

Deep Learning for Early Detection of Depression in University Students: A Mental Health Perspective

By
Mst. Fariha Akter
ID: 211-15-3978

FINAL YEAR DESIGN PROJECT REPORT

This Report Presented in Partial Fulfillment of the Requirements for the Degree of Bachelor of Science in Computer Science and Engineering

Supervised by

Mr. Partha Dip Sarkar
Lecturer
Department of Computer Science and Engineering
Daffodil International University

Co-Supervised by

Mr. Husne Mubarak
Lecturer
Department of Computer Science and Engineering
Daffodil International University



DAFFODIL INTERNATIONAL UNIVERSITY
Dhaka, Bangladesh

September 17, 2025

APPROVAL

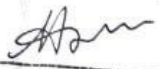
This Project titled **"Deep Learning for Early Detection of Depression in University Students: A Mental Health Perspective"**, submitted by **Mst. Fariha Akter**, ID No: **211-15-3978** to the Department of Computer Science and Engineering, Daffodil International University has been accepted as satisfactory for the partial fulfillment of the requirements for the degree of B.Sc. in Computer Science and Engineering and approved as to its style and contents. The presentation has been held on **17 September, 2025**.

BOARD OF EXAMINERS



Dr. S.M Aminul Haque (SMAH)
Professor and Associate Head
Department of Computer Science and Engineering
Faculty of Science & Information Technology
Daffodil International University

Chairman



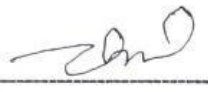
Ms. Nazmun Nessa Moon (NNM)
Associate Professor
Department of Computer Science and Engineering
Faculty of Science & Information Technology
Daffodil International University

Internal Examiner



Md. Abbas Ali Khan (AAK)
Assistant Professor
Department of Computer Science and Engineering
Faculty of Science & Information Technology
Daffodil International University

Internal Examiner



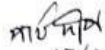
Dr. Md. Zulfiker Mahmud (ZM)
Professor
Department of Computer Science and Engineering
Jagannath University

External Examiner

DECLARATION

We hereby declare that this project has been done by us under the supervision of Mr. Partha Dip Sarkar, Lecturer Department of Computer Science and Engineering, Daffodil International University. We also declare that neither this project nor any part of this project has been submitted elsewhere for the award of any degree or diploma.

Supervised by:

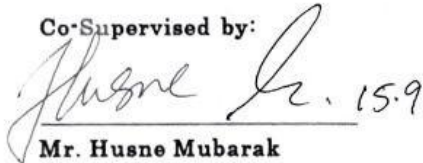

15/09/25

Mr. Partha Dip Sarkar

Lecturer

Department of Computer Science and
Engineering Daffodil International
University

Co-Supervised by:

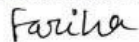

15.9

Mr. Husne Mubarak

Lecturer

Department of Computer Science and
Engineering Daffodil International
University

Submitted by:


15-09-25

Mst. Fariha Akter

Student ID: 211-15-3978

Department of Computer Science and
Engineering
Daffodil International University

ACKNOWLEDGEMENTS

This work would not have been possible without the support and contributions of many individuals over the past two semesters. I am deeply grateful to everyone who has assisted us in one way or another.

First, i express our heartfelt thanks and gratefulness to the almighty for His divine blessing making it possible for us to complete the Final Year Design Project (FYDP) successfully.

I am grateful and wish our profound indebtedness to Mr. Partha Dip Sarkar, Lecturer, Department of Computer Science and Engineering, Daffodil International University, Dhaka, Bangladesh. Deep knowledge and keen interest of my supervisor in the field of Machine learning to carry out this project. His endless patience, scholarly guidance, continual encouragement, constant and energetic supervision, constructive criticism, valuable advice, reading many inferior drafts, and correcting them at all stages have made it possible to complete this project.

I would like to express our heartfelt gratitude to the Head of the Department of Computer Science and Engineering, for his kind help in finishing our project and also to other faculty members and the staff of the Department of Computer Science and Engineering, Daffodil International University.

I would like to thank our entire course-mates at Daffodil International University, who took part in this discussion while completing the coursework.

Finally, I must acknowledge with due respect the constant support and patience of our parents.

ABSTRACT

Nowadays, depression has spread at an alarming rate among students. There are many reasons for this such as lack of sleep, academic decline, loneliness, limited mental support, personal, family or financial problems, etc. In this thesis, I used 1977 datasets where there were 15 universities. Out of which 9 were public and 6 were private. The data was surveyed using PHQ-9 method. It has 6 levels. My proposed method is ANN- artificial neural network. Through which I got 97.98% accuracy. my dataset is a tabular form such as mixed with demographic, academic, phycological. Poor concentration, lack of sleep, and feelings of low self-worth were more influential predictors of depression than demographic and academic information. For better understanding, i used explainable ai (LIME), which full form is Local Interpretable Model Agnostic Explanation. with the help of this method, we can easily understand the reason of the depression which is causing more effectively. I also preprocess the dataset for cleaning the data and understanding for human language to machine language like labeling hot encoding etc. I also used PCA-IG feature selection method for feature selection and changed the hyperparameter of each model for knowing which one is more suitable and will give us more accuracy. The results can help shape targeted mental health programs, guide policy decisions, and inspire future research that tracks students over time and involves more universities—helping address mental health needs in settings where resources are limited.

Keywords : Depression, Deep Learning, Explainable AI, University Students, Bangladesh, LIME

Table of Contents

Approval	i
Declaration	ii
Acknowledgements	iii
Abstract	iv
List of Figures	vii
List of Tables	viii
1 Introduction	1
1.1 Introduction.....	1
1.2 Motivation	2
1.3 Objectives	2
1.4 Methodology	3
1.5 Project Outcome.....	3
1.6 Organization of the Report	4
2 Background	2
2.1 Introduction.....	5
2.2 Literature Review	6
2.2.1 Related Research.....	8
2.3 Gap Analysis	11
2.4 Summary	12
3 Research Methodology	3
3.1 Methodology	13
3.1.1 Overview	13
3.1.2 Proposed Methodology	14
3.2 Detailed Methodology and Design.....	15
3.2.1 Data Collection & Preprocessing.....	15
3.2.2 Feature Selection.....	16
3.2.3 Machine Learning and Deep Learning Model Selection.....	17
3.2.4 Hyperparameter Tuning.....	20
3.2.5 Evaluation Matrices.....	20
3.2.6 Use of Explainable AI.....	21
3.3 Project Plan	22

3.4	Task Allocation.....	23
3.5	Summary	23
4	Implementation and Results	4
4.1	Environment Setup	24
4.2	Comparative Analysis	25
4.3	Results and Discussion	28
4.4	Summary	32
5	Engineering Standards and Design Challenges	5
5.1	Compliance with the Standards.....	33
5.1.1	Software Standards.....	33
5.1.2	Hardware Standards	33
5.1.3	Communication Standards.....	33
5.2	Impact on Society, Environment and Sustainability	34
5.2.1	Impact on Life.....	34
5.2.2	Impact on Society & Environment.....	34
5.2.3	Ethical Aspects	34
5.2.4	Sustainability Plan.....	35
5.3	Project Management and Financial Analysis.....	35
5.4	Complex Engineering Problem.....	35
5.4.1	Complex Problem Solving.....	35
5.4.2	Engineering Activities.....	36
5.5	Summary	37
6	Conclusion	10
6.1	Summary	38
6.2	Limitation	38
6.3	Future Work	39
	References	40

List of Figures

Figure 3.1	Proposed Methodology for our research study.....	14
Figure 3.2	Dataset.....	15
Figure 3.3	Feature Importance Plots for PCA-IG Based Model.....	16
Figure 4.1	Confusion matrix of the performance of our proposed ANN model.....	27
Figure 4.2	ROC curve demonstrating the performance of our proposal model.....	28
Figure 4.3	Correlation heatmap for Our Study's Dataset.....	29
Figure 4.4	LIME Tabular Plot.....	30
Figure 4.5	LIME Feature Importance Plot.....	30

List of Tables

Table 2.1	Summary of Literature Reviewed.....	06
Table 3.1	Hyperparameter Settings of All ML Models.....	20
Table 3.2	Project plan.....	22
Table 3.3	Task allocation.....	23
Table 4.1	Model Result before preprocessing.....	24
Table 4.2	Models Result After Preprocessing.....	26
Table 5.1	Budget Plan (within BDT 50,000).....	35
Table 5.2	Mapping with complex engineering problem.....	35
Table 5.3	Mapping with knowledge Profile.....	36
Table 5.4	Mapping with complex engineering activities.....	36

Chapter 1

Introduction

This chapter explores the connection between depression and Bangladeshi university students using deep learning and explainable AI. I uncover meaningful patterns through machine learning, while XAI helps make these results easier to understand. The chapter also covers the study's main goals and what the findings could mean for students' academic success and personal well-being.

1.1 Introduction

Depression has become a serious global problem with devastating effects on academic performance, individual health, and eventual psychiatric health. For the case of Bangladesh, the problem takes a special precedence because students are experiencing additional sources of stress such as high competition at a professional level, social pressures from culture and society as a whole, and limited access to psychiatric attention [1] [2]. All these pressures push young students towards a risk of depression. Social stigma and ineffective screening tools also lead to a delay of early recognition and consequently many students are left to their own resources to battle their condition [3]. Some of these features like sleep problems, and inability of focusing their academic performance as well as their quality of life. Some of research using advanced computer methods to find early signs of depression in university students in Bangladesh so that they can provide mental health support. [4]. Some software can help detect with depressive symptoms by examining a variety of types of information such as attention of students, their grades, and mental health situation [5]. With the help of deep learning we can detect easily but there is no explanation. So, we used explainable AI (XAI) methods are applied. Algorithms such as Local Interpretable Model-agnostic Explanations (LIME) assist with explaining why models making any result ,they explain all with mental health [6] [7]. The study utilizes data from 1,977 students from 15 major private and government universities of Bangladesh. Basic information such as age, gender, university and department attended, study year, and CGPA are included. Also included are depressive symptoms such as sleep issues, low self-esteem, and difficulty concentrating. The aim of the study is to identify relevant factors pertaining to the depression of students through the analysis of this extensive dataset. With the expected results assisting the development of improved mental health programs and revealing how technology can positively impact mental health care tremendously, the study utilizes deep learning techniques to create a predictor of depression level from university students of Bangladesh with a target of increasing awareness of early signs and providing rapid relief. Ensuring correct results and simplicity of understanding, the model utilizes explainable AI techniques offering insight into how variations of the differing factors will impact depression risk. With this potentially resulting in improved mental health programs at university and improved compassionate learning environment, it can encourage a higher level of support and a better learning environment.

1.2 Motivation

Depression among university students growing day by day. The rate of depression is increasing rapidly. There are lots of reason like low self-confidence, self-worth, low sleep, academic or social pressure, lacking of mental support which harming their well-being also causing a student choose death or suicide over fighting. So here my study will help them for early detection of depression that can help them for saving their life. Also, with the help of ANN custom model with 1977 complexed dataset from 15 different top university with PHQ-9 method, using PCA-IG feature selection and Explainable ai that can make us understand more fluently. My aim to help student reduce mental pain also both academically and phytologically.

1.3 Objectives

The main purpose of my study is to detect depression early and can prevent easily. My goal to train and create a customize model with changing the hyper parameter for detect depression with the help of 1977 dataset which includes academic, demographic, phycological information. My thesis aims to provide significant protection from depression, help univarsity students growing well and lead a better academic and mental health with the proposed customized model ANN - Artificial neural network, PCA-IG method and explainable ai (LIME). Here is some following objective of this thesis:

- To create a trustworthy deep learning model that can accurately predict depression levels among Bangladeshi university students, allowing for early detection and timely intervention.
- To identify the primary characteristics associated with depression by analyzing data from 1,977 students at 15 leading universities, including demographic, academic, and psychological features.
- To increase the predictive model's clarity and utility by using explainable AI (XAI) methodologies like LIME, so that the results are transparent and provide practical assistance for dealing with student mental health issues.
- To provide data-driven recommendations to inform university mental health policies, improve resource allocation, and promote student well-being in academic environments.

1.4 Methodology

My thesis is to detect depression in University student of Bangladesh using deep learning and explained ai. That start with gathering dataset from University student likely 1977 dataset which belong from 15 University where 9 is public and 6 is from private and it have academic, demographic, phycological information of each student. During preprocessing, missing data is handled, and categorical information is converted into a usable format through techniques like label and one-hot encoding. The dataset is further improved through feature engineering. To focus on the most important factors, like sleep problems and trouble concentrating—the study uses a combined method called PCA–IG for feature selection, which reduces unnecessary data and highlights key predictors. Various machine learning and deep learning models, including Random Forest, Support Vector Classifier, K-Nearest Neighbors, LightGBM, Logistic Regression, and an Artificial Neural Network, are then carefully tuned to deliver accurate and reliable predictions of depression severity.

To make the models easier to understand, I use explainable AI methods, especially LIME, which help show exactly how different features impact the predictions. Our process follows a clear step-by-step approach, starting with collecting and cleaning the data, then training the models, fine-tuning their settings, and finally evaluating their performance using measures like accuracy, precision, recall, F1-score, and AUC. I split the data into training (80%), validation (10%), and testing (10%) sets to build and test the models fairly. By combining powerful computational techniques with explainable AI, this approach not only aims for accurate predictions but also makes sure the results are clear and useful for people like educators and mental health professionals, helping them intervene early and support student well-being.

1.5 Project Outcome

This study sets out to create a highly accurate and easy-to-understand deep learning model that predicts depression levels among university students in Bangladesh. The research examines data from numerous universities and identifies key features that significantly impact students' mental health, such as difficulties with attention, sleep disorders, and low self-esteem. With a current machine learning approach known as an Artificial Neural Network (ANN), which reached more than 98% accuracy with careful data preparation, the model can classify depression into a scale of five from mild to severe. The powerful tool enables universities and mental health clinicians to identify at-risk students at an early stage and offer timely assistance beneficial both to academic performance and overall well-being.

The major contribution of the current work is using explainable AI (XAI) methods, namely LIME, in order to shed light on and justify the model's predictions. Interesting information for counselors, teachers, and politicians is offered through an example of how personal variables, e.g., symptoms of mental health or academic pressure, elevate the risk of depression. The dataset indicates that psychological symptoms often have a higher impact compared with demographical or academic ones, allowing mental health programs to aim at high-impact issues. It further increases the transparency of the model's results beyond raising confidence in the results through making the results more actionable in practice.

It has the potential of truly making a difference to the delivery of mental health service, not just at Bangladeshi universities but beyond. With the use of cloud-computing infrastructures like Google Colab, the model becomes scalable and economical and therefore can be adjusted and used with a number of datasets and practice settings and

ensuring it remains useful with the course of time. The insights from this study can help shape university policies, improve how mental health resources are distributed, and reduce stigma by raising awareness through clear, data-driven evidence. In the end, this work aims to create a healthier, more supportive academic environment that enhances students' well-being and tackles a pressing public health issue, especially in places where resources are limited.

1.6 Organization of the Report

This report is organized into several chapters, each designed to guide the reader through the various stages of the research.

Chapter 1 opens the study by clearly explaining the research problem, objectives, and the methods used to explore how depression impacts university students. It sets the report's context as an overarching one and emphasizes the importance of this report and the potential of enhancing the well-being of students.

Chapter 2 presents a clear summary of current research with a summary of notable studies regarding depression and the impact of depression on students' academic performance. It paves the way for the research by referring to notable findings of prior research and indicating how the current research contributes towards and beyond prior research with new ideas.

Chapter 3 shows the research main work or methodology. It shows how I collect the data, how I created the survey questionnaire, and how I used deep learning methods. Also explained every detail information about the cleaning and processing and how results were more understandable of explainable AI (XAI) tools

Chapter 4 shows the Results and Discussion. It reports the fundamental findings of the deep learning and machine learning models. It analyzes the correlations of each of the factors with depression. It also places the findings of this research in comparison with previous research and highlights their implications and surprises.

Chapter 5 presents a clear picture of the project. It discusses technical standards and design issues encountered during the project. It refers to issues of gathering data, applying models in practice, and applying machine learning to real life for the prediction of depression. The chapter also discusses how such issues were resolved and presents useful lessons and directions of future research in this direction.

Lastly, Chapter 6 concludes the paper with a summary of findings and discussion of their greater significance. It dwells upon methods of further research and how the research may be utilized in practical life as well as reducing depression among university students. Generally, the report has a coherent and systematic order designed to lead the readers from the introduction through the chapters and end with the last suggestions and remarks.

Chapter 2

Background

In this chapter I have shown the review of the existing literature on depression, highlighting how it affects students' overall well-being. I also compare with the previous research on the significance of mental health also find out all the key point of my reviewed literature and noted at how machine learning and deep learning techniques have been applied in similar studies for better understanding.

2.1 Introduction

Depression in university students is an increasing global issue. It affects grades as well as health and future mental health. It is a particularly challenging issue in Bangladesh because of high academic pressures and social stresses and a lack of adequate access to mental health support. Many students endure mental health issues in silence due to stigma. These can cause severe symptoms such as persistent depression, sleep issues, and difficulties concentrating that may not become apparent until they severely impact daily life. Fortunately, innovative developments of machine learning and deep learning can assist in discovering early indicators of depression through examining intricate data. These can assist in compensating for shortcomings of mental health screenings. These literature reviews examine prior literature regarding depression among university students with an indication of how emerging technologies were utilized to discover and predict depression in countries like Bangladesh.

Research on mental health is increasingly relying on explainable AI (XAI) methods such as Local Interpretable Model-agnostic Explanations (LIME). It helps make sure that the models predicting depression are not just accurate but also easy to understand for people like teachers and healthcare workers. Many past studies have used machine learning methods such as Random Forest, Support Vector Machines, and Neural Networks to predict depression with pretty good accuracy—often over 80%. However, many of these studies don't focus enough on making their results easy to interpret or tailored to specific cultural contexts. In Bangladesh, where mental health support is limited, researchers are starting to look at how these technologies can help overcome issues like stigma and lack of resources. This review brings together important findings from both global and local studies, points out gaps like the need for more real-world testing and attention to cultural differences, and explains how the current research aims to improve early depression detection for university students in Bangladesh.

2.2 Literature Review

A literature review is basically a thorough look at all the important research and studies done on a specific topic. It helps summarize key ideas, methods, and any missing pieces in the existing knowledge. By reviewing these earlier works, researchers can make sure their own study builds on what's already known while also bringing something new and valuable to the field. In short, it's a way to see how their work fits into the bigger picture.

Table 2.1: Summary of Literature Reviewed.

Author (s)	Year	Title	Methodology	Key Findings
Lin Luo et al.	2025	Predictors of depression among Chinese college students: a machine learning approach	Random Forest Algorithm (RFA)	Suicidal ideation, anxiety, and sleep quality are top predictors; gender differences in risk patterns observed.
Nicola Meda	2023	Frequency and machine learning predictors of severe depressive symptoms and suicidal ideation	Multiple regression, Random Forest	Economic worry strongly linked to depression; demographic factors poorly predicted worsening mental health over time.
Md. Iqbal Hossain Nayan	2022	Comparison of ML algorithms for predicting depression and anxiety among Bangladeshi students	Logistic Regression, RF, SVM, etc.	RF best predicted depression (89% accuracy), SVM best for anxiety (91.49% accuracy); higher prevalence among women.
Minji Gil	2022	ML models predicting depression risk and family/individual factors in Korean students	Sparse Logistic Regression, SVM, RF	RF showed best performance; self-perceived mental health and family factors significant predictors.
Qi Qiang	2023	Identifying risk factors for depression changes in college students using ML	Logistic Regression, RF, SVM, KNN	SVM performed best; baseline depression and parents' emotional expression important predictors for depression changes.
Christo El Morr	2023	Predictive ML models for depression, anxiety, and stress among Lebanese university students	LR, MLP, SVM, RF, XGBoost, AdaBoost, etc.	RF, NB, and AdaBoost were top performers; self-rated health and age key predictors.

Mohammed A. Mamun	2024	Predictive modeling of comorbid depression and anxiety using GIS and ML	XGBoost, GBM, CatBoost, GIS analyses	Past-year suicidal ideation strongest predictor; spatial disparities identified; high prevalence of comorbidity.
Nasirul Mumenin	2021	Screening depression among students using GHQ-12 and ML	16 ML models including Extremely Randomized Trees	ET model achieved 90.26% accuracy; 60% prevalence of depression; complex interplay of socio-demographic and stress factors.

2.2.1 Related Research

Luo et al. [8] used Random Forest method for detecting depression. It shows us comprehensive cross-sectional study including over 10,000 Chinese college students. They show us that the best reason is psychological variables, including anxiety, suicidal thoughts, and sleep quality. Academic stress, body mass index (BMI), physical fitness, and social media etc. This shows us substantial gender differences: BMI was more significant for female students, while physical fitness had a bigger effect on depression risk for male students. The mental health problem for college students, as seen by its accuracy of 87.5% and AUC of 0.927, also with 1,388 participants, Meda et al. [9] found out severe depression symptoms and suicide thoughts in university students.

Using multiple regression and machine learning techniques, they found that about 20% of students experienced severe mental health issues. Economic worry was strongly linked to depression both at baseline and follow-up. While their random forest model accurately predicted students maintaining well-being, it struggled to predict worsening symptoms. The study highlighted that demographic factors were poor predictors of mental health outcomes, emphasizing the need for further research incorporating lived experiences to better identify students at risk.

Hossain et al. [10] conducted a study during the first wave of the COVID-19 pandemic to predict depression and anxiety among Bangladeshi university students using multiple machine learning algorithms. Surveying 2,121 students, the study found that severe depression and anxiety were more prevalent among women. Among the models tested, Random Forest achieved the highest accuracy (89%) for predicting depression, while Support Vector Machine performed best for anxiety prediction with 91.49% accuracy. The study recommends using Random Forest and SVM algorithms for effective mental health prediction in this population.

Gil et al. [11] used machine learning to predict depression risk in Korean college students, finding that random forest performed best. Key factors included students' self-perceived mental health, neuroticism, attachment style, family cohesion, and parents' mental health. The study highlights ML's potential to improve early depression detection and support. Another author [12] conducted a systematic review evaluating machine learning models used to detect depression, anxiety, and stress among undergraduate students. Analyzing 48 studies, they found that most models demonstrated strong accuracy—often above 70%—with depression detection reaching up to 99.1%. However, nearly all studies relied on internal validation, raising concerns about overestimated performance. The review highlights the promise of machine learning in mental health assessment but calls for more external validation to ensure real-world applicability.

Qiang et al. [13] used machine learning to predict changes in depression levels among college students, analyzing psychological and demographic factors collected over two time points. Their study found that support vector machines (SVM) performed best, achieving accuracies of 89.4% for predicting negative changes and 91.9% for positive changes in depression. Parents' emotional expressions and baseline depression levels were important predictors, underscoring the impact of family dynamics on students' mental health

El Morr et al. used machine learning algorithms [14] that shows us stress, anxiety, and depression in Lebanese university students during the COVID-19 epidemic. In this research they told, with AUC values ranging from 72.96% to 78.27%, Random Forest, Naïve Bayes, and AdaBoost were the best at predicting stress, anxiety, and depression. Self-rated health for depression and age for stress and anxiety were significant predictors. From that we can understand machine learning can give us more accuracy with the help of strategies like data augmentation and more comprehensive classification approaches will use in future studies to detect depression.

Mamun et al. shows us depression and anxiety among prospective university students using a combination of GIS and machine learning techniques. [15] Analyzing data from 1,485 participants, they identified key risk factors including female gender, living in urban areas, family structure, and past-year suicidal thoughts. Machine learning models such as XGBoost and GBM showed strong predictive performance, with past-year suicidal ideation emerging as the most important predictor. GIS analysis shows the differences in mental health aspects, policies to better support students' mental well-being.

Mumenin et al. [16] also published a paper with Bangladeshi university students for depression using machine learning. they also used General Health Questionnaire-12 (GHQ-12). It shows almost 60% of students suffered from depression after analysing data from 804 individuals and the accuracy of 90.26%, the Extremely Randomized Tree (ET) method performed well from other 16 models. This paper shows how social information and work-related stress impact students' mental health and shows machine learning method is more effective for detecting depression and anxiety.

Iparraguirre-Villanueva et al. [17] used machine learning for predicting depression such as logistic regression, K-nearest neighbor, and decision tree etc. they analyses from 787 people, they discovered that logistic regression was the most effective. It's accuracy rate of 77%.it have level of depression like the high prevalence of moderate depression, the low rates of treatment seeking, and depression, especially among male students. From this research we can understand that machine learning is more effective for detecting depression with accurate result.

Mumenin et al. [18]proposed DDNet, a robust hybrid machine learning model combining multiple classifiers and a meta-learner to detect depression among university students. They tested in two datasets, after using DDNet, they got accuracies near 99%. The study emphasized model interpretability using SHAP and reliability through uncertainty estimation with Monte Carlo Dropout, highlighting DDNet's potential as an effective, transparent, and trustworthy tool for early depression detection in educational environments.

Siraji et al. [19] studied depression detection among 444 Bangladeshi university students using eight machine learning models. Feature selection and extraction methods like RFE and PCA improved accuracy by 3–15%. The best results came from SparsePCA with CatBoost, effectively classifying depression levels. They found 44% had no depression, 25% mild-to-moderate, and 31% severe-to-extreme symptoms.

Chowdhury et al. [20] examined anxiety, depression, and insomnia among 1,250 Bangladeshi university students using tree-based machine learning models like Decision Tree, Random Forest, and XGBoost. The study found high prevalence rates: 83.3% for moderate to severe anxiety, 84.7% for depression, and 76.5% for insomnia. Key risk factors included social media addiction, age, academic performance, smoking, family income, and

chronotype. With the strongest predictive performance, the XGBoost model demonstrated its promise for university settings' mental health screening and intervention planning.

Using machine learning classification techniques, Munir et al. [21] investigated depression prediction among university students in Bangladesh. They used backward elimination and Pearson's chi-squared test to identify important predictive variables after gathering data from social media questionnaires. They tested six distinct machine learning classifiers and most accuracy giver model was stacking classifier that used 24 features. This study shows Bangladeshi mental health research and explain how machine learning can be used to detect depression. Siddiqua et al. [22] used a phycological questionnaire with both personal and clinical questions to create an AI-based depression method for university students in Bangladesh.it shows us, they divided depression into three categories—normal, moderate, and severe—using ten machines learning and two deep learning models. Their Random Forest model outperformed CNN and Gradient Boosting models and accuracy of 98.08%. They also used explainable ai for explanations and make it more trustworthy but those are limited.

This paper shows using machine learning models, such as logistic regression, decision trees, random forests, neural networks, K-Nearest Neighbor, and a hybrid stacking model, Rahman et al. [23] shows us the detection of the stages of anxiety and depression among Bangladeshi university students. The hybrid model, it have five classifiers using Linear Discriminant Analysis as a meta-classifier, every models, achieving 99% accuracy for depression prediction and 97% accuracy for anxiety prediction. Siddique et al. [24]used machine learning models to predict depression among Bangladeshi university students, comparing algorithms like logistic regression, KNN, and random forest. They also developed an Android app to support students' mental health, showing how technology can aid early detection and intervention.

2.3 Gap Analysis

Despite numerous studies applying machine learning (ML) and deep learning (DL) techniques to detect depression and related mental health issues among university students, several critical gaps persist. Most existing research focuses heavily on achieving high predictive accuracy using various ML algorithms but often overlooks the integration of explainability and interpretability. Although some studies incorporate explainable AI methods to provide model transparency, a standardized and consistent approach to ensure trustworthiness and actionable insights for mental health professionals remains lacking. This gap limits the practical adoption of these models in real-world settings where understanding *why* a prediction was made is crucial.

Another significant gap is that there is insufficient attention dedicated to social and cultural issues special to developing country students. Not many studies attend closely to the influence of social norms, cultural stigma, and restricted access to mental health resources on depression and prediction model performance, although many attend closely to demographics and academic characteristics. Utilization of the tools may vary across communities due to a lack of locally congruent culture models. Data-based models may become more transferable and impactful at a real-world level if they integrate qualitative understanding and attend closely to the local context.

There are also problems like differing features, class imbalances, and low-quality data that are not properly addressed in most works. Even though most of the feature selection and engineering techniques are established and ample datasets are provided, often there isn't a proper comparison and description of methods. Moreover, studies scarcely use external or longitudinal validation; it often doesn't use internal validation. But this raises the question of how the models perform under differing circumstances and longitudinally and this is quite important because students' risk factors and mental health symptoms are always changing.

Current research often investigates anxiety and depression separately or concurrently and that itself is a problem. Comorbid conditions of stress, insomnia, and suicidal behaviors are often studied separately and are paid less attention. Few studies adopt a thorough methodology that combines several aspects of mental health into a single framework for prediction. By doing this, more comprehensive insights may be provided, and more successful interventions to promote students' general wellbeing may be created. While some research includes the development of mobile applications or cloud-based platforms, many stop short of addressing usability, scalability, and ethical considerations such as privacy and data security. Moreover, user acceptance and stakeholder involvement in the design and implementation of AI-driven mental health tools are rarely examined. Addressing these issues is essential to bridge the gap between theoretical model development and effective real-world application in university mental health support systems.

2.3 Summary

The literature review reveals that machine learning and deep learning models, such as Random Forest and Artificial Neural Networks, achieve high accuracy (often over 80%) in detecting depression among university students, identifying key predictors like suicidal ideation and sleep issues. Bangladeshi studies highlight gender and socio-economic factors, with some models reaching near-perfect accuracy. However, gaps include limited model interpretability, lack of external validation, and insufficient focus on socio-cultural contexts, scalability, and related mental health conditions, emphasizing the need for transparent, culturally sensitive, and scalable solutions for early depression detection in resource-limited settings.

Chapter 3

Research Methodology

This chapter outlines the research strategy, methodology, and procedures employed to investigate the relationship between depression and overall mental health among university students. It describes the methods used for data collecting, analysis, and machine learning models that are used to forecast outcomes related to mental health.

3.1 Methodology

In this section, first I started my methodology with collecting data collection and then I preprocessed the dataset for better analyzing, make the dataset clean for using with labeling, encoding, one hot encoding etc. For better accuracy. I then chose the best machine learning models after comparing with others models and changed the hyperparameter to get the best results. Also, I used best feature selection method, it's called PCA-IG from the help of a research study. Then I used explainable AI (XAI) methods analytical tools for better understanding and explanation of the reason causing depression.

3.1.1 Overview

This chapter presents a clear description of how the research considered the interrelation of depression and mental health among Bangladeshi university students. The process commenced with data collection and data preparation to make it clean and neat and ready for analysis. Then with the help of the best feature selection for my complex dataset PCA-IG. In this way the model gives us more accurate result after preprocessing and feature selection. Then with the help of explainable ai called LIME this research also explains all the reason and work for better understanding. It gives us more trusting reason and real-life comparison. Such a technique doesn't dwell only on proper forecasting; it also dwells on discovering and explaining associations of depression with other variables. With this, the study aspires to make results easier to understand, more accessible, and usable in real-world applications.

In this chapter, I present a clear and coherent description of how the research has been conducted. With the use of a workflow diagram, each step of the research process from the choice of features and dataset analysis through data collection and data preparation can be viewed. Each stage of the process itself is described fully, detailing the use of machine learning algorithms, how results were evaluated and what changes were made in order to optimize the performance of the model. I also detail the datasets fully and demonstrate how they are relevant to the research itself and how preparation methods were used effectively to yield proper and high-quality results.

3.1.2 Proposed Methodology

I have explained the study design here. First is data collection and preprocessing, feature engineering to improve the dataset, I choose suitable machine learning models, and took steps like changing their hyperparameters, those are some steps that are used for methodology. An explainable AI (XAI) framework was used to help evaluate and explain the model predictions once the models were finalized, which improved the transparency and actionability of the results. Figure 1 presents a clear and concise visualization of the methodology, capturing the core of the proposed approach.

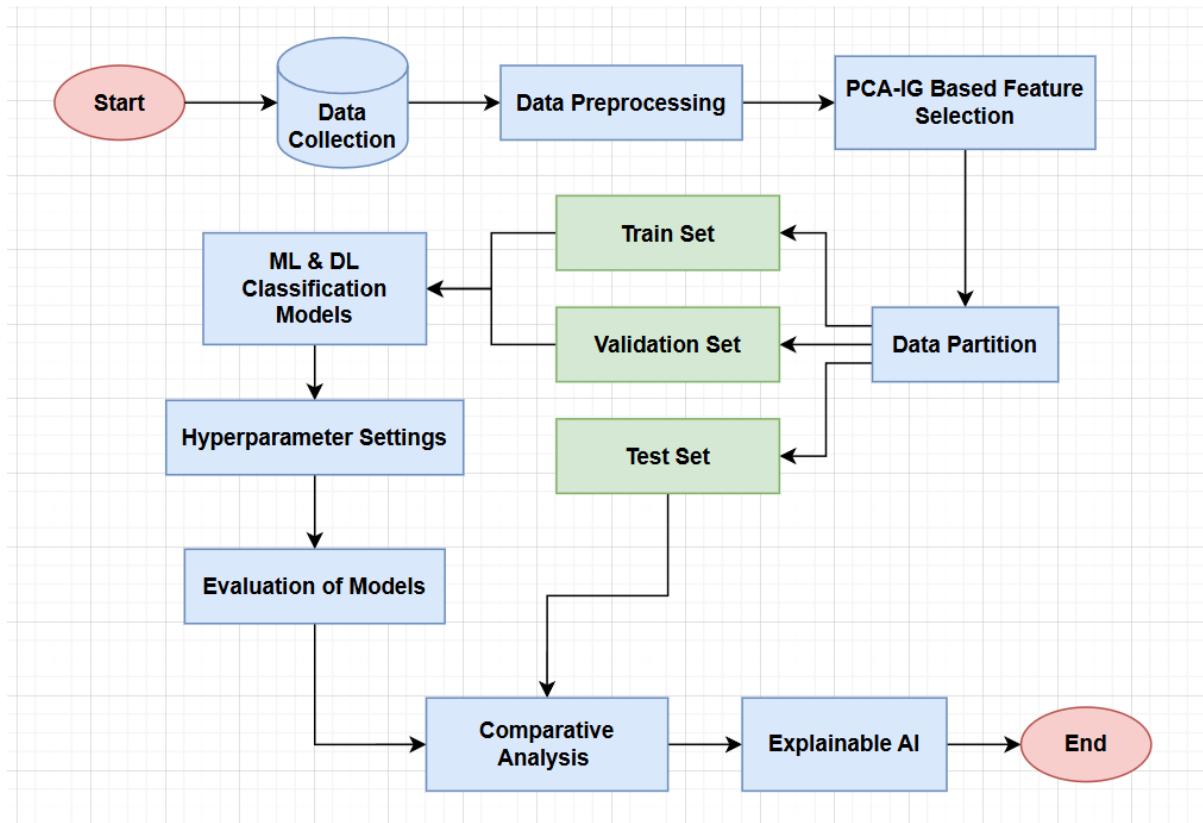


Figure 3.1: Proposed Methodology for our research study.

3.2 Detailed Methodology and Design

Section 3.2 offers a detailed overview of the methodology and design framework used in this study. It summarizes the entire research process, outlining the key procedures and approaches applied. The following subsections then break down each of these components, explaining them in greater detail.

3.2.1 Data Collection & Preprocessing

Our analysis is based on a well-curated dataset from the study “A Comprehensive Standardized Dataset on Mental Health Problems (MHPs) of University Students. [25]” It includes information on 1,977 students from 15 of Bangladesh’s top-ranked universities, comprising 9 public and 6 private institutions. The dataset was carefully reviewed by five experienced academics, each with 10 to 20 years of expertise in teaching and research. To assess anxiety, stress, and depression, three established psychometric evaluation models were applied systematically. For this study, we focused specifically on depression, aiming to identify the key factors affecting students and highlight the most significant elements linked to depressive symptoms.

	A	B	C	D	E	F	G	H	I	J	K	L	M	N	O	P	Q	R	S	T	U
1	1. Age	2. Gender	3. Universi	4. Departm	5. Academ	6. Current CGPA		1. In a sem	2. In a sem	3. In a sem	4. In a sem	5. In a sem	6. In a sem	7. In a sem	8. In a sem	9. In a sem	Depressio	Depressio			
2	18-22	Female	Independe	Engineerin	Fourth Yea	2.50 - 2.99		1	2	1	1	2	1	1	1	1	11	Moderate Depression			
3	18-22	Male	Independe	Engineerin	First Year c	3.80 - 4.00		1	1	1	1	1	1	1	1	1	9	Mild Depression			
4	18-22	Male	Independe	Engineerin	First Year c	3.00 - 3.39		2	0	2	3	2	2	2	2	2	16	Moderately Severe Depression			
5	18-22	Male	Independe	Engineerin	First Year c	3.40 - 3.79		1	1	1	1	1	1	1	1	1	9	Mild Depression			
6	18-22	Male	Independe	Engineerin	First Year c	3.40 - 3.79		1	1	1	1	1	1	1	1	1	9	Mild Depression			
7	18-22	Male	Independe	Engineerin	First Year c	3.80 - 4.00		2	2	2	2	2	2	2	2	2	18	Moderately Severe Depression			
8	18-22	Male	Independe	Engineerin	First Year c	3.80 - 4.00		0	0	1	0	0	0	1	1	0	3	Minimal Depression			
9	18-22	Male	Independe	Engineerin	Third Year c	2.50 - 2.99		2	3	3	0	0	3	3	0	0	14	Moderate Depression			
10	18-22	Female	Independe	Engineerin	First Year c	3.80 - 4.00		2	2	2	2	2	2	2	2	2	18	Moderately Severe Depression			
11	18-22	Female	Independe	Engineerin	First Year c	3.80 - 4.00		1	3	3	3	1	3	3	2	0	19	Moderately Severe Depression			
12	18-22	Male	Independe	Engineerin	First Year c	3.40 - 3.79		2	2	2	2	2	2	2	2	2	18	Moderately Severe Depression			
13	18-22	Male	Independe	Engineerin	First Year c	3.80 - 4.00		1	3	3	3	1	3	3	3	3	23	Severe Depression			
14	18-22	Male	Independe	Engineerin	First Year c	3.80 - 4.00		1	1	1	2	0	1	1	1	1	9	Mild Depression			
15	18-22	Female	Independe	Engineerin	First Year c	3.80 - 4.00		3	1	2	2	1	3	3	2	3	20	Severe Depression			
16	18-22	Female	Independe	Engineerin	First Year c	3.80 - 4.00		2	2	2	2	2	2	2	2	2	18	Moderately Severe Depression			
17	18-22	Male	Independe	Engineerin	First Year c	2.50 - 2.99		2	2	3	3	3	3	1	0	3	20	Severe Depression			
18	18-22	Female	Independe	Engineerin	Third Year c	2.50 - 2.99		2	2	2	2	2	2	2	2	2	18	Moderately Severe Depression			
19	18-22	Female	Independe	Engineerin	First Year c	3.80 - 4.00		1	1	2	2	2	0	1	0	0	9	Mild Depression			
20	18-22	Female	Independe	Engineerin	First Year c	3.40 - 3.79		2	0	0	0	0	1	0	0	0	3	Minimal Depression			
21	18-22	Female	Independe	Engineerin	First Year c	3.40 - 3.79		1	1	0	1	1	0	1	1	0	6	Mild Depression			
22	18-22	Male	Independe	Engineerin	First Year c	3.80 - 4.00		0	0	0	0	0	0	0	0	0	0	No Depression			
23	18-22	Female	Independe	Engineerin	First Year c	3.80 - 4.00		2	1	0	1	0	0	0	0	0	4	Minimal Depression			
24	18-22	Male	Independe	Engineerin	First Year c	3.80 - 4.00		3	2	3	3	3	2	3	2	0	21	Severe Depression			
25	18-22	Female	Independe	Engineerin	First Year c	3.80 - 4.00		1	3	0	3	3	0	0	1	1	12	Moderate Depression			
26	23-26	Male	Independe	Engineerin	Fourth Year c	2.50 - 2.99		0	1	1	3	3	0	2	3	3	16	Moderately Severe Depression			
27	18-22	Male	Independe	Engineerin	First Year c	3.00 - 3.39		1	1	2	1	1	1	0	1	2	10	Moderate Depression			
28	18-22	Male	Independe	Engineerin	First Year c	2.00 - 2.99		0	2	0	2	0	1	0	0	0	7	Mild Depression			

Figure 3.2: Dataset Sample.

The dataset comprises both demographic and academic attributes, including Age, Gender, University, Department, Academic Year, and Current CGPA, alongside psychometric assessment items related to depressive symptoms. These assessment items capture the frequency of experiences such as loss of interest or pleasure, feelings of sadness or hopelessness, sleep disturbances, fatigue, appetite changes, low self-worth, difficulty concentrating, noticeable changes in movement or speech, and thoughts of self-harm. Each symptom is measured on a scale contributing to the Depression Value (maximum score of 27), which is further categorized into Depression Levels: Minimal (1–4), Mild (5–9), Moderate (10–14), Moderately Severe (15–19), and Severe (20–27).

After data collection, the dataset was preprocessed by addressing missing values and encoding categorical variables using both label encoding and one-hot encoding techniques. The original dataset contained 1,977 records and 18 columns, which increased to 22

columns after encoding. The key categorical variables included Age (5 categories), Gender (3 categories), University (15 categories), Department (12 categories), Academic Year (5 categories), and Current CGPA (6 categories). The target variable, Depression Level, was scaled and categorized into six distinct classes for analysis.

3.2.2 Feature Selection

After preprocessing the data, I identified the most impactful features using the PCA–IG method. This approach combines Principal Component Analysis (PCA) and Information Gain (IG) to make feature selection more effective. PCA reduces redundancy by transforming correlated features into a smaller set of uncorrelated components while keeping as much variance as possible. IG then measures how much each feature contributes to predicting the target, allowing us to rank them by importance. By applying PCA first to clean and condense the data, and then IG to pinpoint the most relevant features, this method ensures we focus on information that is both meaningful and non-redundant—ultimately improving model performance and interpretability. Figure 3.3 presents the feature importance plots for the PCA–IG-based model.

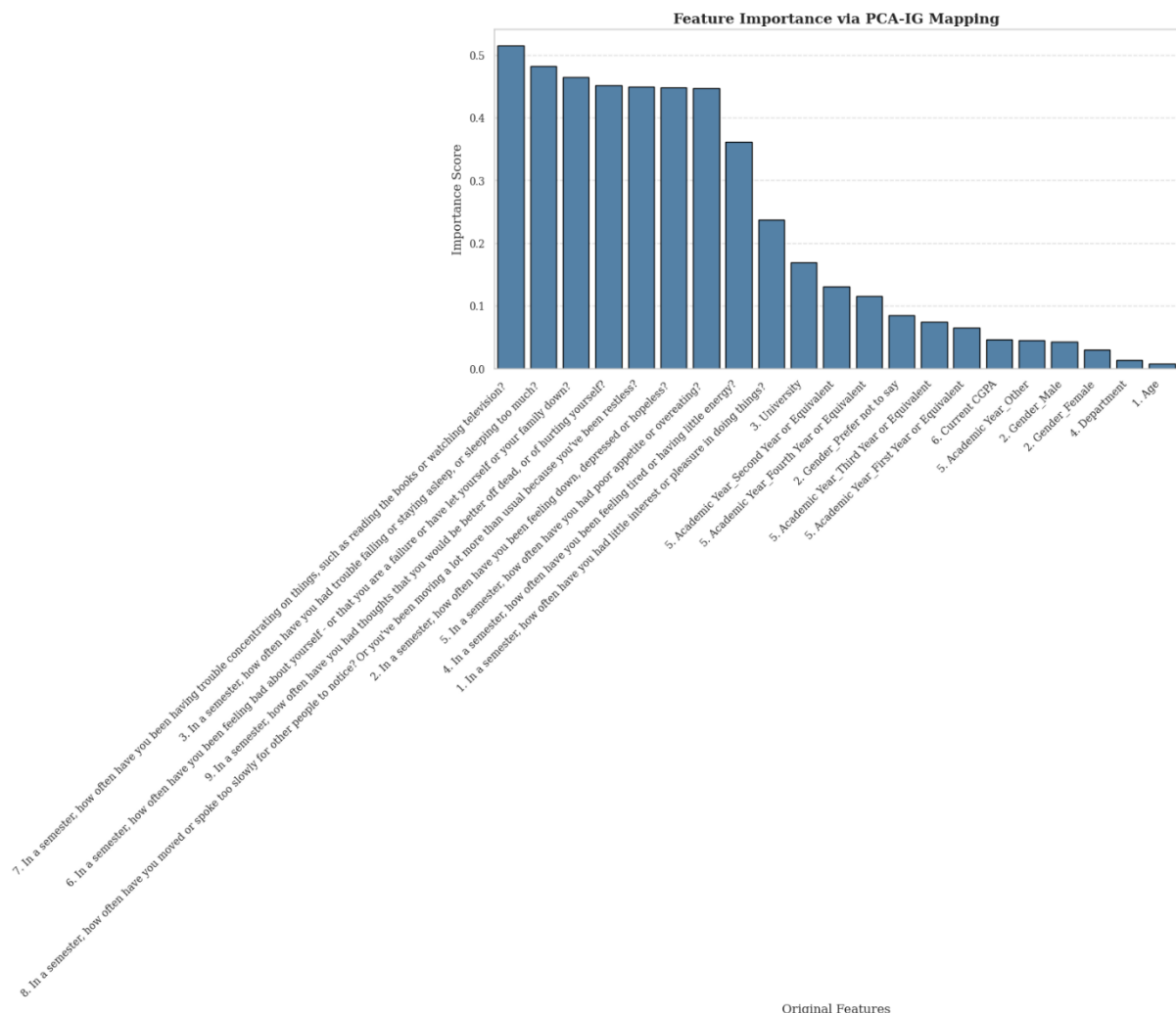


Figure 3.3: Feature Importance Plots for PCA-IG Based Model.

Figure 3.3 presents the feature importance scores mapped from the PCA–IG method to

the original features. The highest impact was seen in concentration difficulties during the semester (score: 0.516), followed closely by sleep problems (0.482) and feelings of low self-worth or failure (0.465). Suicidal thoughts (0.452) and noticeable changes in movement or restlessness (0.450) also ranked highly, along with feelings of depression or hopelessness (0.448) and appetite changes (0.447). Fatigue and low energy had a moderate importance (0.361), while loss of interest or pleasure scored lower (0.237). Among demographic and academic variables, the university attended (0.169) and specific academic years—second year (0.131), fourth year (0.116), and third year (0.074)—showed some influence. Gender categories and current CGPA had comparatively lower scores, ranging from 0.087 to 0.030, with age and department contributing minimally (below 0.015). These results highlight that symptom-related features play a much stronger role in predicting depression levels than demographic or academic variables. The dataset was divided into three subsets: training, testing, and validation, to guarantee a balanced and robust model preparation procedure. Specifically, 80% of the data was earmarked for training, with an additional 20% left out for validation. The remaining 20% of the original dataset was set aside for testing, resulting in a well-organized data split for best model performance.

3.2.3 Machine Learning and Deep Learning Model Selection

After completing data preprocessing and feature engineering, I selected a set of high-performing machine learning and deep learning models to ensure the best possible results. The models used in this study include Random Forest, Support Vector Classifier (SVC), K-Nearest Neighbors (KNN), LightGBM (LGBM), Logistic Regression, and an Artificial Neural Network (ANN). These algorithms are widely recognized for their strong capabilities in classification tasks. In the following sections, I provide a detailed description of each model, highlighting their unique features and strengths in effectively handling classification problems.

Random Forest Classifier:

Random Forest is an ensemble learning method that builds multiple decision trees during training and combines their predictions to improve accuracy and reduce overfitting. Each tree is trained on a random subset of the training data (bootstrapping) and considers a random subset of features when splitting nodes. For classification tasks, the final output is determined by majority voting among all trees.

Let $h_1(x)$ be the prediction of the t^{th} decision tree for input x , and T be the total number of trees. The Random Forest classification rule is:

$$\hat{y} = \text{mode}\{h_1(x), h_2(x), \dots, h_T(x)\} \dots \dots \dots (i)$$

Where:

$\hat{y}$ = final predicted class

mode = most frequent class among the predictions

$h_t(x)$ = output of tree t

Random Forest improves generalization by averaging across many de-correlated trees, making it robust to noise and less prone to overfitting compared to a single decision tree.

Support Vector Classifier (SVC)

The Support Vector Classifier is a supervised learning algorithm based on the principles of Support Vector Machines (SVM). It aims to find the optimal hyperplane that separates data points of different classes with the maximum margin. Only the data points lying closest to the hyperplane, called support vectors, determine its position and orientation. SVC can handle both linear and non-linear classification by using kernel functions (e.g., linear, polynomial, RBF) to map the data into higher-dimensional spaces where separation is easier.

The decision function for classification is:

$$f(x) = \text{sign}(w^T x + b) \dots \dots \dots (ii)$$

Where:

w = weight vector defining the orientation of the hyperplane

b = bias term

x = input feature vector

$\text{sign}(\cdot)$ = determines the predicted class (+1 or -1)

The optimization problem is:

$$\min(w, b) \frac{1}{2} \|w\|^2 \quad \text{subject to} \quad y_i(w^T x_i + b) \geq 1, \forall i \dots \dots \dots (iii)$$

Where y_i are the class labels. The goal is to maximize the margin $\frac{2}{\|w\|}$ while ensuring correct classification of the training data (or with minimal errors in the soft-margin case).

K-Nearest Neighbors (KNN)

K-Nearest Neighbors (KNN) is a straightforward algorithm that classifies a new data point by looking at its k closest neighbors in the dataset and assigning it the most common class among them. It doesn't "learn" a model in advance—it simply remembers the data and makes predictions based on how similar the new point is to existing ones, using distances like Euclidean or Manhattan to measure closeness.

For a query point x , the predicted class $\hat{y}$ is:

$$\hat{y} = \text{mode}\{y_i \mid x_i \in N_k(x)\} \dots \dots \dots (iv)$$

Where:

$N_k(x)$ = set of k nearest neighbors of x

y_i = class label of neighbor x_i

$\text{mode}(\cdot)$ = most frequent class among neighbors

Euclidean Distance Formula (most common in KNN):

$$d(x, x_i) = \sqrt{\sum_{j=1}^n (x_i - x_{ij})^2} \dots\dots\dots v$$

Where n is the number of features.

LGBM (Light Gradient Boosting Machine)

LGBM is a fast, efficient gradient boosting framework that builds decision trees leaf-wise rather than level-wise, which helps reduce loss more quickly. It works by training an ensemble of trees sequentially, where each new tree focuses on correcting the errors of the previous ones.

The model minimizes a differentiable loss function $L(y, y)$ using gradient descent, with each tree's contribution weighted by a learning rate:

$$F_m(x) = F_{m-1}(x) + \eta \cdot T_m(x) \dots\dots\dots vi$$

Here, $F_m(x)$ is the model after m iterations, η is the learning rate, and $T_m(x)$ is the prediction from the m-th tree.

Logistic Regression

Logistic Regression is a simple yet powerful algorithm used for binary and multi-class classification. Instead of predicting a continuous value, it predicts the probability that a given input belongs to a certain class. It uses the logistic (sigmoid) function to map any real-valued number into a range between 0 and 1:

$$P(y = 1 | x) = \frac{1}{1 + e^{-(\beta_0 + \beta_1 x_1 + \beta_2 x_2 + \dots + \beta_n x_n)}} \dots\dots\dots vii$$

Here, $P(y = 1 | x)$ is the predicted probability, and $\beta_0, \beta_1, \dots, \beta_n$ are model coefficients learned from the data. Classification is done by applying a threshold (commonly 0.5) to this probability.

An Artificial Neural Network (ANN)

ANN is a machine learning model inspired by the way the human brain processes information. It consists of layers of interconnected “neurons,” where each neuron takes inputs, applies weights, sums them up, passes the result through an activation function (like sigmoid, ReLU, or tanh), and sends the output to the next layer.

For a single neuron, the computation is:

$$y = f(\sum_{i=1}^n (w_i x_i + b)) \dots\dots\dots viii$$

Here x_i are inputs, w_i are weights, b is a bias term, f is the activation function. ANNs learn by adjusting weights and biases through training to minimize prediction errors.

3.2.4 Hyperparameter Tuning

After completing the data preprocessing, I chose the best models for the task. I then finetuned the hyperparameters to improve their performance. Table 2 shows the hyperparameter settings employed in our models to improve predicted accuracy.

Table 3.1: Hyperparameter Settings of All ML Models.

Model Name	Hyperparameter Setting
Random Forest Classifier	n_estimators=100, criterion='gini'
Support Vector Classifier (SVC)	kernel='rbf', random_state=42
K-Nearest Neighbors (KNN)	n_neighbors=5, weights='uniform', algorithm='auto'
LGBM Classifier	boosting_type='gbdt', num_leaves=31
Logistic Regression	max_iter=1000, solver='lbfgs'
Artificial Neural Network (ANN)	Input: Dense(256, activation=swish) → BN → Dropout(0.3) → Dense(256, activation=swish) → BN → Dropout(0.2) → Dense(256, activation=swish) → BN → Dropout(0.1) → Output: Dense(num_classes, activation='softmax'); Optimizer: Adam(learning_rate=0.1); Loss: categorical_crossentropy; Batch size: 128; Epochs: 300; Validation split: 0.1; Callbacks: EarlyStopping & ReduceLROnPlateau

Table 2.2 shows the key hyperparameters used for all the machine learning and deep learning models in this study. For Random Forest, I set 100 trees and used the Gini criterion to split nodes. The Support Vector Classifier used an RBF kernel and a fixed random seed for consistency. K-Nearest Neighbors considered 5 neighbors with uniform weighting, letting the algorithm choose the best search method automatically. LGBM was run with its default settings, including 31 leaves per tree. Logistic Regression was given a higher iteration limit of 1000 to ensure convergence. The ANN had three layers of 256 units each, using the swish activation function, with batch normalization and dropout to help prevent overfitting. It was trained with the Adam optimizer set at a learning rate of 0.1, and we used early stopping and learning rate reduction to keep the training stable and efficient.

3.2.5 Evaluation Matrices

A critical component of classification jobs is assessing the model's performance, which demonstrates how well the model can forecast results. Classification algorithms are usually evaluated using a range of important indicators, each of which offers a distinct perspective on the prediction capabilities and limitations of the model.

Accuracy:

The percentage of successfully predicted cases out of all predictions is known as accuracy, and it quantifies a classification model's total correctness.

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

Precision:

The precision or Positive Predictive Value shows how many of the predicted positive cases were actually positive.

$$\text{Precision} = \frac{TP}{TP+FP} \dots\dots\dots \text{x}$$

Recall

The model properly detection of the number of true positive is recall, which is sometimes we called it true positive rate.

$$\text{Recall} = \frac{TP}{TP+FN} \dots\dots\dots \text{xi}$$

F1 Score:

The balance between Precision and Recall by taking the harmonic mean of the two measurements is called f1 score.

$$F1 = 2 * \frac{\text{Precision} * \text{recall}}{\text{precision} + \text{recall}} \dots\dots\dots \text{xii}$$

Where, TP = True Positives, TN = True Negatives, FP = False Positives, FN = False Negatives

The ROC curve shows the classification model's significance between distinct classes across a range of thresholds, it is the source of the AUC or Area Under the Curve. In a word, it shows us how well the model detects true positives while reducing false positives. The probability that the model will rate a positive instance higher than a negative one is targeted by the AUC value. This model randomly receives a score of about 0.5, This is a perfect model receives a score of 1. This statistic is particularly helpful for assessing model performance on imbalanced datasets.

3.2.6 Use of Explainable AI

Explainable AI (XAI) means the better understanding and understandability of artificial intelligence systems. Deep learning networks and other complex AI models frequently behave like "black boxes," making it hard to understand how machines make decisions or make predictions. But the XAI help this problem by explaining two key ideas: interpretability and transparency. Transparency shows the inner workings of a model and the data it uses, in the other hand, interpretability explains how a model comes at its conclusions. There are many explainable ai methods like as SHAP (SHapley Additive Explanations) and LIME (Local Interpretable Model-agnostic Explanations) in explaining how many types of features affect the model's predictions. So that we can have the trust and guarantee of the work and can know how it is mainly working and which feature is causing depression, also it explains why and technologies by making AI easier to understand, it allowing people to trust their decisions from the technology.

LIME (Local Interpretable Model-agnostic Explanations):

The explainable ai that I am using in my model is LIME. It is a powerful tool for increasing the understandability of my AI models. I am using LIME for testing how nicely different variables impact the model's predictions and can explain its decision-making process. This work will help to understand the complex model outputs into understandable explanations. It is making it easier to communicate our findings and build confidence in the AI results. The LIME explanatory model is given bellow:

$$\hat{g}(x) = \operatorname{argmin}_{g \in G} GL(f, g, \pi_x') + \Omega(g) \dots\dots\dots(xiii)$$

In this explainable ai method, LIME, the explanation model $\hat{g}(x)$ refer to complex model behavior locally form a set G of possible models, often linear. After using a loss function $L(f, g, \pi_x')$, the explanation model's simplicity and understandability.

3.3 Project Plan

Section 3.3 presents the project plan for our research, with Table 3.2 providing a clearer overview.

Table 3.2: project plan

Phase	Tasks	Timeline
Phase 1: Research	Conduct literature review and finalize research questions.	1-2 Weeks
Phase 2: Data Collection	collect data from Internet	1 week
Phase 3: Model Development	Select and train machine learning models.	2 Weeks
Phase 4: Analysis	Analyze results and interpret findings.	1 Week
Phase 5: Explainable AI Integration	integrating LIME for model explainability. Analyzing feature importance.	1 week
Phase 6: Documentation	Prepare the final report and supporting materials.	2 Weeks
Phase 7: Presentation	Create slides and present findings.	1 Week
Phase 8: Review and Feedback	Incorporate feedback and finalize deliverables.	1 Week

3.4 Task Allocation

Table 3.3: Task Allocation

Task	Description
Research and Background Study	Reviewed studies and identified research gaps.
Data Collection and Preparation	Designed surveys, collected data, and cleaned it for analysis.
Building the Model	Selected machine learning models, optimized them, and tested accuracy.
Analysis and Results	Interpreted data and used tools like LIME to explain results.
Report Writing	Compiled findings into a detailed and well-organized report.
Presentation	Created slides and presented the findings effectively.
Quality Checks	Reviewed work at each step to ensure accuracy and quality.

3.5 Summary

This chapter outlines the research methodology, from data collection to analysis and refinement. A workflow diagram displays the entire process, including feature selection, dataset exploration, machine learning and deep learning application, and iterative method refining. The dataset is reviewed to highlight its significance, and pretreatment phases and algorithms are described to improve transparency and clarity in research procedures.

Chapter 4

Implementation and Results

The fourth chapter concentrates on the application of the recommended methodology and gives the analytical results. It provides details on how to install the model, assess its performance, and compare the results.

4.1 Environment Setup

For environment setup, I am using Google Colab. It is a cloud-based platform that significantly increase the research process. I have used it to build, execute, and optimize our machine learning models. Except the free version, I have been able to create a powerful and inbuilt computing environment. Because of this, I have been able to conduct extensive experiments without being operated by hardware limitations, which often hinder progress. Google Collab's user-friendly design has made it easy to get started, allowing me to focus on exploring new ideas and methods rather than solving technical problems. I have greatly welcomed Google Collab's collaboration tools, which have made it easy to share code and datasets and fostered a strong sense of teamwork throughout our projects. We have been able to collaborate in real time, share ideas, gain knowledge from each other's perspectives, and improve our models as a team. In addition to increasing our productivity, this collaborative setting has improved our conversations and added joy to the learning process. I have coded using Python 3 and the scikit-learn package for machine learning-related work. This powerful combination improved the caliber and impact of our work. It helps us to confidently handle the difficulties of our project.

4.2 Comparative Analysis

4.2.1 Performance Analysis Before Preprocessing

I have run many models with my dataset for choosing accurate model and chose the top six in it after comparing the accuracy to predict depression among university students in Bangladesh. The final prediction outcomes of these models on the testing dataset before preprocessing is given in Table 4.1.

Table 4.1: Models Result Before Preprocessing.

Model	Accuracy (%)	Weighted Precision (%)	Weighted Recall (%)	Weighted F1-Score (%)
Random Forest Classifier	85.61	85.62	85.61	85.00
Support Vector Classifier (SVC)	89.90	90.29	89.90	89.00
K-Nearest Neighbors (KNN)	71.97	72.99	71.97	71.00
LightGBM Classifier (LGBM)	86.36	86.68	86.36	86.00
Logistic Regression	95.96	95.81	95.96	96.00
Artificial Neural Network (ANN)	94.44	94.53	94.44	94.00

In the Table 4.1 shows the performance of six machine learning models accuracy before data preprocessing where Logistic Regression performed the best, the accuracy of 95.96%, followed closely by the Artificial Neural Network at 94.44%. Both the Support Vector Classifier and LightGBM also performed well, scoring above 86% . The Random Forest results, with accuracy and precision around 85%, while K-Nearest Neighbors had the lowest performance, just under 72%. Overall, these results shows us the effectiveness of the models, with Logistic Regression and ANN emerging as the most high accuracy for predicting depression levels among Bangladeshi university students.

4.2.2 Performance Analysis After Preprocessing

Before selecting the best six models , I have tested lots of models from open source to detect depression in college students in Bangladesh. Following preprocessing, Table 3 displays the models' performance on the test data.

Table 4.2: Models Result After Preprocessing.

Model	Accuracy (%)	Weighted Precision (%)	Weighted Recall (%)	Weighted F1-Score (%)
Random Forest (RFC)	85.86	86.00	86.00	86.00
Support Vector (SVC)	89.90	90.29	90.00	89.00
K-Nearest Neighbor (KNN)	73.23	75.80	73.23	72.00
LightGBM (LGBM)	86.36	86.68	86.00	86.00
Logistic Regression	97.73	97.94	97.73	97.83
Artificial Neural Network (ANN)	97.98	97.97	97.98	97.98

After data preprocessing, the performance of the six models is summarized in Table 4.2. Logistic Regression and the Artificial Neural Network (ANN) achieved the highest accuracy, at 97.73% and 97.98%, respectively, demonstrating their strong predictive capabilities on the processed dataset. Both models also showed excellent weighted precision, recall, and F1-scores, all above 97%, indicating balanced and reliable classification across all categories. LightGBM and the Support Vector Classifier (SVC) performed well too, with accuracies of 86.36% and 89.90%, respectively, and robust weighted metrics, showing effective handling of the data. Random Forest delivered solid results with an accuracy of 85.86% and balanced weighted precision, recall, and F1-scores around 86%. The K-Nearest Neighbor (KNN) model, while achieving a reasonable accuracy of 73.23%, lagged behind the other models, particularly in weighted F1-score, highlighting some limitations in classification after preprocessing. Overall, these results emphasize the superior performance of the ANN model for this classification task once the data had been processed.

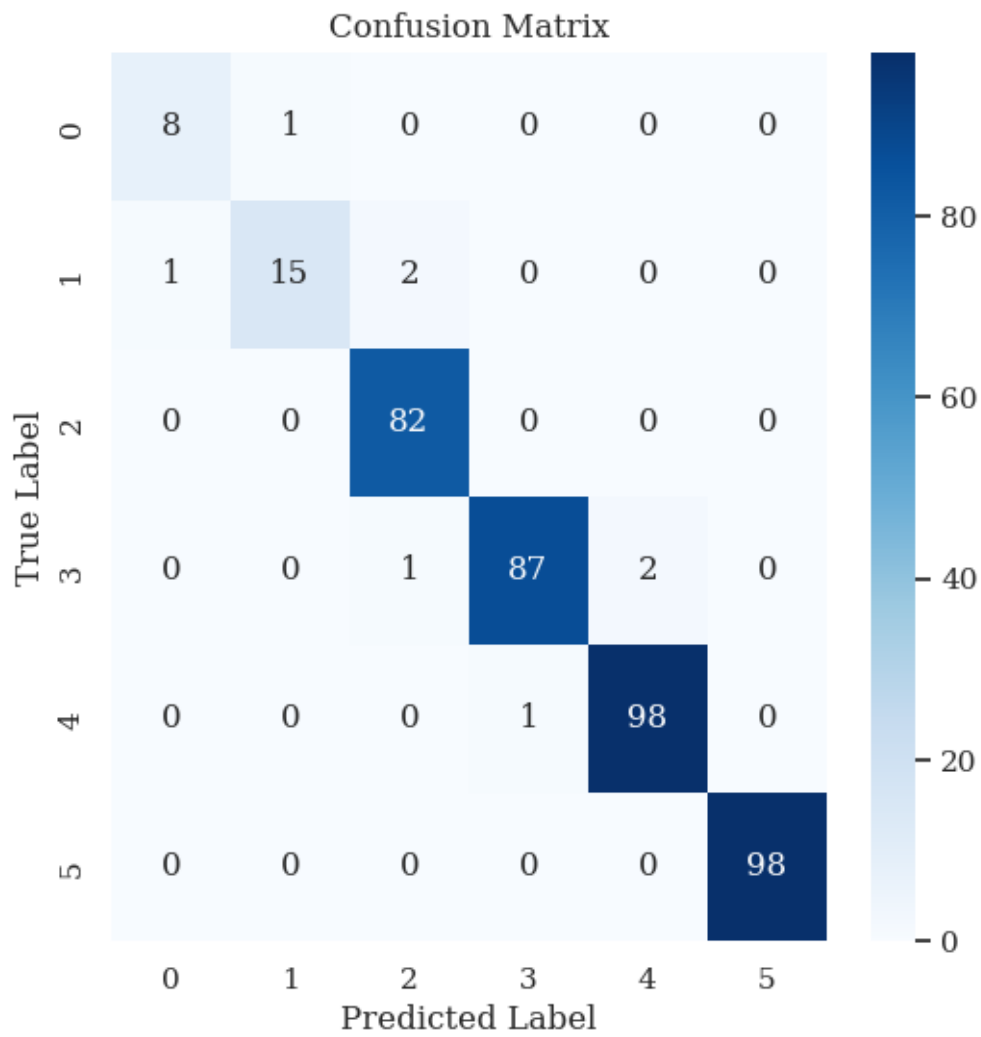


Figure 4.1: Confusion matrix illustrating the performance of our proposed ANN model.

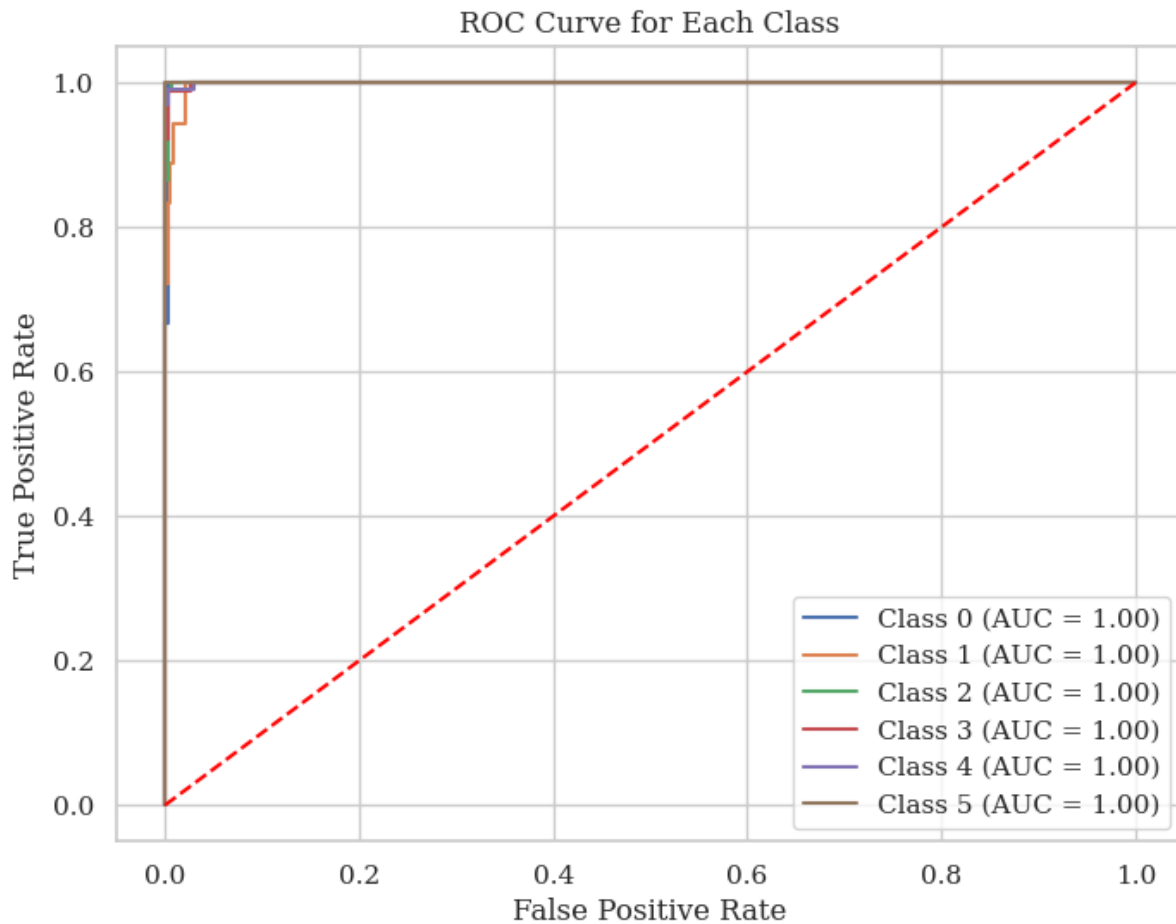


Figure 4.2: ROC curve demonstrating the performance of our proposed model.

4.3 Results and Discussion

I go into greater detail about the results of our investigation and offer a comprehensive analysis of them in this part. The findings are thoroughly explained and described in the next subsections, which also include an assessment of the performance indicators and their relevance to our study question.

4.3.1 EXPLORATORY DATA ANALYSIS (EDA)

I began with exploratory data analysis (EDA) to gain a better understanding of the features in each dataset and ensure they aligned with our research goals. EDA is an important step because it helps reveal the data's distribution, identify any outliers, and uncover hidden patterns. These insights guided more effective feature engineering, informed smarter model selection, and ensured cleaner data overall. This process ultimately contributed to stronger models with improved accuracy and a deeper understanding of the dataset. Figure 4.3 shows the correlation heatmap, highlighting the relationships between different features in our data.

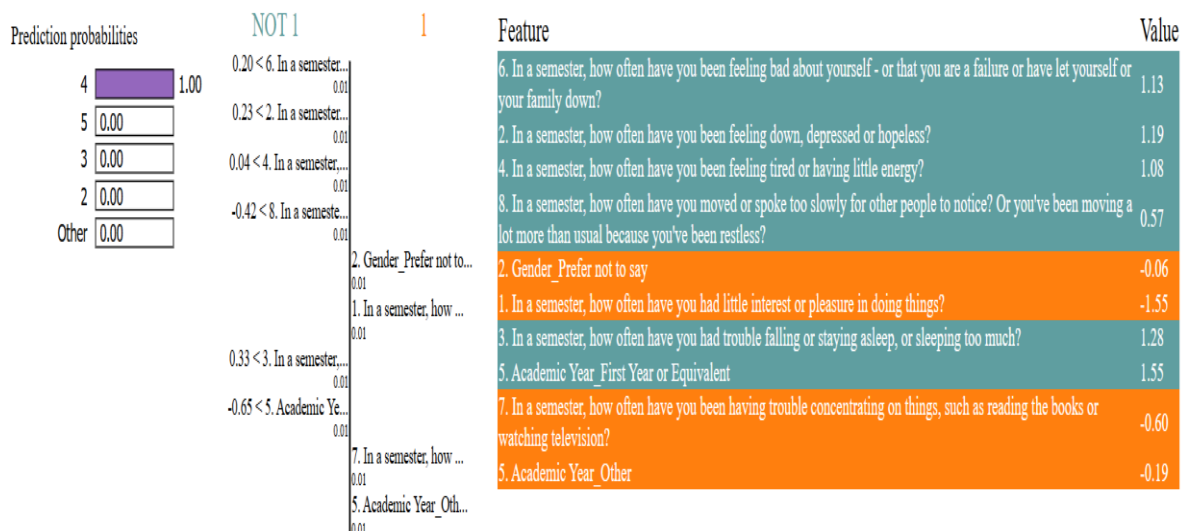


Figure 4.4: LIME Tabular Plot.

This LIME tabular plot explains the key features influencing the prediction of the depression class labeled "1" for a specific instance. The prediction probability for class 4 is very high at 1.00, while all other classes have zero contribution. The features positively impacting the prediction (highlighted in blue) include: "Academic Year_First Year or Equivalent" with a value of 1.55, "In a semester, how often have you had trouble falling or staying asleep, or sleeping too much?" at 1.28, "In a semester, how often have you been feeling down, depressed or hopeless?" at 1.19, "In a semester, how often have you been feeling bad about yourself or that you are a failure or have let yourself or your family down?" at 1.13, "In a semester, how often have you been feeling tired or having little energy?" at 1.08, and "In a semester, how often have you moved or spoke too slowly or have been restless?" at 0.57. Conversely, features negatively influencing the prediction (highlighted in orange) include "In a semester, how often have you had little interest or pleasure in doing things?" with a value of -1.55, "In a semester, how often have you been having trouble concentrating on things?" at -0.60, "Gender_Prefer not to say" at -0.06, and "Academic Year_Other" at -0.19. This balance of positive and negative contributions from these features shapes the model's prediction outcome.

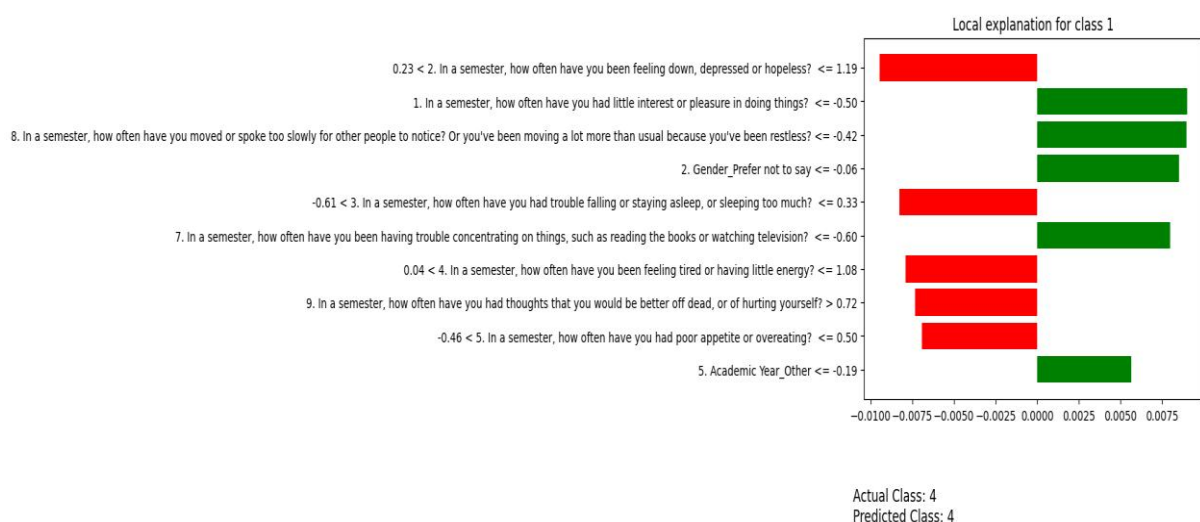


Figure 4.5: LIME Feature Importance Plot.

The LIME feature importance plot for class 1 shows the impact of various features on the model's prediction for an instance with an actual and predicted class of 4. The most significant positive contributor to class 1 is "In a semester, how often have you had thoughts that you would be better off dead, or of hurting yourself?" with a weight of 0.72, indicating a strong association with this class. Other notable positive contributors include "In a semester, how often have you been feeling tired or having little energy?" (0.1) and "In a semester, how often have you had trouble falling or staying asleep, or sleeping too much?" (0.33). Negative contributors that reduce the likelihood of class 1 include "In a semester, how often have you had little interest or pleasure in doing things?" (-0.5), "In a semester, how often have you moved or spoke too slowly for other people to notice? Or you've been moving a lot more than usual because you've been restless?" (-0.42), and "In a semester, how often have you had poor appetite or overeating?" (-0.5). Less impactful features include "Gender-Prefer not to say" (-0.06) and "Academic Year-Other" (-0.19), suggesting minimal influence on the prediction.

4.3.3 Discussion

This study utilizes a carefully curated dataset of 1,977 university students from 15 top-ranked Bangladeshi universities to explore factors influencing depression. The dataset combines demographic and academic information with detailed psychometric assessments based on established models, capturing a range of depressive symptoms. Depression severity is categorized across five levels, enabling a nuanced analysis aimed at identifying key contributors to student depression and highlighting the most impactful factors linked to depressive symptoms. The feature importance analysis reveals that depressive symptoms, particularly concentration difficulties, sleep problems, and feelings of low self-worth, are the strongest predictors of depression levels among students. In contrast, demographic and academic factors such as university, academic year, gender, CGPA, age, and department have much less influence. This underscores that symptom-related features carry greater weight in understanding and predicting depression, supported by a carefully balanced data split to ensure robust model training and evaluation.

The comparative analysis of model performances before and after data preprocessing reveals a clear improvement across all six machine learning algorithms. Initially, Logistic Regression led with 95.96% accuracy, closely followed by the Artificial Neural Network (ANN) at 94.44%, while models like Support Vector Classifier, LightGBM, and Random Forest demonstrated moderate effectiveness, and K-Nearest Neighbors lagged behind. Post-preprocessing, both Logistic Regression and ANN significantly enhanced their accuracy, surpassing 97%, with ANN slightly outperforming Logistic Regression. These two models also achieved the highest weighted precision, recall, and F1-scores, reflecting their balanced and robust classification capabilities. Other models, including LightGBM and SVC, showed modest gains, maintaining strong but lower accuracy and metric scores, while KNN remained the least effective overall. Overall, the Artificial Neural Network emerged as the best-performing model, demonstrating superior predictive power and reliability for classifying depression levels in Bangladeshi university students after data preprocessing.

The LIME explanation highlights the key features driving the model's prediction of a high depression class for a specific student, showing how certain symptoms like sleep problems, feelings of hopelessness, low self-worth, fatigue, and restlessness strongly increase the likelihood of depression, while other factors such as low interest, concentration difficulties, and some demographic variables slightly reduce it. This balance of positive and negative feature impacts provides clear insight into the model's decision-making for

that individual prediction. The LIME feature importance plot illustrates how specific depressive symptoms, particularly suicidal thoughts, fatigue, and sleep disturbances, strongly drive the model's prediction toward class 1, while other symptoms like lack of interest, noticeable changes in movement, and appetite issues decrease its likelihood. Demographic factors have minimal impact, emphasizing that symptom-related features predominantly influence the prediction for this instance.

4.4 Summary

After preprocessing the data and analyzing, I found out that the main feature for depression are concentrating, sleep problems, and low self-esteem, far more so than age or academic year. After data preparation and cleaning, the Artificial Neural Network emerged as the most accurate machine learning model when tested other models. We were able to understand how particular symptoms impacted each student's predictions by using techniques such as LIME, which explained every information and reason that is increased the trust of the model selection.

Chapter 5

Engineering Standards and Design Challenges

This chapter shows us the thesis's technological and cooperative frameworks also the engineering standards and design challenges I ran during my study of thesis. It focusses on the instruments, systems, and procedures that we use to create, run, and improve our machine learning models.

5.1 Compliance with the Standards

In this section, I explained how the project established engineering standards so that the models I created were reliable, efficient, and scalable. There were three main areas of emphasis to achieve high-quality results: software, hardware, and communication standards.

5.1.1 Software Standards

I built our machine learning and deep learning models using Python 3, a language generally known to the research community. It gives us lots of tools to use in our model and algorithms, the scikit-learn package. we can use it for building and training the thesis models. I make the code simple, efficient by using to best practice for coding. This approach seamlessly shows the teamwork and maintains transparency by ensuring that the work is not only easy to understand but also easy to share and extend in this research.

5.1.2 Hardware Standards

Google Colab, a cloud-based platform that provided the processing power needed to train our models, served as the computational backbone of the project. Having access to fast hardware like GPUs allowed me to conduct complex research without spending a lot of money on equipment. Despite several shortcomings, the free version of Colab provided enough resources to support our work, allowing for cost-effective model optimization and experimentation.

5.1.3 Communication Standards

My project has effective communication standards and Google Colab's real-time collaboration tools makes it more useful. I was able to collaborate and maintain all the instruction in my thesis because of sharing code, data, and outcomes. It has made providing and receiving feedback simple and quick, which improved workflow nicely. It makes us to improve our work.

5.2 Impact on Society, Environment and Sustainability

This section shows us the impact on society, environment and sustainability of my thesis work also it focuses on how it might affect the long run.

5.2.1 Impact on Life

My deep learning model for depression prediction holds great promise for improving mental health services, especially in underdeveloped countries like Bangladesh. This model can facilitate rapid diagnosis and intervention by providing a reliable and accessible early detection tool, thereby encouraging better depression management and prevention. Ultimately, this could improve the quality of life of many people, especially in areas where healthcare resources are lacking.

5.2.2 Impact on Society & Environment

By increasing the accuracy of depression predictions, the software would be able to benefit the mental health care system by better allocation of resources and healthcare services. Better predictions enable better decision-making among medical professionals, which eventually decreases the medical facility's workload. We also minimized our dependence on tangible hardware by employing cloud-based networks like Google Colab, which promotes environmental sustainability through lower energy usage and less e-waste.

5.2.3 Ethical Aspects

Ethical considerations were the top priority during the project, with particular emphasis on the fairness of predictions made by the model and patient data privacy. I made our deep learning model open, unbiased, and reliable by adhering to standard industry ethics protocols. Our solution was based on the core principles of protecting personal information and prediction fairness so that our model produced morally and fairly sound results.

5.2.4 Sustainability Plan

It is because our project is scalable and flexible that it is sustainable. I designed the project so that it is maintainable and updatable with the least number of resources possible using a cloud architecture. The model is a sustainable, long-term way of forecasting depression since it is also capable of adapting and getting better with time based on new information.

5.3 Project Management and Financial Analysis

Our activities related to money, particularly how we spent money and resources, are highlighted in this section. We were able to keep expenses minimal by using Google Colab's free version, which saved money for more important tasks like data collection and model enhancement. In case extra processing capacity is needed in the future during the project, I have put in contingency funds.

Table 5.1: Budget Plan (within BDT 50,000)

Category	Details	Estimated cost (BDT)
Cloud Computing	Google Colab, Pro/Pro+ (GPU access)	10,000/year
Data Storage & Backup	Google Drive	6000/year
Software & Tools	Open source frameworks	Free
Internet & Utilities	Reliable internet for cloud training	12,000
Documentation & Reports	Printing, diagrams, formatting, professional reports	5,000
Miscellaneous Expenses	Extra costs (unexpected needs)	7,000
Total Estimated Cost		40,000-45,000

5.4 Complex Engineering Problem

It has increased the model performance, preparing data, and guaranteeing effective system integration. It was the hardest engineering challenges my thesis has faced. I utilized a structured method in addressing these, dividing each of the issues into more smaller tasks. It has improved the models also shows us the real-life scenario.

5.4.1 Complex Problem Solving

In this section, provide a mapping with problem solving categories. For each mapping add subsections to put rationale (Use Table 5.1). For P1, I need to put another mapping with Knowledge profile and rational thereof.

Table 5.2: Mapping with Complex Engineering Problem.

EP1 Dept of Knowledge	EP2 Range Of Conflicting Requirements	EP3 Depth of Analysis	EP4 Familiarity of Issues	EP5 Extent of Applicable Codes	EP6 Extent Of Stake- holder Involvement	EP7 Interdependence
✓	✓		✓			✓

Mapping with Knowledge Profile

This section is designed to map the overall problem and EP1 (*multiple between K3, K4, K5, K6, K8 for attaining EP1*) to the Knowledge Profile.

Table 5.3: Mapping with knowledge Profile.

K1 Natural Science	K2 Mathem atics	K3 Engineeri ng Fundame ntals	K4 Special ist Knowle dge	K5 Enginee ring Design	K6 Enginee ring Practice	K7 Comprehe nsion	K8 Resear ch Literat ure
		✓	✓		✓		✓

K3 (Engineering Fundamental) : Explain the engineering fundamentals, especially in the context of software engineering and data science.

K4 (Specialist Knowledge) : The thesis explained the knowledge in XAI and machine learning.

K6 (Engineering Practice) : The engineering principles in the design and implementation of my proposed model is the practical application.

K8 (Research Literature) : It makes connection between the research literature in XAI and machine learning.

5.4.2 Engineering Activities

In this section, provide a mapping with engineering activities. For each mapping add subsections to put rationale (Use Table 5.3).

Mapping with Complex Engineering Activities

This section is designed to map the overall problem and EA's (*multiple*).

Table 5.4s: Mapping with Complex Engineering Activities.

EA1 Range of re- sources	EA2 Level of Interaction	EA3 Innovation	EA4 Consequences for society and environment	EA5 Familiarity
✓	✓	✓	✓	✓

EA1: Integrating Emotional Intelligence and Sleep Data. My study have tested 518 data points to predict sleep quality using emotional intelligence and demographic data. It is showing the relationship between mental and physical health.

EA2: Collaborative Multidisciplinary Effort. Combining expertise from data science and psychology, the project utilized Explainable AI for clear, actionable insights.

EA3: AI Insights on Sleep Quality. Machine learning models revealed emotional intelligence dimensions like clarity and repair as critical for sleep improvement.

EA4: Tackling Academic Stress and Health. My thesis can help both academically and mentally.

EA5: Usable and Innovative Approach. With the help of explainable ai and many tools, my thesis is now more easy for understanding.

5.5 Summary

A comprehensive review of the engineering standards, the challenges of design, and the procedures that governed the project were provided in Chapter 5. It also brought out how critical it is to get the right tools and platform, like Google Colab, in order to conduct successful and productive model building. It also addressed more generic ethical and social issues and how what we do has repercussions on individuals and societies. Project execution was considered vastly determined by cost factors. Last but not least, it went through the intricate engineering issues encountered during the creation of the depression prediction model and demonstrated how an application of systematic methods towards problem resolution became a requirement in order to transcend the issues and meet with success.

Chapter 6

Conclusion

The Chapter 6 shows with a summary of the thesis and the report's conclusions and their problems. It also shows us the future work with limitation of this thesis for more further research of this sector to detect and prevent the depression among students.

6.1 Summary

After preparing the data, my research proposed the model the Artificial Neural Network (ANN) to detect depression levels among 1,977 students at 15 of Bangladesh's best university. The ANN achieved roughly 98% accuracy. It shows us the best feature or reason of depression, such as difficulties with concentration, sleep issues, and low self-esteem, and they were more effective than demographics or academics feature. They attained this through extensive use of data in the dataset, such as psychometric tests and demographics and academics. Explainable AI (XAI) such as LIME used to clarify the model's conclusions and indicate how certain symptoms of detection. This works helps to make reliable effectiveness in mental health. For accuracy and reliability, I tried changing of the model's parameters, intelligent feature choice through PCA-IG, and data cleaning. Generally, results indicate how AI can detect early indicators of depression and elicit rapid assistance as well as improved mental health policies at university. My thesis presents a practical, scalable, and transparent solution towards a significant public health problem through a combined use of deep learning and XAI.

6.2 Limitation

The forecasts of this study regarding depression levels among Bangladeshi university students were fairly reliable, but certain crucial limitations should be considered. First, the tests were based on self-reporting and could be biased due to social stigma or a lack of awareness regarding mental health symptoms, although it covered 1,977 records of students from 15 universities. Secondly, the research utilized largely cross-sectional data, making it challenging to evaluate fluctuating patterns of mental health and monitor fluctuations of depression from time to time. Moreover, while the models were evaluated internally and were reliable as per their internal assessment, their use with other groups or settings may be restricted until their external validation. Also, the dataset consisted of only students from well-established universities (six private and nine state-run), and the results may not reflect the data of students from small or rural university institutions, which may impact the findings. Even though explainable AI (XAI) methods such as LIME assist with results understanding, currently they are also not fully standardized and may impact the reliability of insight findings useful to mental health professionals. Thirdly, ethical issues such as ensuring data privacy protection and preventing possible outcomes and result during model prediction while I tested the data must in order to keep the results fair, reliable, and culturally applicable in the context of Bangladesh,

6.3 Future Work

Future research should consider the current limitations and expand on what this work has established. Researchers can use data from several years to track longitudinal changes in students' mental health, which will help them gauge the progression of depression and the success of treatment. To make the model more generalizable and to demonstrate alternative socio-economic and cultural scenarios, the dataset should include students from a wide range of schools, such as rural and less well-known Bangladeshi colleges. Testing the model on alternative datasets from other countries or regions would also provide external validation and make the model more useful and robust. To better understand student well-being, future research could aim to investigate other mental health issues, such as anxiety, stress, and insomnia, under the same framework. In the future work if we can get real-time feedback systems for better accuracy and combine it with explainable ai like SHAP or LIME for more accurate understanding and explaining with more human thinking based, then it will help the mental health consultant. We can also create an easily manageable interface, such as a website or mobile app for use then it would help to apply the concept to real-life situations, especially in universities students. Finally, more steps should take for one's information to keep safe and confidential. So, we can protect data privacy for fairness and trust, especially sensitive country like our Bangladesh.

References

- [1] M. Z. a. A.-M. M. a. E. M. A. a. A. M. H. a. K. A. a. U. N. I. a. D. P. K. Karim, "Understanding mental health challenges and associated risk factors of post-natural disasters in Bangladesh: a systematic review," *Frontiers in Psychology*, p. 1466722, 2024.
- [2] S. Arafat, "Mental Health in Bangladesh," *From Bench to Community*, p. 2024, 2024.
- [3] F. a. E. F. W. a. F. I. H. a. M. M. a. F. F. Wohid, "Work Life Balance and Its Influence on Physical and Mental Health Among Female Teachers of Public University in Bangladesh," *Asia Pacific Journal of Medical Innovations*, vol. 1, pp. 68--75, 2024.
- [4] M. A. a. S. M. R. U. Z. a. R. F. a. E. S. a. H. M. a. S. A. a. A. E. a. N. N. a. I. T. T. a. R. S. a. o. Alam, "Implications of Big Data Analytics, AI, Machine Learning, and Deep Learning in the Health Care System of Bangladesh: Scoping Review," *Journal of medical Internet research*, vol. 26, p. e54710, 2024.
- [5] I. I. a. B. R. K. a. A. A. a. R. P. V. a. T. S. Ria, "Depressive Symptoms Among Adolescents in Bangladesh," *International Journal of Mental Health and Addiction*, vol. 22, pp. 75--91, 2024.
- [6] G. H. a. S. R. I. a. S. I. a. K. M. R. Al Masud, "Effective depression detection and interpretation: Integrating machine learning, deep learning, language models, and explainable AI," *Array*, vol. 25, p. 100375, 2025.
- [7] S. a. A. S. a. T. S. A. a. M. R. F. A. a. S. T. E. a. H. M. S. A. a. A. M. A. Sahrin Jisha, "Mental Health Analysis: ML And Explainable AI Predict Depression Among Bangladeshi University Students," *Proceedings of the 3rd International Conference on Computing Advancements*, pp. 491--497, 2024.
- [8] P. o. d. a. C. c. s. a. m. l. approach, "Predictors of depression among Chinese college students: a machine learning approach," *BMC Public Health*, vol. 25, p. 470, 2025.
- [9] N. a. P. S. a. R. P. a. V. F. a. N. C. Meda, "Frequency and machine learning predictors of severe depressive symptoms and suicidal ideation among university students," *Epidemiology and Psychiatric Sciences*, vol. 32, p. e42, 2023.
- [10] M. I. H. a. U. M. S. G. a. H. M. I. a. A. M. M. a. Z. M. A. a. H. I. a. R. M. M. a. R. R. a. M. M. I. H. Nayan, "Comparison of the Performance of Machine Learning-based Algorithms for Predicting Depression and Anxiety among University Students in Bangladesh," *Asian Journal of Social Health and Behavior*, vol. 5, pp. 75-84, 2022.
- [11] M. a. K. S.-S. a. M. E. J. Gil, "Machine learning models for predicting risk of depression in Korean college students: Identifying family and individual factors," *Frontiers in Public Health*, vol. 10, p. 1023010, 2022.
- [12] B. L. a. C. P. { a. H. S. a. D. G. B. a. R. M. a. C. S. C. a. S. A. T. a. B. H. M. T. a. S. P. C. d. a. {Schaab, "How do machine learning models perform in the detection of depression, anxiety, and stress among undergraduate students? A systematic review," *Cadernos de Sa{\'u}de P{\'u}blica*, vol. 40, p. e00029323, 2024.
- [13] Q. a. H. J. a. C. X. a. G. W. a. Y. Q. a. W. Z. a. L. Z. a. Z. Y. a. Q. L. Qiang, "Identifying risk factors for depression and positive/negative mood changes in college students using machine learning," *Frontiers in Public Health*, vol. 13, p. 1606947, 2025.
- [14] C. a. J. M. a. B.-H. I. a. H. S. a. A. D. a. R. M. a. H. R. El Morr, "Predictive Machine Learning Models for Assessing Lebanese University Students' Depression, Anxiety, and Stress During COVID-19," *Journal of Primary Care \& Community Health*, vol. 15, p. 21501319241235588, 2024.
- [15] M. A. a. A.-M. F. a. H. M. E. a. J. M. H. I. a. L. M. H. a. M. N. B. a. I. T. a. T. M. K.

- a. C. T. B. K. a. S. N. P. a. o. Mamun, "redictive Modeling of Comorbid Depression and Anxiety Symptoms Among Prospective University Students: A GIS-Based and Machine Learning Study," *Psychological Reports*, p. 00332941251358203, 2025.
- [16] N. a. H. A. K. a. H. M. A. a. D. P. P. a. D. M. N. a. R. M. M. H. a. H. A. a. B. M. R. a. H. M. S. Mumenin, "Screening depression among university students utilizing GHQ-12 and machine learning," *Heliyon*, vol. 10, 2024.
- [17] O. a. P.-M. C. a. E.-H. A. a. T.-C. C. Iparraguirre-Villanueva, "Machine Learning Models to Classify and Predict Depression in College Students," *International Journal of Interactive Mobile Technologies*, vol. 18, no. 14, p. 148, 2024.
- [18] N. a. Y. M. A. a. A. M. O. a. M. M. M. a. H. M. A. Mumenin, "DDNet: A Robust, and Reliable Hybrid Machine Learning Model for Effective Detection of Depression Among University Students," *IEEE Access*, 2025.
- [19] M. I. a. R. A. A. a. N. M. M. a. A. M. M. A. a. F. F. a. K. L. I. a. A. A. Siraji, "Impact of mobile connectivity on students' wellbeing: Detecting learners' depression using machine learning algorithms," *Plos one*, vol. 18, p. e0294803, 2023.
- [20] A. H. a. R. D. a. R. M. S. Chowdhury, "Predicting anxiety, depression, and insomnia among Bangladeshi university students using tree-based machine learning models," *Health Science Reports*, vol. 7, p. e2037, 2024.
- [21] U. B. a. K. M. S. a. I. U. I. a. S. F. H. Munir, "Machine Learning Classification Algorithms for Predicting Depression Among University Students in Bangladesh," *Proceedings of the Third International Conference on Trends in Computational and Cognitive Engineering: TCCE 2021*, pp. 69--80, 2022.
- [22] R. a. I. N. a. B. J. F. a. K. R. a. M. S. Siddiqua, "AIDA: Artificial intelligence based depression assessment applied to Bangladeshi students," *Array*, vol. 18, p. 100291, 2023.
- [23] A. a. A. A. A. a. F. S. a. M. M. S. Rahman, "Prediction of Depression and Anxiety on University Students in Bangladesh Using Machine Learning," *2023 5th International Conference on Sustainable Technologies for Industry 5.0 (STI)*, pp. 1--6, 2023.
- [24] A. B. a. C. M. Siddique, "A machine learning-based approach to predict university students' depression pattern and mental healthcare assistance scheme using Android application," *International Journal of Data Analysis Techniques and Strategies*, vol. 14, pp. 122--139, 2022.
- [25] M. a. R. A. a. A. L. a. F. K. a. K. R. H. a. K. M. R. a. H. M. S. a. U. M. F. Syeeda, "A comprehensive standardized dataset on mental health problems (mhps) of university students," *measurement*, vol. 8, p. 9, 2024.

211-15-3978

ORIGINALITY REPORT

21%

SIMILARITY INDEX

15%

INTERNET SOURCES

17%

PUBLICATIONS

13%

STUDENT PAPERS

PRIMARY SOURCES

1

dspace.daffodilvarsity.edu.bd:8080

Internet Source

1%

2

"Brain Informatics", Springer Science and Business Media LLC, 2025

Publication

1%

3

www.frontiersin.org

Internet Source

1%

4

Submitted to Universidad Internacional de la Rioja

Student Paper

<1%

5

Submitted to Georgia State University

Student Paper

<1%

6

Dr. M M Mahbubul Syeed, Ashifur Rahman, Laila Akter, Kaniz Fatema et al. "A comprehensive standardized dataset on Mental Health Problems (MHPs) of University Students", Open Science Framework, 2024

Publication

<1%

7

arxiv.org

Internet Source

<1%

8

Submitted to University of South Florida

Student Paper

<1%

9

inass.org

Internet Source

<1%

10

Thangaprakash Sengodan, Sanjay Misra, M Murugappan. "Advances in Electrical and

<1%