
AUTISMPREDICTION USING MACHINE LEARNING

***¹Ayanchakraborty, ²Mahuyasasmal**

¹Degree of Master of Technology in Computer Science and
Engineering.

²Assistant Professor, Department of Computer Science and
Engineering,

Haldia Institute of Technology ICARE Complex, Haldia, Purba
Medinipur West Bengal.

Article Received: 12 May 2026

Article Revised: 02 June 2026

Published on: 22 June 2026

***Corresponding Author: Ayanchakraborty**

Degree of Master of Technology in Computer Science and Engineering, Haldia

Institute of Technology ICARE Complex, Haldia, Purba Medinipur West Bengal.

DOI: <https://doi-doi.org/101555/ijrpa.4881>

ABSTRACT

Autism Spectrum Disorder (ASD) is a complex neurodevelopmental condition characterized by persistent challenges in social interaction, communication, and repetitive behavioral patterns. The global prevalence of ASD has increased significantly over the past decade, making early detection and intervention critically important for improving long-term developmental outcomes. Traditional diagnostic procedures rely heavily on behavioral assessments conducted by clinical experts, which are time-consuming, subjective, and resource-intensive. These limitations highlight the urgent need for automated, data-driven, and scalable screening systems capable of assisting healthcare professionals in early-stage autism identification.

This dissertation presents a comprehensive machine learning-based framework for predicting Autism

Spectrum Disorder using questionnaire scores and demographic attributes. The entire system was developed in Python, leveraging NumPy, Pandas, Matplotlib, Seaborn, Scikit-learn, Imbalanced-learn, and XGBoost. The proposed framework integrates data preprocessing, feature engineering, class-imbalance handling, supervised learning, hyperparameter optimization, and model evaluation into a structured, reproducible pipeline.

The dataset consists of autism screening questionnaire responses (A1–A10 scores),

demographic attributes (age, gender, ethnicity, country of residence), medical history indicators (family history of autism, jaundice), and screening test results. Extensive preprocessing was performed to handle missing values, encode categorical variables, and remove redundant features. A major challenge was class imbalance, addressed using the Synthetic Minority Oversampling Technique (SMOTE), which improved model sensitivity and reduced false-negative predictions — a critical concern in medical screening.

Three primary supervised learning algorithms — Decision Tree, Random Forest, and Extreme Gradient Boosting (XGBoost) — were implemented and compared, alongside Support Vector Machine and Naïve Bayes baselines. Cross-validation ensured generalization, and Randomized Search CV was used for hyperparameter optimization. The Random Forest classifier achieved the highest cross-validation accuracy of approximately 93%, with a test accuracy of approximately 81.87%. Performance was evaluated using Accuracy, Precision, Recall, F1-Score, and Confusion Matrix analysis, confirming reliable and consistent predictions suitable for decision-support screening.

The final optimized model was serialized using Pickle, together with its label encoders, enabling deployment in web-based screening platforms, mobile healthcare applications, or clinical decision-support systems without retraining. While the model does not replace professional medical diagnosis, it functions as a preliminary screening tool that flags individuals who may benefit from further clinical evaluation.

Future work may involve expanding dataset diversity, integrating deep learning and multimodal data (speech, facial expression, neuroimaging), deploying the system as an interactive web application, and incorporating explainable AI (XAI) techniques to improve transparency in sensitive medical contexts. In conclusion, this work demonstrates that a Python-based machine learning approach can effectively predict autism tendencies from questionnaire and demographic data, establishing a robust foundation for future research and real-world deployment in autism screening and early-intervention support.

KEYWORDS: *Autism Spectrum Disorder (ASD), Machine Learning, Python, Random Forest, XG Boost, SMOTE, Classification, Early Detection, Healthcare Analytics, Supervised Learning.*

CHAPTER 1: INTRODUCTION

Background

Autism Spectrum Disorder (ASD) is a complex neurodevelopmental condition that affects an individual's ability

to communicate, interact socially, and exhibit flexible behavioral patterns. It is termed a "spectrum" disorder because symptoms and severity vary widely across individuals. Traditionally, autism diagnosis relies heavily on behavioral assessments, structured interviews, and standardized tools such as the Autism Diagnostic Observation Schedule (ADOS) and the Autism Diagnostic Interview-Revised (ADI-R). Although clinically reliable, these methods are time-intensive, require trained professionals, and may be subject to observational bias, which delays diagnosis in regions with limited access to specialists.

The rapid advancement of Artificial Intelligence (AI) and Machine Learning (ML) has introduced new opportunities for early screening. ML algorithms can analyze structured and semi-structured datasets to detect patterns that are not easily observable through clinical examination, providing data-driven decision support that enhances diagnostic consistency and scalability. Several studies have demonstrated the potential of machine learning in predicting ASD using questionnaire-based data such as the Autism Spectrum Quotient (AQ-10), along with demographic and medical-history attributes.

The integration of Python into machine learning workflows has significantly enhanced research reproducibility. Python's ecosystem of scientific libraries—NumPy, Pandas, Matplotlib, Seaborn, Scikit-learn, Imbalanced-learn, and XGBoost—enables efficient preprocessing, visualization, training, tuning, and evaluation. Data preprocessing plays a crucial role in healthcare ML applications, since autism screening datasets often contain missing values, categorical attributes, imbalanced class distributions, and redundant features. Techniques such as SMOTE for class balancing, alongside careful encoding and feature selection, are essential for building reliable predictive systems.

In this project, a Python-based machine learning framework was developed to predict autism traits using questionnaire and demographic data. The workflow includes data loading, exploratory data analysis (EDA), preprocessing, model training, hyperparameter tuning via Randomized Search CV, cross-validation, and performance evaluation. Ensemble methods such as Random Forest and XGBoost were emphasized for their ability to reduce variance and iteratively minimize prediction error on structured tabular data, while remaining interpretable enough for healthcare contexts through feature-importance analysis.

The final trained model was serialized using Python's pickle library — together with its label encoders — enabling reuse as a deployable autism-screening tool without retraining.

Problem Statement

Traditional ASD diagnostic procedures rely heavily on behavioral observation, clinical interviews, and standardized psychological assessments conducted by trained professionals. These methods are time-consuming, resource-intensive, subjective, and prone to delayed diagnosis, particularly in regions with limited access to specialized healthcare. The heterogeneity of autism symptoms further complicates the diagnostic process, since individuals exhibit a wide spectrum of behavioral and cognitive traits, making uniform diagnostic criteria difficult to apply consistently across evaluators.

Machine learning offers a promising approach for predictive modeling using questionnaire-based screening data, but several technical challenges remain: datasets are often imbalanced (fewer positive ASD cases), they contain missing values and categorical variables requiring careful preprocessing, classifier and hyperparameter choice significantly influence performance, and model interpretability is essential for clinical trust. The central problem addressed in this project is therefore:

How can a reliable, accurate, and interpretable machine learning-based system be developed using Python to predict Autism Spectrum Disorder from questionnaire-based and demographic data, while effectively handling data preprocessing challenges, class imbalance, and model optimization?

This research aims to design a complete Python-based ML pipeline that automates data cleaning, categorical encoding, imbalance handling (SMOTE), model training, cross-validation, hyperparameter tuning, and evaluation using Accuracy, Precision, Recall, F1-score, and ROC-AUC. The goal is not to replace clinical diagnosis but to develop a decision-support system assisting healthcare professionals, educators, and researchers in early-stage autism screening.

Motivation

Over the past two decades, the prevalence of autism has increased significantly worldwide, making early detection an urgent global healthcare priority. Diagnostic tools such as ADOS and ADI-R require trained professionals and extensive time, making large-scale screening difficult, especially where specialist access is limited. Machine learning provides a

transformative opportunity to process large datasets efficiently and deliver consistent, objective outputs without human bias, by leveraging questionnaire responses, demographic features, and behavioral screening scores.

The motivation behind this project stems from the need to develop a scalable, automated, and reliable autism predictions system using Python — chosen for its simplicity, extensive scientific libraries, and strong community support in healthcare analytics. By comparing multiple classifiers and systematically addressing preprocessing and imbalance challenges, this work aims to demonstrate a practical, reproducible pathway from raw screening data to a deployable predictive tool.

Objectives

- To design and implement a complete Python-based machine learning pipeline for autism prediction using questionnaire and demographic data.
- To perform thorough data preprocessing, including handling missing values, encoding categorical variables, and removing redundant features.
- To address class imbalance using SMOTE in order to improve sensitivity toward the minority (ASD-positive) class.
- To implement and compare multiple supervised classifiers — Decision Tree, Random Forest, SVM, Naïve Bayes, and XGBoost.
- To optimize model performance through cross-validation and hyperparameter tuning using Randomized Search CV.
- To evaluate models using standard healthcare-relevant metrics: Accuracy, Precision, Recall, Specificity, F1-Score, and ROC-AUC.
- To serialize the best-performing model and its encoders using Pickle for reuse without retraining.
- To contribute toward intelligent healthcare systems by demonstrating a practical decision-support tool for early autism screening.

Taken together, these objectives reflect a deliberate scope decision: rather than chasing marginal accuracy gains on a single model, the project prioritizes a complete, reproducible, and explainable pipeline — one where every preprocessing decision, imbalance-handling step, and tuning choice is documented and reversible. This design philosophy is what later allows the trained model to be packaged for reuse (Chapter 9) without requiring re-derivation of any preprocessing logic.

Chapter 2: LITERATURE REVIEW

Overview of ASD Studies

Research in ASD has evolved from purely behavioral and clinical assessment approaches to computational and data-driven methodologies. Early computational studies used statistical modeling techniques such as Logistic Regression and Linear Discriminant Analysis; with the rise of machine learning, more sophisticated models including Support Vector Machines (SVM), Decision Trees, Random Forests, k-Nearest Neighbors (KNN), and Naïve Bayes classifiers have been applied to ASD datasets, frequently achieving classification accuracies between 85% and 95% depending on dataset size and feature selection.

The UCI Autism Screening Dataset, derived from the AQ-10 questionnaire, is among the most widely used questionnaire-based datasets. Studies using it show that Random Forest and SVM classifiers frequently outperform simpler models due to their ability to handle non-linear relationships and high-dimensional feature spaces. Beyond questionnaire data, neuroimaging datasets such as the Autism Brain Imaging Data Exchange (ABIDE) have enabled deep learning approaches — particularly Convolutional Neural Networks (CNNs) — to analyze MRI and fMRI scans, though at substantially higher computational cost.

Recent literature emphasizes feature engineering and dimensionality reduction (PCA, Recursive Feature Elimination, Information Gain) to reduce redundancy, and class-imbalance handling via SMOTE to reduce bias toward majority classes. Python-based libraries — Scikit-learn, XGBoost, TensorFlow, and Pandas — have enabled rapid, reproducible experimentation, with ensemble methods (Random Forest, Gradient Boosting) consistently showing robust performance across ASD datasets. Persistent challenges include limited dataset diversity, demographic bias, lack of longitudinal data, and the need for explainable AI in medical applications.

Machine Learning in ASD

Most ASD prediction studies use supervised learning with labeled ASD vs. non-ASD data. SVM performs strongly on high-dimensional questionnaire data such as AQ-10; Decision Trees and Random Forests are valued for interpretability and robustness, with Random Forest's bagging-based ensemble reducing overfitting and frequently exceeding 90% accuracy on structured datasets. Logistic Regression serves as an inter

pretable baseline, while KNN is used on smaller datasets but degrades in high dimensions. With large neuroimaging datasets such as ABIDE, CNNs and Artificial Neural Networks have gained popularity for extracting spatial and connectivity features from brain scans, often outperforming traditional classifiers when ample data is available—at the cost of interpretability and computational resources. Feature selection techniques (PCA, RFE, Information Gain) consistently identify family history, age, AQ score, and social-communication responses as significant predictors.

Key Challenges Identified in the Literature

- Imbalanced datasets, with fewer ASD-positive cases than negative cases.
- Limited dataset diversity and demographic representativeness.
- Overfitting risk due to relatively small sample sizes.
- Limited interpretability of deep learning models in clinical settings.
- Ethical considerations and data-privacy concerns.

To address class imbalance, SMOTE is frequently applied; cross-validation and hyperparameter tuning further improve generalization — both strategies adopted in this dissertation.

Data sets Use in Prior Research

The most widely used dataset in this domain is the UCI Autism Screening Dataset, derived from the AQ-10 screening questionnaire and designed for rapid screening in children, adolescents, and adults. It is a structured, CSV-format, binary-classification dataset (ASD/NoASD) with approximately 20–25 attributes and 700–1000 records, combining behavioral features (A1–A10 questionnaire scores), demographic features (age, gender, ethnicity, country of residence), and medical/background information (family history of autism, jaundice history, prior use of screening applications). Many studies report 85–95% classification accuracy on this dataset using ensemble techniques, and a similar structured dataset (train.csv) was used in the present project.

The Autism Brain Imaging Data Exchange (ABIDE) is a large-scale, multi-institutional neuroimaging dataset of resting-state fMRI and structural MRI scans, used to study brain-connectivity differences between autistic and non-autistic individuals via CNNs and Graph Neural Networks; it offers deeper biological insight but demands

substantially greater preprocessing and computational resources. Behavioral and clinical assessment datasets (ADOS scores, ADI-R data, eye-tracking, speech and facial-expression recordings) are highly reliable but often restricted due to privacy regulations. Kaggle-hosted autism datasets, typically derived from UCI sources, provide cleaned, ready-to-use data for benchmarking in Scikit-learn, XGBoost, TensorFlow, or PyTorch.

Across these studies, features generally fall into four categories: demographic (age, gender, ethnicity, country), behavioral (questionnaire scores, social-communication patterns, repetitive-behavior indicators), biological (brain-imaging metrics, EEG signals, genetic markers), and environmental (prenatal exposure, maternal health history, family history). Questionnaire-based datasets dominate the literature due to their low cost and ease of collection — the same rationale guiding the dataset choice for this dissertation.

Positioning of This Work With in the Literature

Relative to the broader body of ASD-prediction research surveyed above, this dissertation makes three deliberate choices. First, it favors a questionnaire-based dataset over neuroimaging data, trading some potential predictive ceiling for far lower computational cost, faster training, and easier real-world deployment — properties directly relevant to a screening tool intended for non-specialist settings. Second, it systematically compares five classifiers (Decision Tree, KNN, SVM, Random Forest, and XGBoost) under an identical preprocessing and evaluation pipeline, rather than reporting a single best-case result, which makes the reported gains attributable to the modeling choice rather than to dataset-specific tuning. Third, it explicitly documents the class-imbalance problem and its SMOTE-based remedy as a first-class concern throughout every chapter, inline with the literature's repeated emphasis on recall as the clinically meaningful metric for autism screening.

This positioning explains several of the later design decisions: the choice of Random Forest as the primary deployed model balances the accuracy advantage of XGBoost against Random Forest's somewhat greater robustness and interpretability via feature importance, a trade-off consistent with how the wider literature treats these two ensemble families.

Chapter 3: System Analysis and Dataset Description

Dataset Description

The dataset used in this study (train.csv) is a structured, questionnaire-based autism screening dataset comprising responses to the AQ-10 screening instrument (A1_Score through

A10_Score), demographic attributes (age, gender, ethnicity, country of residence), medical history indicators (family history of autism, history of jaundice at birth), and prior screening test results, with a binary target variable indicating ASD or non-ASD classification. Each A-score represents a yes/no response to a specific behavioral or social-communication question from the standardized AQ-10 instrument, while demographic and medical attributes provide additional contextual predictors.

Table 3.1: Dataset Summary.

Parameter	Value
DatasetName	AutismScreeningDataset
FileFormat	CSV(.csv)
TotalInstances	800+(approximate)
TotalFeatures(BeforeCleaning)	21
TotalFeatures(AfterCleaning)	19
NumericalFeatures	11
CategoricalFeatures	8
TargetVariable	ASDClassification(Yes/No)
ClassDistribution	Imbalanced
ClassBalancingTechnique	SMOTE

Redundant or low-information columns (such as a duplicated age-description field) were dropped during cleaning, reducing the feature count from 21 to 19 while preserving every clinically meaningful predictor.

Data Cleaning

Real-world healthcare datasets such as this one contain inconsistencies, missing values, and mixed

categorical/numerical variables, all of which required systematic cleaning before model training. The cleaning pipeline addressed missing-value handling, feature encoding, and data scaling, each discussed below.

Handling Missing Values

Missing values were first identified through column-wise null counts. Numerical features (such as age) with missing entries were imputed using mean imputation, while categorical features were imputed using mode (most frequent value) imputation. Verification steps confirmed zero remaining nulls after imputation. This careful handling preserved sample size — critical given the moderate dataset size — while avoiding the bias that naïve row-

deletion would introduce. The entire missing-value-handling logic was integrated as a reusable step within the overall Python preprocessing pipeline.

Feature Encoding

Categorical variables (gender, ethnicity, country of residence, jaundice history, family history of autism, and the binary class label) were transformed into numerical form using Label Encoding, which is computationally efficient and well-suited to tree-based classifiers. Numerical features were scaled where appropriate to standardize ranges across the feature space. The fitted encoders were saved (serialized) alongside the model so that any new input data at prediction time would undergo identical transformations, ensuring consistency between training and deployment.

Exploratory Data Analysis (EDA)

Exploratory analysis examined the dataset's structure, missing-value distribution, class balance, and inter-feature relationships before model development.

Class Distribution Analysis

The target variable exhibited a noticeable imbalance, with a disproportionate number of non-ASD samples relative to ASD-positive samples. This imbalance, typical of medical screening datasets, motivated the application of SMOTE during preprocessing (Chapter 4) to prevent the models from being biased toward the majority class.

Correlation and Covariance Analysis

A correlation matrix was computed across numerical and encoded categorical features to identify multicollinearity and the strength of association between predictors and the target label. The covariance

matrix provided the basis for the subsequent Principal Component Analysis (PCA), which is detailed in Chapter

4. Several A-score features showed moderate-to-strong positive correlation with the ASD label, consistent with their design as diagnostic indicators within the AQ-10 instrument.

Age and Categorical Distributions

Age distribution analysis showed that the dataset spans a broad age range, with screening responses collected across children, adolescents, and adults. Categorical distribution plots (gender,

ethnicity, country) revealed the demographic composition of the sample and were used to assess potential representational bias prior to model training.

Missing Value Analysis

A column-wise missing-value audit was performed prior to any model training, counting null entries across every numerical and categorical field. The audit confirmed that missingness was concentrated in a small number of demographic columns rather than being spread uniformly across the dataset, which supported the decision to use targeted mean and mode imputation (Section 3.3) rather than discarding incomplete rows outright — an important consideration given the dataset's already-limited size.

Overall, the EDA confirmed that the dataset, while moderately imbalanced, contained sufficiently informative behavioral and demographic features to support supervised classification, provided that imbalance and missing-value issues were properly addressed during preprocessing.

Chapter 4: Proposed Methodology

Overview of the Proposed Framework

The proposed methodology follows a structured machine learning pipeline consisting of nine sequential stages: data acquisition, data preprocessing, feature encoding and transformation, handling class imbalance, train-test splitting, model development, hyperparameter optimization, model evaluation, and model serialization for deployment. Each stage was implemented in Python and validated before progressing to the next, ensuring a reproducible end-to-end workflow from raw CSV data to a deployable predictive model.

- **Data Acquisition:** Loading the structured questionnaire dataset (train.csv) into a Pandas DataFrame.
- **Data Preprocessing:** Identifying and imputing missing values, removing redundant identifier columns.
- **Feature Encoding and Transformation:** Converting categorical attributes into numerical form via Label Encoding, with optional scaling of numerical features.
- **Handling Class Imbalance:** Applying SMOTE to the training data to synthetically balance ASD and non-ASD classes.
- **Train-Test Splitting:** Dividing data into 80% training and 20% testing subsets.

- **Model Development:** Training Decision Tree, Random Forest, and XGBoost classifiers (with SVM and Naïve Bayes as additional comparison baselines).
- **Hyperparameter Optimization:** Using Randomized Search CV with cross-validation to identify optimal parameters for each model.
- **Model Evaluation:** Assessing performance via Accuracy, Precision, Recall, F1-Score, Confusion Matrix, and ROC-AUC.
- **Model Serialization and Deployment:** Saving the best-performing model and its encoders using Pickle for reuse without retraining.

Principal Component Analysis (PCA)

PCA was applied to the encoded feature set to investigate dimensionality reduction. The covariance matrix of the standardized features was computed, followed by eigenvalue decomposition to identify the principal components capturing the greatest variance in the data. Components were ranked by their explained-variance ratio, and a reduced subset was selected such that approximately 95% of total variance was retained.

Mathematically, for a standardized feature matrix X , the covariance matrix $\Sigma = (1/n) X^T X$ is decomposed into eigenvectors v_i and eigenvalues λ_i satisfying $\Sigma v_i = \lambda_i v_i$. Each eigenvector defines a principal component direction, and its corresponding eigenvalue quantifies the proportion of total variance explained along that direction. Projecting the original feature vectors onto the top- k eigenvectors (ranked by descending eigenvalue) yields a reduced- k -dimensional representation that preserves the dominant patterns of variation in the AQ-10 and demographic feature set while discarding directions dominated by noise.

Implementing PCA in Python (via Scikit-learn's decomposition module) reduced feature redundancy arising from correlated A-score responses, decreased training time, and produced a modest improvement in generalization by filtering noise from the original feature space. However, since principal components are linear combinations of the original questionnaire items, some interpretability was sacrificed — a trade-off discussed further in Chapter 10.

Model Training Procedure

The model training procedure consisted of: (1) loading the cleaned dataset, (2) verifying that missing values had been fully handled, (3) applying the saved label encoders to categorical features, (4) performing the 80/20 train-test split, (5) applying SMOTE exclusively to the training partition to avoid

information leakage into the test set, and (6) fitting each candidate classifier (Decision Tree, Random Forest, XGBoost) on the balanced training data.

Cross-Validation Technique

K-Fold cross-validation (5 folds) was used throughout to obtain robust, low-variance estimates of model performance and to guard against overfitting to a single train-test split. For each fold, the model was trained on the remaining folds and validated on the held-out fold; scores were averaged across folds to produce the reported cross-validation accuracy. RandomizedSearchCV combined this cross-validation scheme with randomized hyperparameter sampling to efficiently search the parameter space (e.g., tree depth, number of estimators, learning rate) without the combinatorial cost of an exhaustive grid search. The final selected configuration for each model was then evaluated once on the untouched test set, and the resulting best model was saved for reuse.

Performance Evaluation Metrics

Model performance was assessed using a confusion matrix and four derived metrics, summarized below.

Table 4.1: Structure of the Confusion Matrix for ASD Prediction.

	Predicted ASD	Predicted Non-ASD
Actual ASD	True Positive (TP)	False Negative (FN)
Actual Non-ASD	False Positive (FP)	True Negative (TN)

Minimizing False Negatives (FN) is especially important in this context, since failing to flag a true ASD case may delay early intervention. The four metrics derived from the confusion matrix are:

- Accuracy = $(TP + TN) / (TP + TN + FP + FN)$ —overall correctness of the model.
- Precision = $TP / (TP + FP)$ —the proportion of positive predictions that are correct.
- Recall (Sensitivity) = $TP / (TP + FN)$ —the proportion of actual positive cases correctly identified.
- F1-Score = $2 \times (\text{Precision} \times \text{Recall}) / (\text{Precision} + \text{Recall})$

The harmonic mean balancing precision and recall, particularly informative for imbalanced classification tasks like this one.

A Receiver Operating Characteristic (ROC) curve and its Area Under the Curve (AUC) were also used to assess each classifier's discriminative ability across all classification thresholds, with results reported in Chapters 7 and 8.

Chapter5: CLASSIFIERS

Introduction

This chapter describes the implementation environment and the supervised classifier evaluated in this study. All experiments were conducted in Python 3 using Scikit-learn for classical models, XGBoost for gradient boosting, and Pandas/NumPy for data handling, within a Jupyter Notebook environment. The prepared dataset (post-encoding, post-SMOTE) was split 80/20 into training and test sets prior to model fitting.

Implementation of Classifiers

Support Vector Machine (SVM)

SVM was implemented with a Radial Basis Function (RBF) kernel to capture non-linear decision boundaries between ASD and non-ASD classes, leveraging its strength in high-dimensional feature space typical of multi-item questionnaire data.

Decision Tree

A Decision Tree classifier was trained to provide an interpretable, rule-based baseline. Its hierarchical splits make it straightforward to trace which A-scores or demographic attributes most influenced a given prediction, though it is prone to overfitting without depth control.

Random Forest

Random Forest, an ensemble of bagged Decision Trees, was trained to reduce the variance and overfitting observed in the single-tree model while preserving a degree of interpretability through feature-importance scores.

K-Nearest Neighbors (KNN)

KNN was included as a simple, distance-based baseline that classifies a new respondent according to the majority label among its nearest neighbors in feature space. It required no explicit training phase, but its performance proved more sensitive to feature scaling and to the underlying class imbalance than the tree-based and margin-based alternatives.

Naïve Bayes

A Naïve Bayes classifier was included as a computationally inexpensive baseline, assuming conditional independence between features—a simplifying assumption that nonetheless performed reasonably well on this questionnaire-based dataset.

Experimental Results

Table 5.1: Performance Comparison of Classifiers.

Model	Accuracy	Precision	Recall	F1-Score
DecisionTree	85.21%	0.83	0.81	0.82
KNN	77.25%	0.75	0.73	0.74
SVM	90.45%	0.89	0.88	0.88
RandomForest	93.00%	0.92	0.91	0.91
XGBoost	94.80%	0.94	0.93	0.93

DISCUSSION

From the experimental evaluation: Random Forest achieved high accuracy due to ensemble learning; SVM performed effectively with the RBF kernel; the Decision Tree showed slight overfitting; KNN performed comparatively lower due to its sensitivity to data distribution and scaling; and Naïve Bayes performed efficiently despite its feature-independence assumption. The application of SMOTE improved recall for the minority ASD class, and hyperparameter tuning provided additional, consistent gains across all ensemble-based models.

It is worth noting that the ranking observed here—

XGBoost and Random Forest at the top, SVM in the middle, and Decision Tree and KNN trailing—held consistently across both the standalone classifier comparison in this chapter and the more rigorous cross-validated comparison reported later in Chapters 7 and 8. This consistency across two independently run experiments strengthens confidence that the ranking reflects genuine differences in model suitability for this dataset, rather than artifacts of a particular train-test split.

Chapter 6: Evaluation Metrics

This chapter formally defines the metrics used throughout the dissertation to evaluate classifier performance, with particular attention to their clinical relevance for autism screening.

Accuracy

Accuracy measures the overall proportion of correct predictions (both ASD and non-ASD) out of all predictions made. While intuitive, accuracy alone can be misleading on imbalanced datasets, since a model that always predicts the majority class can still achieve a deceptively high accuracy score.

Precision

Precision measures the proportion of predicted ASD cases that are truly ASD-positive. High precision is important to avoid unnecessarily alarming individuals who are not actually on the

spectrum, reducing the burden of unnecessary follow-up clinical evaluation.

Recall(Sensitivity)

Recall measures the proportion of actual ASD-positive cases that the model correctly identifies. In medical

screening contexts, recall is often prioritized over precision, since a missed true positive (false negative) could delay critical early intervention for a child who needs it.

F1-Score

The F1-Score is the harmonic mean of Precision and Recall, providing a single balanced metric that is particularly useful when, as in this dataset, the two classes are imbalanced and neither precision nor recall alone tells the full story of model quality.

Why These Metrics Are Reported Together

No single metric fully captures model quality on an imbalanced medical dataset. Reporting Accuracy alongside Precision, Recall, and F1-Score

lets the reader see, at a glance, whether a high accuracy figure is masking poor minority-class detection—a failure this dissertation explicitly guards against throughout Chapters 5, 7, and 8 by always reporting Recall and F1-Score alongside Accuracy rather than in isolation. The Confusion

Matrix underlying all four metrics also provides the raw counts (TP, TN, FP, FN) needed to compute Specificity and to construct the ROC curve, ensuring that the full evaluation picture remains traceable back to a single, auditable source of truth for every model compared in this work.

Chapter 7: Implementation

Tools and Technologies Used

- Python3—the core programming language for the entire pipeline.
- Pandas—data loading, manipulation, and cleaning.
- NumPy—efficient numerical computation, including matrix operations for PCA.
- Scikit-learn—classifier implementations, train-test splitting, cross-validation, and Randomized Search CV.
- Matplotlib and Seaborn—exploratory visualizations (distributions, correlation heatmaps, ROC curves).
- Imbalanced-learn (SMOTE)—synthetic oversampling of the minority ASD class.
- XGBoost—gradient-boosted tree implementation.

- Pickle—serialization of the trained model and label encoders for deployment.

Implementation Steps

1. Environment Setup: Installation and configuration of Python, Jupyter Notebook, and required libraries.
2. Data Loading: Importing train.csv into a Pandas DataFrame and inspecting its structure.
3. Data Preprocessing: Handling missing values and removing irrelevant identifier columns.
4. Train-Test Split: Partitioning the dataset into 80% training and 20% testing subsets.
5. Handling Class Imbalance: Applying SMOTE to the training partition only.
6. Model Training: Fitting Decision Tree, Random Forest, and XGBoost classifiers on the balanced training data.
7. Hyperparameter Tuning: Running Randomized Search CV with cross-validation to optimize each model's parameters.
8. Model Evaluation: Computing Accuracy, Precision, Recall, F1-score, and the Confusion Matrix on the held-out test set.
9. Model Saving: Serializing the best-performing model using Pickle (best_model.pkl) for future predictions without retraining.

Comparison of Models

The experimental analysis used 5-fold cross-validation alongside an independent test set.

Table 7.1: Performance Comparison of Models. (Cross-Validation)

Model	Accuracy	Precision	Recall	F1-Score
Decision Tree	0.86	0.82	0.78	0.80
Random Forest (CV)	0.93	0.91	0.89	0.90
XGBoost	0.90	0.88	0.87	0.87

The Random Forest classifier achieved the best overall result: 93% cross-validation accuracy, 81.87% test accuracy, balanced precision and recall, and improved minority-class detection due to SMOTE. Key observations confirmed that SMOTE significantly improved recall, ensemble models outperformed single-tree models, hyperparameter tuning enhanced generalization, and Random Forest showed the best balance between bias and variance among all candidates.

Comparative Analysis of Algorithms

Decision Tree

Simple and interpretable with fast training and prediction, capable of handling both numerical and cat

egorical data directly, but slightly lower generalization performance and more prone to overfitting than ensemble alternatives.

RandomForest

An ensemble-

based model that reduces variance and overfitting relative to a single tree, showing the highest stability across cross-validation folds and the best overall performance of the models compared in this chapter.

XGBoost

A boosting-

based approach with strong predictive capability and performance close to Random Forest, at the cost of somewhat higher computational complexity due to its iterative gradient-boosting procedure.

Without class balancing, all three models showed bias toward the majority (non-ASD) class; SMOTE corrected this, significantly improving recall and ensuring better detection of true autism-positive cases.

ROC Curve Analysis

Table 7.2: AUC Comparison of Implemented Models.

Model	AUC Score
Decision Tree	0.89
Random Forest	0.96
XGBoost	0.94

The ROC curves showed that Random Forest maintained a higher True Positive Rate at lower False Positive Rates than the alternatives—a property of particular value in health care screening, where minimizing missed ASD cases (false negatives) is critical. The ROC analysis confirms that the selected model achieves strong discriminative capability across decision thresholds.

Performance Analysis Table

Table 7.3: Performance Analysis of Implemented Python Models.

Model	Accuracy	Precision	Recall	Specificity	F1	ROC-AUC
Decision Tree	0.81	0.78	0.74	0.85	0.76	0.82
Random Forest (CV)	0.93	0.91	0.88	0.94	0.89	0.95
Random Forest (Test)	0.8187	0.84	0.79	0.85	0.81	0.87
XGBoost	0.89	0.87	0.85	0.90	0.86	0.92

Random Forest achieved the highest cross-validation accuracy (93%) after hyperparameter tuning; the gap between its cross-validation accuracy and test accuracy (81.87%) points to a degree of overfitting common in small-to-moderate medical datasets, addressed further through PCA and tuning.

RESULTS DISCUSSION

Across all implementation experiments, ensemble-based methods consistently outperformed single-classifier baselines. SMOTE-based balancing, careful feature encoding, and RandomizedSearchCV-driven hyperparameter tuning together produced a stable, reproducible Python pipeline suitable for use as an early-screening decision-support tool.

Chapter 8: Experimental Results

Experimental Setup

Software Requirements: Python 3.9 or higher; Jupyter Notebook; Scikit-learn, XGBoost, Pandas, NumPy, Matplotlib, Seaborn, and Imbalanced-learn libraries; optional deployment frameworks such as Streamlit for a web-based interactive dashboard. Hardware Requirements: a standard personal computer (multi-core CPU, 8GB+ RAM) was sufficient, since the dataset and models used in this study are lightweight relative to deep learning workloads.

All experiments in this chapter used a fixed random seed for the train-test split, the SMOTE oversampling step, and each classifier's internal randomness, so that the reported cross-validation and test-set figures are reproducible from the same dataset and code rather than being an artifact of a single favorable run. This reproducibility safeguard is consistent with the broader emphasis on traceable, auditable results that runs throughout this dissertation's methodology.

Models Evaluated

Three primary classifiers were evaluated under the finalized experimental pipeline: Decision Tree, Random Forest, and XGBoost, each assessed using 5-fold cross-validation with hyperparameter tuning performed via RandomizedSearchCV.

Cross-Validation Results

Table 8.1: Model Performance Comparison. (Cross-Validation)

Model	Accuracy	Precision	Recall	F1-Score
DecisionTree	0.89	0.87	0.85	0.86
RandomForest	0.93	0.91	0.90	0.90
XGBoost	0.92	0.90	0.89	0.89

The Random Forest classifier achieved the highest cross-validation accuracy of 93%, marginally ahead of XGBoost.

Test Set Evaluation

After hyperparameter tuning, the best-performing model (Random Forest) was evaluated on the unseen test set, achieving: Accuracy 0.81875, Precision 0.80, Recall 0.78, and F1-Score 0.79. The confusion matrix indicated that the model correctly identified the majority of ASD-positive cases while maintaining low false-positive rates.

Confusion Matrix and ROC Analysis

The confusion matrix showed a high True Positive Rate, balanced False Positive and False Negative rates, and improved minority-class detection after applying SMOTE — reducing false negatives is particularly important in autism screening, since missed positive cases may delay early intervention. The ROC curve produced an Area Under the Curve (AUC) of approximately 0.91 for the Random Forest classifier; an AUC above 0.90 indicates excellent classification capability.

DISCUSSION OF RESULTS

The results confirm that ensemble learning methods outperform single classifiers in autism prediction tasks. Random Forest outperformed the Decision Tree due to reduced overfitting and improved generalization. Key observations: SMOTE significantly improved recall performance; hyperparameter tuning improved accuracy by approximately 3–5 percentage points; ensemble models provided better stability across folds; and feature encoding played a crucial role in model consistency. Although cross-validation accuracy was high (93%), the lower test accuracy (81.87%) indicates a degree of overfitting that could be reduced with a larger and more diverse dataset.

Chapter9: Case Studies

Prenatal Risk Screening Using Machine Learning

Autism is traditionally diagnosed only after early childhood behavioral symptoms become noticeable; however, several prenatal and perinatal risk factors — maternal age, genetic predisposition, pregnancy complications, premature birth, and environmental exposures — have been associated with increased ASD likelihood in medical research. This case study explores extending the ML framework to assist in prenatal risk screening using demographic and maternal health data, with the objective of developing a risk-assessment system that could help identify high-risk pregnancies for closer monitoring.

Candidate prenatal risk features include maternal age, family history of autism, preterm-birth indicators, birth weight, prenatal infection history, medication exposure during pregnancy, smoking or alcohol exposure, and genetic markers where available. The same Python-based ML pipeline structure used for the primary AQ-10-based model — preprocessing, SMOTE-based imbalance handling, supervised classification (Random Forest, SVM, XGBoost, Logistic Regression), and standard evaluation metrics — was proposed for this extended use case. In a representative trial using this framework, a Random Forest model achieved a cross-validation accuracy of 92–94%, with Precision 0.91, Recall 0.93, and ROC-AUC 0.95, successfully classifying high-risk cases with minimal false negatives.

Python Deployment of the Autism Prediction System

While research models are typically trained within Jupyter notebooks, real-world use requires deployment into a reusable, accessible format. Python provides several suitable deployment pathways: Flask for a lightweight web API, Streamlit for an interactive dashboard, FastAPI for a high-performance API layer, and Pickle/Joblib for model persistence. This project's entire pipeline — from data ingestion through to the trained, serialized model — was implemented end-to-end in Python, making it directly portable to any of these deployment options.

Python-based deployment of this kind offers several advantages: scalability for integration into hospital information systems, straightforward integration with electronic health records, real-time prediction capability, cloud deployment options, and the possibility of a user-friendly graphical interface for non-technical end users such as caregivers or educators. As a practical extension of this dissertation, the author also implemented an interactive web-based ASD screening application (HTML/JavaScript front-end calling the Anthropic Claude API)

that replicates the trained model's prediction logic through a guided, multi-phase questionnaire interface, illustrating one concrete deployment path from the research pipeline to an end-user-

facing tool. Such as systems support early monitoring, assist pediatricians and gynecologists or other relevant specialists, enables preventive intervention strategies, and reduces dependency on subjective evaluation alone — while explicitly not replacing clinical diagnosis.

Worked Example: A Single Prediction Walk through

To illustrate how the deployed model consumes a new respondent's answers, consider a representative walkthrough. The system first collects the ten AQ-10 screening responses (each coded 0 = No, 1 = Yes), covering items such as noticing small sounds others miss, focusing on the whole picture versus fine detail, ease of multitasking, ease of returning to a task after interruption, reading between the lines in conversation, recognizing when a listener is bored, inferring a character's intentions while reading, collecting detailed information about categories of things, inferring feelings from facial expressions, and difficulty making new friends.

The system then collects background details: age, gender, ethnicity, jaundice history at birth, family history of autism, country of residence, prior use of a screening application, and a numeric prior screening score, together with an indicator of whether the respondent is filling out the form for themselves or on behalf of someone else. In one representative trial run, a 24-year-old male respondent of Asian ethnicity from India, with a positive family history of autism but no jaundice history, a prior screening score of 5, and no prior use of a screening app, produced a final classification of “does NOT have Autism (No ASD).” This walkthrough demonstrates how the saved label encoders and trained model combine at inference time to convert raw form responses into a single, immediately interpretable screening outcome.

Chapter 10: DISCUSSION

Analysis of Classifier Strengths

The Decision Tree classifier offered clear interpretability through its rule-based hierarchical structure and was effective at identifying dominant behavioral features, but remained sensitive to overfitting without depth control. Random Forest outperformed the single tree by aggregating many trees through bagging, significantly improving generalization and reducing variance, while still providing feature-importance rankings for interpretability. SVM performed strongly in the high-dimensional

feature space typical of multi-item questionnaire data when paired with an RBF kernel, though at greater computational cost on larger datasets. Naïve Bayes offered fast, low-cost classification despite its feature-independence assumption, performing reasonably well on this questionnaire-based data. XGBoost demonstrated the strongest predictive power overall via iterative gradient-boosting optimization, effectively capturing complex nonlinear relationships. Overall, the ensemble-based methods (Random Forest and XGBoost) showed the highest predictive reliability among all classifiers tested.

Impact of Principal Component Analysis (PCA)

PCA reduced feature redundancy among the correlated AQ-10 questionnaire responses, retained approximately 95% of total variance using a reduced number of components, decreased training time accordingly, and produced a slight improvement in generalization performance. By transforming correlated features into orthogonal principal components, PCA reduced noise and helped prevent overfitting; however, excessive dimensionality reduction introduced a minor interpretability cost, since principal components are linear combinations of the original questionnaire items rather than directly meaningful behavioral indicators. On balance, PCA enhanced computational efficiency while preserving classification performance.

Interpretation of Experimental Results

Across the evaluation metrics used — Accuracy, Precision, Recall, F1-score, Specificity, and ROC-AUC — high accuracy (above 90% in cross-validation) indicated strong overall correctness, high recall confirmed that autistic individuals were being correctly identified at an acceptable rate, high precision kept false-positive predictions low, and ROC-AUC values close to 1 confirmed strong class separability. Random Forest and XGBoost achieved the best balance between precision and recall, yielding the highest F1-scores, and confusion matrix analysis confirmed that false negatives were minimized — a critical property for any medical prediction system. SMOTE-based class balancing further improved minority-class detection performance throughout.

Strengths and Limitations of the Proposed Model

Strengths of the proposed framework include a robust preprocessing pipeline (missing-value handling, categorical encoding, and class balancing improving reliability), systematic comparison across multiple candidate classifiers, dimensionality reduction via PCA for

improved efficiency, strong predictive accuracy from the ensemble models, scalability afforded by the Scikit-learn/XGBoost implementation, and reproducibility guaranteed by serializing the trained model with Pickle.

Limitations included dependence on the quality and size of the underlying dataset, the introduction of possible response bias inherent to questionnaire-based features, the inability of LabelEncoder to handle previously unseen categorical values without retraining, the system's role strictly as a screening aid rather than a diagnostic replacement, and limited demographic diversity in the available training data potentially reducing generalizability to broader populations.

Ethical and Practical Considerations

Because the system is intended to operate in a healthcare-adjacent context, several ethical and practical considerations were kept in mind throughout development. First, the model's output is framed strictly as a screening signal rather than a diagnosis, and any deployment of the system should present results alongside a clear recommendation to seek professional evaluation rather than as a standalone verdict. Second, since the training data's demographic coverage is not exhaustive, predictions for individuals from underrepresented countries, ethnicities, or age groups should be interpreted with additional caution, and any production deployment would benefit from ongoing monitoring for performance drift across subgroups. Third, because questionnaire responses are self-reported (or reported by a caregiver), the system inherits whatever response bias is already present in the AQ-10 instrument itself; this is a limitation of the underlying screening tool rather than of the machine learning layer built on top of it, but it is an important caveat for anyone interpreting the model's output.

Comparison with Existing Systems

Compared to conventional ADOS/ADI-R-based diagnosis, the proposed model provides faster automated screening, reduces dependence on subjective human judgment, and offers objective, data-driven decision support. Compared to previously published ASD-prediction research, this work's incorporation of SMOTE improved class-imbalance handling, RandomizedSearchCV-based hyperparameter tuning enhanced performance beyond default configurations, and the ensemble techniques used here outperformed several single-classifier approaches reported in earlier studies—consistent with the broader literature trend favoring ensemble methods for ASD classification.

Overall Discussion

Taken together, the results demonstrate that Python-based machine learning techniques can provide reliable autism prediction using structured questionnaire and demographic data. Ensemble classifiers combined with PCA-based dimensionality reduction and SMOTE-based data balancing meaningfully enhance predictive stability and generalization. While not intended to replace clinical diagnosis, the resulting system can function as an intelligent screening tool that assists healthcare professionals in early detection and decision-making, and establishes a solid technical foundation for future integration with real-time web applications, cloud-based healthcare platforms, and explainable AI systems.

Chapter 11: Conclusion and Future Work

CONCLUSION

The primary objective of this project was to design and develop a machine learning-based predictive system for autism screening using Python. The system leverages questionnaire-based and demographic data to classify individuals into autistic and non-autistic categories, with the entire pipeline implemented using NumPy, Pandas, Scikit-learn, Matplotlib, Seaborn, Imbalanced-learn, and XGBoost to ensure robustness, scalability, and reproducibility.

The project followed a systematic machine learning workflow: data collection and understanding, data preprocessing, handling class imbalance via SMOTE, model training across multiple algorithms, hyperparameter optimization via Randomized Search CV, evaluation using accuracy, precision, recall, F1-score, confusion matrix, and cross-validation, and finally model export via Pickle for future reusability. Experimental results demonstrated that ensemble-based methods, particularly Random Forest, achieved superior performance compared to other models, with strong cross-validation accuracy and satisfactory generalization on the held-out test dataset. The use of SMOTE significantly improved recall values—especially important in medical diagnosis, where minimizing false negatives is paramount.

From a technical standpoint, this project demonstrates that Python-based machine learning frameworks can effectively assist in autism screening, with a modular implementation structure that supports easy adaptation,

retraining, and deployment. From a societal standpoint, the work contributes to the broader advancement of AI-driven healthcare solutions: although the system does not replace clinical diagnosis, it can function as a decision-support tool assisting healthcare professionals, educators, and parents in early screening. The project successfully fulfills its stated objectives by designing a complete end-to-end autism prediction pipeline, comparing multiple classification algorithms, optimizing performance through tuning and class balancing, and developing a reusable, deployable predictive system in Python.

More broadly, the dissertation shows that careful engineering of a relatively conventional machine learning pipeline — rigorous preprocessing, principled imbalance handling, systematic cross-validated comparison, and disciplined model serialization — can produce a screening tool whose performance is competitive with, and whose implementation is considerably simpler than, deep-learning-based alternatives reported elsewhere in the literature. That simplicity is itself a contribution: a Random Forest model that runs on a standard laptop and ships as two small .pkl files is far easier for a small clinic, school, or independent researcher to actually adopt than a neuroimaging-based deep network requiring specialized hardware and data-collection infrastructure.

Limitations of the Present Work

- **Dataset Dependency** — model performance depends heavily on the quality and diversity of the dataset used; limited demographic representation may reduce generalizability.
- **Questionnaire-Based Features Only** — the current implementation relies on behavioral and demographic features, excluding neuroimaging or genetic data.
- **Handling Unseen Categories** — Label Encoder may raise errors if unseen categorical values appear during prediction.
- **Not a Clinical Diagnostic Tool** — the system serves as a predictive aid and cannot substitute professional medical diagnosis.
- **Static Model** — the model requires retraining as new data distribution emerges over time.

Future Work

Future extensions of this work may include integrating deep learning architectures — Artificial Neural Networks, Convolutional Neural Networks for neuroimaging data, and Recurrent Neural Networks for sequential behavioral analysis — to capture more complex nonlinear relationships as larger datasets become available. Multi-modal

data integration, combining structured questionnaire data with speech, facial-expression, EEG, or neuroimaging inputs, represents another promising direction, alongside deployment as a fully interactive web or mobile application (using frameworks such as Flask or Streamlit) and the incorporation of explainable AI (XAI) techniques to improve transparency and clinical trust in model predictions.

In the nearer term, three concrete next steps follow directly from the limitations identified in Section 11.2. First, retraining the encoders to gracefully handle unseen categorical values (for example, by mapping any unrecognized country or ethnicity label to a designated “other” category at inference time) would remedy the most brittle failure mode of the current Label Encoder-based pipeline. Second, expanding the training dataset with additional, demographically diverse samples—even without changing the modeling approach—would likely narrow the gap between the 93% cross-validation accuracy and the 81.87% test accuracy observed in Chapter 8, since that gap is consistent with a moderate-sized dataset rather than with a flaw in the chosen algorithms. Third, wrapping the existing pickled model and encoders in a lightweight Flask or Streamlit service, rather than leaving them as notebook artifacts, would convert this dissertation’s pipeline into something a collaborating clinician or researcher could realistically try without needing to read the underlying Python code—anatural complement to the standalone web-based screening interface already built as part of this project (Chapter 9).

Final Remarks

In conclusion, this dissertation successfully demonstrates that a Python-based machine learning approach can effectively predict autism tendencies using questionnaire and demographic data. The structured workflow, balanced data handling, optimized model selection, and reproducible implementation collectively establish a robust foundation for future research and real-world deployment in autism screening and early-intervention support systems.

Appendix: Implementation Configuration and Publication

Implementation Configuration Summary

The entire system was developed in Python following a modular structure spanning data preprocessing, model training, evaluation, and model export. Table A.1 summarizes the development environment and system configuration used throughout the project.

Table A.1: Implementation Configuration Summary

Component	Specification
Programming Language	Python 3.x
IDE/Environment	Jupyter Notebook/VS Code
Operating System	Windows 10/11
Processor	Intel Core i5 or higher
RAM	8GB (minimum recommended)
Core Libraries	NumPy, Pandas, Matplotlib, Seaborn
ML Libraries	Scikit-learn, XGBoost
Imbalance Handling	Imbalanced-learn (SMOTE)
Model Persistence	Pickle (.pkl files)
Hyperparameter Tuning	Randomized Search CV
Evaluation Metrics	Accuracy, Precision, Recall, F1-Score
Cross-Validation	K-Fold Cross-Validation

Hyperparameter tuning targeted parameters such as the number of estimators, maximum tree depth, learning rate (for XGBoost), and minimum samples required to split a node. The best-performing configuration for each model was selected based on cross-validation accuracy, and the final artifacts—`best_model.pkl` and `encoders.pkl`—were saved for direct reuse in any downstream deployment.

Software Work flow Summary

1. Import required Python libraries.
2. Load the dataset from the source CSV file.
3. Perform exploratory data analysis (EDA).
4. Preprocess the data (cleaning and encoding).
5. Split the dataset into training and testing sets.
6. Apply SMOTE for class balancing on the training partition.
7. Train multiple candidate classification models.
8. Perform hyperparameter tuning via Randomized Search CV.
9. Evaluate each model using standard performance metrics.
10. Save the best-performing model and its encoders for future predictions.

Journal Publication

The research presented in this dissertation was published as a peer-reviewed journal article, summarized in Table A.2.

Table A.2: Journal Publication Details.

Field	Detail
PaperTitle	AutismPredictionUsingMachineLearning
Authors	AyanChakraborty,MahuyaSasmal
Journal/ ISSN	IJSRED,ISSN2581-7175
PublishedIssue	Volume9,Issue3
YearofPublication	2026
UniqueID	IJSRED-V9I3P128

The published paper reproduces the dissertation's score experimental comparison across Decision Tree, KNN, SVM, Random Forest, and XGBoost classifiers, and reaches the same central conclusion: ensemble learning methods, particularly Random Forest and XGBoost, provide the strongest, most generalizable performance for autism screening from structured questionnaire and demographic data, while remaining computationally lightweight enough for real-world deployment.

REFERENCES

1. Rajesh Prasad et al., "Autism Spectrum Disorder Prediction Using Machine Learning and Design Science," Int. J. Exp. Res. Rev., Vol. 39 (2024).
2. Yang Ding, Heng Zhang and Ting Qiu, "Deep learning approach to predict autism spectrum disorder: a systematic review and meta-analysis."
3. Rajesh Prasad, Farida Musa, Hadiza Muhammad Ahmad, Santosh Kumar Upadhyay and Birendra Kumar Sharma, "Autism Spectrum Disorder Prediction Using Machine Learning and Design Science."
4. Suman Raj, Sarfaraz Masood, "Analysis and Detection of Autism Spectrum Disorder Using Machine Learning Techniques," Procedia Computer Science 167 (2020), pp. 994–1004.
5. Shyam Sundar Rajagopalan, Yali Zhang, Ashraf Yahia, Kristiina Tammimies, "Machine Learning Prediction of Autism Spectrum Disorder From a Minimal Set of Medical and Background Information."
6. Shanthi Selvaraj, Poonkodi Palanisamy, Summia Parveen, and Monisha, "Autism Spectrum Disorder Prediction Using Machine Learning Algorithms," ICCVBIC 2019.
7. Muhammad Shoaib Farooq, Rabia Tehseen, Maidah Sabir and Zabihullah Atal, "Detection of autism

- spectrumdisorder(ASD)inchildrenandadultsusingmachinelearning,”DOI:10.1038/s41598-023-35910-1.
8. KoushikChowdhury,MirAhmadIraj,“PredictingAutismSpectrumDisorderUsingMachineLearning Classifiers,” DOI: 10.1109/RTEICT49044.2020.9315717.
 9. AmeerS.Jaradat,MohammadWedyan,SajaAlomari,MalekMahmoudBarhoush,“UsingMachine Learning to Diagnose Autism Based on Eye Tracking Technology,” DOI: 10.3390/diagnostics15010066.
 10. KanimozhiA,DhanasriA,“AutismSpectrumDisorderPredictionbyFacialRecognitionUsing Deep Learning,” ISSN: 2320-2882.
 11. Khandaker Mohammad Mohi Uddin, Hasibur Rahman, Mahadi Hasan, Fatema Akter, Suman Chandra Das,“Amachinelearningapproachtopredictautismspectrumdisorder(ASD)forbothchildrenandadults using feature optimization,” ISSN 2220-8879.
 12. M.ShuaibQureshi,MdBilalQureshi,J.Asghar,F.Alam,“PredictionandAnalysisofAutism Spectrum Disorder Using Machine Learning Techniques,” DOI: 10.1155/2023/4853800.
 13. S.Yuvaraj,S.Sugavanaesh,C.TharunKumar,G.Saran,R.Subha,“IdentificationofAutismSpectrum Disorder (ASD) in Adults through Various Machine Learning Algorithms,” DOI: 10.1109/IDCIoT59759.2024.10467893.
 14. SiddhantDudhane,SrijanMohan,PrajaktaSahu,UshasukhanyaS,“AutismSpectrumDisorder (ASD)
 15. ScreeningandDetectionusingMachineLearningTechniques,”DOI:10.1109/ICNWC60771.2024.10537301.
 16. D.AarthiandS.Kannimuthu,“AComprehensiveAnalysisofAutismSpectrumDisorderUsing Machine Learning Algorithms: Survey,” DOI: 10.1007/978-981-99-7216-6_20.
 17. LongjieJin,HualeiCui,PeiyuanZhangandChunquanCai,“Earlydiagnosticvalueofhomevideobased machine learning in autism spectrum disorder: a meta-analysis,” DOI: 10.1007/s00431-024-05837-4.