



Anxiety and Depression Mental Health Factors: Predicting At-Risk Individuals using Supervised Learning Techniques

Kathleen Nicole D. Oliva

College of Information and Communication Technology,
South East Asian Institute of Technology, Inc.
Philippines

ABSTRACT

Mental health is a crucial component of human well-being, yet conditions such as anxiety and depression continue to affect individuals globally, often remaining undetected due to stigma and limited access to care. This study applies supervised machine learning techniques to predict individuals at risk of anxiety and depression using the Anxiety and Depression Mental Health Factors dataset sourced from Kaggle. The dataset consists of 1,200 survey responses encompassing demographic, lifestyle, medical, and psychosocial variables. Data preprocessing involved normalization, factor conversion, and the creation of binary “High” and “Low” risk labels based on a threshold score of 12, aligned with established clinical measures such as PHQ-9 and GAD-7. Three supervised learning models—Random Forest, Decision Tree, and Naive Bayes—were developed using an 80:20 train-test split and 5-fold cross-validation. Among these, the Random Forest model performed best, particularly in predicting depression ($F1 = 0.473$, $Recall = 0.452$, $AUC = 0.545$). Anxiety prediction, however, exhibited weaker performance across all models, indicating potential limitations in feature diversity or data balance. Feature importance analysis identified stress level, sleep hours, financial stress, and social support as significant predictors of mental health risk. The findings suggest that while machine learning models show promise in supporting early mental health detection, their accuracy depends heavily on data quality and feature representation. Future work may incorporate more comprehensive datasets and hybrid model approaches to strengthen predictive performance and support data-driven mental health interventions.

Key Words: Mental Health, Supervised Learning, Machine Learning, Anxiety, Depression, Random Forest, Decision Tree, Naive Bayes, Predictive Modeling

1. INTRODUCTION

1.1 Problem Background

Mental health is a state of well-being that allows people to handle everyday stress, recognize their abilities, learn and work productively, and contribute to their communities (World Health Organization, 2022). It's not just about the absence of mental illness—it's about how well someone can manage life's ups and downs, grow into their potential, and stay connected. Because of this, mental health is considered both a basic human right and a key factor in personal and societal development.

However, despite the increasing awareness, mental health issues like anxiety and depression continue to take a toll on individuals and communities around the world. People struggling with these conditions often face social isolation, lower productivity, and poor physical health. Research has shown that mental health disorders do not arise from a single source, they're often shaped by a mix of things like genetics, substance use, poverty, violence, inequality, and environmental stress (Kirkbride et al., 2024).

The role of stigma further complicates early detection. In the Philippines, a study on nursing students found that although personal mental health stigma levels were relatively low, it still significantly influenced their willingness to seek help (Tagufa et al., 2023). International studies supported this, where in UK, young adults delayed seeking

support due to feelings of shame, fear of judgment, and a preference for self-reliance, even when experiencing severe symptoms (Salaheddin & Mason, 2016). These barriers reveal that individuals at risk may remain undiagnosed or unsupported—not because of lack of need, but due to limited awareness, emotional discomfort, or perceived inaccessibility of care.

The challenge lies in identifying at-risk individuals early to enable timely intervention. This project aimed to explore how predictive analytics can be used to understand mental health data better and identify patterns that signal a higher risk for anxiety and depression. Using supervised learning models, the goal was to predict who might be struggling based on various factors—demographic, lifestyle, medical, and psychosocial—to support early intervention and, hopefully, better outcomes.

1.2 Dataset Description

This dataset utilized in this study, titled “Anxiety and Depression Mental Health Factors,” was sourced from Kaggle. It contains 1,200 survey responses collected from individuals reporting on various lifestyle, demographic, and psychological factors related to mental health. The dataset is a mix of numerical, categorical, and binary data, which includes the following features: 1) demographics: age, gender, education, and employment status, 2) lifestyle factors: sleep patterns, physical activity, and social support, 3) mental health metrics: anxiety score, depression score, stress level, 4) medical history: family history of mental illness, chronic illnesses, medication use, 5) coping strategies: therapy, meditation, substance use, and 6) additional factors: financial and work-related stress, self-esteem, life satisfaction, and feelings of loneliness. The dataset is designed for mental health analysis, predictive modeling, and research on the impact of various factors on mental well-being.

2. METHODOLOGY

This chapter outlines the procedures followed in preparing the dataset and developing the predictive models. The methodology is divided into two major stages: data preprocessing and model building.

2.1 Data Preprocessing Steps

The dataset went through a series of preprocessing steps to make sure it was clean, consistent, and ready for analysis. Since the dataset has no missing values, the researcher focused on data transformation and preparation for machine learning. Column names were standardized for consistency and ease of reference. The categorical variables included in the dataset such as Gender, Education Level, Employment Status, and several mental health-related indicators were all converted to factors to properly represent their qualitative nature in the models. A separate copy of the dataset with raw scores were kept for visualization purposes.

Binary classification labels for anxiety and depression were created by setting a threshold score of 12, which is above the ~75th percentile of the dataset, categorizing individuals into “High” and “Low” risk groups. This cutoff aligns with the principles behind the PROMIS T-score metric, where a score of 50 represents the mean of the reference population and 10 the standard deviation. In this metric system, higher scores indicate a greater presence of the measured concept (HealthMeasures, 2017)—in this case, anxiety or depression symptoms. Setting the threshold above this level effectively separates individuals exhibiting clinically significant symptoms, categorizing them into “High” and “Low” risk groups for targeted analysis. Furthermore, this approach is informed by widely used clinical screening tools such as the PHQ-9 and GAD-7, where thresholds of 10 to 15 are typically used to identify moderate to severe cases. However, recent studies caution that optimal cutoff values may vary depending on the population and setting, and overly rigid thresholds may limit screening accuracy in diverse samples (Snijkers et al., 2021). Thus, the 12-point cutoff was selected as a flexible yet meaningful boundary to reflect elevated risk without directly replicating clinical diagnoses.

Moreover, the raw and anxiety depression scores were excluded from the training dataset to avoid data leakage, which could lead to an artificially inflated model performance. These raw scores act as direct indicators or outcomes of mental illness severity, rather than explanatory factors influencing mental health status. Removing them ensured that the model learned from the underlying features and risk factors, which made the predictions more meaningful.

2.2 Model Building Techniques

Three supervised machine learning models were used to predict anxiety levels: Random Forest, Decision Tree, and Naive Bayes. These models were chosen not only for their ability to handle classification problems well, but also because they’ve been widely used in similar research studies related to mental health (e.g., Pandit et al., 2023; Qasrawi et al., 2022).

Before training the models, the dataset was split into a 80:20 ratio for training and testing sets respectively. Stratified sampling was used to maintain the balance of the binary labels, ensuring the models would be evaluated fairly on unseen data.

All models were trained using 5-fold cross-validation to ensure consistency and reliability in evaluation. Once trained, they were used to make predictions on the test set, which allowed comparison on how well each model performed in classifying individuals into “High” and “Low” risk groups.

3. RESULTS AND DISCUSSION

3.1 Model Performance Evaluation

The performance of the machine learning models was evaluated using standard classification metrics: accuracy, precision, recall, F1-score, and the Area Under the Curve (AUC). These numbers indicate whether the model actually understood which cases were serious enough to flag—especially for high anxiety or depression scores.

For the prediction of anxiety levels, it can be seen in Table 3.1.1 that the Decision Tree achieved the highest overall accuracy at 55.2%, yet came with a recall of 0.000, which clearly indicates a complete failure to identify high-anxiety cases. This meant that the accuracy is somewhat misleading, as the model’s predictive usefulness in this context is essentially null. The Random Forest model produced a slightly lower accuracy (54.0%) but yielded the best F1-score (0.267) and a modest AUC of 0.504. Naive Bayes, on the other hand, while achieving the lowest recall (0.019) and F1-score (0.035), had a comparable AUC of 0.505. Altogether, the results suggest that none of the models performed significantly better than random chance in identifying at-risk individuals for anxiety—reflecting a possible imbalance in the dataset or a lack of distinct patterns in the features used for prediction.

Table 3.1 Performance Evaluation Analysis

MODEL	ACCURACY	F1	PRECISION	RECALL	AUC
Anxiety					
Random Forest	0.540	0.267	0.465	0.187	0.503
Decision Tree	0.552	NaN	NaN	0.000	0.500
Naive Bayes	0.536	0.035	0.250	0.019	0.505
Depression					
Random Forest	0.517	0.473	0.495	0.452	0.545
Decision Tree	0.492	0.299	0.441	0.226	0.481
Naive Bayes	0.508	0.272	0.468	0.191	0.578

In contrast, the models exhibited relatively improved performance in predicting depression levels. Random Forest again emerged as the top-performing algorithm, with an accuracy of 51.7%, a recall of 0.452, and the highest F1-score (0.473). Although not exceptional as other related studies, these metrics indicate a comparatively more balanced capability to detect high-depression cases. The model’s AUC of 0.545 further supports its relative strength. Naive Bayes also showed promising results, achieving an AUC of 0.578 and an accuracy of 50.8%. It produced moderately average results in terms of precision and recall, which contributed to a reasonable F1-score of 0.272. Meanwhile, the Decision Tree model is weak in terms of classification ability, with lower precision (0.441), recall (0.226), and F1-score (0.299).

3.2 Visualization Outputs

To better understand the patterns and underlying relationships in the dataset, visualizations were produced using key psychosocial and engineered variables. The figures below provide descriptive insights into how anxiety and depression risks are distributed across key factors. The researcher recognized 4 groups of individuals: 1) low anxiety and low depression, 2) low anxiety and high depression, 3) high anxiety and low depression, and 4) high anxiety and high depression.

As shown in Figure 3.2.1, individuals in the high-risk groups demonstrated notably lower coping scores compared to those in low-risk groups. This pattern indicates that people with elevated anxiety and depression levels tend to have weaker coping mechanisms, which may increase their susceptibility to mental health challenges.

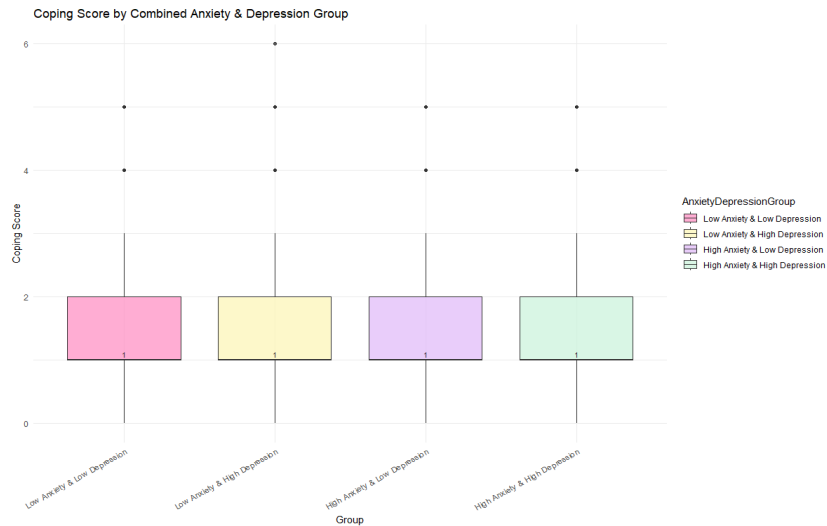


Figure 3.2.1 Coping Score by Combined Anxiety and Depression Group Box Plot

The box plot in Figure 3.2.2 reveals that overall stress levels rise in parallel with mental health risk. Respondents who fall under both high anxiety and high depression categories exhibit the greatest stress scores, underscoring the strong relationship between perceived stress and overall mental well-being.

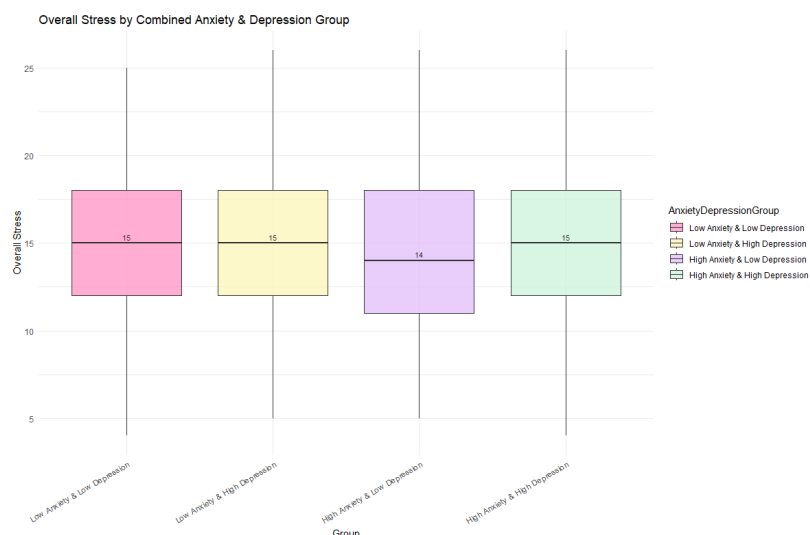


Figure 3.2.2 Overall Stress by Combined Anxiety and Depression Group Box Plot

Figure 3.2.3 highlights the most influential predictors identified by the Random Forest model for anxiety classification. Variables such as Age, Sleep Hours, Physical Activity Hours, Stress Level, and Financial Stress were among the top contributors. This suggests that anxiety symptoms in the dataset are primarily influenced by lifestyle and situational stress factors.

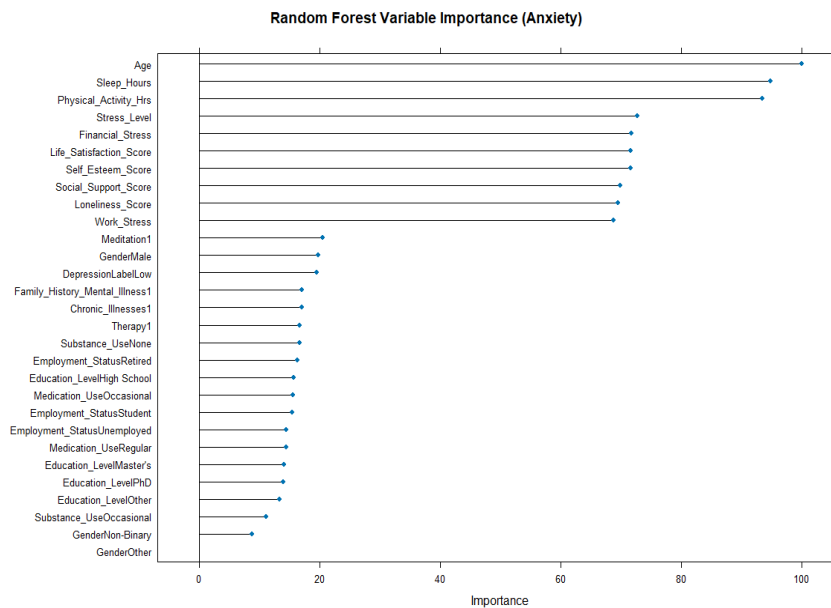


Figure 3.2.3 Random Forest Variable Importance of Anxiety

In contrast, Figure 3.2.4 presents the most significant predictors of depression. The model identified Sleep Hours, Age, Physical Activity Hours, Life Satisfaction Score, Loneliness Score, and Stress Level as key factors. These findings emphasize that depression risk is closely tied to psychosocial well-being and perceived connectedness.

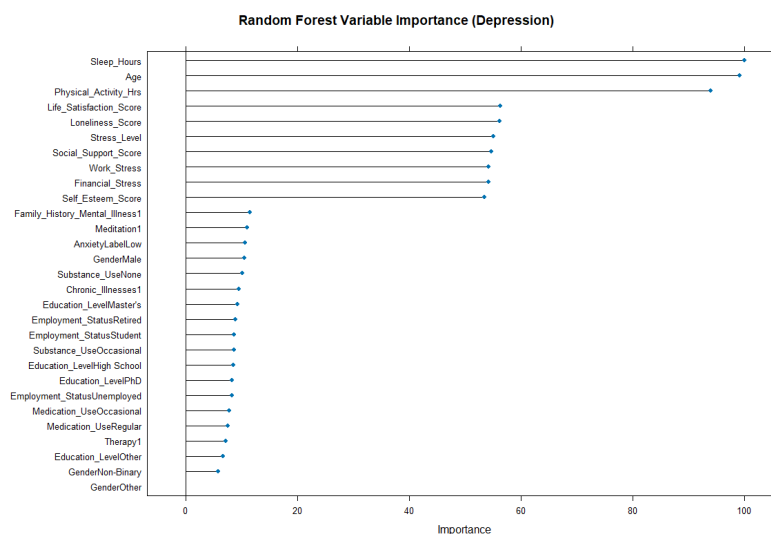


Figure 3.2.4 Random Forest Variable Importance of Depression

The ROC curve shown in Figure 3.2.5 compares the classification performance of the three supervised learning models for depression prediction. Among them, the Random Forest achieved the best balance between sensitivity and specificity, demonstrating modest predictive capability compared to Decision Tree and Naive Bayes.

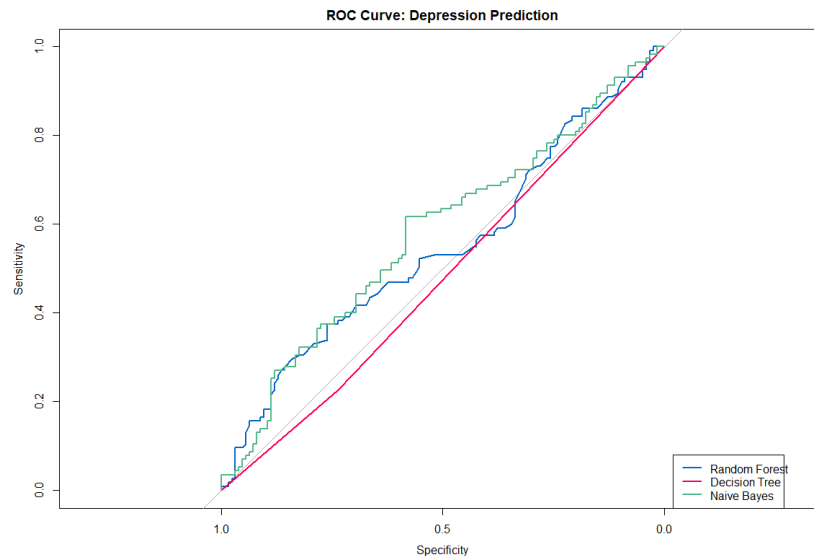


Figure 3.2.5 ROC Curves of Depression Prediction

3.3 Interpretation of Results

This study aimed to assess the effectiveness of supervised learning models in predicting anxiety and depression risks based on demographic, lifestyle, and psychosocial features. Among the models tested, Random Forest delivered the most balanced results in terms of accuracy and overall classification performance. However, the outcomes were notably better for depression prediction compared to anxiety. For depression, the Random Forest model achieved an F1-score of 0.473 and a recall of 0.452, indicating a fairly balanced ability to identify true positive cases while managing false predictions. In contrast, anxiety prediction produced relatively weaker results, with the highest recall being only at 0.187. Decision Tree, despite achieving the highest accuracy for anxiety (55.2%), recorded a recall of 0.000, which means it failed to detect any actual high-risk cases. These limitations were evident in the confusion matrices, particularly for Naive Bayes and Decision Tree, both of which struggled with sensitivity in detecting high-anxiety cases.

The Random Forest feature importance plots shed light on the variables most influential in the classification process. For anxiety prediction, the top contributors included Age, Sleep Hours, Physical Activity Hours, Stress Level, and Financial Stress (Figure 3.2.3). These suggest that anxiety risk in this dataset is more tied to lifestyle and demographic factors than to emotional or coping-related indicators. For depression, key predictors were Sleep Hours, Age, Physical Activity Hours, Life Satisfaction Score, Loneliness Score, and Stress Level (Figure 3.2.4), emphasizing the role of psychosocial wellbeing and connectedness in identifying depressive tendencies. Similar findings reported that Random Forest and Support Vector Machine (SVM) outperformed other models when applied to datasets that included physical, mental, and social health indicators (Qasrawi et al., 2022; Pandit et al., 2023). Both studies also noted the underperformance of Naive Bayes in mental health prediction, which aligns with the current study's results.

The integration of non-clinical features in this study also aligns with another study, where it was discussed that generalized anxiety disorder (GAD) and major depressive disorder (MDD) could be effectively predicted using routine demographic and biomedical data without relying on traditional psychiatric markers (Nemesure et al., 2021). Valuable variables such as SES, life satisfaction, and stress were features that also played a critical role in the models used in the current study. Meanwhile, a study emphasized the importance of evaluating models using more than just accuracy, especially when working with imbalance datasets (Norouzi & Machado, 2024). It was discussed that recall and F1-score act as crucial metrics, which echoed this study's approach and helped explain the gap in model performance between anxiety and depression classification.

Findings in demographic group suggest that mental health risks were more prevalent among older individuals (particularly those aged 50+), while younger respondents tended to fall into low-risk groups. It was also revealed that females exhibit a higher rate of comorbid high-risk categories, and individuals with lower education levels or unstable employment were also more likely to belong to higher-risk groups. High-risk respondents consistently had lower coping scores, wellbeing, and lifestyle quality, alongside higher levels of overall stress (Figures 3.2.1 to 3.2.2). These

relationships are supported by prior research, emphasizing the role of behavioral and environmental factors, such as stress, life satisfaction, and support in the development of mental health disorders (Nemesure et al., 2021; Iyortsuun et al., 2023).

Moreover, while fewer respondents in the dataset reported chronic illnesses or a family history of mental disorders, patterns in the data indicated that those who did were proportionally more represented in high-risk groups. These findings reflect broader research that highlights the predictive strength of medical and hereditary factors in mental health classification (Pandit et al., 2023; Qasrawi et al., 2022). Random Forest and SVM continue to be the most reliable algorithms in diagnosing a variety of mental health conditions, particularly in small-to-medium datasets where interpretability and generalization are key (Iyortsuun et al., 2023).

One important limitation of the present study is the relatively poor performance of all models in predicting anxiety. Past research has shown that anxiety symptoms tend to be more complex, internalized, and context-sensitive compared to depression (Nemesure et al., 2021; Iyortsuun et al., 2023), which made them more difficult to capture using general survey data. This observation also points to the challenge of class imbalance, where models are not equally trained to detect less frequently occurring outcomes. While accuracy is a common metric, it can be misleading in such situations, and greater focus should be placed on sensitivity, specificity, and the F1-score when evaluating performance (Norouzi & Machado, 2024).

Nonetheless, the results affirm that with proper feature selection, data preprocessing, and careful evaluation, supervised learning methods like Random Forest can uncover meaningful patterns from survey-based data. Although this study does not claim clinical-level diagnostic accuracy, it shows potential for scalable, non-invasive mental health screening, especially for depressive symptoms.

4. CONCLUSION AND RECOMMENDATION

4.1 Summary of Key Findings

This study explored the use of supervised learning techniques: Random Forest, Decision Tree, and Naive Bayes, to predict individuals at risk of anxiety and depression based on survey data from Kaggle (Kumar, 2025), which covered demographic, lifestyle, medical, and psychosocial variables. Results show that while all models struggled to classify anxiety cases effectively, Random Forest outperformed the others in terms of balanced metrics for both disorders. For depression, Random Forest achieved the highest F1-score (0.473) and recall (0.452), indicating moderate success in identifying individuals with higher depression scores. In contrast, anxiety classification produced weaker results, with a maximum recall of only 0.187 and limited ability to detect true positive cases, particularly by the Decision Tree model, which failed to identify any high-risk anxiety cases.

Variable importance plots revealed that both emotional and behavioral indicators played key roles in prediction, especially for depression. Visual analysis supported these patterns, showing that high-risk individuals tended to report lower coping and wellbeing scores, higher stress, and weaker social support. Additionally, mental health risks were more prevalent among older adults, females, unemployed individuals, and those with chronic illnesses or a family history of mental disorders.

One notable limitation of the study is the relatively small dataset ($n = 1,200$), which may have limited the models' ability to generalize and capture complex mental health patterns—particularly for anxiety. This aligns with observations from prior research (Iyortsuun et al., 2023; Nemesure et al., 2021), which cautions against overreliance on smaller sample sizes due to the risk of overfitting and class imbalance. The absence of temporal, clinical, or narrative data also limited the depth of the models' contextual understanding. Nevertheless, the study demonstrates that supervised learning can provide meaningful early insights into mental health risks, especially for depression, when paired with appropriate feature selection and evaluation metrics.

4.2 Recommendations for Business Decision-Making and Model Improvements

The findings of this study present several implications for business and organizational decision-making, particularly in contexts such as schools, workplaces, and community health programs. Early identification of individuals at-risk for anxiety and depression can inform the development of targeted wellness strategies, preventive mental health interventions, and more responsive resource allocation. Even though the models used in this study are

not diagnostic tools, they offer valuable insight into how data-driven approaches may support timely detection and action. Businesses, for instance, may use similar predictive frameworks to improve employee wellbeing and initiatives or personalize mental health outreach in human resource settings.

In terms of technical improvements, one of the most critical recommendations is to expand the dataset size and diversity. A larger, more diverse dataset can help mitigate the limitations posed by the current sample of 1,200 individuals and reduce the likelihood of overfitting. This is especially important when attempting to model complex disorders like anxiety, which tend to be more variable and context-sensitive than depression. Future studies could also integrate clinically validated psychometric instruments such as the GAD-7 and PHQ-9 to strengthen the reliability of target labels and better align prediction with clinical benchmarks.

Additionally, addressing class imbalance through techniques like SMOTE (Synthetic Minority Over-sampling Technique) or cost-sensitive learning could improve recall scores and enhance the model's sensitivity in identifying high-risk cases, particularly for anxiety. Expanding the range of features to include real-time behavioral data, temporal tracking, and even qualitative indicators such as journal entries or mood logs may also provide deeper insight into subtle psychological patterns. These improvements can help models learn more contextually rich and dynamic representations of mental health.

Finally, it is recommended to explore hybrid modelling strategies by combining traditional supervised learning models with deep learning approaches (Iyortsuun et al., 2023). Merging ML and DL techniques has shown promising results, particularly when working with symptomatically complex disorders or when additional predictive precision is required. However, such enhancements would demand access to larger, high-quality datasets to prevent overfitting and ensure generalizability. In conclusion, while the current models provide a strong foundational framework for predictive mental health analysis, further refinement through thoughtful feature engineering and architectural design will be essential for real-world application.

ACKNOWLEDGEMENT

The researcher would like to express heartfelt gratitude to the Lord Almighty for the knowledge, strength, and wisdom bestowed throughout the course of this study. Deep appreciation is extended to Ms. Mary Joy M. Fernandez, MIT, LPT for her continuous guidance, support, and invaluable feedback that shaped the direction of this research, and to Mr. Reginald S. Prudente, MIT for his encouragement and belief in the researcher's capability.

The researcher is also sincerely thankful to their family for always keeping them in their prayers and for providing constant emotional and spiritual support. Gratitude is likewise extended to the individual who made the dataset used in this study publicly available, which made this project possible.

Above all, this work is also dedicated to the researcher themselves—for enduring the challenges, for showing resilience, and for continuing to believe even in moments of doubt. This study stands as a personal milestone of growth, faith, and perseverance.

REFERENCES

- HealthMeasures. (2017). *PROMIS*. HealthMeasures. <https://www.healthmeasures.net/score-and-interpret/interpret-scores/promis>
- Iyortsuun, N. K., Kim, S.-H., Jhon, M., Yang, H.-J., & Pant, S. (2023). A Review of Machine Learning and Deep Learning Approaches on Mental Health Diagnosis. *Healthcare*, 11(3). <https://doi.org/10.3390/healthcare11030285>
- Kirkbride, J., Anglin, D., Colman, I., Dykxhoorn, J., Jones, P., Patalay, P., Pitman, A., Sonesson, E., Steare, T., Wright, T., & Griffiths, S. L. (2024). The social determinants of mental health and disorder: evidence, prevention and recommendations. *World Psychiatry*, 23(1), 58-90. <https://doi.org/10.1002/wps.21160>
- Kumar, A. (2025, March 14). *Anxiety and Depression Mental Health Factors Dataset*. Kaggle. https://www.kaggle.com/datasets/ak0212/anxiety-and-depression-mental-health-factors?select=anxiety_depression_data.csv
- Nemesure, M., Heinz, M., Huang, R., & Jacobson, N. (2021). Predictive modeling of depression and anxiety using electronic health records and a novel machine learning approach with artificial intelligence. *Scientific Reports*, 11(1980). <https://www.nature.com/articles/s41598-021-81368-4>

- Norouzi, F., & Machado, B. L. M. 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*. <https://ijadseh.com/index.php/ijadseh/article/view/19/23>
- Pandit, M., Azwaan, M., Wani, S., Ibrahim, A. A., Abdulkhaleq, R., Abdulghafor, A., & Gulzar, Y. (2023). Examining Factors for Anxiety and Depression Prediction. *International Journal on Perceptive and Cognitive Computing*, 9(1). <https://journals.iium.edu.my/kict/index.php/IJPC/article/view/368/219>
- Qasrawi, R., Polo, S. P. V., Al-Halawa, D. A., Hallaw, S., & Abdeen, Z. (2022). Assessment and Prediction of Depression and Anxiety Risk Factors in Schoolchildren: Machine Learning Techniques Performance Analysis. *JMIR Form Res*, 6(8). <https://doi.org/10.2196/32736>
- Salaheddin, K., & Mason, B. (2016). Identifying barriers to mental health help-seeking among young adults in the UK: a cross-sectional survey. *British Journal of General Practice*, 66(651). <https://doi.org/10.3399/bjgp16X687313>
- Snijders, J., van den Oever, W., Weerts, Z., Vork, L., Mujagic, Z., Leue, C., Hesselink, M., Kruimel, J., Muris, J., Bogie, R., Masclee, A., Jonkers, D., & Keszthelyi, D. (2021). Examining the optimal cutoff values of HADS, PHQ-9 and GAD-7 as screening instruments for depression and anxiety in irritable bowel syndrome. *Neurogastroenterology & Motility*, 33(12). <https://doi.org/10.1111/nmo.14161>
- Tagufa, G., Magobo, A. K. L., Andal, D. J., Eisma, D. J., Tanhueco, H., Macalinga, H., Manahan, I. C., Tumolva, M. 2., Opur, Y. G., & Pobre, R. (2023). Mental Health Stigma and Help-Seeking Behavior Among Manila Central University. *Philippine Scientific Journal*, 56(1). <https://www.ejournals.ph/article.php?id=26323>
- World Health Organization. (2022, June 17). *Mental health*. World Health Organization (WHO). <https://www.who.int/news-room/fact-sheets/detail/mental-health-strengthening-our-response>