debt collector" reflect psychological pressure and dissatisfaction with payment and collection systems experienced by people using online credit service providers.

However, the lexicon-based approach used in classification has limitations, particularly in detecting context-dependent meanings, irony, and sarcasm typical in social media language. Therefore, future research should focus on developing deep learning-based methods such as LSTM and BERT, which can capture more complex contexts and semantics. Given its accuracy, stability, and efficiency in processing social media data, the Linear SVM model is recommended as a primary strategy for developing public opinion classification systems, particularly in the Indonesian digital financial services industry.

4. CONCLUSION

This study successfully developed and implemented a machine learning approach to categorize the sentiment of Indonesian public opinion regarding online credit services. The research analyzed 25,068 comments from Twitter and YouTube using comprehensive data preprocessing, lexicon-based labeling, TF-IDF feature extraction, and classification using multiple models. Among the tested algorithms, including Support Vector Machines with linear, polynomial, and RBF kernels, and Logistic Regression, the Linear SVM with 90.17% accuracy and 0.902 F1-score showed consistent performance across all sentiment categories. Sentiment distribution analysis revealed primarily negative attitudes toward online credit services, particularly related to aggressive collection practices, excessive interest rates, and psychological harassment from debt collection agencies. WordCloud visualization confirms these results with terms such as "Pinjol" (online loan), "debt collector", "DC" (debt collector), and "terror". The lexicon-based approach provided efficient labeling methods but showed limitations in interpreting context-dependent nuances, irony, and linguistic ambiguity. Future research would benefit from incorporating deep learning architectures such as LSTM and BERT to improve



semantic understanding and labeling accuracy. Furthermore, extending data collection to other platforms and demographic variables would allow for more comprehensive public opinion analysis. This study contributes to developing opinion analysis systems and provides valuable insights to regulatory authorities and service providers, helping them better understand and address consumer concerns through evidence-based approaches.

REFERENCES

- [1] M. Idris and Mussalimun, "Sentiment Analysis of Google Play Store Reviews using Support Vector Machines," *International Journal of Applied Information Systems*, vol. 12, no. 42, pp. 48–53, 2024. [Online]. Available: www.ijais.org
- [2] K. Ahmad, "Analisis Sentimen Pinjaman Online Akulaku dan Kredivo dengan metode Support Vector Machine (SVM)," *Jurnal Mandalika Literature*, vol. 4, no. 4, pp. 323–332, 2023, doi: 10.36312/jml.v4i4.2045.
- [3] T. A. Nurdin, M. B. Alexandri, W. Sumadinata, and R. Arifianti, "Sentiment Analysis of User Preference for Old Vs New Fintech Technology Using SVM and NB Algorithms," *Management Systems in Production Engineering*, vol. 31, no. 4, pp. 373–380, 2023, doi: 10.2478/mspe-2023-0041.
- [4] A. A. Ilham, E. Warni, and Pahrul, "Soft Voting Classifier With Optimized Weight Using Particle Swarm Optimization on Sentiment Analysis for Online Credit and Loan Application Reviews," *International Journal of Innovative Computing, Information and Control*, vol. 20, no. 2, pp. 359–372, 2024, doi: 10.24507/ijic.20.02.359.
- [5] N. Ranti, M. Hanif, K. H. Hanif, C. Nisa, S. Informasi, and B. Kaltara, "Perbandingan Algoritma Regresi Logistik, Support Vector Machine, dan Gradient Boosting Pada Analisis Sentimen Data Komentar Siswa," *IKOMTI*, vol. 4, no. 2, pp. 27–32, 2023. [Online]. Available: <http://ejournal.uhb.ac.id/index.php/IKOMTI>
- [6] E. Budianita, E. P. Cynthia, A. Pranata, and D. Abimanyu, "Pendekatan berbasis Machine Learning dan Leksikal Pada Analisis Sentimen," in *Seminar Nasional Teknologi Informasi, Komunikasi dan Industri*, pp. 99–104, 2022. [Online]. Available: <https://ejournal.uin-suska.ac.id/index.php/SNTIKI/article/view/19137>
- [7] M. S. Islam, E. Sultana, S. R. Faizabadi, and J. Uddin, "Challenges and future in deep learning for sentiment analysis: a comprehensive review and a proposed novel hybrid approach," *Artificial Intelligence Review*, vol. 57, no. 3, pp. 1–47, 2024, doi: 10.1007/s10462-023-10651-9.
- [8] H. M. Jelodar, Y. Wang, C. Yuan, X. Feng, X. Jiang, Y. Li, and L. Zhao, "Latent Dirichlet allocation (LDA) and topic modeling: models, applications, a survey," *Multimedia Tools and Applications*, vol. 78, no. 11, pp. 15169–15211, 2019, doi: 10.1007/s11042-018-6894-4.
- [9] M. Taboada, J. Brooke, M. Tofiloski, K. Voll, and M. Stede, "Lexicon-Based Methods for Sentiment Analysis," *Computational Linguistics*, vol. 37, no. 2, pp. 267–307, 2011. [Online]. Available: <https://www.proquest.com/docview/896181231/>
- [10] K. Du, F. Xing, R. Mao, and E. Cambria, "Financial Sentiment Analysis: Techniques and Applications," *ACM Computing Surveys*, vol. 56, no. 9, pp. 1–41, 2024, doi: 10.1145/3649451.
- [11] N. Wang, H. Wang, Y. Jia, and Y. Yin, "Explainable recommendation via multi-task learning in opinionated text data," in *Proceedings of the 41st International ACM SIGIR Conference on Research and Development in Information Retrieval*, pp. 165–174, 2018, doi: 10.1145/3209978.3210010.
- [12] A. Gandhar, S. Gandhar, S. B. Kumar, A. Rehalia, P. Priyadarshi, and M. Tiwari, "A Comparative Analysis of Machine Learning Algorithms for Sentiment Analysis in Indian Social Media," *Library Progress International*, vol. 44, no. 3, pp. 8139–8145, 2024.
- [13] D. H. Wahid and A. SN, "Peringkasan Sentimen Esktraktif di Twitter Menggunakan Hybrid TF-IDF dan Cosine Similarity," *IJCCS (Indonesian Journal of Computing and Cybernetics Systems)*, vol. 10, no. 2, pp. 207–218, 2016, doi: 10.22146/ijccs.16625.
- [14] I. A. Fahrezi, R. Rudiman, and N. A. Verdikha, "Analisis Sentimen Twitter atas Isu Hak Angket Menggunakan Pembobotan TF-IDF dan Algoritma SVM," *Sci-Tech Journal*, vol. 3, no. 2, pp. 179–192, 2024, doi: 10.56709/stj.v3i2.526.
- [15] S. Khairunnisa and S. Al Faraby, "Pengaruh Text Preprocessing terhadap Analisis Sentimen Komentar Masyarakat pada Media Sosial Twitter (Studi Kasus Pandemi)," *Jurnal Media Informatika Budidarma*, vol. 5, no. 2, pp. 406–414, 2021, doi: 10.30865/mib.v5i2.2835.
- [16] M. Y. Rifai, A. Fadlil, and Sunardi, "Optimizing Text Preprocessing for Accurate Sentiment Analysis on E-Wallet Reviews," *Journal of Information and Communication Technology*, vol. 7, no. 2, pp. 42–50, 2023, doi: 10.21070/jicte.v7i2.1650.
- [17] D. A. K. Khotimah and R. Sarno, "Sentiment analysis of hotel aspect using probabilistic latent semantic analysis, word embedding and LSTM," *International Journal of Intelligent Engineering and Systems*, vol. 12, no. 4, pp. 275–290, 2019, doi: 10.22266/ijies2019.0831.26.
- [18] I. Zulfa and E. Winarko, "Sentimen Analisis Tweet Berbahasa Indonesia Dengan Deep Belief Network," *IJCCS (Indonesian Journal of Computing and Cybernetics Systems)*, vol. 11, no. 2, pp. 187–198, 2017, doi: 10.22146/ijccs.24716.
- [19] D. S. Utami and A. Erfina, "Analisis Sentimen Pinjaman Online di Twitter Menggunakan Algoritma Support Vector Machine (SVM)," in *SISMATIK (Seminar Nasional Sistem Informasi dan Manajemen Informatika)*, vol. 1, no. 1, pp. 299–305, 2021.
- [20] M. Iqbal, M. Afdal, and R. Novita, "Implementasi Algoritma Support Vector Machine Untuk Analisa Sentimen Data Ulasan Aplikasi Pinjaman Online di Google Play Store," *MALCOM Indonesian Journal of Machine Learning and Computer Science*, vol. 4, no. 4, pp. 1244–1252, 2024, doi: 10.57152/malcom.v4i4.1435.
- [21] M. A. Ardiansyah, M. Alamsyah, and M. F. Arif, "Analisis Sentimen Twitter Tentang Pinjaman Online di Indonesia Menggunakan Metode Random Forest," *Jurnal Teknologi Informasi dan Ilmu Komputer*, vol. 8, no. 3, pp. 567–574, 2021.
- [22] R. R. Santoso, R. Megasari, and Y. A. Hambali, "Implementasi Metode Machine Learning untuk Analisis Sentimen," *Jurnal Aplikasi dan Teori Ilmu Komputer*, vol. 3, no. 2, pp. 85–97, 2020. [Online]. Available: <https://ejournal.upi.edu/index.php/JATIKOM>
- [23] F. Z. Tala, "A Study of Stemming Effects on Information Retrieval in Bahasa Indonesia," University of Amsterdam, Netherlands, 2003.
- [24] A. Ahmad Irfa, Adiwijaya, and M. Syahrul Mubarak, "Klasifikasi Topik Berita Berbahasa Indonesia Menggunakan k-Nearest Neighbor," in *Proceedings of Engineering*, vol. 5, no. 2, pp. 3631–3640, 2018. [Online]. Available: <https://core.ac.uk/download/pdf/299923375.pdf>