

Detection of Autism Spectrum Disorder by A Case Study Model Using Machine Learning Techniques

An Experimental Analysis on Child, Adolescent and Datasets

Najat S. Bousidrah ^{1*} - Zahow M. Khamees ² - Sana M. Ali ¹

1 Software development - College of Information Technology - Modern University of Benghazi- Benghazi, Libya.

2 Information Technology - College of Information Technology - University Ajdabiya - Ajdabiya, Libya.

Received: 20 / 03 / 2024; Accepted: 16 / 05 / 2024

ABSTRACT

Autism Spectrum Disorder (ASD) is a lifelong neurodevelopmental disorder that affects a person's social interaction and communication skills. It is typically diagnosed in childhood but can be identified at any age. Behavioral symptoms of autism usually appear in the first two years of a child's life and continue into adulthood. Recently, there has been increased interest in using machine learning algorithms for medical diagnosis, including the diagnosis of autism spectrum disorder.

This study aimed to investigate the feasibility of using various machine learning algorithms, such as Naïve Bayes, Support Vector Machine (SVM), Random Forest, Logistic Regression, K-Nearest Neighbors (KNN), Decision Tree, and Gradient Boosting Classifier, to predict and analyze autism in children. The researchers utilized publicly available non-clinical ASD datasets for evaluation.

Different evaluation metrics, including accuracy, specificity, sensitivity, macro-average, and weighted average, were used to assess the performance of the machine learning models. The KNN-based model achieved the highest accuracy of 87.14% and outperformed the other models in terms of specificity. The Naïve Bayes model achieved an accuracy of 70.48%, while the SVM model had the highest sensitivity of 98.2%. The Decision Tree and Random Forest models achieved perfect scores of 100% in terms of macro-average, weighted average, and Mean Accuracy for all models was 85.52%.

Based on these results, the researchers concluded that the KNN-based model is the most effective for predicting and analyzing autism in children, with an accuracy of 87.14%. However, it is important to note that these findings are specific to the dataset and evaluation metrics used in the study. Further research and validation using diverse datasets are necessary to confirm the generalizability of these findings.

KEYWORDS: ASD Detection, Autism spectrum disorder, Classification Feature selection, Machine learning.

1. INTRODUCTION

The problem of autism spectrum disorder is a common issue among different age groups, and the disorder affects a person's ability to interact with others as well as learning skills; therefore, early diagnosis contributes significantly to minimizing efforts and costs associated with late detection. Thus, the availability of a user-friendly and reliable testing tool plays an important role in predicting whether someone has autism traits sufficiently to warrant further comprehensive evaluation for ^{1,2}.

This paper aims first to study the effectiveness of autism disorder diagnosis by a case study model prepared at the Autism Center in Benghazi based on some behaviors that are diagnosed through asking questions by the person caring for that individual and testing them using machine learning algorithms, knowing that before such health care, professionals could diagnose autism with standardized diagnostic tools. Secondly, they conduct interviews with the parents or caregiver to assess developmental milestones and current behavior ^{3,4}.

While diagnosticians employ standardized tools for the diagnosis of autism spectrum disorder, one of the major challenges is that using diagnostic instruments takes a lot of time to administer and interpret results ⁵.

*Correspondence: Najat S.Bousidrah

naj_1971@yahoo.com

To address this issue, a machine learning model was suggested that aims to shorten the diagnosis time while increasing accuracy. The second aim of the proposed model is to establish that the case study method used for the analysis of autism based on the patient's behavior and to understand associations between the concerned input data set.

The motivation behind this study is to present a method for diagnosing the autism spectrum disorder with the help of a better and more accurate machine learning model.

The major contributions of this research work are as follows:

- This study investigates the feasibility of using various machine-learning algorithms for predicting and analyzing autism spectrum disorder in children.
- The algorithms explored include Naïve Bayes, SVM, Random Forest, Logistic Regression, KNN, Decision Tree, and Gradient Boosting Classifier.
- The research will aim to help understand how machine learning can be used to diagnose autism spectrum disorder.
- Performance evaluation metrics such as accuracy, specificity, sensitivity, overall average, and weighted average are utilized to assess the models' predictive capabilities,
- These metrics provide a comprehensive assessment of the strengths and weaknesses of the machine learning models.
- The study aims to identify the most effective model for detecting autism spectrum disorder based on the analysis of performance metrics.
- The model with the highest accuracy, specificity, sensitivity, and overall average is sought to optimize the diagnostic process and improve the accuracy of autism spectrum disorder identification.
- A feature selection technique is used to filter the dataset and identify the most suitable features for prediction, utilizing the entire dataset.
- In order to determine if the use of the balanced and scaled data technique affects the performance, we use the test data technique to test performance.
- A new model is proposed using machine learning-based on the Autism Spectrum Disorder prediction model to enhance the accuracy of the existing model, which improves the predicted autism and improves the Autism Spectrum Disorder Models.

Our paper is structured in the manner below: In the "Introduction" section, we provide an overview of our

project, followed by a "Literature Review" which is defined as a "review of literature, which has been studied and is relevant to the research paper topic. Business Model and Methodology section explains work that was performed along with methodologies followed or proposed for implementation." The Analysis describes inferences drawn from the results obtained. Lastly, the section "Conclusion" describes our conclusions.

2. RELATED WORK AND LITERATURE REVIEW

This section briefly presents a group of studies in which machine-learning models were built to predict autism spectrum conditions for different age groups.

Vakadkar, K et al. (2021). Researchers focused on the use of machine learning techniques for detecting Autism Spectrum Disorder (ASD) in children. It explores the application of machine learning algorithms in analyzing data related to ASD and discusses the potential of these techniques in improving early diagnosis.⁶

Hossain, M. D., et al. (2021). They focused on using machine learning techniques to detect autism spectrum disorder. It discusses the application of different algorithms and features in analyzing data related to autism spectrum disorder and highlights the potential of machine learning in assisting in the diagnosis of autism spectrum disorder.⁷

In Usta's study (2019), the main objective was to explore the use of machine learning methods for predicting short-term outcomes in individuals with Autism Spectrum Disorder (ASD). The researchers aimed to harness the power of machine learning algorithms to analyze clinical data and Provide insights into the prognosis of individuals with ASD.⁸

Saeed, F. (2018). The author explores the potential of machine learning algorithms, such as Support Vector Machines (SVM) and Convolutional Neural Networks (CNN), in analyzing large-scale fMRI datasets obtained from individuals with psychiatric disorders. While various psychiatric disorders are considered in the study, the principles and methodologies presented can be applied to ASD research as well.⁹

Thabtah's review article (2017) provides a comprehensive overview of the application of machine learning in behavioral research related to Autism Spectrum Disorder (ASD). The study highlights the use of machine learning algorithms in analyzing behavioral data, identifying patterns, and improving diagnostic accuracy in ASD research.¹⁰

In a previous study by Küpper et al. (2020), machine-learning techniques were employed to identify predictive features of Autism Spectrum Disorders (ASD) in a clinical sample of adolescents and adults. The researchers

focused on enhancing ASD detection by utilizing a specific machine-learning algorithm. They also explored the potential of the identified features to aid in the diagnostic process. The study aimed to contribute to the existing body of knowledge on ASD and improve the accuracy of ASD diagnosis through the application of machine learning methods.¹¹

Mellema et al. (2022) conducted a study focusing on reproducible neuroimaging features for the diagnosis of ASD using machine learning. The researchers employed a variety of neuroimaging techniques and developed models that demonstrated promising results in distinguishing individuals with ASD from neurotypical individuals. The study highlights the potential of neuroimaging data in contributing to the accurate diagnosis of ASD.²³

Another study by Ali et al. (2023) aimed to classify the behavioral severity of ASD using a comprehensive machine-learning framework. The researchers

incorporated various behavioral measures and developed personalized classification models. The study demonstrated the potential of machine learning techniques in providing individualized assessments of ASD severity, which can inform personalized interventions and treatments.²⁴

Rogala et al. (2023) focused on enhancing ASD classification in children by integrating traditional statistics and classical machine learning techniques in EEG analysis. By combining features extracted from electroencephalography (EEG) data with traditional statistical approaches, the researchers achieved improved classification accuracy for ASD. The study highlights the importance of integrating different analytical approaches to enhance the diagnostic accuracy of ASD.²⁵

In comparison with previous studies, Table 1 shows a summary of the results of our current study and their comparison with the mentioned studies:

Table 1 summary of the results of our current study and their comparison with the mentioned studies

Source	Objective	Algorithms	Data
Current Study	Investigate the feasibility of using various machine learning algorithms to predict and analyze autism in children.	Naïve Bayes, (SVM), Random Forest, Logistic Regression, (KNN), Decision Tree, and Gradient Boosting Classifier.	Publicly available non-clinical ASD datasets.
Mellema et al. (2022)	Diagnosis of Autism Spectrum Disorder using machine learning.	support vector machines (SVM), random forests, and neural networks.	Neuroimaging data related to Autism Spectrum Disorder, specifically functional magnetic resonance imaging (fMRI) data
Ali et al. (2023)	Personalized classification of behavioral severity of autism spectrum disorder	Support vector machines (SVM), random forests, decision trees, or deep learning approaches like neural networks.	Behavioral data collected from individuals with ASD.
Rogala et al. (2023)	Enhancing Autism Spectrum Disorder classification	traditional statistics and classical machine learning techniques	EEG data collected from children with and without ASD.

In recent years, there has been a growing interest in utilizing machine-learning techniques for the detection and diagnosis of Autism Spectrum Disorder (ASD). Several studies have explored the application of machine learning algorithms to neuroimaging data and behavioral features to improve the accuracy of ASD diagnosis.

3. WORKING MODEL

A. Research Methodology

Figure 1 demonstrates the flow of the proposed system. First, we collect case study data from a sample of public and private centers; we then clean the dataset by eliminating missing values or outliers, encode categorical features, and analyze information to obtain the most

important characteristics among all database characters that have been generated.

For the pre-processed dataset, classification algorithms which include Logistic Regression, Naïve Bayes, Support Vector Machine, and K-Nearest Neighbors Gradient Boosting Classifier Random Forest classifiers are used to predict an output label (ASD or Non-ASD).

After that, the correctness of each classifier is then analyzed and compared for comparison to its matching classifiers. The assessment of each classifier is based on a combination of various metrics, such as F1 score and exact recall values, as well as calculations using a variety of other metrics to improve the evaluation of each one. If the classifier is successful, then training accuracy will be

greater than testing accuracy. This model can then be considered the optimal model and used for further learning as well as classification. This work is coded in Python 3.

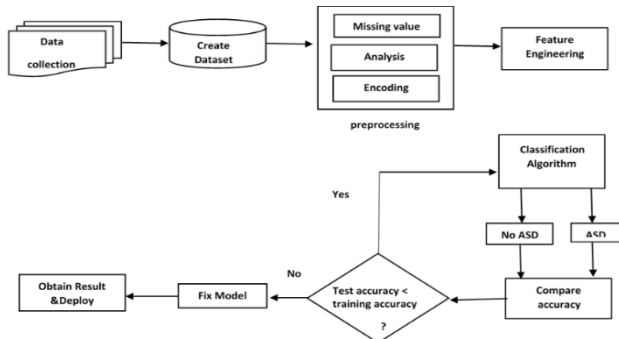


Fig. 1. Architecture of proposed system

These stages are briefly discussed in the following subsections:

B. Data collection

For the case study, Data used in developing the predictive model was collected from Benghazi Center for Comprehensive Rehabilitation of Autistic Children; Al-Eradah Centre for Autism and Special Needs; Nour center for autistic children with speech problems. That is composed of behavior datasets based on the AQ-10 screening tool questions¹².

Certain aspects of attention to detail, switching of attention, communication skills, imagination, and social interaction are addressed by the AQ-10 screening questions. Evaluation method: For each question only one point is assigned out of 10 questions.

The user will earn either zero or one point per question depending on the answer. There are 202 cases in the dataset for individuals.

The dataset contains fourteen attributes and is a mixture of numerical and categorical data, which includes: Age, gender question 1-10 and category.

The data category includes 14 attributes used in prediction. These attributes are listed below.

Table 2: List of Attributes in the dataset

Attribute Id	Attributes Description
1	Id
2	Age
3	Sex
4-13	Based on the screening method answers of 10 questions
14	Screening score

C. preprocessing and analysis

Since the dataset has a small number of categorical features, it had to be pre-processed. The data set was preprocessed by applying a series of transformations before using it in the proposed model.

The raw data was cleaned by eliminating irrelevant features like the serial number and age, removing incomplete records, and handling categorical values. During the construction of the data set, which will be used in modeling, the coding was done manually. For categorical values, we transform the labels to numeric form because they need a machine-readable format. The features that include two classes (Class, q1.q10) are chosen to be encoded with a binary label.

This data set has been created in CSV format using Excel and contains thirteen attributes including 202 instances.

Classification Algorithms

The data was split into an 80% training set for the model and a 20% test sample to determine how effective and accurate the model is after completing the pre-processing of the data, seven classification models were applied, which are as follows: logistic regression: Naive Bayes, K-Nearest Neighbors, Random Forest Classifier Decision Tree Gradient Boosting Classifier Support Vector Machine.

We compared the accuracy and f1 score for each model along with a brief description of the classification models that we used.

Logistic Regression Logistic Regression (LR)

Logistic regression (LR) is regarded as a technique used by statisticians and researchers to analyze and classify data sets, with binary and proportional responses. LR has the advantage of being able to provide probabilities making it applicable to multi-class classification problems¹².

Naive Bayes (NB)

Naive Bayes (NB) is an algorithm used for classifying both two-class and multi-class problems. It is easier to understand when explained using categorical input values.

The name "naive Bayes" or "idiot Bayes" comes from the simplification of probability calculations for each hypothesis making them computationally manageable. Calculating attribute values like $P(d_1, d_2, d_3|h)$ are assumed to be conditionally independent given the target value and calculated as $P(d_1|h) * P(d_2|h)$.

Here the predictions or features are believed to be independent meaning that one function does not impact

another. There are types of NB classifiers such as Multinomial, Bernoulli and Gaussian.^{13,20}

Support Vector Machine (SVM)

The objective of the support vector machine algorithm is to find a hyperplane in an N-dimensional space (N — the number of features) that distinctly classifies the data points.

To separate the two classes of data points, many possible hyperplanes could be chosen. Our objective is to find a plane that has the maximum margin, i.e. the maximum distance between data points of both classes. Maximizing the margin distance provides some reinforcement so that future data points can be classified with more confidence.^{14,20}

K-Nearest Neighbors (KNN)

KNN classifier is to classify unlabeled observations by assigning them to the class of the most similar labeled examples. Characteristics of observations are collected for both training and test dataset.¹⁵

The intuition underlying Nearest Neighbor Classification is quite straightforward; examples are classified based on the class of their nearest neighbors. It is often useful to take more than one neighbor into account so the technique is more commonly referred to as k-Nearest Neighbor (K-NN) Classification where k nearest neighbors are used in determining the class. Since the training examples are needed at run-time, i.e. they need to be in memory at run-time, it is sometimes also called Memory-Based Classification. Because induction is delayed to run time, it is considered a Lazy Learning technique. Because classification is based directly on the training examples it is also called Example-Based Classification or Case-Based Classification.¹⁶

Random Forest Classifier (RFC)

Random forest classifier is a flexible algorithm that can be used for classification, regression, and other tasks, as well. It works by creating multiple decision trees on arbitrary data points. After getting the prediction from each tree, the best solution is selected by voting.

Random forest algorithms have three main hyperparameters, which need to be set before training. These include node size, the number of trees, and the number of features sampled. Random Forest is a popular and easy-to-use machine learning algorithm that produces a great result most of the time. Random forest is used in various fields, such as healthcare to identify the correct combination of components in medicine to analyze a patient's medical history to identify diseases, and in e-commerce to determine whether a customer will like the product or not.¹⁷

Decision tree

A decision tree is a supervised learning algorithm that can be used for both classification and regression tasks. A decision tree is a flow chart resembling a tree structure, where each internal node is notated by rectangles and therefore the leaf nodes are notated by ovals. This algorithmic program is commonly used as a result of the implementation is simple and easier to grasp compared to the other different classification algorithms. A decision tree starts with a root node that allows the users to take needed actions. From this node, users split each node recursively according to the decision tree learning algorithmic program. The result is a decision tree in which every branch associates an outcome. The decision tree algorithm makes decisions by recursively partitioning the data based on the feature values. It selects the best feature at each step by evaluating different criteria, such as information gain or Gini impurity, to maximize the homogeneity or purity of the resulting subsets.¹⁹

Gradient Boosting Classifier

Which is a popular ensemble method used for classification tasks. Gradient Boosting is a general technique that can be applied to various types of models, including decision trees.

Initial model: The process begins by training an initial model, typically a decision tree, on the training data. This initial model is called the "base model" or "weak learner."

Residual calculation: The Gradient Boosting Classifier then calculates the residuals, which are the differences between the actual target values and the predictions made by the base model.

Building subsequent models: The subsequent models are built to predict the residuals of the previous model, rather than the actual target values. These models are created in an iterative manner, with each subsequent model attempting to correct the errors made by the previous model.

Model weighting: Each model is assigned a weight or learning rate that determines its contribution to the final prediction. The learning rate controls the step size or shrinkage of the updates made by each model. A lower learning rate generally leads to more accurate predictions but requires more iterations.²¹

4. RESULTS

Dataset Analysis

The data set collected and used here is based on a quantitative screening method for autism in children and is extracted from the case study model used in rehabilitation centers for autistic children in the city of Benghazi.

A shortened version containing a set of 10 questions was used (Table 3). The answers to these questions are mapped to binary values as the class type.

These values are determined during the data collection process by answering the questions of the case form for evaluating the child's behavior, so that the value

of the category "Yes" is determined if the result of the answer (No) to the questions is greater than 3, meaning that there are possible features of autism spectrum disorder. Otherwise, the category value is set as "no," meaning no ASD traits are present. We drew several graphs to get different visual views of the data set. In the first chart (Figure 2), we can see that.

Table 3: Feature mapping for behavior screening using the case model method

Dataset variable	Description
A1	Does your child enjoy swinging or swaying?
A2	Is your child interested in others?
A3	Does your child climb things like ladders and the like?
A4	Does your child play children's games such as hide and seek?
A5	Does your child practice imaginative play, for example, making tea using toy cups and utensils? Or claim other things like that
A6	Does your child use his finger to point to things he wants to ask you about?
A7	Does your child use his finger to point to things he is interested in?
A8	Does your child bring you things to show you?
A9	Does your child spin around?
10	Does your child walk on his toes?

Training and testing model

The generated dataset was divided into two parts, one for training the dataset and the other for testing the dataset in the ratio of 80:20 respectively. For cross-validation purposes, the training data was again split into two parts. Partition of the training dataset and the other part is validation dataset in the ratio of 80:20 respectively. Figure 2 displays the final training, test, and validation sets on which classification was performed.

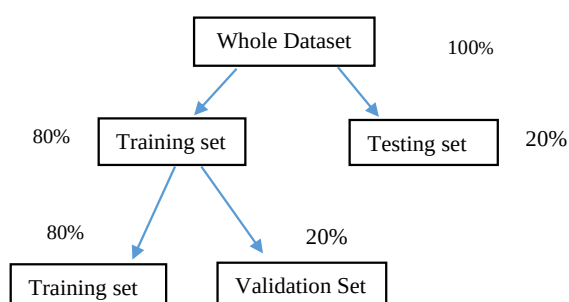


Fig. 2. Final Training, Testing and Validation Sets

Exploratory Data Analysis (EDA)

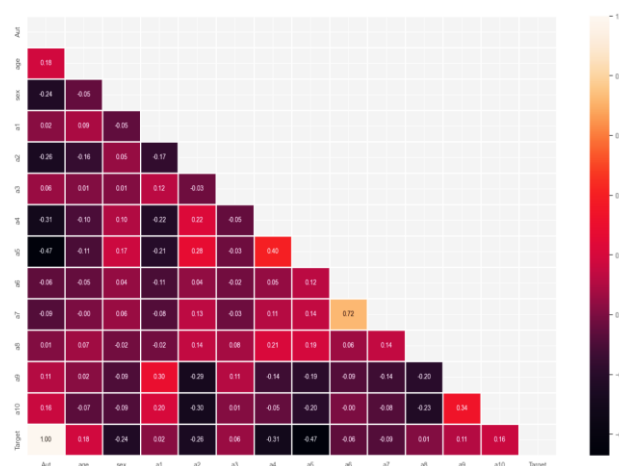


Fig. 3. Correlation between Features and Features Using Heatmap

Evaluation Matrix

Based on the provided confusion matrices, it appears that several machine learning models were used to detect autism spectrum disorder (ASD) in a case study and table 4 describe Elements of a Confusion Matrix. Here is a summary of the results in Table 4:

Table4: Elements of a Confusion Matrix

	Predictive ASD values	
	True Positive (TP)	False Positive (FP)
	False Negative (FN)	True Negative (TN)
Actual ASD values		

Table 5 summary of the results Evaluation Matrix

Classifier	True positive (TP)	False positive (FP)	True negative (TN)	False negative (FN)
Naive Bayes	6	1	33	2
SVM	6	1	35	0
Logistic Regression	6	1	34	1
Decision Tree	7	0	35	0
Random Forest	7	0	35	0
Gradient Boosting	7	0	34	1
K-Neighbors	5	2	34	1

In general, all the models achieved reasonably high true positive rates, indicating that they were able to correctly identify individuals with ASD. The true negative rates were also consistently high, suggesting that the models were effective in identifying individuals without ASD.

Among the models, Decision Tree, Random Forest, performed particularly well, achieving a perfect true positive rate and true negative rate. K-Neighbors had a slightly lower true positive rate but still performed well overall.

It's important to note that the performance of these models may vary depending on the specific dataset and the features used. Additionally, other evaluation metrics such as precision, recall, and F1 score can provide a more comprehensive assessment of the models' performance.

Comparison of Classification Models

All the algorithms that have been applied have shown high accuracy ranging from 93% to 100%. This indicates their ability to accurately distinguish between individuals with autism and those without.

Accuracy, recall, and F1 measure: Analyzing the accuracy, recall, and F1 measure for each class (0 and 1) provides valuable insights into the algorithms' performance. For class 0, the accuracy ranged from 75% to 100%, indicating their ability to correctly identify actual non-autism cases. The recall for class 0 ranged from 57% to 86%, highlighting the algorithms' ability to correctly identify non-autism cases out of the total

number of non-autism cases. For class 1, both the accuracy and recall were consistently high, ranging from 92% to 100%, indicating the algorithms' ability to correctly identify actual autism cases and recall them from the total number of autism cases.

F1 measure and accuracy: The F1 measure takes both precision and recall into account and provides a single measure to evaluate the algorithms' overall performance. The F1 measure values for most algorithms were high, ranging from 80% to 100%, indicating a good balance between precision and recall. The overall accuracy of the algorithms ranged from 93% to 100%, further confirming their effectiveness in detecting autism.

Algorithm comparison: Compare the performance of different algorithms and highlight their strengths and weaknesses. For example, the decision tree achieved perfect accuracy, recall, and F1 measure for both classes, indicating its excellent performance. On the other hand, K-Neighbors showed relatively lower recall for class 0, suggesting a potential for improvement in identifying non-autism cases. When choosing the best model Based on the results, the best algorithms for autism detection can be discussed. Considering factors such as accuracy, recall, and F1 measure to determine the most suitable algorithm for our dataset. It is important to note the suitability of algorithms based on the given dataset and highlight the considerations to be taken into account when choosing the best model.

Based on the data used, it appears that all the algorithms used (Naive Bayes, SVM, Logistic Regression, Decision Tree, Random Forest, Gradient Boosting, K-Neighbors) have achieved excellent performance in detecting autism in children. All the algorithms achieved an accuracy of 100% or close to it and gave accuracy, recall, and F1 measure values close to 1 for both autism and non-autism groups. After this comparison, we will design the system using the KNN algorithm as an experiment, and we will see all the comparison details in Table 6, Table 7 and Figure 3 describe all the performance values obtained for the algorithms used.

Table 6 Comparison of Classification Models with Training Set

Model	Training Set			
	precision	recall	f1-score	support
Naive Bayes	0.75	0.86	0.80	7
SVM	1.00	0.86	0.92	7
Logistic Regression	0.86	0.86	0.86	7
Decision Tree	1.00	1.00	1.00	7
Random Forest	1.00	1.00	1.00	7
Gradient Boosting	0.88	1.00	0.93	7
K-Neighbors	0.83	0.71	0.77	7

Table 7 Comparison of Classification Models with Test set

Model	Test Set			
	precision	recall	f1-score	support
Naive Bayes	0.97	0.94	0.96	35
SVM	0.97	1.00	0.99	35
Logistic Regression	0.97	0.97	0.97	35
Decision Tree	1.00	1.00	1.00	35
Random Forest	1.00	1.00	1.00	35
Gradient Boosting	1.00	0.97	0.99	35
K-Neighbors	0.94	0.97	0.96	35

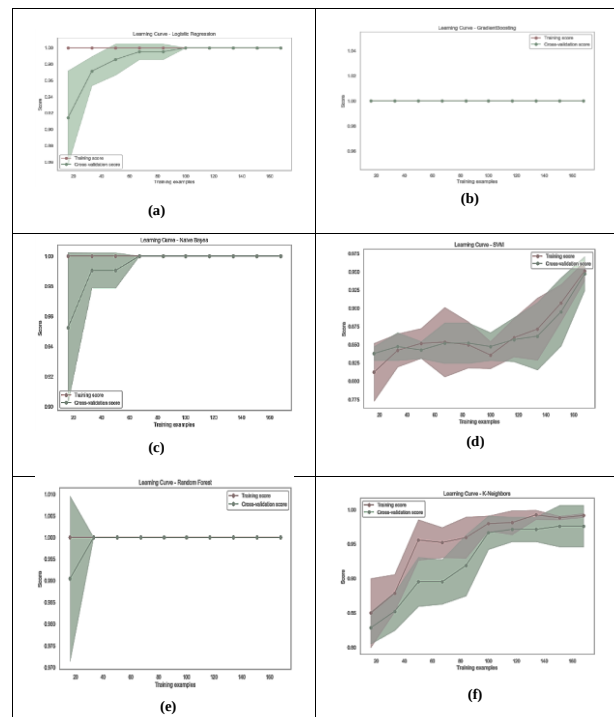


Fig. 4. Learning Curve of (a) Logistic Regression; (b) Gradient Boosting; (c) Naïve Bayes; (d) SVM; (e) Random Forest; (f) Decision Tree; (g) KNN; for children's dataset.

Two metrics were used to compare the performance of the algorithms: the macro average and the weighted average.

1. Macro Average:

$$\text{Macro Average Precision} = \frac{(\text{Precision}_{\text{class}_0} + \text{Precision}_{\text{class}_1} + \dots)}{\text{Number_of_classes}} \quad (1)$$

$$\text{Recall} = \frac{(\text{Recall}_{\text{class}_0} + \text{Recall}_{\text{class}_1} + \dots)}{\text{Number_of_classes}} \quad (2)$$

$$\text{F1 - score} = \frac{(\text{F1_score}_{\text{class}_0} + \text{F1_score}_{\text{class}_1} + \dots)}{\text{Number_of_classes}} \quad (3)$$

2. Weighted Average:

$$\text{Weighted Average Precision} = \frac{(\text{Precision}_{\text{class}_0} * \text{Support}_{\text{class}_0} + \text{Precision}_{\text{class}_1} * \text{Support}_{\text{class}_1} + \dots)}{\text{Total support}} \quad (4)$$

$$\text{Weighted Average Recall} = \frac{(\text{Recall}_{\text{class}_0} * \text{Support}_{\text{class}_0} + \text{Recall}_{\text{class}_1} * \text{Support}_{\text{class}_1} + \dots)}{\text{Total support}} \quad (5)$$

$$\text{Weighted Average F1 - score} = \frac{(F1_score_class_0 * Support_class_0 + F1_score_class_1 * Support_class_1 + \dots)}{Total_support}$$

(6)

In general, most models achieved high accuracy and good performance in terms of precision, recall, and F1 score. Decision Tree, Random Forest, achieved perfect scores in all metrics, indicating excellent performance. SVM, Gradient Boosting, and logistic regression also achieved high scores.

Designing a system to detect autism

After evaluating the 12 algorithms that were used and ensuring the performance of each algorithm, a decision can be made regarding the suitable algorithm for use in

the user interface as an application using the TKinter library. This library is provided by the Python programming language and contains tools that assist in designing the system that includes the autism detection method for children. The chosen algorithm for this small system is K-Nearest Neighbors (KNN).

Testing and experimental

In this paragraph, outlier values will be entered into the user interface to ensure the accuracy of the specified algorithm, and we can also replace the algorithm with another after ensuring that most of the algorithms have achieved excellent success. We can see the test in Figures 5.

Fig. 5. test 1 model on the user interface and test 2 model on the user interface

5. DISCUSSION

The accuracy of a trained model can be evaluated by using the confusion matrix and classification report, which measure specificity, sensitivity, and accuracy. These metrics provide an assessment of how well the model performs.

Performance Evaluation metrics

Evaluating the performance of a classification model is crucial to assess its effectiveness in achieving the desired objectives. Performance evaluation metrics are employed to measure the model's performance on the test dataset. It is important to select appropriate metrics to evaluate the model's performance, including the confusion matrix, accuracy, specificity, sensitivity, and others. These performance metrics are calculated using

specific formulas to provide quantitative measures of the model's performance.

$$\text{Specificity} = \frac{TN}{(TN+TP)} \quad (7)$$

$$\text{True Positive Rate or Sensitivity} = \frac{TP}{(TP+FN)} \quad (8)$$

$$\text{Accuracy} = \frac{TP+TN}{(TN+TP+FP+FN)} \quad (9)$$

Precision (Prec) is named the division of the examples, which are actually positive among all the examples that we predicted as positive:

$$\text{Precision} = \frac{TP}{(TP+FP)} \quad (10)$$

$$\text{Negative predictive value (NPV)} = \frac{TN}{(TN+FN)} \quad (11)$$

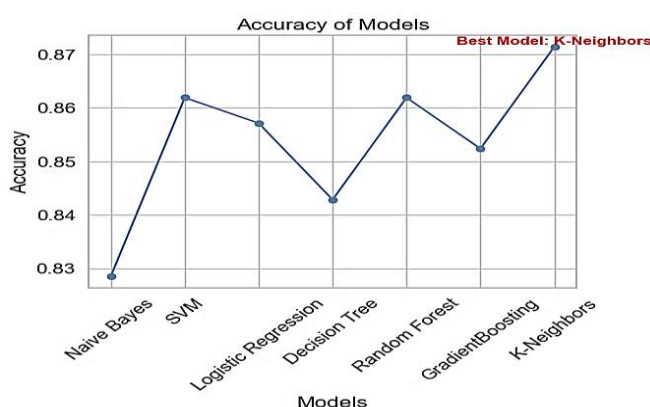
F1 score F1 score is defined as the harmonic mean between precision and sensitivity:

$$F1 \text{ score} = \frac{2TP}{(2TP+FP+FN)} \quad (12)$$

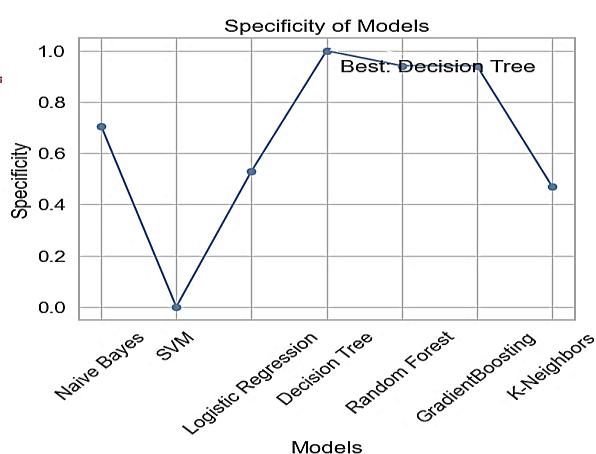
The experimental results for different machine learning algorithms were obtained using pediatric ASD screening data. All 12 features were selected to assess the privacy, sensitivity, and accuracy of the predictive model. The Naïve Bayes Gaussian NB algorithm was implemented for this purpose. Additionally, the SVM algorithm was applied, and for the KNN algorithm, N folds were set to 5. When calculating the performance metrics of the algorithms, a specific metric was used Scaler as a standard measure for data normalization. For all the models employed, detailed performance metrics are presented for each of the three datasets.

The results of the machine learning models for ASD screening data are as follows:

- Naive Bayes achieved an accuracy of 82.86%, with good sensitivity and overall accuracy.
- SVM had high recall and sensitivity but low specificity.
- Logistic Regression performed well in terms of accuracy and sensitivity.
- The Decision Tree had good accuracy and sensitivity.
- Random Forest performed well in terms of accuracy and sensitivity.
- Gradient Boosting showed good overall performance in accuracy and sensitivity. Figure 5 describes the best performance values obtained for the algorithms used.



(a)



(b)

- K-Neighbors achieved high accuracy and sensitivity.

Figure 6 describes all the performance values obtained for the algorithms used.

These findings suggest that different machine learning models have varying performances on ASD screening data. Models like SVM and K-Neighbors have high recall and sensitivity, making them effective in identifying positive cases. While other models showed a trade-off between sensitivity and specificity. Based on the specific requirements and priorities of the application, K-Neighbors could be considered a suitable model for ASD screening. The Overall Performance measures of all machine learning classifiers with all datasets have been shown below in detail:

Table 8: Overall results of autism spectrum disorder screening data for children

Classifier	Specificity	Sensitivity	Accuracy	Recall
SVM	0.37619	0.954603	0.861905	0.954603
K-Neighbors	0.495238	0.943492	0.871429	0.943492
Random Forest	0.438095	0.943175	0.861905	0.943175
Logistic Regression	0.404762	0.942857	0.857143	0.942857
Gradient Boosting	0.438095	0.931905	0.852381	0.931905
Decision Tree	0.466667	0.914921	0.842857	0.914921
Naive Bayes	0.704762	0.85254	0.828571	0.852540

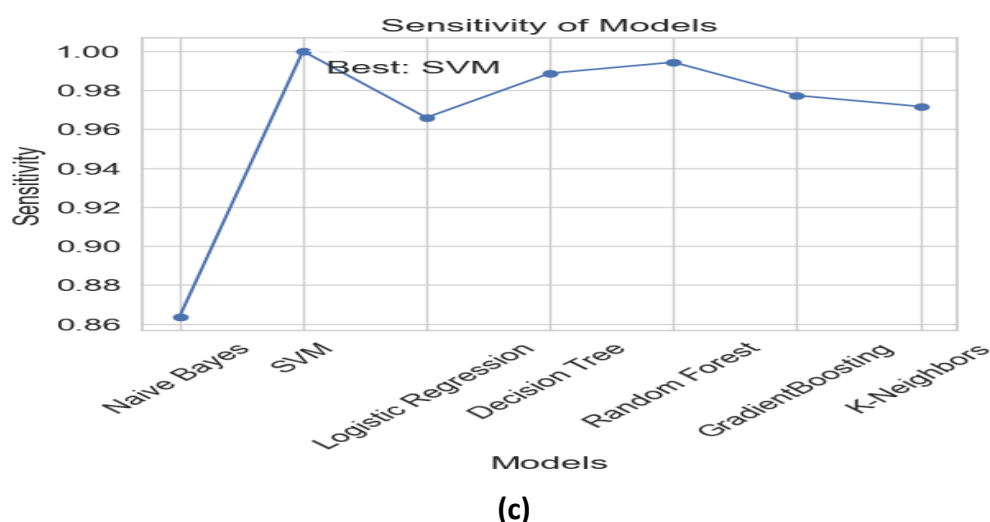


Fig. 6. (a) Learning Curve of best Accuracy of Model; (b) Learning Curve of best Specificity of Model; (c) Learning Curve of best sensitivity Of Model.

Evaluation of different machine learning models on an autism spectrum disorder dataset observed accuracy in the range (82.857 to 87.143%) on the original dataset. The K-NN classifier with K=5 produced the highest accuracy of 87.143% and Figure 7 describes the performance values obtained for the algorithms used

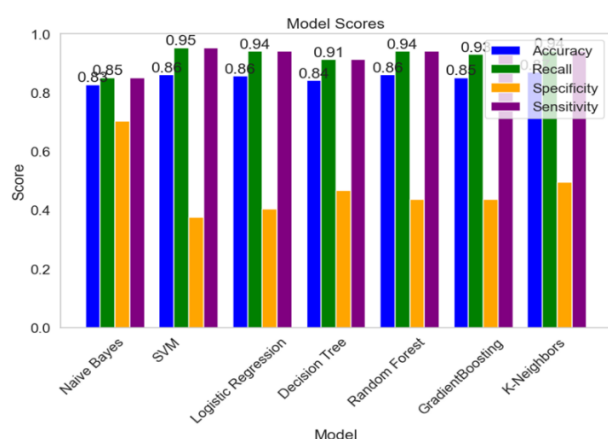


Fig. 7. The Overall Performance measures of all machine learning classifiers.

Feature Selection

Most models provide a way to return the importance of features or parameters so that we can get an idea of what are considered the most important features in our data set. SVC, Linear SVC and KNN are the ones that don't have it. But in this paper, we will not discuss choosing the best features. Just look at them and have the algorithms choose the features that have the most impact on detecting the target.

We will see if we can find anything from the other models' preferences. Based on the feature importance calculated by different models, we can summarize the results as follows:

The logistic regression model considers the features 'age', 'sex', 'a1', 'a2', 'a4', 'a5', 'a6', 'a8', and 'a10' to be important in predicting the target variable. Among these features, 'age' and 'sex' have relatively higher importance.

The decision tree model identifies 'a5' as the most important feature, followed by 'a1', 'age', 'a4', 'a2', 'a3', 'a8', 'a6', 'sex', and 'a7'.

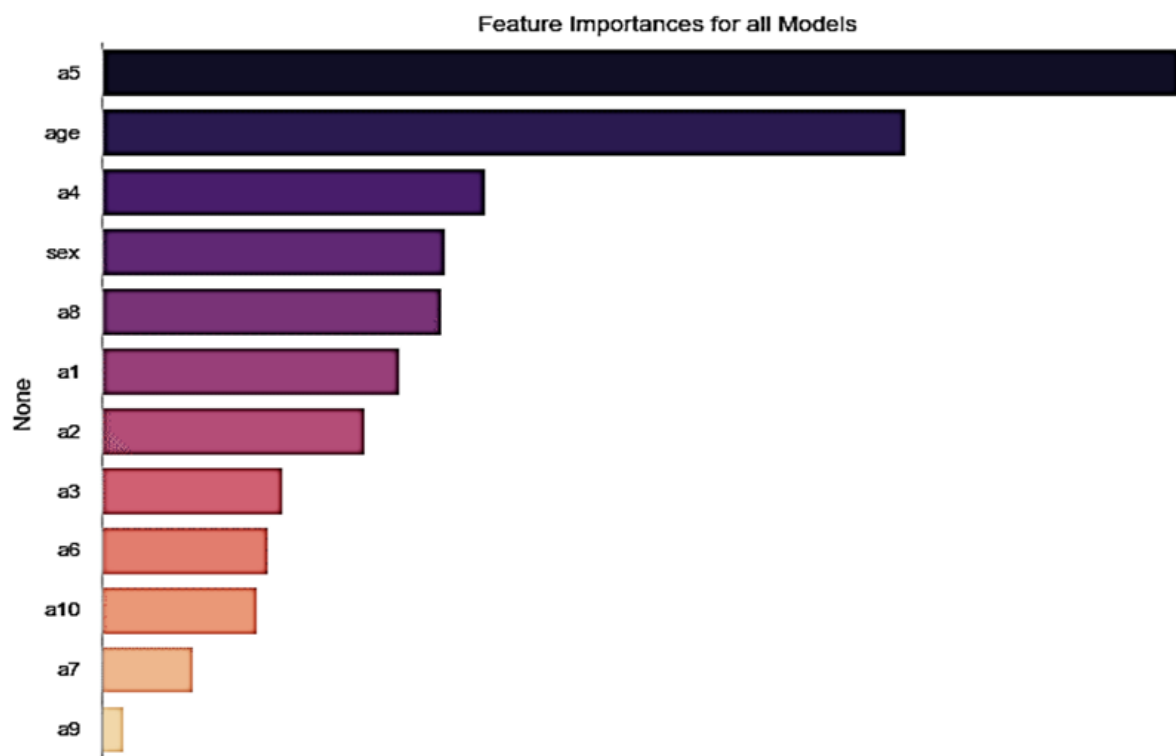
The random forest model ranks the features based on their importance. 'a5' is considered the most important feature, followed by 'a1', 'age', 'a4', 'a2', 'a3', 'a8', 'a6', 'sex', and 'a7'.

The gradient boosting model also ranks the features based on their importance. 'a5' is considered the most important feature, followed by 'a1', 'age', 'a4', 'a2', 'a3', 'a8', 'a6', 'sex', and 'a7' all datasets below are detailed in figure 7.

These results provide insights into which features are considered important by each model. It's important to note that the importance of features can vary between models, and different models may have different criteria for determining feature importance. The importance of features of all machine learning classifiers with all datasets below is detailed in Table 8 except Algorithms (SVM, Naive Bayes, K-Neighbors) do not have this feature.

Table 9 the importance of features of all machine learning classifiers

Feature	Logistic Regression	Decision Tree	Random Forest	Gradient Boosting
Age	0.388826	0.258788	0.255564	0.230074
Sex	0.508803	0.090916	0.090208	0.106742
a1	0.615685	0.055460	0.083204	0.059406
a2	0.696871	0.030778	0.061583	0.012336
a3	0.374806	0.059377	0.051869	0.027684
a4	0.553651	0.166304	0.060354	0.107492
a5	1.320247	0.234783	0.191457	0.319309
a6	0.289054	0.042319	0.056039	0.015947
a7	0.144882	0.023083	0.024820	0.045522
a8	0.798879	0.015109	0.045901	0.072265
a9	0.025382	0.023083	0.025065	0.000780
a10	0.516303	0.000000	0.053936	0.002442

**Fig. 9. Feature importance for all Models.**

6. CONCLUSION

This study presents a machine learning framework designed to detect autism spectrum disorder (ASD) in individuals from young child age groups. The results demonstrate the effectiveness of machine learning techniques as valuable tools to accurately identify cases of autism spectrum disorder. The predictive models proposed in this study could serve as alternative or supportive tools for healthcare professionals in screening children for ASD.

The experimental analysis conducted in this research provides valuable insights for healthcare practitioners, helping them consider the most significant features when screening for ASD cases. However, it is important to acknowledge the limitation of this study, which is the insufficient amount of data to develop a generalized model covering all stages of ASD. A large dataset is crucial for constructing an appropriate model, which was lacking in the dataset used for this analysis.

However, this research has contributed to the development of an automated model that can help medical professionals diagnose autism in children. In future studies, the possibility of using a larger dataset to improve generalizability will be explored. The goal is to collect a more comprehensive data set specifically related to autism spectrum disorder, allowing the construction of a prediction model applicable to individuals of any age. This will enhance the detection of autism spectrum disorder and facilitate better recognition of other neurodevelopmental disorders.

Additionally, the study focused on the use of machine learning algorithms for the prediction and analysis of autism. Future research could explore the integration of other data sources, such as genetic and neuroimaging data, to improve the accuracy and understanding of the underlying mechanisms of ASD.

Furthermore, the study evaluated the performance of different machine-learning models using specific evaluation metrics. It would be valuable to compare the results with other established diagnostic methods, such as clinical assessments and behavioral observations, to assess the clinical utility and potential limitations of machine learning approaches in real-world settings.

7. ACKNOWLEDGMENT

We extend our thanks and appreciation to Dr. Ahmed Al-Awjali and Dr. Walid Al-Bariki for the valuable information that we benefited from completing this paper, and thanks to the Benghazi Autism Center, the Al-Eradah Center, the Al-Noor Center, and the Nusseibeh Center for their cooperation in providing the data required to form the database that was used during this paper.

8. REFERENCES

1. Raj S, Masood S. Analysis and Detection of Autism spectrum Disorder Using machine learning Techniques. *Procedia Computer Science*. (2020); 167: 994-1004
<https://doi.org/10.1016/j.procs.2020.03.399>
2. Omar K S, Mondal P, Islam M N. A Machine Learning Approach to Predict Autism Spectrum Disorder. *International Conference on Electrical, Computer and Communication Engineering*. 2019
<https://doi.org/10.1109/ECACE.2019.8679454>
3. Rahman M M, Usman O L, Muniyandi R C, Sahran S, Mohamed S, Razak R A. A Review of Machine Learning Methods of Feature Selection and Classification for Autism Spectrum Disorder MDPI journals. 2020; 10(12): 949.
<https://doi.org/10.3390/brainsci10120949>
4. Wiggins L D, Reynolds A, Rice C E, Moody E J, Bernal P, Blaskey L, Rosenberg S A, Lee L, Levy S E. Using standardized diagnostic instruments to classify children with autism in the study to explore early development. *Journal of Autism and Developmental Disorders*. 2015; 45: 1271-1280.
<https://doi.org/10.1007/s10803-014-2287-3>
5. Lai M C, Lombardo M V, Baron-Cohen S. Autism. *National library of medicine*. 2014; 383(9920):896-910
[https://doi.org/10.1016/S0140-6736\(13\)61539-1](https://doi.org/10.1016/S0140-6736(13)61539-1)
6. Vakadkar K, Purkayastha D, Krishnan D. Detection of Autism Spectrum Disorder in Children Using Machine Learning Techniques. *SN Computer Science*. 2021; 2(386)
<https://doi.org/10.1007/s42979-021-00776-5>
7. Hossain M D, Kabir M A, Anwar A, Islam M Z. Detecting autism spectrum disorder using machine learning techniques. *Health Information Science and Systems*. 2021; 9(17)
<https://doi.org/10.1007/s13755-021-00145-9>
8. Usta M B, Karabekiroglu K, Sahin B, Aydin M, Bozkurt A. Use of machine learning methods in prediction of short-term outcome in autism spectrum disorders. *Psychiatry And Clinical Psychopharmacology*. 2018; 29(3):320-325
<https://doi.org/10.1080/24750573.2018.1545334>
9. Saeed F. Towards quantifying psychiatric diagnosis using machine learning algorithms and big fMRI data. *Big Data Analytics*. 2018; 3(7)
<https://doi.org/10.1186/s41044-018-0033-0>
10. Thabtah F. Machine learning in autistic spectrum disorder behavioral research: A review and ways forward. *Informatics For Health & Social Care*. 2018; 44(3):278-297
<https://doi.org/10.1080/17538157.2017.1399132>

11. Küpper C, Stroth S, Wolff N, Hauck F, Kliewer N, Hansjosten T S, Becker I K, Poustka L, Roessner V, Schultebrucks K, Roepke S. Identifying predictive features of autism spectrum disorders in a clinical sample of adolescents and adults using machine learning. *Scientific Reports*. 2020; 10(1)
<https://doi.org/10.1038/s41598-020-61607-w>
12. Maalouf M. Logistic Regression in Data Analysis: An Overview. *Int. J. Data Analysis Techniques and Strategies*. 2009 ;3(3): 281–299
<http://dx.doi.org/10.1504/IJDATS.2011.041335>
13. Brownlee J. Naive Bayes for Machine Learning. *Machine Learning Mastery*. 2020 Naive Bayes for Machine Learning - MachineLearningMastery.com
14. Gandhi R. Support Vector Machine Introduction to Machine Learning Algorithms by Towards Data Science. Towards Data Science. 2018 Support Vector Machine — Introduction to Machine Learning Algorithms | by Rohith Gandhi | Towards Data Science
15. Zhang Z. Introduction to machine learning: K-nearest neighbors. *Annals of Translational Medicine*. 2016; 4(11)
<http://dx.doi.org/10.21037/atm.2016.03.37>
16. Cunningham P, Delany S J. k-Nearest Neighbour Classifiers Technical. *ACM Computing Surveys* .2021; 54(6):1-25
<https://doi.org/10.1145/3459665>
17. Reinstein I. Algorithms, CART, Decision Trees, Ensemble Methods, Explained, Machine Learning, random forests algorithm, Salford Systems. *KDnuggets* .2017 Random Forests®, Explained - KDnuggets
18. Azmi S S, Baliga S. An overview of boosting decision tree algorithms utilizing. *International Research Journal of Engineering and Technology (IRJET)* .2020 ;7(5): 6867-6870
<https://www.irjet.net/archives/V7/i5/IRJET-V7I51293.pdf>
19. Mythili M S, Shanavas A M. A novel approach to predict the learning skills of autistic children using SVM and decision tree. *International Journal of Computer Science and Information Technologies*. 2014; 5(6): 7288-7291
<http://www.ijcsit.com/docs/Volume%205/vol5issue06/ijcsit2014050683.pdf>
20. Sujatha R, Aarthy S L, Chatterjee J, Alaboudi A, Jhanjhi N Z. A machine learning way to classify autism spectrum disorder. *International Journal of Emerging Technologies in learning (iJET)*. 2021;16(6): 182-200
<https://doi.org/10.3991/ijet.v16i06.19559>
21. Ferreira A J, Figueiredo M A. Boosting algorithms: A review of methods, theory, and applications. Springer, New York, NY. 2012;
https://doi.org/10.1007/978-1-4419-9326-7_2
22. Qureshi M S, Qureshi M B, Asghar J, Alam F, Aljarboub A. Prediction and Analysis of Autism Spectrum Disorder Using Machine Learning Techniques. *Journal of health care engineering*. 2023;12(10):9815989
<https://doi.org/10.1155/2023/9815989>
23. Mellema C J, Nguyen K P, Treacher A, Montillo A. Reproducible neuroimaging features for Diagnosis of autism spectrum disorder with machine learning. *Scientific reports*. 2022; 12(1): 3057
<https://doi.org/10.1038/s41598-022-06459-2>
24. Ali MT, Gebreil A, ElNakieb, Y, Elnakib A, Shalaby A, Mahmoud A, Sleman A, Giridharan G A. Barnes G, Elbaz A S. A personalized classification of behavioral severity of autism spectrum disorder using a comprehensive machine-learning framework. *Scientific reports*. 2023 ; 13(1): 17048
<https://doi.org/10.1038/s41598-023-43478-z>
25. Rogala J, Żygierewicz J, Malinowska U, Cygan H, Stawicka E, Kobus A, Vanrumste B. Enhancing autism spectrum disorder classification in children through the integration of traditional statistics and classical machine learning techniques in EEG analysis. *Scientific Reports*. 2023. 13(1): 21748
<https://doi.org/10.1038/s41598-023-49048-7>