



An Improved Hybrid Machine Learning-Based Approach for Depression Detection on Social Media Posts

Wadzani Aduwamai Gadzama^{1*} Mohammed Ali Kawo² & Lucy Bulus Dalhatu³
^{1&3}Department of Computer Science, Federal Polytechnic, Mubi, Adamawa State, Nigeria

²Department of Computer Science, Federal University Gusau, Zamfara State, Nigeria

Corresponding Author: ask4gadzama@gmail.com

Abstract

Depression is one of the most common diseases these days due to many economic and financial problems that befell individuals and families, especially in Nigeria. It is the major reason for anxiety and sleeping disorders. It can have a significant impact on a person's daily functioning and their ability to participate in their work, school, meals, sleep and play. Depression, if left untreated or not given early attention, may sometimes lead to self-harm and suicide. The study evaluates the effectiveness of SVM, RF, DT, K-NN, and NB machine learning models. It proposes a hybrid SVM-RF model to detect likely depressed posts on Facebook and X (Twitter) using users' tweets and comments. Datasets were collected via Facebook using the Facebook_scraper library and the Sentiment140 dataset on the Kaggle website. Datasets were split into 80:20 for model training and testing. Python 3.7 and Google Colab were used to analyse the machine learning algorithms. This study showed that SVM outperformed other techniques, such as RF, DT, K-NN, and NB, achieving 91% accuracy in identifying depressive content. The outcome also

shows that the hybrid SVM-RF model achieves 93% better performance in capturing and understanding more complex patterns of depression-related language, as measured across several metrics. The outcome of the study will help psychologists, policymakers, and other concerned stakeholders in society, as well as the Nigerian government, to identify vulnerable individuals who are at risk of experiencing depression among social media users.

Keywords: Depression, Detection, Machine Learning, Social Media, and Algorithm

Introduction

Depression is among the most important public health problems of the 21st century and is a major contributor to the global burden of disease (Shen et al., 2017). A recent study by the World Health Organisation states that 1 in 8 people worldwide has a mental illness, and more than 300 million people worldwide are affected by depression each year (Gadzama et al., 2024). In addition to its significant effect on quality of life and productive functioning within society, the risk of suicide is a major concern, as almost 80% of people who attempt or die by suicide have a known history of mental illness (Alsagri & Ykhlef, 2020a). Despite its high prevalence, there remains a large treatment gap of over 70% of people in the early stages of depression never access professional psychological services, and often the depression gets worse (Shen et al., 2017). Depression can be detected using various criteria and scales, including the widely employed Diagnostic and Statistical Manual of Mental Disorders (DSM), which details specific signs of depression (Alkasem et al., 2026).

Recent years have seen the emergence of new social media platforms that enable the large-scale understanding and monitoring of mental health. Communities like Twitter (now X) and Facebook have emerged as a common forum for people to share their thoughts, emotions, and personal experiences, creating a wealth of user-generated content that can be used to capture both

everyday life and mental state (Shen et al., 2017; Coppersmith et al., 2014). This digital footprint provides a wealth of non-invasive information for researchers and clinicians to work with, in addition to the information gathered from traditional clinical assessments (Mowery et al., 2017). Analysed data using advanced computational methods is especially timely, because the negative emotions spread of one person to another in social networks can aggravate mental health problems (Alsagri & Ykhlef, 2020). Accordingly, there has been an increasing interest in using social media to detect, estimate and monitor the occurrence of diseases, as well as to find a potential window for early intervention to improve mental health outcomes (Mowery et al., 2017).

Machine Learning (ML) is playing a crucial role in this journey, with the potential to detect depressive symptoms across a range of data types, including text, audio, and visual sources. The main benefits of machine learning techniques are their remarkable algorithmic scalability and modelling flexibility (Wei et al., 2025). In particular, support vector machines (SVM), random forest (RF), decision trees (DT), K-nearest neighbours (K-NN) and Naïve Bayes (NB) are the machine learning models commonly used to detect linguistic cues for depression in text data from social media posts, medical records, and patient interviews (Vasha et al., 2023; Kour & Gupta, 2022). These models generally process word usage, emotional tone and topic modeling to distinguish depressive and non-depressive text. For example, these methods have been effective for identifying depressed language on social media platforms such as Twitter and Reddit and for analysing electronic health records (Kumbhar et al., 2021; Jayanthi et al., 2022). Yet, although these contributions, it is still a great challenge to develop accurate and reliable models with early detection.

In Nigeria, for instance, the problem is more pronounced in areas where socio-economic factors such as economic hardship, unemployment, marital difficulties, abuse, use of drugs and alcohol, internal conflict, and prolonged illness contribute to the development of depression in individuals

(Vasha et al., 2023). While these various factors are major contributors to the current mental health crisis, the current research, which already has improved the accuracy of the models, is largely lacking in solutions to prevent depression or allow for timely intervention using Twitter (X) and Facebook datasets. The need for early detection of depression is critical for timely intervention (Fisher et al., 2026). Although ML algorithms can be used to analyse large data sets and find patterns, concerns have been raised about the accuracy, reliability and interpretability of the algorithms when applied to the unique, non-structured character of social media posts.

This study aims to fill these gaps by proposing an effective hybrid machine learning algorithm to detect depressive messages in English language tweets and Facebook posts at an early stage using Twitter (X) and Facebook datasets. This study aims to assess the performance of the current machine learning algorithms and introduce an enhanced machine learning-based algorithm for depression detection from social media. Specifically, the goals are to compare the effectiveness of the current SVM, RF, DT, K-NN, and NB methods for detecting depressed posts in a Twitter and Facebook dataset, and to develop and analyse a hybrid model to determine its effectiveness in detecting depressed individuals on Twitter and Facebook.

This work is based on the available literature. For instance, Kumbhar et al. (2021) reported that Decision Trees performed at a high level of accuracy of 98.55% on Twitter data, and Adarsh et al. (2023) designed an ensemble model using SVM and KNN with an accuracy of 98.05% for the classification of depression from suicidal ideation using Twitter data. Some other studies have used more than one classifier, in which Naïve Bayes and NBTree have similar accuracy (97.31%) in certain circumstances (Govindasamy & Palanichamy, 2021), while Random Forest has better accuracy than others in other cases (Angskun et al., 2022). The hybrid methods have also been investigated; for example, Skaik and Inkpen (2020) used SVM, LR, and XGBoost separately, but still achieved F1-scores as high as 0.961. For non-English data, Vasha et al. (2023) used SVM on Bangla Facebook comments,

achieving an accuracy of 77%. The work has incorporated deep learning techniques, with Gupta et al. (2022) showing that LSTM models are superior to baseline classifiers for detecting depression, even with imbalanced data. However, these developments have been accompanied by numerous studies that highlight problems of overfitting, feature selection, and the need for cross-platform generalisation (Alsagri & Ykhlef, 2020; Chatterjee et al., 2021). Chiong et al. (2021) showed that models developed using Twitter data can be transferred to other platforms, such as Facebook and Reddit, demonstrating cross-platform generalizability. These studies demonstrate the promise of ML for depression detection and show that improved hybrid models that can predict depression at an early stage and in a wide range of social media platforms are needed.

Methodology

This section provides an in-depth discussion of the methodology used to detect depression in social media posts in this study. The methodology includes a research design, data collection methods, data preprocessing and feature engineering techniques, baseline SVM, RF, DT, K-NN, and NB machine learning models, the proposed SVM-RF hybrid machine learning architecture, and evaluation metrics used to assess the model's performance. The ultimate objective is to provide a repeatable, clear way to build a reliable, precise depression-detection system from textual data on Twitter and Facebook.

Research Design

The study is a quantitative, experimental research design with classification using supervised machine learning. The research is organised to see how well traditional SVM, RF, DT, K-NN, and NB machine learning algorithms perform versus the proposed SVM-RF hybrid model in terms of classifying depressive from non-depressive social media posts. The design includes a structured approach to data acquisition, preprocessing, feature extraction, model training, hyperparameter tuning, cross-validation, and

model evaluation. The outcome variable or dependent variable is a binary classification of posts (depressive vs. non-depressive), and the independent variables are the textual features extracted from the posts.

Data Collection and Description

For this study, two popular social media platforms, Twitter (currently X) and Facebook, were used to collect a dataset of textual posts. The datasets were combined and preprocessed. These platforms are chosen because they have large user populations and provide a great deal of self-disclosed user information, which is rich in emotions. Data collection is done through scraping, API-based data extraction and publicly available depression-related corpora, in a mix.

Twitter Data Collection

Tweets are collected using the Twitter API targeting English-language tweets. A query-based approach is used to detect potentially depressive content, which includes depression-related keywords and phrases based on previous studies (e.g., depressed, hopeless, suicidal, worthless, sad, lonely) and phrases commonly used to express negative affect. The dataset has been balanced by including both depressive tweets (tweets that include the keywords and were manually annotated) and randomly sampled non-depressive tweets (posts containing no depression-related term). Posts from varied geographical areas and dates are included to achieve data representativeness.

Facebook Data Collection

Data are collected from public pages and groups discussing mental illnesses, as well as from publicly accessible depression datasets. Only public posts are shown because of privacy restrictions. The collection process uses the Facebook Graph API and follows ethical best practices for collecting user information. Like Twitter posts, the Facebook posts are classified as depressive or non-depressive based on the presence of depression related keywords and then manually annotated by trained annotators.

Data Composition

The combined dataset comprises 3,980 posts, with half belonging to the depressive class and the other half to the non-depressive class, thereby avoiding the class imbalance that could affect the model's performance. The data is then split into 80% for training and 20% for validation to tune the hyperparameters, and the 10-fold cross-validation method is used to make the model generalisable.

Data Preprocessing

A detailed pipeline for preprocessing raw textual data is developed to prepare it for machine learning analysis. This is an important step to eliminate noise and to enhance the input features. The following are the steps taken for preprocessing.

- i. **Text Normalisation**

All text is converted to lowercase to make it uniform and reduce feature dimensionality.

- ii. **Noise Removal**

Special characters, numbers, punctuation, hyperlinks, emojis, and non-ASCII characters are stripped out because they are not particularly helpful for depression detection.

- iii. **Tokeniser**

The cleaned text is split into Natural Language Processing (NLP) tokens (words) using the Natural Language Toolkit (NLTK) tokeniser.

- iv. **Stop Words Removal**

Common stop words, such as "the," "is," "am," and "are," that are not informative for the depression meaning are removed with the predefined English stop word list in NLTK.

- v. **Stemming and Lemmatisation**

Stemming (Porter Stemmer) and Lemmatization (WordNet Lemmatizer) were used to reduce word inflections to their root forms. Lemmatization is more linguistically correct, because it converts words to their dictionary base forms (e.g., "running" to "run").

vi. **Abbreviations and slang**

Common abbreviations and slang terms used on social media are expanded to full words to help readers understand the content (e.g., "u" becomes "you" and "lol" becomes "laugh out loud").

Feature Extraction

Feature extraction converts preprocessed text data into numerical features suitable for machine learning algorithms. Two feature extraction methods are used: TF-IDF (Term Frequency-Inverse Document Frequency) for traditional machine learning models, and word embeddings for the deep learning models.

TF-IDF Vectorization

For baseline ML models (SVM, RF, DT, K-NN, NB), TF-IDF was used to convert the text into a vector space model. TF-IDF assigns weights to terms based on how many times they appear in a document relative to the rest of the collection (Chiong et al., 2021), helping identify which terms are more discriminative for classification. The TF-IDF matrix is formed from unigrams and bigrams to capture single-word and two-word-phrase context.

Word Embeddings

Pre-trained Word embeddings are used in the proposed hybrid deep learning model to represent the semantic and syntactic relationships among the words. In particular, we use the 300-dimensional GloVe Word embeddings, trained on 2 billion tweets (Pennington et al., 2014). These embeddings are dense vectors associated with each Word, in which words with similar meanings share similar vectors, and they represent an improved representation over TF-IDF.

Model Development

In this study, the SVM, RF, DT, K-NN, and NB machine learning models, as well as a proposed hybrid SVM-RF machine learning model, are tested. The development process consists of training, tuning and testing each model to find the best approach in detecting depression. The following 5 machine learning models are implemented as baseline models and compared with the proposed hybrid model. In selecting these models, the ones widely used in previous work on depression detection are chosen.

Support Vector Machine (SVM): A linear SVM classifier is implemented with a linear kernel and optimised for high-dimensional sparse data, such as TF-IDF vectors. A grid search optimises the regularisation parameter (C). It is a supervised machine learning classifier (Rehmani et al., 2024).

Random Forest (RF): An ensemble classifier consisting of multiple decision trees is implemented. Random Forest improves predictive performance by constructing multiple decision trees across different data and feature subsets (Luo et al., 2025; Mortier et al., 2017). The model is trained with many trees (n_estimators) and a maximum depth, both optimised to avoid overfitting.

Decision Tree (DT): A single decision tree classifier is used, with the parameter controlling the model's complexity set to a value. It is used in this study due to its interpretability and ease of use (Song et al., 2025).

K-Nearest Neighbors (K-NN): A distance-based classifier that assigns posts to the majority class of their k nearest neighbors in the feature space is implemented. K is a parameter optimised via cross-validation. KNN works well for examining proximity-based relationships in the data (Norouzi et al., 2024).

Naïve Bayes (NB): It is a straightforward algorithm that relies on Bayes' theorem, assuming that the attributes are conditionally independent given the class (Yadav & Bhutia, 2025). It is a simple yet effective technique for predicting depression. The multinomial Naïve Bayes is ideal for discrete features, such as Word count and TF-IDF values.

Proposed Hybrid SVM-RF Model

Integrating SVM with RF machine learning frameworks for depression identification yields a robust hybrid SVM-RF model that capitalises on the strengths of both systems. The SVM classifier learns an ideal separating hyperplane that distinguishes between depressive and non-depressive text in a high-dimensional feature space. Simultaneously, the Random Forest classifier constructs numerous decision trees and aggregates their predictions via majority voting, thereby enhancing classification resilience and mitigating overfitting. The results from the SVM and RF classifiers are then integrated in a decision-level fusion module using a weighted voting or stacking approach. This hybridisation leverages the synergistic advantages of both algorithms, allowing the framework to generate more precise and dependable predictions than either classifier functioning alone. The unified classifier ultimately classifies each social media post into either Depressed or Non-depressed categories. The hybrid SVM-RF model, by combining SVM and RF, can effectively capture both the contextual significance and temporal dynamics of depressive symptoms in textual data, thereby improving the precision of depression detection in user-generated content on Facebook and Twitter. Figure 3 below illustrates the suggested hybrid SVM-RF methodology for detecting depression.

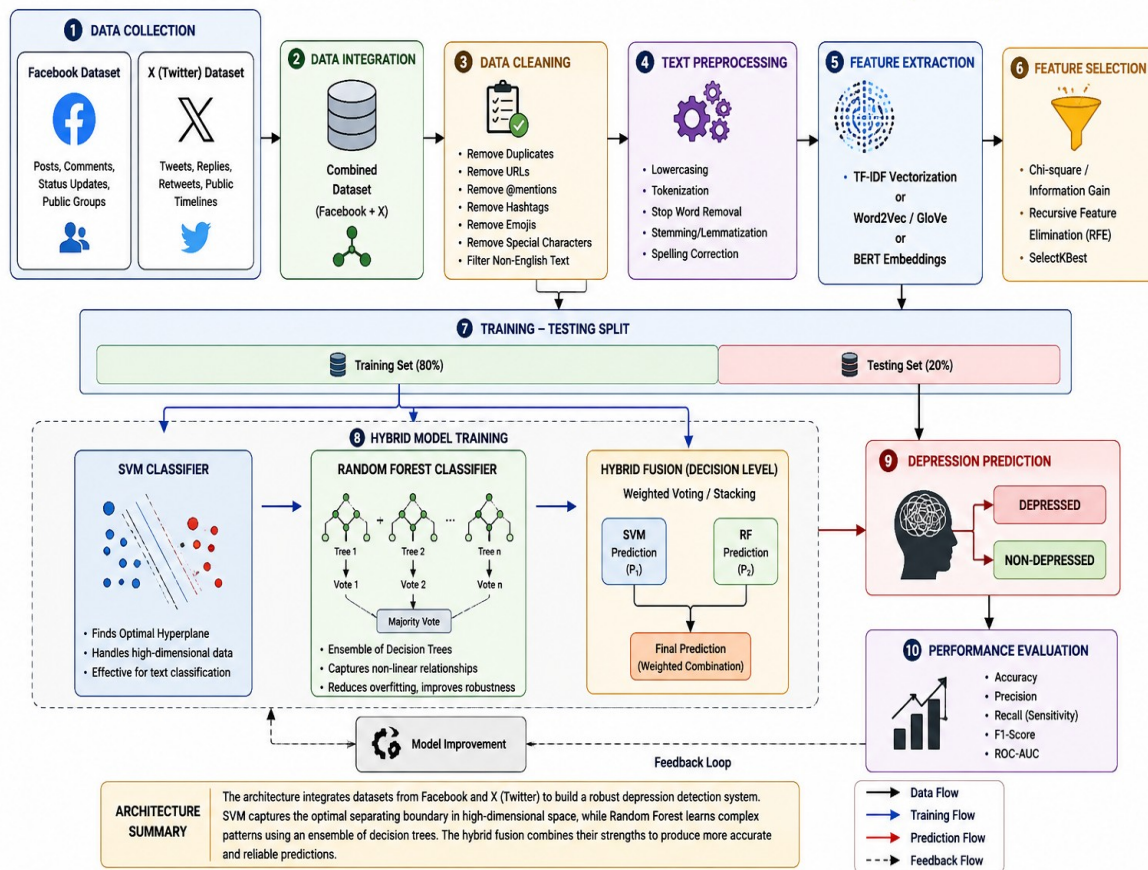


Figure 1: Proposed Hybrid SVM - RF Architecture for Depression Detection

Figure 1 illustrates the proposed hybrid architecture of Support Vector Machine SVM and RF for the identification of depression utilising integrated textual datasets from Facebook and X (Twitter). The framework begins by gathering and consolidating data from various social media platforms, then performs data cleansing and text preprocessing to eliminate extraneous elements and standardise the textual material. The transformed data is converted to numerical formats using feature extraction methods such as TF-IDF, which subsequently enable the selection of the most pertinent features to enhance classification efficacy. The refined dataset is subsequently partitioned into training and test subsets, and the SVM and Random Forest classifiers are trained concurrently. Their prediction are integrated using a decision-level fusion approach, such as weighted voting or stacking, to

leverage the strengths of both classifiers and improve predictive accuracy. Ultimately, the hybrid model categorises each post as either depressed or non-depressed, and its efficacy is assessed using conventional metrics such as accuracy, precision, recall, and F1-score. This architecture provides a robust, effective framework for identifying depression in social media data by integrating the complementary strengths of SVM and Random Forest classifiers.

Model Training and Validation

The Training Procedure: The baseline models are trained on the training data, and the proposed hybrid model is trained on the training data. The feature vectors for traditional ML models are the TF-IDF vectors. In the hybrid deep learning model, preprocessed sequences of fixed length are fed into the embedding layer.

Cross-validation: To obtain robust, generalisable performance, 10-fold stratified cross-validation is performed on the training set. It partitions the data into 10 equal sections, using 9 for training and 1 for validation. This is repeated 10 times, and the average performance is presented. Stratification ensures that each fold has the same distribution of data classes as the entire dataset.

Hyperparameter Tuning: Grid search and random search techniques are used to find the best hyperparameter settings in each model. In the hybrid approach, hyperparameters such as the number of filters, kernel size, LSTM units, dropout rate, learning rate, etc., are manually tuned and optimised using automated search techniques (e.g., Keras Tuner).

Evaluation Metrics

All models are tested using a wide range of metrics to provide an overall sense of each model's classification performance. In real-world scenarios, metrics other than accuracy are emphasised because of the imbalanced nature of such datasets to ensure reliable evaluation. These metrics were taken into account, including all four significant dimensions, i.e., true

positives (TPs), false positives (FPs), false negatives (FNs), and true negatives (TNs), as shown below.

Accuracy: This measure counts the number of depression cases correctly identified, meaning it considers only correctly classified cases. It is considered the percentage of correctly predicted values by the model (Alsagri & Ykhlef, 2020). The total number of instances the model forecasts, divided by the number of instances it correctly predicts as depressed, is the accuracy metric, which can vary by method, as illustrated in Equation (1).

$$Accuracy = \frac{TP + TN}{TP + FP + FN + TN}$$

(1)

F1-Scores: The F1-score considers precision and recall; it is the harmonic mean of precision and recall (Nadeem et al., 2016). The F1-score, also known as F1-Measure, is a better measure of depression than accuracy, as it accounts for cases that were not correctly classified. The F-measure evaluates the harmonisation of two factors based on precision and recall, since they measure different features, as shown in Equation (2).

$$F1\text{-Score} = 2 * \frac{Precision * Recall}{Precision + Recall}$$

(2)

Precision: The proportion of correctly identified positive depression cases among the total number of predicted positive depression cases (Alsagri & Ykhlef, 2020; Nadeem et al., 2016). Generally speaking, for precision prediction, it is to be expected that it will be favourable. Precision is the proportion of people classified as depressed who are actually predicted to be depressed—this could be evaluated with the formula in Equation (3).

$$Precision = \frac{TP}{TP + FP}$$

(3)

Recall: It is the true positive cases that have been correctly identified from all positive cases in the dataset (Alsagri & Ykhlef, 2020). Recall measures how many positive depression cases the model correctly predicted out of all positive cases. The equation (4) for the recall percentage is shown below:

$$Recall = \frac{TP}{TP + FN}$$

(4)

	Actually Positive (1)	Actually Negative (0)
Predicted Positive (1)	True Positives (TPs)	False Positives (FPs)
Predicted Negative (0)	False Negatives (FNs)	True Negatives (TNs)

Figure 2: Confusion Matrix (Draelos, 2019)

Table 1: Detail of TPs, TNs, FPs, and FNs (Nadeem et al., 2022)

True Positives (TP)	Depression instances that are positive and determined as positive
True Negatives (TN)	Instances that are negative and are determined as negative/non-depressed
False Positives (FP)	Instances that are negative and are determined as positive/depressed
False Negatives (FN)	Depression instances that are positive but are determined as negative

Simulation Environment

The machine learning algorithms will be analysed using Python 3.7 and Google Colab. Python 3.7 is a Python version. It is a high-level programming language that most researchers use for their research. It is highly recommended for AI-based work, and it is very popular among new-generation programmers because it is easy to learn and understand. Google Colab is a free, open-source platform for Python. Google Colaboratory is a free, web-based Jupyter notebook environment used to train machine learning and deep learning models on central processing units (CPUs), graphics processing units (GPUs), and tensor processing units (TPUs).

Results and Discussion

This section delineates the experimental setups and evaluation criteria employed in the studies. This study used Twitter and Facebook datasets to evaluate the performance of SVM, RF, DT, K-NN, and NB machine learning algorithms. Additionally, it showcases and analyses the findings from experiments performed using the developed hybrid SVM-RF model utilising the datasets. The implemented models were trained and tested. The research obtained data-acquisition metrics from the confusion matrix. The study employed a confusion matrix to consolidate a positioning model comprising many sub-metrics. The sub-metrics are precision, accuracy, F1 score and recall from the confusion matrix. Figure 3 below compares the performance of the proposed hybrid model with other models.

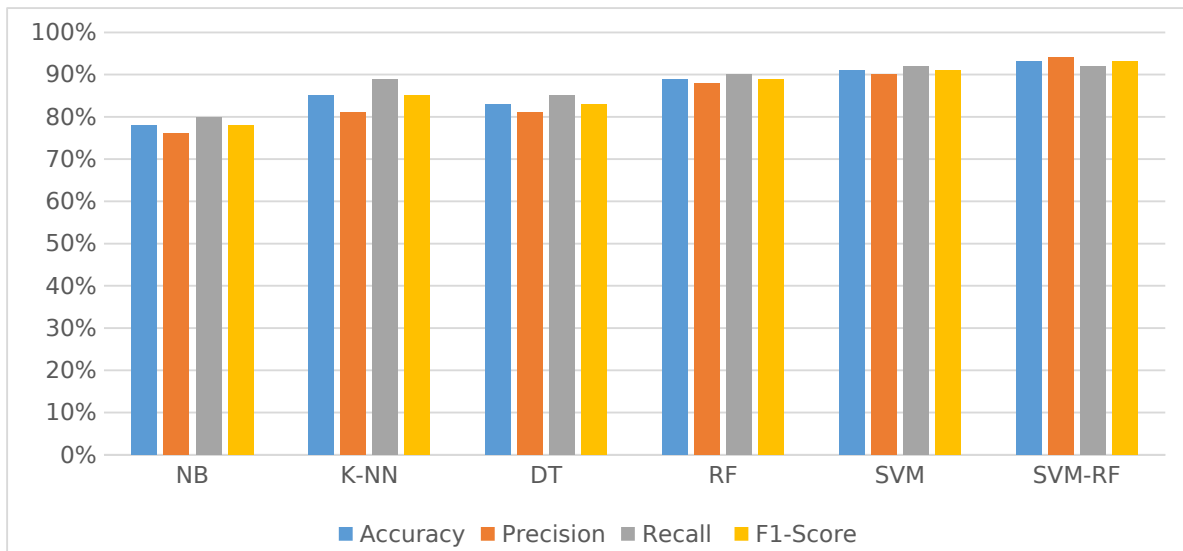


Figure 3: Comparison of the Proposed Hybrid SVM-RF Model and the Existing Models

Figure 3 presents a comparative analysis of the SVM, RF, DT, K-NN, NB, and SVM-RF models for depression detection, evaluating accuracy, precision, recall, and F1-score. This evaluation uses datasets from posts on Facebook and X (Twitter). Table 2 below provides a comparative analysis of the SVM, RF, DT, K-NN, NB, and the hybrid SVM-RF model performance.

Table 2: Classification Report for the Models

Model	Accuracy	Precision	Recall	F1-Score
NB	78%	76%	80%	78%
K-NN	85%	81%	89%	85%
DT	83%	81%	85%	83%
RF	89%	88%	90%	89%
SVM	91%	90%	92%	91%
SVM-RF	93%	94%	92%	93%

Table 2 presents the classification performance of five standalone machine learning algorithms. The NB, K-NN, DT, RF, and SVM algorithms, along with the proposed hybrid SVM-RF model for depression detection, were evaluated using the combined Facebook and X (Twitter) datasets. The evaluation was conducted using four widely adopted performance metrics: accuracy, precision, recall, and F1-score.

With 78% accuracy, 76% precision, 80% recall, and 78% F1-score, NB was the least effective standalone classifier. NB's lower precision suggests a higher rate of false positives, despite its relatively good recall. This suggests that NB's simplifying assumption of feature independence limits its efficacy in modelling the complex linguistic features of depression-related social media posts.

The K-NN classifier enhanced classification efficacy, achieving 85% accuracy, 81% precision, 89% recall, and 85% F1-score. The comparatively elevated recall signifies that K-NN effectively recognised the majority of depressive posts; nonetheless, its diminished accuracy implies that it misclassified a significant quantity of non-depressive posts as depressive. This behavior is expected, as K-NN is fundamentally distance-based and thus less effective in high-dimensional textual feature spaces.

The DT model achieved 83% accuracy, 81% precision, 85% recall, and an F1-score of 83%. Despite DT offering comprehensible classification guidelines, its performance was lower than that of K-NN, RF, and SVM, suggesting that a single decision tree may fail to adequately capture the intricate associations in the combined Facebook and X datasets and is more prone to overfitting.

The RF classifier achieved significant improvements, with 89% accuracy, 88% precision, 90% recall, and 89% F1-score. The ensemble learning approach utilised by RF combines the forecasts of numerous decision trees, therefore diminishing variance and enhancing resilience. Its equitable performance across all assessment criteria demonstrates that RF successfully identified varied linguistic patterns linked to depressive expressions on both Facebook and X.

Among the individual models, the SVM demonstrated the best performance, achieving 91% accuracy, 90% precision, 92% recall, and 91% F1-score. These outcomes reveal SVM's exceptional capacity to process high-dimensional, sparse textual features commonly generated by text vectorisation approaches such as TF-IDF and Word embeddings. The

increased recall signifies that SVM successfully detected people with depressive inclinations while preserving a reasonably minimal false-positive rate.

The proposed hybrid SVM-RF model attained the highest overall efficacy, with 93% accuracy, 94% precision, 92% recall, and 93% F1-score. In contrast to the independent SVM model, the integrated method enhanced accuracy by 2 percentage points, precision by 4 percentage points, and F1-score by 2 percentage points, while preserving the same recall. These results indicate that integrating the discriminative power of SVM with the ensemble-learning advantages of Random Forests improves the model's proficiency at distinguishing between depressive and non-depressive social media content. The hybrid model's enhanced accuracy suggests reduced false-positive identifications, thereby increasing its reliability for real-world depression screening.

The classification results indicate that the proposed SVM-RF hybrid model outperformed all individual machine learning methods across most evaluation metrics. The results suggest that incorporating complementary classifiers can enhance predictive accuracy when analysing diverse textual data from Facebook and X. Thus, the hybrid model offers a more resilient and precise framework for prompt identification of depression on social media platforms, making it an appropriate method for advanced mental health surveillance systems.

Conclusion

The findings of this research indicate that the proposed hybrid Support Vector Machine-Random Forest (SVM-RF) model offers a superior approach for identifying depression by integrating Facebook and X (formerly Twitter) datasets compared to the individual machine learning algorithms assessed. The hybrid model achieved the highest overall performance, with 93% accuracy, 94% precision, 92% recall, and 93% F1-score, demonstrating its exceptional capability to differentiate between depressive and non-depressive social media posts. These results validate that combining the complementary

strengths of SVM and Random Forests improves classification performance and enhances the resilience of automated depression identification systems. Thus, the suggested hybrid framework serves as a dependable and effective means for facilitating early detection of depression and prompt mental health intervention through the utilisation of social media data.

Recommendations

Based on the findings of this study, the following recommendations are proposed.

1. Facebook and X (Twitter) companies can use the proposed hybrid SVM-RF model in this study to integrate machine learning-based depression-detection systems into their digital platforms.
2. The Nigerian government should make policies that will allow the identification of victims of depression on social media platforms.
3. Government and non-government organisations should introduce campaigns that will enlighten the public on the dangers of depression on individual health.
4. Researchers and developers should consider the proposed Hybrid SVM-RF model for automated depression detection on social media, as it demonstrated improved predictive performance over the individual SVM, RF, DT, K-NN, and NB classifiers.
5. Future studies should validate the proposed hybrid model using larger, more diverse datasets collected across multiple social media platforms.

Reference

- Adarsh, V., Arun Kumar, P., Lavanya, V., & Gangadharan, G. R. (2023). Fair and Explainable Depression Detection in Social Media. *Information Processing and Management*, 60(1), 103168. <https://doi.org/10.1016/j.ipm.2022.103168>
- Alkasem, H., Alsalamah, A., Alhussan, L., Alzahrani, N., & Aba-alkhail, S. (2026). Machine learning methods for detecting depression in Arabic tweets : A comprehensive performance analysis with enhanced evaluation metrics. *Discover Artificial Intelligence*, 6(120), 1-24.
- Alsagri, H. S., & Ykhlef, M. (2020a). Machine Learning-based Approach for Depression Detection in Twitter Using Content and Activity Features. In *IEICE Transactions on Information and Systems* (Vol. E103D, Number 8). Institute of Electronics, Information and Communication, Engineers, IEICE. <https://doi.org/https://doi.org/10.1587/transinf.2020EDP7023>
- Alsagri, H. S., & Ykhlef, M. (2020b). Machine Learning-Based Approach for Depression Detection in Twitter Using Content and Activity Features. *IEICE Transactions on Information and Systems*, E103D(8), 1825-1832. <https://doi.org/10.1587/transinf.2020EDP7023>
- Angskun, J., Tipprasert, S., & Angskun, T. (2022). Big Data Analytics on Social Networks for Real-Time Depression Detection. *Journal of Big Data*, 9(1). <https://doi.org/10.1186/s40537-022-00622-2>
- Chatterjee, R., Gupta, R. K., & Gupta, B. (2021). Depression detection from social media posts using multinomial naïve theorem. *IOP Conference Series: Materials Science and Engineering*, 1022(1). <https://doi.org/10.1088/1757-899X/1022/1/012095>
- Chiong, R., Budhi, G. S., Dhakal, S., & Chiong, F. (2021). A Textual-Based Featuring Approach for Depression Detection Using Machine Learning Classifiers and Social Media Texts. *Computers in Biology and Medicine*, 135. <https://doi.org/10.1016/j.combiomed.2021.104499>
- Coppersmith, G., Dredze, M., & Harman, C. (2014). Quantifying Mental Health Signals in Twitter. In A. for C. Linguistics (Ed.), *Workshop on Computational Linguistics and Clinical Psychology: From Linguistic Signal to Clinical Reality* (Number 222886, pp. 51-60).
- Dhas, J. T. M. (2020). Python 3.7.1 Vol-1. In *Introduction, Control flow, Data Structure, Module, I/O, Errors and Exceptions, Classes, Standard Libraries, Virtual Environments and Packages* (Number March, p. 131). <https://docs.python.org/3/tutorial/index.html>
- Draeos, R. (2019). *MACHINE LEARNING*. Measuring Performance: The Confusion Matrix. <https://glassboxmedicine.com/2019/02/17/measuring-performance-the-confusion-matrix/>

- Fisher, H., Jaffe, N. M., Pidvirny, K., Tierney, A. O., Vaidean, M. S., Dongre, P., & Webb, C. A. (2026). Language-Based Detection of Depression with Machine Learning: Systematic Review and Meta-Analysis. *Npj Digital Medicine*, 9(1), 273.
- Gadzama, W. A., Gabi, D., Argungu, M. S., & Suru, H. U. (2024). The use of machine learning and deep learning models in detecting depression on social media: A systematic literature review. *Personalised Medicine in Psychiatry*, 45-46, 100125. <https://doi.org/10.1016/j.pmip.2024.100125>
- Govindasamy, K. A. L., & Palanichamy, N. (2021). Depression Detection Using Machine Learning Techniques on Twitter Data. *Proceedings - 5th International Conference on Intelligent Computing and Control Systems, ICICCS 2021*, (Iciccs), 960-966. <https://doi.org/10.1109/ICICCS51141.2021.9432203>
- Gupta, S., Goel, L., Singh, A., Prasad, A., & Ullah, M. A. (2022). Psychological Analysis for Depression Detection from Social Networking Sites. *Computational Intelligence and Neuroscience*, 2022. <https://doi.org/10.1155/2022/4395358>
- Jayanthi, K., Deepika, K., Rupanjali, K., Tharun, K., Chowdary, D., & Sai Eswar, K. (2022). Depression Detection using Machine Learning Algorithms. *International Journal of Advanced Research in Science, Communication and Technology (IJARSCT)*, 2(1). <https://doi.org/10.48175/568>
- Kour, H., & Gupta, M. K. (2022). An Hybrid Deep Learning Approach for Depression Prediction from User Tweets Using Feature-Rich CNN and Bi-Directional LSTM. *Multimedia Tools and Applications*, 81(17), 23649-23685.
- Kumbhar, M., Dube, R., Barbade, M., Kulkarni, M., Konda, M., & Konkati, M. (2021). Depression Detection Using Machine Learning. *International Conference on Smart Data Intelligence (ICSMDI 2021)*. <https://ssrn.com/abstract=3851975>
- Luo, L., Yuan, J., Wu, C., Wang, Y., Zhu, R., Xu, H., Zhang, L., & Zhang, Z. (2025). Predictors of depression among Chinese college students: a machine learning approach. *BMC Public Health*, 25(1). <https://doi.org/10.1186/s12889-025-21632-8>
- Mortier, P., Cuijpers, P., Kiekens, G., Auerbach, R. P., Demyttenaere, K., Green, J. G., Kessler, R. C., Nock, M. K., & Bruffaerts, R. (2017). The prevalence of suicidal thoughts and behaviours among college students: A meta-analysis. *Psychological Medicine*, 48(4), 554-565. <https://doi.org/10.1017/S0033291717002215>
- Mowery, D., Bryan, C., & Conway, M. (2017). Feature Studies to Inform the Classification of Depressive Symptoms from Twitter Data for Population

- Health. *ArXiv Preprint*, 0–4. <http://arxiv.org/abs/1701.08229>
- Nadeem, A., Naveed, M., Satti, M. I., Afzal, H., Ahmad, T., & Kim, K. (2022). *Depression Detection Based on Hybrid Deep Learning SSCL Framework Using Self-Attention Mechanism: An Application to Social Networking Data*. 1–28.
- Norouzi, F., Modes, B. L., & Machado, S. (2024). Predicting Mental Health Outcomes: A Machine Learning Approach to Depression, Anxiety, and Stress. *International Journal of Applied Data Science in Engineering and Health*, 1, 1–7. <https://ijadseh.com>
- Rehmani, F., Shaheen, Q., Anwar, M., Faheem, M., & Bhatti, S. S. (2024). Depression detection with machine learning of structural and non-structural dual languages. *Healthcare Technology Letters*, 11(4), 218–226. <https://doi.org/10.1049/htl2.12088>
- Shen, Tg., Jia, J., Nie, L., Feng, F., Zhang, C., Hu, T., Chua, T.-S., & Zhu, W. (2017). Depression Detection via Harvesting Social Media: A Multimodal Dictionary Learning Solution. *JProceedings of the Twenty-Sixth International Joint Conference on Artificial Intelligence (IJCAI-17) Depression*, (26), 3838–3844. <https://doi.org/10.55529/jmhib.26.1.7>
- Skaik, R., & Inkpen, D. (2020). Using Twitter Social Media for Depression Detection in the Canadian Population. *ACM International Conference Proceeding Series*, 109–114. <https://doi.org/10.1145/3442536.3442553>
- Song, Y. L. Q., Chen, L., Liu, H., & Liu, Y. (2025). Machine learning algorithms to predict depression in older adults in China: a cross-sectional study. *Frontiers in Public Health*, 12(3). <https://doi.org/10.3389/fpubh.2024.1462387>
- Vasha, Z. N., Sharma, B., Esha, I. J., Al Nahian, J., & Polin, J. A. (2023). Depression Detection in Social Media Comments Data Using Machine Learning Algorithms. *Bulletin of Electrical Engineering and Informatics*, 12(2), 987–996. <https://doi.org/10.11591/eei.v12i2.4182>
- Wei, Y., Qin, S., Liu, F., Liu, R., Zhou, Y., Chen, Y., Xiong, X., Zheng, W., Ji, G., Meng, Y., Wang, F., & Zhang, R. (2025). Acoustic-based machine learning approaches for depression detection in Chinese university students. *Frontiers in Public Health*, 13. <https://doi.org/10.3389/fpubh.2025.1561332>
- Yadav, K., & Bhutia, D. T. (2025). Application of machine learning models for predicting depression among older adults with non-communicable diseases in India. *Scientific Reports*, 15(1), 1–26. <https://doi.org/10.1038/s41598-025-18053-3>
- Zhao, Y., Liang, Z., Du, J., Zhang, L., Liu, C., & Zhao, L. (2021). *Multi-Head Attention-Based Long Short-Term Memory for Depression Detection*



From Speech. 15(August), 1-11.
<https://doi.org/10.3389/fnbot.2021.684037>