

SENTIMENT ANALYSIS OF MENTAL HEALTH REVIEWS USING MACHINE LEARNING ALGORITHMS

Risa Wati¹, Siti Ernawati^{2*}

¹Sistem Infomasi
Universitas Bina Sarana Informatika
risawati.rwx@bsi.ac.id¹

²Sistem Informasi
Universitas Nusa Mandiri
siti.ste@nusamandiri.ac.id²

(*) Corresponding Author

Abstract

Mental health is a significant issue in the modern era due to lifestyle changes, social pressures, and technological advancements that introduce new challenges. These problems affect various aspects of life, including education, employment, social relationships, and overall quality of life. Technological development enables the use of machine learning to automatically classify large amounts of data. This study aims to analyze and compare the performance of Support Vector Machines (SVM), K-Nearest Neighbor (K-NN), Naïve Bayes (NB), Logistic Regression (LR), Decision Tree (DT), and Random Forest (RF) in sentiment classification on mental health issues, while simultaneously contributing to scientific development and supporting the understanding of public psychological conditions. The dataset used in this research was obtained from Kaggle and consists of 20,364 mental health-related reviews in .CSV format, processed using Google Colab with the Python programming language. The data were categorized into two groups—*depression* and *suicidewatch*—and then underwent preprocessing, data splitting into training and testing sets with an 80:20 ratio, and TF-IDF weighting. The results indicate that the SVM algorithm outperforms the other methods. Using an RBF kernel and a C parameter of 15, SVM achieved an accuracy of 72.09%, a precision of 72.11%, a recall of 72.09%, and an F1-score of 72.09%. This study not only provides scientific contributions but also supports efforts to better understand the psychological conditions experienced by society.

Keywords: Sentiment Analysis; Mental Health; Machine Learning

Abstrak

Kesehatan mental merupakan isu penting di era modern akibat perubahan gaya hidup, tekanan sosial, dan kemajuan teknologi yang menimbulkan tantangan baru. Permasalahan ini berdampak pada berbagai aspek kehidupan, termasuk pendidikan, pekerjaan, hubungan sosial, dan kualitas hidup. Perkembangan teknologi memungkinkan pemanfaatan machine learning untuk mengklasifikasikan data dalam jumlah besar secara otomatis. Penelitian ini bertujuan untuk menganalisis dan membandingkan performa metode Support Vector Machines (SVM), K-Nearest Neighbor (K-NN), Naïve Bayes (NB), Logistic Regression (LR), Decision Tree (DT), dan Random Forest (RF) dalam klasifikasi sentimen pada isu kesehatan mental, sekaligus memperluas pengembangan ilmu pengetahuan serta mendukung pemahaman kondisi psikologis masyarakat. Dataset pada penelitian ini diperoleh dari Kaggle dan terdiri atas 20.364 ulasan terkait kesehatan mental dalam format .CSV lalu diolah menggunakan Google Colab dengan bahasa Python. Data dikategorikan menjadi dua, yaitu *depression* dan *suicidewatch*, kemudian melalui tahap preprocessing, pembagian data train dan test dengan rasio 80:20, serta pembobotan TF-IDF. Hasil pengujian menunjukkan algoritma SVM memberikan performa terbaik dibandingkan metode lainnya. Dengan kernel RBF dan parameter C=15, SVM menghasilkan akurasi 72,09%, precision 72,11%, recall 72,09%, dan F1-score 72,09%. Penelitian ini tidak hanya memberikan kontribusi ilmiah, tetapi juga mendukung upaya memahami kondisi psikologis masyarakat.

Kata kunci: Analisis Sentimen; Kesehatan Mental; Machine Learning

INTRODUCTION

In the modern era, mental health has become an important issue due to lifestyle changes, social pressures, and technological advancements that bring new challenges. This problem affects many aspects of life, such as education, employment, social relationships, and overall quality of life (Ilham & Pramusinto, 2023). In 2024, WHO estimated that about 1 in 8 people worldwide, or around 970 million individuals, experience mental health disorders. Among them, depression and anxiety are the most common, affecting 5% and 4% of adults respectively (Huntington Psychological Service, 2024). Mental health is one of the global issues receiving increasing attention, especially in the digital era where people are more active in expressing their feelings and opinions through social media and other digital platforms. This text-based data holds great potential for analysis to better understand the psychological conditions of individuals as well as society at large. However, manual analysis of large-scale data is inefficient, thus requiring computational approaches capable of detecting patterns and sentiments more quickly and accurately.

In line with this, machine learning approaches have been widely used to automatically classify textual data. However, challenges remain because psychological expressions in text are often implicit. Popular algorithms such as Support Vector Machine (SVM), k-Nearest Neighbor (KNN), Naïve Bayes, Logistic Regression, Decision Tree, and Random Forest are frequently selected in text analysis studies due to their stable and effective performance (Wang et al., 2024), and these algorithms are particularly relevant because they have been proven capable of handling high-dimensional data and the complex characteristics of social media text (Mamani-Coaquira & Villanueva, 2024).

Previous studies have shown that SVM consistently delivers higher accuracy than other algorithms in sentiment analysis (Kanugrahan et al., 2024), and can handle complex problems through its kernel-based approach (Daffa et al., 2024). Using kernel functions, SVM can operate in higher-dimensional spaces, allowing data to be mapped into more complex feature spaces and enabling classes that are difficult to separate linearly to be more easily distinguished (Pramesti & Pratiwi, 2023). However, most existing studies have focused on domains such as transportation, courier service applications, and education. Mental health, despite being highly crucial and far-reaching in its impact, remains rarely examined comprehensively,

particularly in relation to comparative evaluations of various machine learning algorithms. In addition, previous research generally has not evaluated sentiment analysis within the context of mental-health-related digital reviews. In this study, the data object is clearly defined as user-generated reviews on mental health topics sourced from Kaggle.

This research gap forms an important foundation for the present study. This study aims to analyze and compare the performance of six machine learning algorithms in sentiment classification for mental-health-related data. The contributions of this study are twofold: (1) an academic contribution in the form of a comparative evaluation of algorithms within the mental health domain, and (2) a practical contribution by utilizing the selected model to assist healthcare professionals, counselors, and technology developers in understanding public sentiment patterns related to mental health issues, thereby supporting efforts to improve the quality of digital mental-health services.

RESEARCH METHODS

This study employs a quantitative approach with an experimental design, adapting the research workflow of text mining and sentiment analysis from similar studies by Ramayanti et al. (Ramayanti et al., 2025), and Nasir et al. (Nasir & Palanichamy, 2022). The original research models were then modified to align with the specific objectives of this study. The methodological stages include data collection, labeling, text preprocessing, dataset splitting, TF-IDF weighting, classification using six machine learning algorithms, and model evaluation through a confusion matrix.

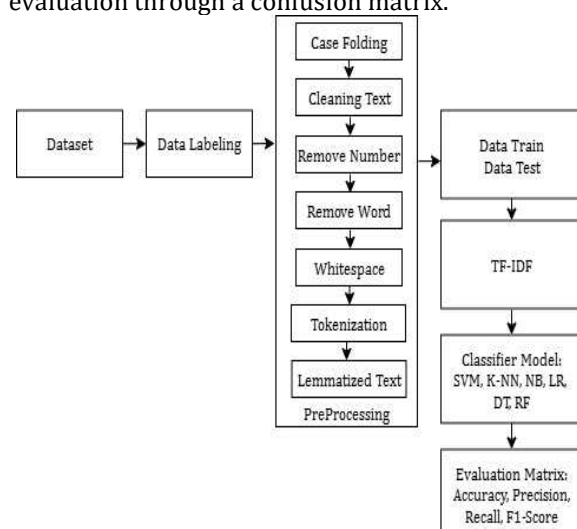


Figure 1 Research Model

The data used in this study are secondary data consisting of 20,364 mental health reviews, obtained from Kaggle at the following link: <https://www.kaggle.com/datasets/mritunjay1708/mental-health-review>. The dataset is stored in .CSV (Comma-Separated Values) format and processed using Google Colab. The data are labeled into two categories, namely *depression* and *suicidewatch*. Figure 2 illustrates the distribution of mental health review sentiment analysis data, with 10,371 reviews categorized as *depression* and 9,993 reviews categorized as *suicidewatch*.

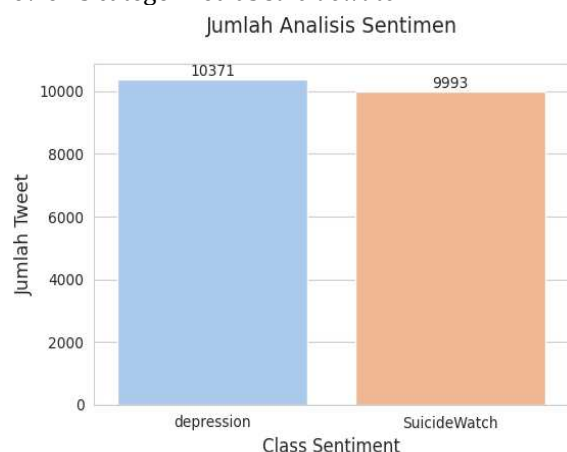


Figure 2 Number of Sentiment Analysis

After the data were collected, data preprocessing was carried out in the form of:

1. Remove Duplicate

Remove Duplicate is the process of eliminating duplicate data to ensure that no redundancy exists within the dataset. This step is carried out by removing reviews that contain identical values (Saputra & Noor Hasan, 2023).

2. Case Folding

Case Folding is a preprocessing step that converts uppercase letters into lowercase letters (Majid et al., 2025). Table 1 presents the results of data preprocessing using case folding.

Table 1. Result of Case Folding

Text	Label	Case Folding
I recently went through a breakup and she said she still wants to be friends so I said I can try doing that but when she talks to me about things it always hurts. I	depression	i recently went through a breakup and she said she still wants to be friends so i said i can try doing that but when she talks to me about things it always hurts. i

just want to lose feelings so all this pain can stop it hurts so much and I cannot even cry about it. I do not want to hurt her because she said she does not want to never speak to me again but I do not know what to do here. When we were together she always hurt me so I do not know why I still love her. I wish we never met it would be much less painful How do I lose feelings?

just want to lose feelings so all this pain can stop it hurts so much and i cannot even cry about it. i do not want to hurt her because she said she does not want to never speak to me again but i do not know what to do here. when we were together she always hurt me so i do not know why i still love her. i wish we never met it would be much less painful how do i lose feelings?

3. Cleaning Text

Cleaning Text aims to remove unnecessary components from the text, such as punctuation, numbers, URL links, emojis, and special characters, leaving only the relevant parts of the text for analysis (Ernawati et al., 2025). Table 2 presents the results of data preprocessing using text cleaning.

Table 2. Result of Cleaning Text

Cleaning Text
i recently went through a breakup and she said she still wants to be friends so i said i can try doing that but when she talks to me about things it always hurts i just want to lose feelings so all this pain can stop it hurts so much and i cannot even cry about it i do not want to hurt her because she said she does not want to never speak to me again but i do not know what to do here when we were together she always hurt me so i do not know why i still love her i wish we never met it would be much less painful how do i lose feelings

4. Remove Number

Remove Number is the process of eliminating numbers in the text that are not relevant to the data being analyzed. Table 3 presents the results of data preprocessing using the remove number step.

Table 3. Result of Remove Number

Remove Number
i recently went through a breakup and she said she still wants to be friends so i said i can try doing that but when she talks to me about things it always hurts i just want to lose feelings so all this pain can stop it hurts so much and i cannot even cry about it i do not want to hurt her because she said she does not want to never speak to me again but i do not know what to do here when we were together she always hurt me so i do not know why i still love her i wish we never met it would be much less painful how do i lose feelings

5. Remove Word

Remove Word is the process of eliminating words with short character lengths. For example, in the sentence *"I recently went through a breakup"*, after applying the remove word process, the words *"I"* and *"a"* are removed. Table 4 presents the results of data preprocessing using the remove word step.

Table 4. Result of Remove Word

Remove Word
recently went through breakup said still wants friends said doing that when talks about things always hurts just want lose feelings this pain stop hurts much cannot even about want hurt because said does want never speak again know what here when were together always hurt know still love wish never would much less painful lose feelings

6. Whitespace

Whitespace is the process of removing unnecessary blank characters or excessive spaces in the text, making it cleaner and more structured. Table 5 presents the results of data preprocessing using the whitespace step.

Table 5. Result of Whitespace

Whitespace
recently went through breakup said still wants friends said doing that when talks about things always hurts just want lose feelings this pain stop hurts much cannot even about want hurt because said does want never speak again know what here when were together always hurt know still love wish never would much less painful lose feelings

7. Tokenization

Tokenization is the process of splitting text into smaller units called tokens, which are generally words or sequences of words (Ernawati et al., 2025). Table 6 presents the results of data preprocessing using the tokenization step.

Table 6. Result of Tokenization

Tokenization
"['recently', 'went', 'through', 'breakup', 'said', 'still', 'wants', 'friends', 'said', 'doing', 'that', 'when', 'talks', 'about', 'things', 'always', 'hurts', 'just', 'want', 'lose', 'feelings', 'this', 'pain', 'stop', 'hurts', 'much', 'cannot', 'even', 'about', 'want', 'hurt', 'because', 'said', 'does', 'want', 'never', 'speak', 'again', 'know', 'what', 'here', 'when', 'were', 'together', 'always', 'hurt', 'know', 'still', 'love', 'wish', 'never', 'would', 'much', 'less', 'painful', 'lose', 'feelings']"

8. Lemmatized Text

Lemmatized text is the process of converting a word into its base form (Majid et al., 2025). Unlike stemming, which only trims word endings, lemmatized text considers context and utilizes a dictionary to ensure that the word is converted into its actual base form. Table 7 presents the results of data preprocessing using lemmatized text.

Table 7. Result of Lemmatized Text

Lemmatized Text
recently went through breakup said still want friend said doing that when talk about thing always hurt just want lose feeling this pain stop hurt much cannot even about want hurt because said doe want never speak again know what here when were together always hurt know still love wish never would much less painful lose feeling

Figure 3 illustrates the 10 most frequently used words in mental health reviews.

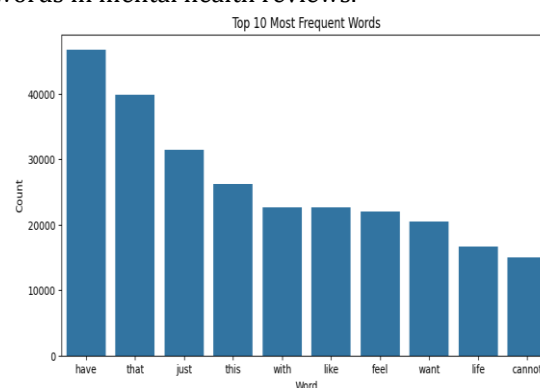


Figure 3 Top 10 Most Frequent Words

Word cloud is a form of visualization that displays a collection of words from a text. In this visualization, the font size is determined proportionally based on the frequency of occurrence, so that words appearing more frequently are displayed in larger font sizes (Kevin et al., 2024). Figure 4 is a word that frequently

appears in mental health reviews with the label depression.

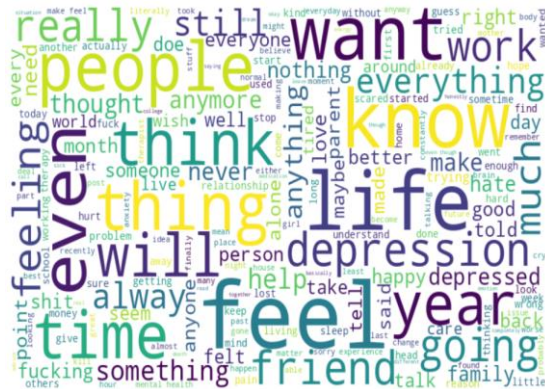


Figure 4 Word Cloud Depression

Figure 4 presents a word cloud visualization of the *depression* category, showing the words that most frequently appear in reviews related to mental health issues. The dominance of words such as “people”, “life”, “feel”, “think”, “want”, and “know” indicates that experiences of depression are closely tied to individuals’ internal struggles regarding the meaning of life and social relationships, while repeated terms like “feel”, “feeling”, and “feelings” emphasize that emotional expression is a central aspect in depression narratives. The appearance of words such as “depression”, “depressed”, and “help” illustrates that many users explicitly express their psychological condition and their need for support, followed by negatively toned words like “alone”, “never”, “anymore”, and “nothing”, which depict a high level of emotional distress. Temporal words such as “year”, “time”, and “month” indicate that depression is often perceived as a long-term condition. In addition, the presence of words like “friend”, “family”, “someone”, and “people” suggests that social relationships play an important role, both as sources of support and as potential triggers of psychological pressure.

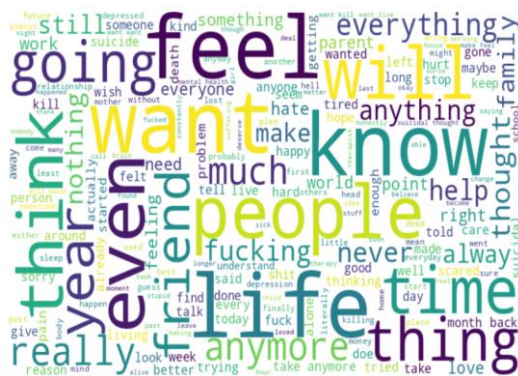


Figure 5 Word Cloud SuicideWatch

Figure 5 presents the word cloud for the *SuicideWatch* category, illustrating high-frequency terms that appear in reviews associated with extreme psychological states. The predominance of words such as “feel,” “want,” “know,” “people,” “life,” and “time” indicates that users frequently articulate intense internal emotional struggles, particularly concerning their desires, thoughts, and life circumstances. The presence of critical terms including “kill,” “suicidal,” “death,” “gone,” and “anything anymore” reflects a heightened level of hopelessness compared to the *depression* category, suggesting expressions indicative of elevated psychological risk. Words such as “alone,” “never,” “tired,” “hurt,” and “nothing” further reinforce sentiments of despair and helplessness commonly reported among individuals under suicide-risk monitoring. Additionally, the appearance of relational terms such as “friend,” “family,” and “parent” suggests that social relationships remain a significant factor, functioning both as potential sources of psychological strain and as emotional support structures. The high frequency of temporal words such as “year,” “time,” and “month” also indicates that these psychological conditions are perceived as persistent and long-standing rather than episodic.

Overall, this visualization underscores the complexity of emotional experiences and the intensity of psychological distress expressed by users, while simultaneously highlighting the urgency of exploring linguistic patterns that signal high-risk states and validating the necessity of machine learning approaches to more accurately interpret critical psychological expressions.

The next stage is splitting the data into two parts, namely training data and testing data with a ratio of 80:20. From this division, 16,266 training data were obtained for variable X and 16,266 training data for variable Y, while 4,067 testing data were obtained for variable X and 4,067 testing data for variable Y. This division aims to allow the model to be trained using the majority of the data, and then tested with data that has never been used in the training process, so that the model's performance and generalization ability can be evaluated more objectively.

After the data were split into training and testing sets, weighting was carried out using TF-IDF. TF-IDF is a statistical measure used to indicate the contribution or importance of a word in a collection of documents (Majid et al., 2025). TF-IDF converts text into numerical vectors so that it can be processed using machine learning. TF-IDF produces values as illustrated in Figure 6.

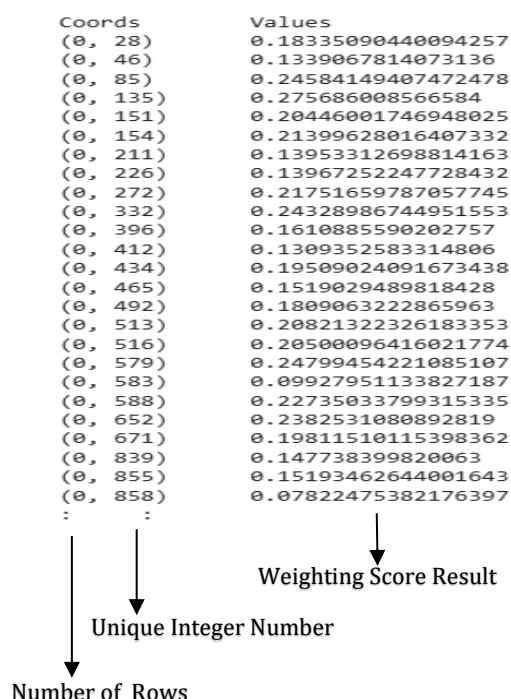


Figure 6 Hasil TF-IDF

Subsequent testing in this study was conducted using six machine learning algorithms, namely Support Vector Machines (SVM), K-Nearest Neighbor (K-NN), Naïve Bayes (NB), Logistic Regression (LR), Decision Tree (DT), dan Random Forest (RF).

1. Support Vector Machines (SVM)

Is recognized as an effective classification algorithm, particularly in addressing both linear and non-linear data separation problems (Ariyo Munandar et al., 2023). The strength of SVM lies in its ability to handle non-linear relationships among features through the application of kernel functions, such as polynomial kernels, Radial Basis Function (RBF), and sigmoid (Anggraini & Alita, 2024).

2. K-Nearest Neighbor (K-NN)

The K-NN algorithm is a supervised learning method used for both classification and regression tasks (Syahril Dwi Prasetyo et al., 2023). The K-Nearest Neighbor (K-NN) algorithm, in contrast, is relatively simple to implement. It classifies data based on the proximity of neighboring data points, with the number of neighbors determined by the user and represented by the parameter k (Isnain et al., 2021).

3. Naïve Bayes (NB)

Naïve Bayes (NB), originally introduced by Thomas Bayes, is a probabilistic algorithm that applies statistical principles grounded in Bayes' Theorem. It has been widely adopted in classification tasks due to its efficiency and

probabilistic interpretability (Ramadhani & Suryono, 2024). Naïve Bayes is also frequently applied in text classification, particularly sentiment analysis, as it can efficiently process large datasets and achieve fairly good accuracy in text analysis (Latifah et al., 2025).

4. Logistic Regression (LR)

Logistic Regression (LR) is a statistical method designed to model the relationship between a dependent variable and one or more independent variables. In machine learning, it is frequently applied to classification problems involving categorical target variables (Ramadhani & Suryono, 2024). Logistic Regression (LR) is employed in machine learning to handle classification problems with target variables consisting of different categories (Kartika Sari et al., 2024).

5. Decision Tree (DT)

Decision Tree is one of the most popular and effective methods for prediction and classification (Cahyaningtyas et al., 2021). A Decision Tree is a classification algorithm that operates by utilizing a tree-structured model (Wiratama Putra & Triayudi, 2022). The Decision Tree algorithm seeks to improve accuracy by pruning tree branches that represent noise in the data (Pramesti & Pratiwi, 2023).

6. Random Forest (RF)

Random Forest is a machine learning method that works by combining multiple decision trees. The final result is determined based on the majority vote among the collection of decision trees (Kanugrahan et al., 2024). Random Forest and Decision Tree are two widely used methods in sentiment analysis. Random Forest works by combining multiple decision trees, while Decision Tree is used to produce more accurate predictions (Huda et al., 2023).

RESULTS AND DISCUSSION

The testing results demonstrated that the model was evaluated using six machine learning algorithms, namely SVM, K-NN, Naïve Bayes (NB), Logistic Regression (LR), Decision Tree (DT), and Random Forest (RF). All algorithms were implemented using the Python programming language.

For the SVM algorithm, experiments were conducted with several types of kernels, including Linear, RBF, Polynomial, and Sigmoid. The results indicated that the RBF kernel with a parameter value of $C = 15$ achieved the best performance, with the highest accuracy of 72.09%. This finding suggests that the RBF kernel is more capable of capturing non-linear patterns within the dataset,

aligning with the data characteristics in this study. These results are consistent with the findings of Liu (Liu et al., 2011) and Ernawati (Ernawati & Wati, 2024), who also reported that the RBF kernel outperforms the Linear kernel when handling complex data distributions.

For the K-NN algorithm, testing was conducted with variations in the number of neighbors. The best result was obtained with $k = 5$, yielding an accuracy of 53.11%. This relatively low value compared to other algorithms may be attributed to K-NN's sensitivity to data distribution and outliers.

The Naïve Bayes (NB) algorithm achieved an accuracy of 70.94%, which proved relatively stable as the independence assumption among features remained sufficiently relevant to the dataset. This performance is competitive compared to Logistic Regression (71.53%) and Random Forest (71.40%).

Meanwhile, the Decision Tree (DT) algorithm recorded an accuracy of only 60.19%, the second-lowest after K-NN. The poor performance of DT may be attributed to overfitting, as the model tends to construct highly specific trees for the training data but fails to generalize adequately to the testing data.

The Logistic Regression (LR) algorithm achieved an accuracy of 71.53%, only slightly lower than SVM. This demonstrates that LR is capable of providing stable performance in binary classification tasks involving textual data. This finding corroborates the work of Ferdiansyah which emphasized that LR remains a strong baseline for comparison with other algorithms.

Meanwhile, Random Forest (RF), which is an extension of DT employing an ensemble approach, yielded better results with an accuracy of 71.40%. Although its performance was slightly lower than that of SVM and LR, RF offers advantages in terms of stability and its ability to handle heterogeneous data. This finding is consistent with the study of Barreñada et al. (2024), which emphasized that RF effectively mitigates the overfitting limitations of a single DT by aggregating predictions from multiple decision trees.

Overall, the evaluation results of the six algorithms are presented in Table 8. SVM demonstrated the best performance with an accuracy of 72.09%, followed by LR (71.53%) and RF (71.40%). These results highlight that more complex algorithms do not necessarily guarantee superior performance; rather, the choice of parameters and the alignment of the algorithm with the data characteristics play a critical role.

Table 8. Research Results

Algoritma	Evaluation Metrics (%)			
	Accuracy	Precision	Recall	F1-Score
SVM	72.09	72.11	72.09	72.09
KNN	53.11	53.11	53.11	52.89
NB	70.94	70.94	70.94	70.93
LR	71.53	71.53	71.53	71.53
DT	60.19	60.20	60.19	60.19
RF	71.40	71.41	71.40	71.40

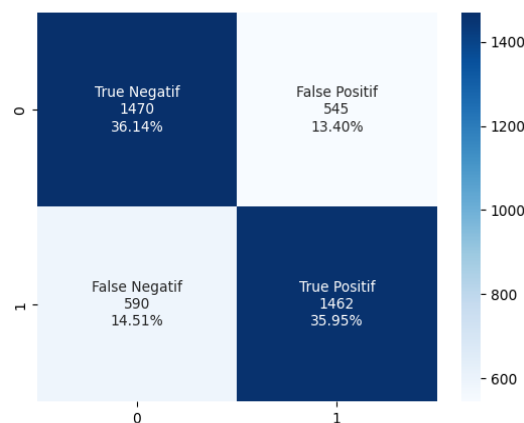


Figure 7 Heatmap Metode SVM

Figure 7 presents the heatmap of the SVM classification results. A total of 36.14% of negative data were correctly predicted as negative, and 35.95% of positive data were correctly predicted as positive, while 13.40% of negative data were misclassified as positive and 14.51% of positive data were misclassified as negative. Despite these errors, the prediction distribution demonstrates a relatively good balance between positive and negative classes.

The AUC Score of the SVM algorithm, shown in Figure 8, reached 78.34%, indicating that the model achieved a fairly good performance.

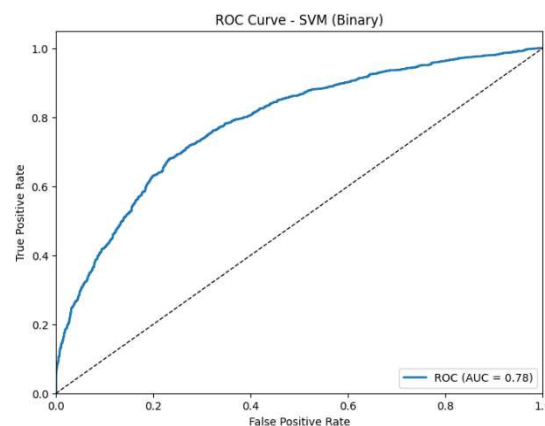


Figure 8 ROC Metode SVM

CONCLUSIONS AND SUGGESTIONS

Conclusion

This study aims to analyze and compare the performance of Support Vector Machines (SVM), K-Nearest Neighbor (K-NN), Naïve Bayes (NB), Logistic Regression (LR), Decision Tree (DT), and Random Forest (RF) methods in sentiment classification on mental health issues. The results show that the SVM algorithm provides the best performance compared to the other algorithms. Using an RBF (Radial Basis Function) kernel and a C parameter of 15, SVM achieves an accuracy of 72.09%, a precision of 72.11%, a recall of 72.09%, and an F1-score of 72.09%. This finding indicates that SVM with this configuration outperforms the other methods in classifying mental health-related review data. In the context of mental health, achieving a high accuracy score is crucial, as misclassification may lead to incorrect interpretations of an individual's emotional state. For example, a negative review misclassified as a positive one may obscure potential signals of stress, anxiety, or depressive symptoms. This aligns with previous studies suggesting that sentiment analysis can serve as an early indicator for identifying emotional patterns and psychological conditions within a population (Chancellor & De Choudhury, 2020). This study demonstrates that machine learning approaches can be utilized to understand public sentiment patterns related to mental health issues. The findings contribute to the development of sentiment analysis research, particularly regarding the application of classification algorithms in the context of mental health. Beyond advancing scientific knowledge, the results may also support practical efforts to better understand the psychological conditions experienced by society.

Suggestion

In future research, hyperparameter tuning using GridSearchCV can be applied to the proposed model to improve accuracy. Furthermore, employing a larger and more diverse dataset from various social media platforms may be considered to enhance the model's generalizability. Deep learning-based NLP methods such as BERT or RoBERTa can also be tested and compared with the traditional algorithms used in this study. From a practical perspective, subsequent research may integrate the model into web- or mobile-based applications to support counseling services or mental health screening more effectively.

REFERENCES

- Anggraini, J., & Alita, D. (2024). Implementasi Metode SVM Pada Sentimen Analisis Terhadap Pemilihan Presiden (Pilpres) 2024 Di Twitter. *Jurnal Informatika: Jurnal Pengembangan IT*, 9(2), 102–111. <https://doi.org/10.30591/jpit.v9i2.6560>
- Ariyo Munandar, A., Farikhin, & Edi Widodo, C. (2023). *Sentimen Analisis Aplikasi Belajar Online Menggunakan Klasifikasi SVM*. 7(1), 77–84.
- Cahyaningtyas, C., Nataliani, Y., & Widiyarsari, I. R. (2021). Analisis sentimen pada rating aplikasi Shopee menggunakan metode Decision Tree berbasis SMOTE. *AITI: Jurnal Teknologi Informasi*, 18(Agustus), 173–184.
- Chancellor, S., & De Choudhury, M. (2020). Methods in predictive techniques for mental health status on social media: a critical review. In *npj Digital Medicine* (Vol. 3, Issue 1). Nature Research. <https://doi.org/10.1038/s41746-020-0233-7>
- Daffa, M., Fahreza, A., Luthfiarta, A., Rafid, M., Indrawan, M., & Nugraha, A. (2024). Analisis Sentimen: Pengaruh Jam Kerja Terhadap Kesehatan Mental Generasi Z. *Journal Of Applied Computer Science And Technology (JACOST)*, 5(1), 2723–1453. <https://doi.org/10.52158/jacost.715>
- Ernawati, S., Frieyadie, & Yulia, E. R. (2025). HyVADSVM: Hybrid VADER-SVM and GridSearchCV Optimization for Enhancing Cyberbullying Detection. *Journal of Robotics and Control (JRC)*, 6(1), 101–113. <https://doi.org/10.18196/jrc.v6i1.24385>
- Ernawati, S., & Wati, R. (2024). Evaluasi Performa Kernel SVM dalam Analisis Sentimen Review Aplikasi ChatGPT Menggunakan Hyperparameter dan VADER Lexicon. In *Jurnal Buana Informatika* (Vol. 15, Issue 1).
- Huda, D. N. I., Prianto, C., & Awangga, R. M. (2023). Analisis Sentimen Perbandingan Layanan Jasa Pengiriman Kurir Pada Ulasan Play Store Menggunakan Metode Random Forest dan Decision Tree. *Jurnal Ilmiah Informatika (JIF)*, 11.
- Huntington Psychological Service. (2024, August 4). *The Latest Mental Health Statistics: What the Numbers Say About the State of Our Minds in 2024*. Huntington Psychological Service.
- Ilham, A., & Pramusinto, W. (2023). Analisis Sentimen Masyarakat Terhadap Kesehatan Mental Pada Twitter Menggunakan Algoritme K-Nearest Neighbor. *SENAFTI*, 2(2).

- Isnain, A. R., Supriyanto, J., & Kharisma, M. P. (2021). Implementation of K-Nearest Neighbor (K-NN) Algorithm For Public Sentiment Analysis of Online Learning. *IJCCS (Indonesian Journal of Computing and Cybernetics Systems)*, 15(2), 121. <https://doi.org/10.22146/ijccs.65176>
- Kanugrahan, G., Hafizh, V., Putra, C., & Ramdhani, Y. (2024). Analisis Sentimen Aplikasi Gojek Menggunakan SVM, Random Forest dan Decision Tree. *Jurnal Infortech*, 6(2), 171–178. <http://ejournal.bsi.ac.id/ejurnal/index.php/infortech>
- Kartika Sari, A., Akhmad Irsyad, Dinda Nur Aini, Islamiyah, & Stephanie Elfriede Ginting. (2024). Analisis Sentimen Twitter Menggunakan Machine Learning untuk Identifikasi Konten Negatif. *Adopsi Teknologi Dan Sistem Informasi (ATASI)*, 3(1), 64–73. <https://doi.org/10.30872/atasi.v3i1.1373>
- Kevin, K., Enjeli, M., & Wijaya, A. (2024). Analisis Sentimen Penggunaan Aplikasi Kinemaster Menggunakan Metode Naive Bayes. *Jurnal Ilmiah Computer Science*, 2(2), 89–98. <https://doi.org/10.58602/jics.v2i2.24>
- Latifah, U., Zuhriya, T. K., & Daniati, E. (2025). Analisis Sentimen Komentar Twitter (X) tentang Kesehatan Mental menggunakan Naive Bayes untuk Mengukur Opini Publik. *Seminar Nasional Inovasi Teknologi*, 9, 2549–7952.
- Liu, Q., Chen, C., Zhang, Y., & Hu, Z. (2011). Feature selection for support vector machines with RBF kernel. *Artificial Intelligence Review*, 36(2), 99–115. <https://doi.org/10.1007/s10462-011-9205-2>
- Majid, A. M., Imelda, K., & Nawangsih, I. (2025). Komparasi Algoritma Klasifikasi Machine Learning dan Penerapan Metode Ensemble Stacking untuk Menganalisis Sentimen Kesehatan Mental. *SKANIKA: Sistem Komputer Dan Teknik Informatika*, 8(2), 293–304. <https://www.kaggle.com/datasets/suchintikasarkar/sentiment-analysis-for-mental-health>
- Mamani-Coaquira, Y., & Villanueva, E. (2024). A Review on Text Sentiment Analysis With Machine Learning and Deep Learning Techniques. In *IEEE Access* (Vol. 12, pp. 193115–193130). Institute of Electrical and Electronics Engineers Inc. <https://doi.org/10.1109/ACCESS.2024.3513321>
- Nasir, A. A. B. M., & Palanichamy, N. (2022). Sentiment Analysis of Covid-19 Tweets by Supervised Machine Learning Models. *Journal of System and Management Sciences*, 12(6), 50–69. <https://doi.org/10.33168/JSMS.2022.0604>
- Pramesti, L. A., & Pratiwi, N. (2023). Analisis Sentimen Twitter Terhadap Program MBKM Menggunakan Decision Tree dan Support Vector Machine. *Journal of Information System Research (JOSH)*, 4(4), 1145–1154. <https://doi.org/10.47065/josh.v4i4.3807>
- Ramadhani, B., & Suryono, R. R. (2024). Komparasi Algoritma Naive Bayes dan Logistic Regression Untuk Analisis Sentimen Metaverse. *JURNAL MEDIA INFORMATIKA BUDIDARMA*, 8(2), 714. <https://doi.org/10.30865/mib.v8i2.7458>
- Ramayanti, I., Hermawan, L., Syakurah, R. A., Stiawan, D., Meilinda, Negara, E. S., Fahmi, M., Ghiffari, A., & Rizqie, M. Q. (2025). Sentiment and Emotion Classification Models Using Hybrid Textual and Numerical Features: A Case Study of Mental Health Counseling. *Journal of Applied Data Sciences*, 6(3), 1482–1494. <https://doi.org/10.47738/jads.v6i3.632>
- Saputra, A., & Noor Hasan, F. (2023). ANALISIS SENTIMEN TERHADAP APLIKASI COFFEE MEETS BAGEL DENGAN ALGORITMA NAIVE BAYES CLASSIFIER. *SIBATIK JOURNAL: Jurnal Ilmiah Bidang Sosial, Ekonomi, Budaya, Teknologi, Dan Pendidikan*, 2(2), 465–474. <https://doi.org/10.54443/sibatik.v2i2.579>
- Syahril Dwi Prasetyo, Shofa Shofiah Hilabi, & Fitri Nurapriani. (2023). Analisis Sentimen Relokasi Ibukota Nusantara Menggunakan Algoritma Naive Bayes dan KNN. *Jurnal KomtekInfo*, 1–7. <https://doi.org/10.35134/komtekinfo.v10i1.330>
- Wang, J., Wei, J., Tian, F., & Wei, Y. (2024). A comparative study of machine learning models for sentiment analysis of transboundary rivers news media articles. *Soft Computing*, 28(23), 13331–13347. <https://doi.org/10.1007/s00500-024-10357-2>
- Wiratama Putra, T., & Triayudi, A. (2022). Analisis Sentimen Pembelajaran Daring menggunakan Metode Naive Bayes, KNN, dan Decision Tree. *Jurnal Teknologi Informasi Dan Komunikasi*, 6(1), 2022. <https://doi.org/10.35870/jti>