

Ensemble Approach for Mental Disorder Severity Prediction

Emechebe Prince Emeka⁽¹⁾ Hambali Moshood Abiola⁽²⁾ Jerry Bala⁽³⁾

⁽¹⁾Department of Computer Science, Federal University Wukari

⁽²⁾Department of Computer Science, Department of Computer Engineering Federal University Wukari

⁽³⁾Department of Computer Science, Taraba State Polytechnic Suntai

Date of Submission: 01-07-2026

Date of Acceptance: 11-07-2026

ABSTRACT

Despite the increasing prevalence of mental health problems around the world, there is still a great need for automated and data-driven systems that aid in the detection and severity assessment of mental disorders. The current research proposes a framework for the prediction of severity of mental disorder based on machine learning ensembles using three different datasets obtained from the areas of clinical records, mental health surveys, and technology workplaces. In total six classifiers were examined, including Logistic Regression, Decision Tree, Support Vector Machine, Random Forest, Gradient Boosting and Voting Ensemble. The performance of these classifiers was measured by the metrics of Accuracy, Precision, Recall, F1 score and AUC-ROC. The experimental findings have shown that the use of ensemble techniques, especially Gradient Boosting and Random Forest, gives consistently better performance in prediction accuracy which ranges from 73.1% to 95.8% based on the dataset properties and feature quality. In addition to the high accuracy, Voting Ensemble offers higher consistency through the combination of the predictions obtained from different base estimators; however, its performance is subject to the level of correlation among base estimators. Comparative experiments showed that dataset properties such as feature separation, dimensionality, and noise levels play a larger role in affecting model performance than the size of the dataset.

Keywords: Mental health prediction, Ensemble learning, Mental disorder severity, classification models, predictive modeling.

I. INTRODUCTION

Mental disorders are one of the most important public health problems around the world; such diseases lead to disability, decreased productivity, and deaths worldwide. The World Health Organization states that depression and

anxiety are among the top reasons of diseases around the globe; in addition, the prevalence of such problems has increased in different populations [1]. Furthermore, burden of disease studies have also confirmed the importance of psychiatric problems on the economy and healthcare system [2]. The prevalence of mental disorders has been found more in conflict areas [3].

Early detection and accurate severity assessment are key to effective intervention; nevertheless, the current process is highly subjective, time-consuming, and does not scale well. It has led to the application of AI/ML-based automated systems for predicting mental well-being and assisting in making decisions [4], [5]. Recent research shows that supervised machine learning models are able to predict mental disorders using structured data of both clinical and behavioral nature [6]. Furthermore, methods like Random Forests and Gradient Boosting provide enhanced predictions by decreasing the variance and achieving better generalizability through ensemble learning techniques [7], [8]. Research has also explored reinforcement learning and other types of artificial intelligence for use in adaptive healthcare decision-making systems despite difficulties with interpretability and generalizability [9], [10].

More research shows that tree-based ensemble algorithms are very efficient for dealing with high-dimensional and non-linear mental health data because of their capability to detect interactions between features and avoid overfitting [8], [11]. On the contrary, linear algorithms such as Logistic Regression and margin-based algorithms such as Support Vector Machines have difficulties dealing with complex distribution of psychological data [6], [12].

Nevertheless, model performance is very dependent on dataset properties such as imbalance of classes, feature separation, and noise instead of just the size of the dataset [13]. This problem is especially important in the field of mental health

analytics due to the existence of overlapping classes in survey data and inconsistent answers [14].

There are many studies conducted in this domain but very few studies that offer an integration of different machine learning approaches on heterogeneous datasets of mental disorders. Most studies have been restricted to a single dataset or to a specific algorithm, which reduces the generality of findings obtained in these studies. Therefore, the current study attempts to propose an approach that will involve the use of multiple machine learning approaches on multiple types of mental health datasets.

II. RELATED WORK

Machine learning is one of the areas that have been greatly used in the field of prediction and decision-making concerning the state of mental well-being of patients. The early global health literature demonstrates an increase in the prevalence of mental diseases in the population and calls for more scalable computational tools to diagnose and treat such conditions [1], [2], [3]. These references create an epidemiological basis for introducing artificial intelligence into mental healthcare. There have been various investigations on the use of machine learning algorithms for mental state categorization. The well-known machine learning algorithms Logistic Regression, Support Vector Machines, and Decision Tree have been extensively used in the prediction of depression and anxiety using clinical and survey-based data [6], [12], [14]. However, the application of these classifiers on non-linear and high dimensional data has proved their poor performance capabilities.

To overcome these shortcomings, ensemble learning approaches have received considerable focus. The use of Random Forest and Gradient Boosting has proved to be successful in terms of better classification performance through the combination of multiple weak learners to minimize variance [7], [8]. XGBoost was proposed by Chen and Guestrin as an efficient and scalable boosting approach with effective predictive power, which has become increasingly popular for health care predictions [7]. Zhang et al. have discussed the effectiveness of using ensemble methods in classifying imbalanced medical data sets [15].

Recent studies further support the superiority of ensembles in mental health prediction. Kaushik et al. have proven that ensemble models are more efficient in mental disorder severity prediction than single classifiers using heterogeneous data sets [14]. Likewise, Ruhollah et al. have suggested ensemble learning methods that are capable of improving the accuracy of prediction in mental

health disorders classification [16]. Tachmazidis et al. have also stressed the significance of applying data balancing along with ensembles in clinical data sets [17].

Alongside conventional machine learning techniques, studies on deep learning and hybrid approaches have been done to predict mental disorders. For example, CNN has been used for identifying depression using electroencephalography signals and achieving high accuracy in features extraction [18]. Additionally, multimodal deep learning that combines social media data and surveys has shown good performance in mental health prediction [19]. Nevertheless, deep learning models are highly dependent on big amounts of data and computational power, which may restrict their use in some clinical settings.

Reinforcement learning is another model which has been proposed for adaptive healthcare systems. The research carried out by Abdellatif et al. and Yu et al. points towards the usefulness of the approach in terms of healthcare [9], [10]. However, even with all these developments, reinforcement learning approaches have not yet been explored for mental health severity predictions because of the difficulties involved.

The most notable problem that has been identified among the literature is generalizability on heterogeneous datasets. Guerreiro et al. have shown that models that have been trained on one mental health dataset tend to perform very poorly on other datasets because of differences in distribution of features and methods of data collection [13].

Also, Organisciak et al. stressed the importance of making machine learning algorithms clinically relevant and externally valid in the domain of mental health care systems [5]. Though many improvements have been achieved, current literature mostly considers one-model analysis and one data set only.

As a result, there is a clear requirement for conducting an integrated and comparative study of the performance of multiple machine learning models, especially ensemble approaches, on diverse data sets in the field of mental health. This paper tries to fill this gap through comparing conventional classification algorithms with ensemble approaches regarding their predictive performance in mental disorders.

III. METHODOLOGY

This research adopts a supervised machine learning framework for mental health condition prediction using both individual classifiers and an ensemble learning strategy. The methodology includes dataset

description, preprocessing procedures, model selection, ensemble design, and evaluation strategy.

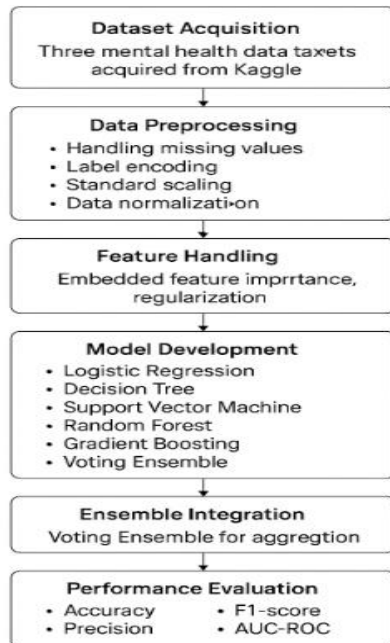


Figure 1. Methodology Workflow

Table 1. Datasets Summary

Dataset Name	Source	No. of Instances	No. of Features	No. of Classes
Mental Disorders Dataset (Clinical)	Kaggle	120	15	4 (Bipolar Type-1, Bipolar Type-2, Depression, Normal)
Mental Health Survey Dataset	Kaggle	5,000	15	2 (Yes / No)
Tech Workers Mental Health Dataset (OSMI)	Kaggle	1,433	52	2 (Yes / No)

IV. MACHINE LEARNING MODELS

These models were chosen to represent both individual learners and ensemble-based approaches for comprehensive comparison.

Logistic Regression (LR): a linear baseline model used for binary and multiclass classification.

Algorithm 1: Logistic Regression for Mental Health Classification

Input: Preprocessed dataset (X, y)

Output: Trained Logistic Regression model and predicted labels

1. Split dataset into training (70%) and testing (30%) sets
2. Initialize Logistic Regression model
3. Train model on (X_{train}, y_{train})
4. Predict labels on X_{test}
5. Evaluate model using accuracy, precision, recall, F1-score, ROC-AUC

Decision Tree (DT): a rule-based model capable of capturing non-linear relationships.

Algorithm 2: Decision Tree Classifier for Mental Health Classification

Input: Preprocessed dataset (X, y)

Output: Trained Decision Tree model and predicted labels

1. Split dataset into training (70%) and testing (30%) sets
 2. Initialize Decision Tree model
 3. Train model on (X_{train} , y_{train})
 4. Predict labels on X_{test}
 5. Evaluate model using accuracy, precision, recall, F1-score, ROC-AUC
-

Support Vector Machine (SVM): a margin-based classifier effective in high-dimensional spaces.

Algorithm 3: Support Vector Machine for Mental Health Classification

Input: Preprocessed dataset (X , y)

Output: Trained SVM model and predicted labels

1. Split dataset into training (70%) and testing (30%) sets
 2. Initialize SVM with selected kernel (e.g., RBF)
 3. Train model on (X_{train} , y_{train})
 4. Predict labels on X_{test}
 5. Evaluate model using accuracy, precision, recall, F1-score, ROC-AUC
-

Random Forest (RF): an ensemble of decision trees that reduces overfitting through bagging.

Algorithm 4: Random Forest for Mental Health Classification

Input: Preprocessed dataset (X , y)

Output: Trained Random Forest model and predicted labels

1. Split dataset into training (70%) and testing (30%) sets
 2. Initialize Random Forest with a specified number of trees
 3. Train model on (X_{train} , y_{train})
 4. Predict labels on X_{test}
 5. Evaluate model using accuracy, precision, recall, F1-score, ROC-AUC
-

Gradient Boosting (GB): a boosting technique that builds models sequentially to minimize error.

Algorithm 5: Gradient Boosting for Mental Health Classification

Input: Preprocessed dataset (X , y)

Output: Trained Gradient Boosting model and predicted labels

1. Split dataset into training (70%) and testing (30%) sets
 2. Initialize Gradient Boosting with chosen hyperparameters
 3. Train model on (X_{train} , y_{train})
 4. Predict labels on X_{test}
 5. Evaluate model using accuracy, precision, recall, F1-score, ROC-AUC
-

Voting Ensemble Classifier: The Voting Ensemble aggregates predictions from individual base models to reduce variance and improve generalization. It combines outputs from Logistic Regression, Decision Tree, SVM, and optionally Random Forest and Gradient Boosting. The final class label is determined based on the combined predictions.

Algorithm 6: Voting Ensemble for Mental Health Classification

Input: Trained base models: Logistic Regression, Decision Tree, SVM (and optionally Random Forest, Gradient Boosting)

Output: Final predicted labels from ensemble

1. Collect predictions or predicted probabilities from all base models on X test
2. Combine predictions using voting (majority vote or averaged probabilities)
3. Select final class label based on ensemble outcome
4. Evaluate ensemble model using accuracy, precision, recall, F1-score, ROC-AUC

V. PROPOSED ENSEMBLE FRAMEWORK

The suggested approach combines different base models together in one Voting Ensemble scheme. The ensemble uses majority voting to integrate predictions made by such models as Logistic Regression, Decision Tree, SVM, Random Forest, and Gradient Boosting. This combination makes it possible to take advantage of the strong sides of different models while avoiding their drawbacks. The use of the ensemble approach works especially well in case of mental health prediction because of the nature of the problem.

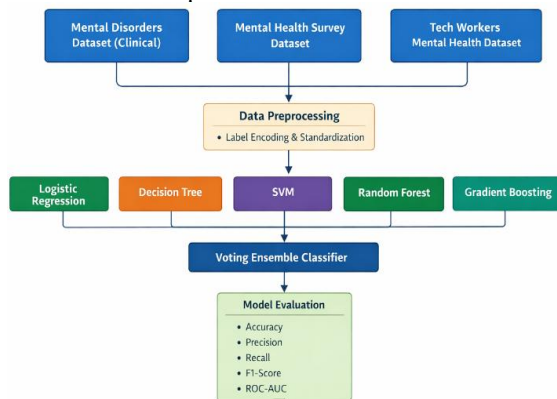


Figure 2. Model of the proposed system

VI. EVALUATION STRATEGY/METRICS

The performance of all models was evaluated based on a combination of the train-test split approach (70/30) and k-fold cross-validation in order to avoid overfitting bias. The cross-validation procedure was used only for the training dataset during model selection and hyperparameter optimization, and the final evaluation was performed on the test dataset.

1. **Accuracy:** Accuracy measures the proportion of correctly predicted instances out of the total instances. It gives an overall assessment of how well the model predicts both positive and negative cases.

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

Where TP is True Positives, TN is True Negatives, FP is False Positives, and FN is False Negatives.

Accuracy gives an overall assessment of the model's correctness.

2. **Precision:** Precision is the ratio of true positive predictions to the total predicted positive instances. It assesses the accuracy of positive predictions, indicating how well the model avoids false positives.

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

3. **Recall (Sensitivity or True Positive Rate):** Recall is the ratio of true positive predictions to the total actual positive instances. It measures the model's ability to correctly identify positive cases and avoid false negatives.

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

4. **ROC Score (Receiver Operating Characteristic Score):** The ROC score is calculated from the ROC curve, which visualizes the trade-off between the true positive rate and the false positive rate. The area under the ROC curve (AUC-ROC) represents the model's ability to distinguish between positive and negative instances. A higher AUC indicates better discrimination.

5. **The F1-score:** is the harmonic mean of precision and recall. It provides a balanced assessment of the model's precision and recall, which is especially useful when classes are imbalanced.

$$F1 = 2 \cdot \frac{Precision \cdot Recall}{Precision + Recall}$$

6. **Confusion Matrix:** The confusion matrix is a table that presents the system's performance by summarizing the predictions against the actual labels. It consists of four elements: true positives (TP), true negatives (TN), false positives (FP), and false negatives (FN).

		Actual Values	
		Positive (1)	Negative (0)
Predicted Values	Positive (1)	TP	FP
	Negative (0)	FN	TN

Figure 3. Confusion Matrix

VII. RESULTS

Models Performance on Mental Disorders Dataset (Dataset 1: N = 120 Samples, 15 Features)

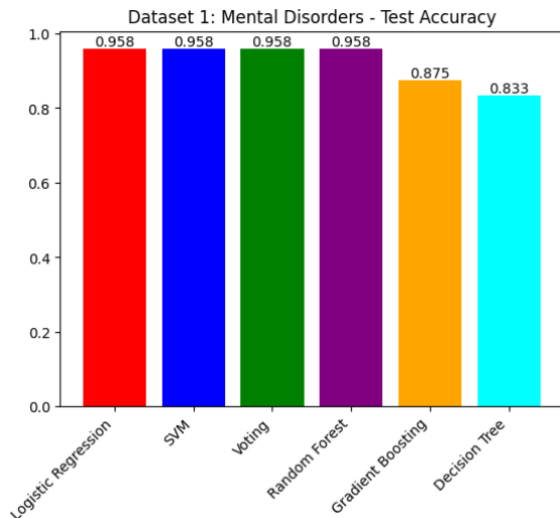


Figure 4. Models Performance on Mental Disorders Dataset

Performance Evaluation of Machine Learning Models on Mental Health Survey Dataset (Dataset 2; N = 5,000 Samples, 15 Features)

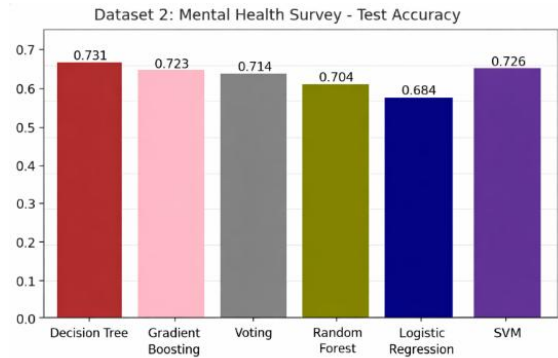


Figure 5. Models Performance on Mental Disorders Dataset

Models Performance on Tech Workers Mental Health Dataset (Dataset 3; N = 1,433 Samples, 52 Features)

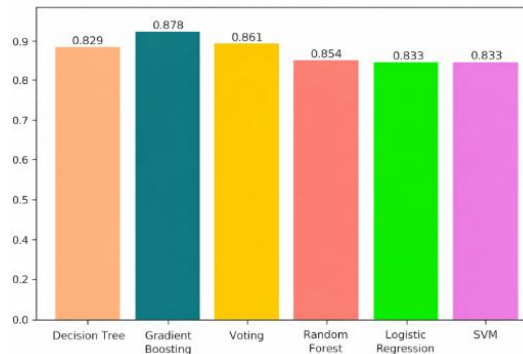


Figure 6: Models Performance on Tech Workers Dataset on Tech Workers Mental Health

Cross Dataset Comparative

Across the three datasets, distinct performance regimes emerge, as shown in Table 4.5. The results demonstrate that model effectiveness varies significantly depending on dataset characteristics such as feature dimensionality, class separability, and noise level. Rather than a single model consistently outperforming others, different algorithms achieve optimal performance under different data conditions

120, 15 features)

2. Mental Health Survey (N = 5,000, 15 features)

3. Tech Workers Mental Health (N = 1,433, 52 features)

Table 2: Cross Dataset Comparative

Dataset	Best Performing Models	Test Accuracy
1. Mental Disorders (Clinical, N =	Decision Tree / Random Forest / Voting Ensemble	95.8%

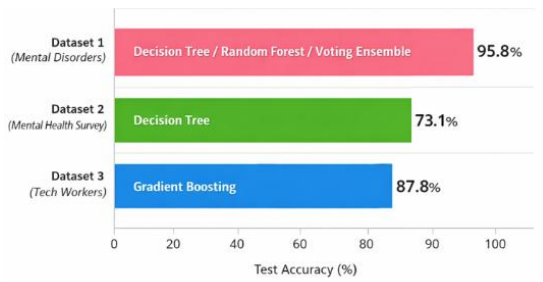


Figure 7: Cross Dataset Comparative

It is evident from the comparative analysis conducted on the three mental health data sets that each of them exhibits different performance levels, thus highlighting the importance of features in the data sets used. It is noteworthy that the Mental Disorders (Clinical) data set had the highest accuracy level of about 95.8% for Decision Tree and Random Forest algorithms.

In contrast, the Mental Health Survey database was the most difficult dataset for the decision tree model, as it attained an accuracy of just 73.1%. The fact that this dataset had a very high number of records but still performed poorly in terms of prediction is mainly attributed to the fact that the data were noisy, subjective, and overlapped.

The Tech Workers Mental Health dataset presented an example of moderate complexity, where Gradient Boosting showed the highest accuracy of 87.8%. This implies that due to a significant amount of structured demographic, occupational, and psychosocial features, boosting algorithms could model more complex feature interactions compared to linear methods.

Across all datasets, ensemble-based methods, including Random Forest, Gradient Boosting, and Voting Ensemble, proved to be effective in all datasets by highlighting the significance of capturing non-linear patterns in mental health prediction. Nonetheless, their dominance was dataset-specific, since no particular algorithm demonstrated a consistent superiority in all data. Logistic Regression and SVM showed stable baselines, but their inability to take into account feature interaction patterns became obvious.

Generally speaking, the results indicate that the quality of datasets and informativeness of features are more important than dataset size when predicting accurately. These lessons have implications in implementing machine learning in the real world to provide mental health services, where ensemble and boosting methods can be used to predict in different scenarios.

VIII. DISCUSSION

The experimental results demonstrate that machine learning models exhibit varying predictive performance across heterogeneous mental health datasets, consistent with findings in prior studies on healthcare predictive modeling [5], [8], [14]. This confirms that dataset characteristics such as feature distribution, noise level, and class separability play a critical role in model performance.

The clinical dataset achieved the highest predictive accuracy across all models, with ensemble-based methods such as Random Forest and Gradient Boosting performing best. This aligns with previous studies showing that tree-based ensembles are highly effective in structured clinical prediction tasks due to their ability to capture non-linear relationships and reduce overfitting [7], [10], [11]. The strong performance across multiple models suggests that the dataset contains highly discriminative and well-structured features that enhance model learning capacity.

On the contrary, the dataset for mental well-being survey performed the poorest even though it was made up of a bigger number of records. It is worth noting that this finding confirms the findings from previous studies showing that mental well-being information can be full of noise, bias, and overlapping classes thereby making it unreliable [13], [14]. The same case applies to behavioral datasets where feature overlap makes classification difficult [5].

The occupational (technology workplace) dataset demonstrated intermediate performance, with Gradient Boosting and Random Forest achieving the highest accuracy. This supports prior studies highlighting the effectiveness of ensemble learning methods in handling heterogeneous and moderately structured healthcare data [8], [10], [16]. The Voting Ensemble provided stable but not superior performance, which is consistent with literature showing that ensemble gains diminish when base learners are highly correlated [8], [15].

Ensemble methods have demonstrated superiority over classifiers such as Logistic Regression and SVM in all data sets. This finding is consistent with literature on how ensemble methods help improve generalization performance through variance reduction and stable decision making in biomedical data sets [7], [10], [17]. Nevertheless, none of the models generalized well in all data sets, which is consistent with the findings of Guerreiro et al. on how cross-domain generalization poses a challenge to mental health prediction [14].

A key finding is that data quality and feature informativeness have a greater impact on predictive performance than dataset size alone. This

observation is strongly supported by prior mental health analytics research, where feature engineering and dataset structure were found to be more influential than sample size in determining model effectiveness [13], [5].

Practically speaking, the outcomes show the promise of using ensemble machine learning methods to assess mental health in clinical and occupational settings. However, similar to Organisciak et al. and Rabipour & Asadi, this application calls for attention to be paid to issues of interpretability, generalizability, and clinical credibility [4], [5].

In general, this demonstrates the strength of the increasing literature that supports the use of ensemble learning for mental health prediction while pointing to the weaknesses of existing models when dealing with noisy real-world data.

IX. CONCLUSION AND FUTURE LOOK

This study offered an approach to building a comparative machine learning model for mental disorders severity predictions based on three diverse datasets for mental health analysis from clinical, survey, and occupational angles. In total, six machine learning algorithms such as Logistic Regression, Decision Tree, Support Vector Machine, Random Forest, Gradient Boosting, and Voting Ensemble were tested within a single experiment setup.

The obtained outcomes clearly show the superiority of ensemble learning techniques over the classical machine learning classifiers on all of the considered datasets. Specifically, the Gradient Boosting and Random Forest performed the best among the others with regard to predictive accuracy. However, while Voting Ensemble had relatively stable performance on all of the datasets, it could not beat the other two.

A key observation from this study is that dataset characteristics, including feature quality, noise level, and class separability, have a greater influence on model performance than dataset size alone. The survey dataset, despite being the largest, produced the lowest predictive accuracy, while the clinical dataset achieved the highest performance due to stronger feature structure.

The findings of this research emphasize the importance of ensemble learning techniques in developing robust mental health prediction systems that can support clinical decision-making, workplace mental health monitoring, and public health analytics.

Further research could involve enhancing the generalization capability of models using techniques such as feature engineering, ensemble-

based deep learning, and domain adaptation approaches. Moreover, applying explainable AI (XAI) algorithms will further improve interpretability and trustworthiness of mental health prediction models.

ACKNOWLEDGMENT

The conference organizer and the anonymous reviewers are acknowledged by the authors for their informative feedback that enhances the manuscript's substance.

REFERENCES

- [1]. World Health Organization, "World mental health forum 2024 report," 2024. [Online]. Available: <https://wfneurology.org>
- [2]. Institute for Health Metrics and Evaluation, "Global burden of anxiety and depression statistics," 2022. [Online]. Available: <https://www.healthdata.org>
- [3]. F. Charlson et al., "New WHO prevalence estimates of mental disorders in conflict settings: A systematic review and meta-analysis," *The Lancet*, vol. 394, no. 10194, pp. 240–248, 2019.
- [4]. S. Rabipour and S. Asadi, "Artificial intelligence applications in mental health: Opportunities and challenges for diagnosis and treatment," *Frontiers in Psychiatry*, vol. 14, 2023.
- [5]. D. Organisciak et al., "Machine learning for mental health: Applications, challenges, and the clinician's role," *Current Psychiatry Reports*, vol. 26, no. 11, pp. 694–702, 2024.
- [6]. A. M. Chekroud et al., "Cross-trial prediction of treatment outcome in depression: A machine learning approach," *The Lancet Psychiatry*, vol. 3, no. 3, pp. 243–250, 2016.
- [7]. T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. ACM SIGKDD*, 2016, pp. 785–794.
- [8]. Y. Zhang and X. Chen, "Machine learning approaches for mental health prediction: A comparative study of classification algorithms," *Journal of Biomedical Informatics*, vol. 140, 2023.
- [9]. K. Vaishnavi et al., "Predicting mental health illness using machine learning algorithms," *Frontiers in Digital Health*, 2025.
- [10]. R. Ruhollah, A. Rahimi, and M. Karimi, "A machine learning framework for mental health disorder prediction using ensemble learning techniques," *Expert Systems with Applications*, vol. 252, 2025.
- [11]. I. Tachmazidis et al., "Mental health prediction using ensemble learning

- approaches with rebalancing technique,” in *IEEE AMLDS*, 2025, pp. 455–460.
- [12]. A. A. Abdellatif et al., “Reinforcement learning for intelligent healthcare systems: A comprehensive survey,” *arXiv*, 2021.
- [13]. J. Yu et al., “Reinforcement learning in healthcare: Applications and challenges,” *Artificial Intelligence in Medicine*, vol. 145, 2023.
- [14]. J. Guerreiro et al., “Transatlantic transferability and replicability of machine-learning algorithms to predict mental health crises,” *npj Digital Medicine*, vol. 7, 2024.
- [15]. T. M. Laursen, M. Nordentoft, and P. B. Mortensen, “Excess early mortality in schizophrenia,” *Annual Review of Clinical Psychology*, vol. 10, pp. 425–448, 2014.
- [16]. M. N. A. Al-Hamadani et al., “Reinforcement learning algorithms and applications in healthcare and robotics: A systematic review,” *Sensors*, vol. 24, no. 8, 2024.
- [17]. M. Niu et al., “HCAG: A hierarchical context-aware graph attention model for depression detection,” in *IEEE ICASSP*, 2021, pp. 4235–4239.
- [18]. Z. Yin et al., “A multimodal deep learning framework for mental health prediction using social media and survey data,” *Computers in Biology and Medicine*, vol. 145, 2022.
- [19]. Y. Yoon et al., “Machine learning in mental health and its relationship with epidemiological practice,” *Frontiers in Psychiatry*, vol. 15, 2024.
- [20]. M. Madububambachu et al., “Machine learning techniques to predict mental health diagnoses: A systematic literature review,” *Clinical Practice & Epidemiology in Mental Health*, vol. 20, 2024.
- [21]. Y. Zhang et al., “Machine learning approaches for mental health prediction: A comparative study of classification algorithms,” *Journal of Biomedical Informatics*, vol. 140, 2023.
- [22]. Z. Zhang et al., “Research and application of XGBoost in imbalanced data,” *Journal of Industrial Information Integration*, vol. 27, 2022.
- [23]. D. Smolyakov, “Ensemble learning: Introduction to ensemble methods in machine learning,” 2017. [Online]. Available: <https://dyakonov.org>
- [24]. Y. Yin et al., “Machine learning techniques to predict mental health diagnoses: A systematic literature review,” *Clinical Practice & Epidemiology in Mental Health*, vol. 20, 2024.
- [25]. M. Yoon et al., “Machine learning in mental health and its relationship with epidemiological practice,” *Frontiers in Psychiatry*, vol. 15, 2024.
- [26]. J. Yu et al., “Reinforcement learning in healthcare: Applications and challenges,” *Artificial Intelligence in Medicine*, vol. 145, 2023.