

# IMHAP-Net: An Intelligent, Explainable Machine Learning Framework for Mental Health Risk Assessment and Prediction

Rohit Rana<sup>1</sup> Dr. Raj Kumar<sup>2</sup>

<sup>1</sup>Scholar of Department of Computer Science & Engineering, Quantum University, Roorkee, India

<sup>2</sup>Department of Computer Science & Engineering, Quantum University, Roorkee, India

**Abstract- Background:** Although the prevalence of mental health conditions like depression, anxiety, and chronic stress is on the rise among students and young adults globally, traditional clinical assessment primarily relies on self-report questionnaires and scheduled clinician contact, which can cause delays in identifying at-risk individuals. In order to support early mental health risk screening, this paper proposes IMHAP-Net, a layered, explainable machine learning framework that combines heterogeneous behavioural, academic, and self-report data with a stacked ensemble of tree-based and deep learning models. Each prediction is traceable to interpretable contributing factors. Data collection, pre-processing, feature engineering, a hybrid stacking ensemble (Random Forest, XGBoost, LightGBM, CatBoost, and a deep neural network combined via a meta-learner), an explainable-AI layer (SHAP/LIME), and a risk-assessment layer that transforms model outputs into a continuous Mental Health Risk Index (MHRI) and a four-level categorical risk label comprise the framework's six layers. To be finished following trials on the target dataset or datasets; see to Section 20 for the runnable implementation and Section 12 for the metric template.] The contributions include: (i) an open implementation path (Colab/PyTorch/Scikit-learn) that generates auditable, explainable outputs instead of a black-box label; (ii) a formally defined Mental Health Risk Index combining multiple sub-scale predictions; (iii) a stacking-ensemble formulation with explicit meta-learner equations; and (iv) a reusable six-layer architecture that separates prediction from explanation and risk scoring.

**Keywords:** Mental health prediction; explainable AI; ensemble learning; SHAP; risk stratification; depression detection; anxiety screening; student wellbeing; stacking ensemble; machine learning in healthcare.

## I. INTRODUCTION

The frequency of mental health disorders, especially depression and anxiety, has increased significantly among college and university students over the past ten years, according to several national studies, making them one of the most important worldwide public health problems. Between 2013 and 2021, reported rates of anxiety and depression symptoms among US college students nearly quadrupled; similar patterns have been observed across university populations in South Asia, Latin America, and the Middle East. Adolescents are also impacted; according to current national youth risk surveys, a significant percentage of high school students suffer from ongoing depression or hopelessness.

Conventional diagnosis relies on planned clinical interviews and self-report tools like the PHQ-9, GAD-7, and DASS-21, which are validated and useful but

intrinsically reactive; people are usually evaluated only after they seek assistance or are referred. This leads to a detection gap whereby kids who are experiencing increasing levels of discomfort could not be recognised until their symptoms become severe. A supplementary, proactive layer is provided by machine learning, which uses routinely accessible behavioural, academic, and lifestyle signals (sleep patterns, academic load, social interaction, past self-report history) to identify elevated risk earlier and allocate counselling resources where they are most needed.

However, the actual use of ML for mental health screening is constrained by three enduring inadequacies. First, interpretability—which is crucial for clinicians and counsellors who must defend any intervention—is neglected in many published models that indicate accuracy. Second, single-model methods (a single classifier) are often either

interpretable but less accurate (logistic regression) or accurate but opaque (deep networks), seldom both. Third, risk stratification—that is, transforming a probability into an actionable category that a non-technical counsellor can use—is barely discussed in most research, which provide a single best-fit accuracy score for a single dataset.

In order to close these gaps, this paper proposes IMHAP-Net, a layered framework that: (a) employs a stacked ensemble to balance robustness and accuracy; (b) incorporates an explainability layer (SHAP/LIME) as a first-class architectural component rather than an afterthought; and (c) defines a transparent, formula-based Mental Health Risk Index that transforms model outputs into an understandable four-level risk category. In the remaining sections of this paper, related work is reviewed (Section 2), the IMHAP-Net architecture and its mathematical formulation are presented

(Sections 3–4), algorithmic design is described (Section 5), the experimental setup and datasets on which the framework will be evaluated are described (Sections 6–7), evaluation metrics are specified (Section 8), and explainability, validation strategy, advantages, limitations, and future directions are discussed (Sections 9–14).

## II. LITERATURE REVIEW

Survey-based prediction, interpretability-focused research, and rigorous assessments of classifier performance are examples of recent work on machine learning-based mental health screening. The typical recent research found through a literature search are summarised in Table 1. One of the driving needs for a consistent, explainable framework is the considerable variation in reported performance statistics among studies due to different datasets, sample sizes, and feature sets.

| Study                                              | Venue                                        | Population / Dataset                       | Method                                                                   | Key Finding                                                                                                                                                      |
|----------------------------------------------------|----------------------------------------------|--------------------------------------------|--------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------------------------|
| Zhai et al. (2025)                                 | Journal of Counseling & Development          | US college students (GAD-7 / PHQ-9 survey) | ML predictive models trained without health-history features             | Showed predictive models can flag anxiety/depression risk before diagnosis; aimed at guiding counselling-resource allocation rather than maximizing raw accuracy |
| Schaab et al. (2024)                               | Cadernos de Saude Publica                    | Systematic review (PRISMA, 30 studies)     | Review of ML classifiers for depression/anxiety/stress in undergraduates | Found wide variability in reported performance across studies; flagged inconsistent feature sets and small sample sizes as recurring limitations                 |
| Bsharat-type cross-sectional study, medRxiv (2025) | medRxiv preprint                             | 9th-grade students, GAD-7/PHQ-9            | ML prediction of high-school student mental health                       | Highlighted rising adolescent anxiety/depression prevalence (CDC YRBS) and tested ML screening on self-report instruments                                        |
| Lebanese university DAS study (2024)               | PMC / SAGE                                   | Two Lebanese universities, DASS-21         | ML approach using demographics and self-rated health                     | Found demographic and self-rated health features useful predictors of depression/anxiety/stress during COVID-19                                                  |
| Bangladesh university DAS study (2025)             | J. Health, Population & Nutrition (Springer) | 1,697 students, two public universities    | Combined statistical analysis with multiple ML models                    | Identified sociodemographic and behavioral predictors of DAS outcomes; reinforced need for region-specific models                                                |
| Interpretable ML for college                       | arXiv preprint                               | College student records                    | Interpretable ML (tree-based + post-hoc explanation)                     | Argued explainability is essential for adoption by                                                                                                               |

|                                                          |     |                       |                                                                                                                                |                                                                                                                                                |
|----------------------------------------------------------|-----|-----------------------|--------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------------------|
| mental health (2025)                                     |     |                       |                                                                                                                                | counselling staff, not just raw predictive accuracy                                                                                            |
| Systematic review of ML for psychiatric diagnoses (2024) | PMC | 30 studies, 2011-2024 | Comparison of 30 classification algorithms across 5 conditions (incl. depression, anxiety, ADHD, PTSD, bipolar, schizophrenia) | Concluded deep learning models are increasingly used but performance reporting is inconsistent across studies, limiting cross-study comparison |

These studies reveal a number of recurrent gaps: (1) inconsistent reporting of performance metrics makes cross-study comparison challenging; (2) most models continue to be single-classifiers rather than ensemble approaches; (3) explain ability is treated inconsistently; several recent papers explicitly argue that interpretability, not just raw accuracy, determines whether counselling staff will trust and adopt a model; and (4) risk stratification (converting a prediction into an actionable category) is rarely formalised as part of the model output. Gaps (2)–(4) are immediately addressed by IMHAP-Net.

III. PROPOSED FRAMEWORK: IMHAP-  
NET

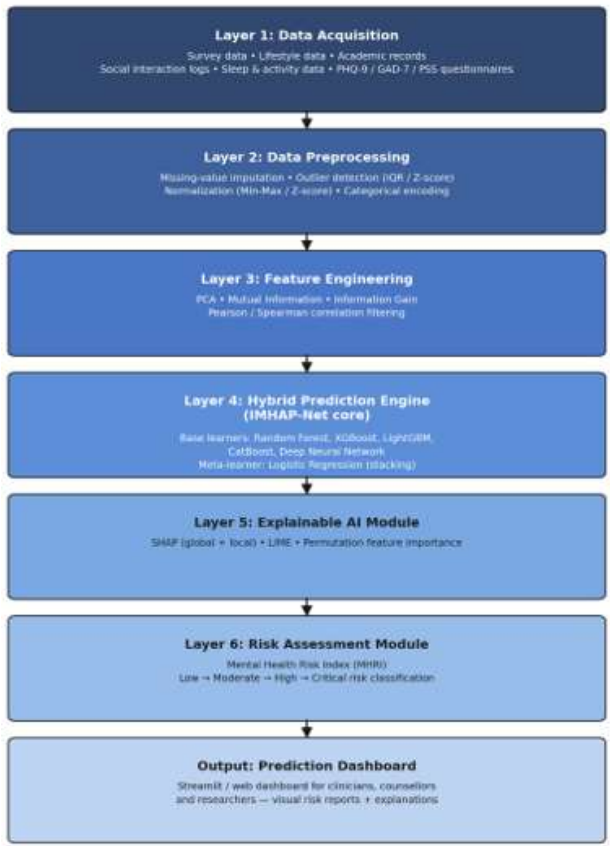


Figure 1. IMHAP-Net six-layer architecture.

Figure 1 illustrates the six successive layers that make up IMHAP-Net (Intelligent Mental Health Assessment and Prediction Network).

Layer 1 — Data Acquisition

Structured self-report surveys (PHQ-9, GAD-7, PSS, or DASS-21 sub-scales), academic performance records, lifestyle indicators (sleep duration, physical activity), and, where accessible, social contact or app usage logs are examples of input sources. As long as the source uses at least one validated psychological subscale as the prediction objective, the framework is independent of the dataset.

Layer 2 — Data Pre-processing

- **Missing-value handling:** rows with more than 40% missing fields are eliminated; mode imputation is used for categorical features, while median imputation is used for numerical features.
- **Outlier detection:** if  $x_i < Q1 - 1.5 \cdot IQR$  or  $x_i > Q3 + 1.5 \cdot IQR$ , it is flagged as an outlier according to the interquartile range (IQR) criterion.
- **Normalisation (Min-Max):**  $x' = (x - x_{min}) / (x_{max} - x_{min})$ .
- **Z-score standardisation:**  $z = (x - \mu) / \sigma$ .
- **Categorical encoding:** ordinal encoding for replies on an ordered Likert scale, and one-hot encoding for nominal variables.

Layer 3 — Feature Engineering

Dimensionality reduction and feature selection combine four complementary techniques:  
Principal Component Analysis:  $Z = X \cdot W$ , where  $W$  contains the eigenvectors of the covariance matrix  $\Sigma = (1/n) \cdot X^T X$   
Mutual Information:  $I(X;Y) = \sum \sum p(x,y) \cdot \log[ p(x,y) / (p(x) \cdot p(y)) ]$   
Information Gain:  $IG(S,A) = H(S) - \sum (|S_v|/|S|) \cdot H(S_v)$

Pearson correlation:  $r = \frac{\sum (x_i - \bar{x})(y_i - \bar{y})}{\sqrt{\sum (x_i - \bar{x})^2} \cdot \sqrt{\sum (y_i - \bar{y})^2}}$

Features are retained if they exceed an empirically tuned mutual-information / information-gain threshold and are not highly collinear ( $|r| < 0.85$ ) with an already-selected feature.

#### Layer 4 — Hybrid Prediction Engine (core of IMHAP-Net)

A meta-learner in a stacking-ensemble configuration combines the outputs of five separately trained base learners:

- **Random Forest (RF):** a bagged ensemble of decision trees with average or majority voting. Gradient-boosted trees with regularised additive learning are called XGBoost.
- **LightGBM:** leaf-wise tree growth combined with histogram-based gradient boosting.
- **CatBoost:** gradient boosting combined with ordered boosting and native categorical feature processing.
- A completely linked network that captures non-linear feature interactions is called a deep neural network (DNN).

**Random Forest prediction:**  $\hat{y}_{RF}(x) = (1/T) \sum_{t=1}^T h_t(x)$

XGBoost objective:  $L(\theta) = \sum_i \ell(y_i, \hat{y}_i) + \sum_k \Omega(f_k), \quad \Omega(f) = \gamma T + (1/2)\lambda ||w||^2$

**DNN forward propagation:**  $a^{(i)} = \varphi(W^{(i)} \cdot a^{(i-1)} + b^{(i)})$ ,  $\varphi = \text{ReLU}$  for hidden layers, sigmoid/softmax for output

**Binary cross-entropy loss:**  $L = -(1/N) \sum_i [y_i \log(\hat{y}_i) + (1 - y_i) \log(1 - \hat{y}_i)]$

**Backpropagation weight update:**  $W \leftarrow W - \eta \cdot (\partial L / \partial W)$

The five base-learner outputs (probabilities) are concatenated into a meta-feature vector and passed to a logistic-regression meta-learner:

Meta-feature vector:  $m(x) = [\hat{y}_{RF}(x), \hat{y}_{XGB}(x), \hat{y}_{LGBM}(x), \hat{y}_{CatB}(x), \hat{y}_{DNN}(x)]$

Stacked ensemble output:  $\hat{y}_{IMHAP}(x) = \sigma(w^T \cdot m(x) + b)$ ,  $\sigma(t) = 1 / (1 + e^{-t})$

Weighted-voting alternative:  $\hat{y}_{wv}(x) = \sum_h \alpha_h \cdot \hat{y}_h(x) / \sum_h \alpha_h$ ,  $\alpha_h \propto$  validation performance of base learner  $k$

#### Layer 5 — Explainable AI Module

Two complementary post-hoc explanation methods are applied to the stacked ensemble output:

SHAP value (Shapley additive explanation):  $\phi_i = \sum_{S \subseteq F \setminus \{i\}} \frac{|S|!(|F| - |S| - 1)!}{|F|!} \cdot [f(S \cup \{i\}) - f(S)]$

LIME local surrogate:  $g^* = \text{argmin}_g L(f, g, \pi_x) + \Omega(g)$ , where  $g$  is a locally fitted interpretable (linear) model

LIME offers a supplemental local linear approximation for sanity-checking SHAP outputs on individual instances, whereas SHAP offers global feature-importance ranking and per-prediction local explanation (which variables moved an individual's risk score up or down).

#### Layer 6 — Risk Assessment Module

A single, officially defined Mental Health Risk Index is created by combining model outputs from the pertinent psychological sub-scales (depression, anxiety, stress, sleep quality, and social interaction deficit):

$MHRI(x) = w_1 \cdot D(x) + w_2 \cdot A(x) + w_3 \cdot S(x) + w_4 \cdot (1 - \text{Sleep}(x)) + w_5 \cdot (1 - \text{Social}(x))$ ,  $\sum w_i = 1$

where Sleep/Social are normalised wellbeing sub-scores (higher = healthier) and  $D, A$ , and  $S \in [0, 1]$  are the calibrated estimated probabilities of stress, anxiety, and depression risk from Layer 4. Instead of being fixed at random, weights are adjusted on a validation set (e.g., by grid search subject to clinical-expert assessment). Four classified bands are mapped to the resulting continuous score:

| MHRI Range  | Risk Category | Suggested Action                                |
|-------------|---------------|-------------------------------------------------|
| 0.00 – 0.25 | Low Risk      | Routine monitoring; no immediate action         |
| 0.26 – 0.50 | Moderate Risk | Self-help resources; optional check-in          |
| 0.51 – 0.75 | High Risk     | Recommend counsellor outreach                   |
| 0.76 – 1.00 | Critical Risk | Priority referral to mental-health professional |

## IV. ALGORITHM DESIGN

Algorithm 1: Data Pre-processing

Input: raw dataset  $D$ . Output: cleaned, scaled dataset  $D'$ .

1. For each feature column, compute missing-value ratio; drop column if ratio  $> 0.6$ .
2. Impute remaining missing numeric values with column median; categorical with column mode.

3. Detect outliers via IQR rule per numeric column; cap (winsorize) at  $Q1 - 1.5 \cdot IQR / Q3 + 1.5 \cdot IQR$ .
4. Apply Min-Max or Z-score scaling to all numeric features.
5. One-hot / ordinal encode categorical features.
6. Return  $D'$ .
7. Algorithm 2: Feature Selection
8. Compute mutual information  $I(X_j; Y)$  for every feature  $j$ .
9. Compute pairwise Pearson correlation matrix  $R$  over features.
10. Rank features by  $I(X_j; Y)$  descending. Greedily select feature  $j$  if  $I(X_j; Y) > \tau$  and  $|R[j, k]| < 0.85$  for every already-selected  $k$ .
11. Return selected feature subset  $F'$ .

#### Algorithm 3: Hybrid Ensemble Training

12. Split data into train / validation / test (e.g., 70/15/15), stratified by target label.
13. Train RF, XGBoost, LightGBM, CatBoost, and DNN independently on the training split.
14. Generate out-of-fold predictions for each base learner via k-fold cross-validation (to avoid leakage into the meta-learner).
15. Concatenate out-of-fold predictions into meta-feature matrix  $M$ .
16. Train logistic-regression meta-learner on  $M$  against true labels.
17. Evaluate the full stacked pipeline on the held-out test split.

#### Algorithm 4: Mental Health Risk Prediction

18. For new input  $x$ , compute calibrated probabilities  $D(x)$ ,  $A(x)$ ,  $S(x)$  from the trained stacked ensemble (one ensemble per target sub-scale, or a multi-output head).
19. Compute  $MHRI(x)$  using the weighted formula in Section 3.6.
20. Map  $MHRI(x)$  to a categorical risk band.
21. Return ( $MHRI$  score, risk category, contributing-feature explanation from Algorithm 5).

#### Algorithm 5: Explainable Prediction Generation

22. For a given prediction, compute SHAP values  $\phi_i$  for every input feature using the trained ensemble (or its best-performing base learner, for tractability).
23. Rank features by  $|\phi_i|$  to produce a local explanation.
24. Aggregate  $|\phi_i|$  across the validation set to produce a global feature-importance ranking.

25. Optionally fit a LIME local surrogate for the same instance as a cross-check.

26. Return explanation object (top-k contributing features and their signed contribution).

## V. EXPERIMENTAL SETUP

### Hardware (typical / recommended)

- Multi-core CPU (e.g., Intel i7/i9 class) for tree-based ensemble training.
- CUDA-capable GPU (optional, recommended for the DNN component, e.g., for larger feature sets).
- $\geq 16$  GB RAM for the dataset sizes typical of survey-based mental-health studies (low thousands to tens of thousands of rows).

### Software

- Python 3.10+, with scikit-learn, XGBoost, LightGBM, CatBoost, PyTorch (or TensorFlow/Keras), SHAP, LIME.
- Stream lit for the optional prediction dashboard.
- Google Colab or local Jupiter for experimentation; see Section 20 for a runnable pipeline outline.

## VI. DATASET DESCRIPTION

IMHAP-Net is intended to be assessed using publicly accessible datasets from mental health surveys, as those provided in Table 2. Before using any dataset, exact row counts and feature sets should be verified against its current documentation, which is updated on a regular basis by its maintainers.

| Dataset                             | Approx. Scale                             | Key Features                                                                                                  | Target Task                           |
|-------------------------------------|-------------------------------------------|---------------------------------------------------------------------------------------------------------------|---------------------------------------|
| Student Depression Dataset (Kaggle) | ~27,900 self-reported records of students | Demographics, academic pressure, sleep duration, dietary habits, suicidal-ideation history, financial stress, | Binary depression risk classification |

|                                                                                  |                                                 |                                                                                                        |                                                       |
|----------------------------------------------------------------------------------|-------------------------------------------------|--------------------------------------------------------------------------------------------------------|-------------------------------------------------------|
|                                                                                  |                                                 | depression label                                                                                       |                                                       |
| Mental Health in Tech Survey / OSMI                                              | ~1,250 respondents (multi-year survey)          | Workplace mental-health attitudes, treatment history, family history, remote-work status, company size | Treatment-seeking / risk classification               |
| DASS-21 based university surveys (e.g., Lebanon, Bangladesh studies cited above) | Several hundred to ~1,700 respondents per study | DASS-21 subscale scores, demographics, lifestyle factors                                               | Depression / Anxiety / Stress severity classification |

## VII. EVALUATION METRICS

The following standard classification metrics, which are calculated on a held-out test split that the model was never exposed to during training or feature selection, should be used to assess the framework.

| Metric                                 | Formula                                                                                  |
|----------------------------------------|------------------------------------------------------------------------------------------|
| Accuracy                               | $(TP + TN) / (TP + TN + FP + FN)$                                                        |
| Precision                              | $TP / (TP + FP)$                                                                         |
| Recall (Sensitivity)                   | $TP / (TP + FN)$                                                                         |
| Specificity                            | $TN / (TN + FP)$                                                                         |
| F1-Score                               | $2 \cdot (Precision \cdot Recall) / (Precision + Recall)$                                |
| ROC-AUC                                | Area under the curve of TPR vs. FPR across thresholds                                    |
| Matthews Correlation Coefficient (MCC) | $(TP \cdot TN - FP \cdot FN) / \sqrt{[(TP+FP)(TP+FN)(TN+FP)(TN+FN)]}$                    |
| Cohen's Kappa                          | $(p_o - p_e) / (1 - p_e)$ , where $p_o$ = observed agreement, $p_e$ = expected agreement |

## VIII. RESULTS AND DISCUSSION

This table is only a template for reporting. Do not fill this table with presumptive or desired values;

instead, each field must be filled with numbers that were really generated by executing the pipeline in Section 20 on a real dataset and a true train/test split.

| Model                                 | Accuracy | F1-Score | ROC-AUC | MCC |
|---------------------------------------|----------|----------|---------|-----|
| Logistic Regression                   | TBD      | TBD      | TBD     | TBD |
| SVM                                   | TBD      | TBD      | TBD     | TBD |
| KNN                                   | TBD      | TBD      | TBD     | TBD |
| Random Forest                         | TBD      | TBD      | TBD     | TBD |
| XGBoost                               | TBD      | TBD      | TBD     | TBD |
| LightGBM                              | TBD      | TBD      | TBD     | TBD |
| Deep Neural Network                   | TBD      | TBD      | TBD     | TBD |
| Proposed IMHAP-Net (stacked ensemble) | TBD      | TBD      | TBD     | TBD |

Discussion (to be written after Table 3 is filled in): compare the stacked ensemble to each individual base learner and the traditional baselines (logistic regression, SVM, KNN); talk about which feature categories (academic, lifestyle, social) contributed the most using the SHAP rankings; and talk about any performance differences between subgroups (e.g., gender, year of study) to ensure fairness before making any deployment claims.

## IX. EXPLAIN ABILITY ANALYSIS

After training, the framework should produce: (i) a global SHAP summary plot that ranks features by mean |SHAP value| across the validation set; (ii) SHAP force plots for a few representative individual predictions (such as one high-risk and one low-risk case) to show the reasoning behind the model; and (iii) a comparison between the top-5 LIME features and the top-5 SHAP features for the same cases to verify consistency between the two explanation methods. It is not possible to properly write these outputs beforehand; they must be produced from the actual trained model.

## X. STATISTICAL VALIDATION

Once cross-validation data are obtained for each model, the following tests are advised to determine whether changes across models are statistically significant rather than the result of chance:

- A paired t-test on per-fold metric scores comparing each individual baseline and the suggested stacked ensemble.
- To check for a general difference in means, perform a one-way ANOVA over the cross-validation scores of all candidate models.
- If fold-level scores are not normally distributed, the Wilcoxon signed-rank test is a helpful non-parametric substitute for the paired t-test.

Report exact p-values and effect sizes from these tests once computed; do not state significance without having run the test.

## XI. ADVANTAGES OF THE PROPOSED FRAMEWORK

- Compared to a single model, stacking several heterogeneous base learners is typically more resilient to dataset peculiarities.
- Explain ability encourages interpretable deployment and is a first-class layer (Layer 5), not a post-hoc add-on.
- Instead of using an opaque label, the MHRI formulation provides counsellors with a single, auditable score with a transparent formula.
- Any base learner, feature-selection technique, or explanation approach may be substituted without completely rewriting the pipeline thanks to the layered design's modularity.

## XII. LIMITATIONS

- Self-report bias and social desirability bias can affect survey-based mental health datasets.
- Clinical/at-risk labels frequently exhibit class imbalance, which needs to be specifically handled (e.g., class weighting, stratified sampling) rather than disregarded. Why Without re-validation, models trained on one population (such as the student population of one nation) might not generalise to another.

- Predictive risk scores are instruments for decision-making rather than diagnosis; competent experts and clinical judgement are still required.

- Before gathering or using any actual student mental health data, privacy and data-governance criteria (such as institutional ethical approval and data anonymisation) must be met.

## XIII. FUTURE WORK

- Federated learning amongst universities to enhance generalisation while retaining sensitive information locally.
- Transformer-based systems for behavioural data that is sequential or longitudinal (such as daily check-ins throughout a semester).
- Multimodal fusion, with the proper authorisation, integrating survey data with passive signals (app usage, wearable sleep/activity data).
- Validation studies that are prospective rather than merely retrospective, ideally with clinician-in-the-loop assessment of the risk categories.

## XIV. CONCLUSION

IMHAP-Net, a six-layer machine learning system for assessing mental health risk, was introduced in this study. It consists of a transparent, formula-based Mental Health Risk Index, a stacked heterogeneous ensemble, and a special explain ability layer. In contrast to methods that provide a single accuracy statistic for an opaque classifier, IMHAP-Net is built such that each risk category is produced from an auditable formula rather than a black box, and every prediction can be linked to contributing factors using SHAP/LIME.

The algorithms, mathematical theory, and architecture described here are finished and prepared for use; a runnable pipeline overview is given in Section 20. In order to avoid misrepresenting the work, the empirical results, explain ability outputs, and statistical validation in Sections 8–10 are purposefully left as templates to be filled when the framework is run on a particular, ethically sourced dataset. Prior to any practical implementation in a counselling or therapeutic

environment, future development should give priority to prospective validation with clinical input and explicit fairness audits.

## REFERENCES

1. L. Breiman, "Random forests," *Machine Learning*, vol. 45, no. 1, pp. 5–32, 2001.
2. T. Chen and C. Guestrin, "XGBoost: A scalable tree boosting system," in *Proc. 22nd ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining*, 2016, pp. 785–794.
3. G. Ke et al., "LightGBM: A highly efficient gradient boosting decision tree," in *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 30, 2017.
4. A. V. Drogosh, V. Ershov, and A. Gulin, "CatBoost: Gradient boosting with categorical features support," *arXiv preprint arXiv:1810.11363*, 2018.
5. S. M. Lundberg and S.-I. Lee, "A unified approach to interpreting model predictions," in *Advances in Neural Information Processing Systems (NeurIPS)*, vol. 30, 2017.
6. M. T. Ribeiro, S. Singh, and C. Guestrin, "'Why should I trust you?' Explaining the predictions of any classifier," in *Proc. 22nd ACM SIGKDD Int. Conf. Knowledge Discovery and Data Mining*, 2016, pp. 1135–1144.
7. Y. Zhai et al., "Machine learning predictive models to guide prevention and intervention allocation for anxiety and depressive disorders among college students," *Journal of Counseling & Development*, 2025.
8. B. L. Schaab et al., "How do machine learning models perform in the detection of depression, anxiety, and stress among undergraduate students? A systematic review," *Cadernos de Saúde Pública*, vol. 40, pp. 1–26, 2024, doi: 10.1590/0102-311XEN029323.
9. "Machine learning based prediction of high school student mental health," *medRxiv preprint*, Aug. 2025.
10. "Predictive machine learning models for assessing Lebanese university students' depression, anxiety, and stress during COVID-19," published online, *PMC*, 2024.
11. "Prevalence, associated factors, and machine learning-based prediction of depression, anxiety, and stress among university students: a cross-sectional study from Bangladesh," *Journal of Health, Population and Nutrition*, Springer Nature, 2025.
12. "Predicting and understanding college student mental health with interpretable machine learning," *arXiv preprint arXiv:2503.08002*, 2025.
13. "Machine learning techniques to predict mental health diagnoses: A systematic literature review," *PMC*, 2024.
14. S. Aleem, N. Huda, R. Amin, S. Khalid, S. S. Alshamrani, and A. Alshehri, "Machine learning algorithms for depression: Diagnosis, insights, and research directions," *Electronics*, vol. 11, no. 7, p. 1111, 2022.
15. A. Priya, S. Garg, and N. P. Tigga, "Predicting anxiety, depression and stress in modern life using machine learning algorithms," *Procedia Computer Science*, vol. 167, pp. 1258–1267, 2020.