

Machine Learning model for Autism Spectrum Disorder Detection using EEG signal

A. Arockia Helen sushma^{1*}, S. PriyaDharshini², S. Nagakumar raj³, A. Kayalvizhi⁴

^{1,2,3,4} Assistant professor, Department of CSE (Cyber Security), SSM Institute of Engineering & Technology, Tamil Nadu, India.

^{1*}Corresponding Author E-mail: helensushma@gmail.com

ABSTRACT: Diagnosing autism spectrum disorder (ASD) in youngsters presents considerable difficulties owing to the disorder's intricacy and unpredictability. Conventional diagnostic methods frequently depend on single-modal data, which can be restrictive and may not produce optimal outcomes. To address these limitations, we present a novel diagnostic methodology that utilizes the integration of electroencephalogram (EEG) data for enhanced ASD detection. The proposed study uses machine learning algorithms, mainly logistic regression (LRC) and extra tree classifier (ETC), to find correlations and new insights from data that is hidden in a latent feature space. Machine learning classifier methods generate optimal feature representations that enhance discriminability and generalization, ultimately improving diagnostic accuracy. The EEG dataset evaluated with ASD indicates that our approach markedly surpasses current methodologies. This innovation has significant potential for equipping physicians with a more objective and dependable tool for diagnosing autism spectrum disorder in children.

Keywords: ASD, EEG, ML, Logistic Regression classifier and Extra tree classifier

1. Introduction

Many symptoms are associated with ASD, the most common of which are impairments in interpersonal interaction and communication. and behaviors that are repetitive or very focused. The idea of a "spectrum" highlights the wide variation in symptoms and how severe they are from one individual to another [1] In order to make an accurate diagnosis of ASD, there are several important steps to take, including developmental monitoring, screening, and thorough examinations. Rendering to the Centers for Disease Control (CDC), ASD is noticed in approximately one child out of every 54 children currently living in the United States. This highlights the essential for early detection and individualized therapies to help individuals impacted [2].

The analysis of ASD relies on the occurrence or lack of specific behaviors, recognizable signs, and growing interruptions. No biological marker or laboratory test is available for the analysis of ASD. The Diagnostic and Statistical Manual (DSM-IV-TR) categorizes ASD as part of the broader classification of Pervasive Developmental Disorders (PDD).

PDD is a diseases marked by deficits in shared social collaboration and statement abilities, alongside the manifestation of stereotyped behaviors, interests, and activities. Pervasive developmental disorders encompass:

- Autism
- Childhood disintegrative disorder (CDD)
- Rett's Syndrome
- PDD-NOS
- Asperger's Disorder

Researchers are closely examining automated algorithms to detect diseases for use in healthcare.

1.1 Literature Survey

Research into and approaches to treating Autism Spectrum Disorder (ASD) have so far garnered a great deal of interest from scientists and medical experts around the world. The root cause of autism spectrum disorder (ASD) is still not well understood, which makes it difficult to create more specific treatments, despite the widespread interest in this illness. While behavioral evaluations and standardized diagnostic scales are helpful in making a clinical diagnosis of ASD, they do nothing to take into account the disorder's molecular complexity [3]. The increasing interest in medical informatics among the data science research community has prominently emphasized disease prediction as a crucial topic of study [4]. As illness prediction leads modern healthcare developments, early intervention becomes feasible, significantly enhancing patient outcomes.

Machine learning (ML) methodologies have arisen as a revolutionary strategy, utilizing complex health data to accurately forecast disease risks. By leveraging historical health records, these machine learning models can predict future outcomes, facilitating more preventive and tailored patient care. Current research is systematically evaluating the influence of supervised machine learning methods on enhancing the precision and dependability of illness prediction models [5–9].

2. System Model

2.1 ML Algorithm

ML is integral to the extensive domain of artificial intelligence (AI), concentrating on the creation of structures which increase their performance via involvement or data acquisition. Artificial intelligence (AI) generally denotes the development of technology that replicates human cognitive functions, including learning, thinking, and decision-making. Machine learning

distinguishes itself by allowing machines to independently identify and adjust to patterns in data, thus executing tasks without explicit human programming for each function. This capability positions machine learning as a crucial catalyst for numerous contemporary AI applications [10].

2.2 Supervised Learning Algorithm

In supervised learning, we train models using input-output pairs, relying on labeled data to guide the prediction of precise outcomes.

To facilitate the algorithm's learning of a mapping function for use in predicting tests, the dataset contains input structures and their matching output labels [11].

In supervised learning, we instruct algorithms to correlate inputs with their corresponding outputs. We annotate both the training and validation datasets. Figure 1 shows the process of Supervised Learning.

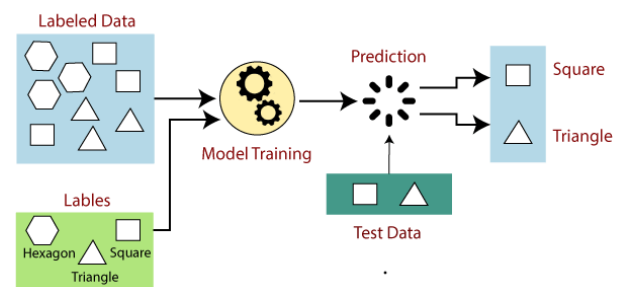


Figure 1: System Model for Working Flow of Supervised Learning Algorithm

Step 1: Begin by identifying the necessary type of training dataset.

Step 2: Obtain the annotated training data.

Step 3: Partition the dataset into three subsets.

Identify the input features in the training dataset that provide adequate information for the model to accurately predict the output. Select an appropriate algorithm for the model, such as a support vector machine or a decision tree.

Step 4: Implement the method on the training dataset, employing validation sets as control parameters as necessary, as they represent subsets of the training data.

Step 5: Assess the accurateness by implementing it on the test dataset. If the model precisely forecasts the result, it is accurate. Several methods are part of supervised learning classification, including logistic regression and extra trees.

2.3 ASD Data set

This study utilized screening data for autism spectrum disorders in teenagers. The dataset comprises 104 observations and 21 variables. The data was gathered by a questionnaire administered via a mobile application, soliciting participants' details, including age, gender, ethnicity, country of residency, presence of jaundice, diagnosis of autism, and ten standardized questions designed to assess the likelihood of ASD.

The list and type of 21 classes in this dataset are shown in Table 1.

Table 1: Type of 21 classes

S.No	Input data sets
1	A1_Score
2	A2_Score
3	A3_Score
4	A4_Score
5	A5_Score
6	A6_Score
7	A7_Score
8	A8_Score
9	A9_Score
10	A10_Score
11	Age
12	Gender
13	Ethnicity
14	Jaundice
15	Autism
16	Country_of_res
17	Used_app_before
18	Result
19	Age_desc
20	Relation
21	Class/ASD

A1-A10 Score: We report the results for attributes A1_Score to A10_Score in binary format, either 0 or 1, based on the respondents' answers to the questions. We examined the data frequency for each attribute to determine the presence of any

sparse classes. We observed a normal distribution of all responses of 0 and 1, with no exceptional question score.

3. Proposed Methodology

3.1 Logistic Regression Classifier

For classification tasks, supervised machine learning techniques like logistic regression predict the likelihood of an instance falling into a specific category. It is a statistical technique used to determine the relationship between two variables.

3.1.1 A machine learning classification system comprises four components:

1. The patterns of inputs are illustrated. The actual input is represented by the value of a feature vector $[x_1, x_2, \dots, x_n]$ for each observation x (i). For this discussions, it refer the value of feature i for input x (j) as x (j) i , or simply x_i ; it also notated as f_i , $f_i(x)$, or $f_i(c, x)$ in the case of multiclass classification problems.
2. The classification function returns some y 's expected class, given $p(y|x)$. In the following sections, it present the sigmoid function and the softmax function which is used to classify data.
3. The objective function is defined for purposes of optimizing learning via the reduction of a loss function that represents errors made while processing training samples. The cross-entropy loss function is defined in this section.
4. Describe an algorithm that is capable of optimizing the objective function. Hence provide an example of the stochastic gradient descent algorithm.

3.1.2 The Logistic regression follows a two-step process which includes:

- Training: Stochastic gradient descent is used to optimize the parameters w and b for the Logistic regression function, which provides the probability of the event occurring.

- Testing: Produce the results of the model. By computing $p(y|x)$ for a given test input x , and return as the output the label of response y .

3.1.3 The Significance of the Sigmoid Function

The primary objective of Binary Logistic Regression is to create a classification model that classify a new observation as either a class or not a class. Hence utilize the sigmoid function to reach this classification model. Analyze a single observation or feature vector ($x_1, x_2, \dots, x_n$) which produce one output response y that either fits in or does not fit. The classification use features which indicate about autism, and $P(y=1|x)$ is a probability that the given document relates to autism and $P(y=0|x)$ is the probability that it does not relate to autism.

To make predictions using logistic regression, a weight vector and a bias (intercept) term are calculated based on a training set of observations. Each weight (w_i) is a real number representing the contribution, or importance of each input feature, (x_i) to the classification decision as to whether or not the instance belongs to the positive class. The bias term (b) is used to represent the contribution of all weighted input features.

Once the weights are learned from the training data, the classification algorithm compute the total evidence (z) for classifying a test case using the following process: for each input feature value (x_i), multiply it by the corresponding weight (w_i); then add all of the products to produce a weighted sum (z) of the input feature contributions.

$$z = \left(\sum_{i=1}^n w_i x_i \right) + b \quad (1)$$

The dot product of two vectors ($a \cdot b$) is defined as the sum of the products of their corresponding components. Boldface notation (b) represents vectors. Consequently, the ensuing expression constitutes an equivalent version of

$$z = \mathbf{w} \cdot \mathbf{x} + b \quad (2)$$

Equation 2 requires z to have a valid probability, specifically within the range of 0 to 1. The weights in the network are real-valued, however, there is no restriction on the output; z range from negative infinity to positive infinity.

In training of a neural network the optimize weights w and b using stochastic gradient descent and the cross-entropy loss function

In test, $p(y|x)$ and compare the results from two outputs ($y=0$, or 1). The output with the highest probability is the correct label.

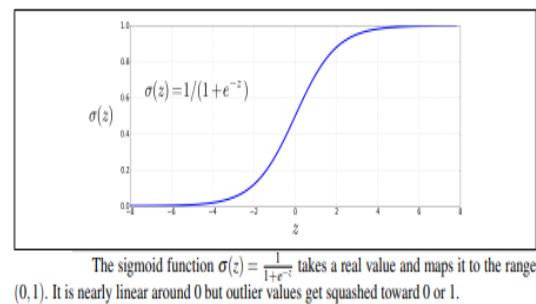


Figure 2: Sigmoid function

The sigmoid function, named for its S-shaped curve, is sometimes referred to as the logistic function, which is the origin of the term logistic regression. Figure 2 visually illustrates the equation of the logistic function, which represents the sigmoid.

$$\sigma(z) = \frac{1}{1 + e^{-z}} = \frac{1}{1 + \exp(-z)} \quad (3)$$

Among its many benefits is the fact that the sigmoid function may convert any real number to the interval (0,1), the exact range needed for a probability. As a result of its flattening at the extremes and near-linear behavior at 0, it successfully compresses outlier values toward 0 or 1.

The total of the weighted attributes, when applied to the sigmoid function, provides a quantity amongst 0 and 1. To calculate the probability of $y=1$, start with the two probabilities, $p(y=1)$ and

$p(y = 0)$, and check for the correct probabilities. Here's how to accomplish this goal:

$$\begin{aligned}
 P(y=1) &= \sigma(\mathbf{w} \cdot \mathbf{x} + b) \\
 &= \frac{1}{1 + \exp(-(\mathbf{w} \cdot \mathbf{x} + b))} \\
 P(y=0) &= 1 - \sigma(\mathbf{w} \cdot \mathbf{x} + b) \\
 &= 1 - \frac{1}{1 + \exp(-(\mathbf{w} \cdot \mathbf{x} + b))} \\
 &= \frac{\exp(-(\mathbf{w} \cdot \mathbf{x} + b))}{1 + \exp(-(\mathbf{w} \cdot \mathbf{x} + b))}
 \end{aligned} \tag{4}$$

The sigmoid function possesses the characteristic

$$1 - \sigma(x) = \sigma(-x) \tag{5}$$

On the other hand, we might have represented the probability of y equaling zero as $\sigma(-(\mathbf{w} \cdot \mathbf{x} + b))$.

Usually, we refer to the score $\text{logit } z = \mathbf{w} \cdot \mathbf{x} + b$ from equation (2) as the logit. It is the inverse of sigmoid function thus $p/(1 - p)$ is the log function of odds.

$$\text{logit}(p) = \sigma^{-1}(p) = \ln \frac{p}{1-p} \tag{6}$$

Classification Utilizing Logistic Regression
The previous section's sigmoid function provides a method for calculating the probability $P(y = 1|x)$ for a given instance x . The decision boundary is as follows,

$$\text{decision}(x) = \begin{cases} 1 & \text{if } P(y=1|x) > 0.5 \\ 0 & \text{otherwise} \end{cases} \tag{7}$$

3.2 Extra Trees Classifier

One ensemble learning technique that produces classification results by combining the outputs of numerous independent decision trees into a single "forest" is the Extremely Randomized Tree Classifier, also recognized as the Extra Trees Classifier. Although it constructs decision trees in a distinct manner, it resembles the Random Forest Classifier significantly. At each split node, Extra Trees casually picks a division of k structures after the complete feature set, utilizing the original

training data to construct each decision tree. We select the optimal feature from this subset using a mathematical metric, usually the Gini Index, to divide the data. This random selection process resulted in the creation of uncorrelated decision trees is the result of this random selection process.

We calculate the total decrease in the mathematical criterion (Gini Index or entropy) for each feature while building the tree to select the best features. To help users choose the best k features for their purposes, this score—also known as the feature's Gini or entropy value—ranks the attributes from most to least important. When compared to other methods, Extra Trees' computational efficiency typically yields faster results [11].

For feature selection, the Extra Trees Classifier has multiple significant benefits:

1. **Resilience to noise and extraneous features:**
By utilizing several trees and prioritizing features based on significance, Extra Trees adeptly handles datasets characterized by numerous characteristics and noise.
2. **Computational efficiency:** The concurrent creation of decision trees expedites the training process, particularly advantageous for high-dimensional data.
3. **Reducing bias:** Choosing feature subsets and split points at random reduces the bias that comes with using a single decision tree, making the evaluation of feature importance fair.
4. **Feature ranking:** Users can prioritize features and understand their impact on the target variable by evaluating them based on their significance.
5. **Dealing with multicollinearity:** Extra Trees reduce the impact of feature correlations by picking feature subsets and splits at random, which is different from methods that depend on clear feature relationships.
6. **Adaptability in feature selection:** Users can modify the threshold for feature inclusion according to feature significance, opting

for either a minimal set or a more extensive group to reconcile model complexity and performance.

7. Enhanced generalization and interpretability: By picking only the most pertinent characteristics, Extra Trees mitigates overfitting and improves the model's interpretability by emphasizing the most significant aspects.

The advantages include an active instrument for feature collection, particularly in complicated datasets where speed and efficiency are paramount.

4. Results and Discussion

The input EEG data set have 800 sample details in the training and 200 test samples are tested with A1 score to A10 score and detected autism

4.1 Target class output

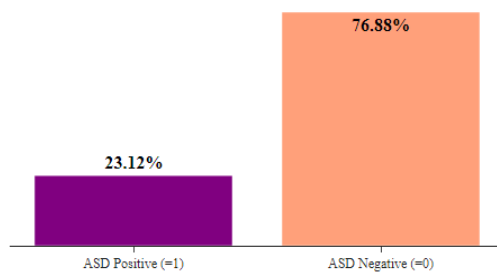


Table 2: Logistic regression classifier -ROC-AUC curve output

INDEX	0	1	2	3	4	MEAN	STD
Fit_time	1.293571949	1.553783	1.325692	1.37479	1.539857	1.417539	0.108771
Score_time	0.015843868	0.005867	0.004108	0.003963	0.003908	0.006738	0.004611
Test_roc_auc	0.919797847	0.863107	0.900022	0.954296	0.939134	0.915271	0.031832
Train_roc_auc	0.924165019	0.935825	0.92812	0.915403	0.916392	0.923981	0.007597

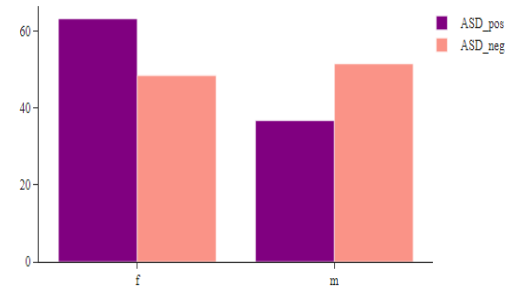
Table 3: Extra tree classifier- ROC-AUC curve output

INDEX	0	1	2	3	4	MEAN	STD
Fit_time	1.049814	1.027649	1.014085	1.64289	1.635304	1.273949	0.29837
Score_time	0.083253	0.067348	0.076027	0.103637	0.106161	0.087285	0.015259
Test_roc_auc	0.907273	0.875192	0.892112	0.946825	0.95935	0.91615	0.032067
Train_roc_auc	0.923094	0.932831	0.926211	0.912959	0.911482	0.921315	0.008078

The above table 2 and table 3 provide ROC-AUC output of mean and standard deviation values of both logistic regression and extra tree classifier for both training and test case.

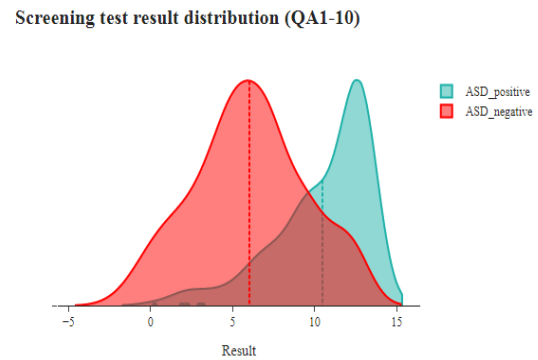
This figure shows 23.12% of ASD positive results and 76.88% negative results.

4.2 Gender of patient result



This bar chart describes patient result of male which has 62% of ASD positive and 50% ASD negative. The female result has 38% of ASD positive and 52% ASD neagative.

4.3 Distribution of input data set: A1-A10 score plot



The ROC-AUC curve is a critical tool in medicine for disease prediction and determining the optimal system model. ROC curves in logistic regression aid in determining the suitable threshold to

classify a new observation as a "failure" (0) or a "success" (1).

1) Optimal ROC curve:

An optimal ROC curve resembles a right angle. When this transpires, the differentiation between 0 and 1 becomes so pronounced that statistical analysis is nearly superfluous. At each point along the curve, either sensitivity or specificity is at 100%, indicating a complete absence of overlap between the two. In this case, both sensitivity and specificity may attain 100%, through the AUC for this curve equaling 1.

2) ROC Curve with No Predictive Power:

Conversely, the least favorable ROC curve manifests as a straight line at a 45-degree angle, signifying a lack of predictive capability. This indicates that the model forecasts results no more accurately than random chance, similar to the act of flipping a coin. In this case, increasing the true positive rate (sensitivity) results in increase, thereby undermining the model's effectiveness.

The AUC for this particular kind is 0.5. The test_ROC_AUC has a mean of 91.6%, and the train_ROC_AUC has a mean of 92.1%, signifying a highly effective predictive model exposed in table 4.

Table 4. *Comparative analysis of ROC-AUC curve*

Sl. No	ROC_AUC score (For machine learning model performance analysis)	Average (Should be greater than 50%)
1	Train_ROC_AUC (in both logistic regression and Extra tree classifier)	91.61 (Best model)
2	Test_ROC_AUC (in both logistic regression and Extra tree classifier)	92.13 (Best model)

5. Conclusion

This research presents a method for detecting autism spectrum disorder with an enhanced and

more accurate machine learning model. Both classifiers utilize the ROC curve, which exceeds conventional prediction techniques. The ROC curve visually evaluates model performance across different thresholds, with the AUC signifying the probability that the method would accurately rank a randomly chosen positive example above a negative one. These classifier algorithms yielded superior results compared to conventional machine learning methods, with an accuracy of 91.6%.

References

1. Catalina Sarmiento; Chloe Lau, Year: 2020, "Diagnostic and Statistical Manual of Mental Disorders, 5th Ed. Hoboken, Nj, Usa: Wiley, pp. 125–129.
2. Guifeng Xu; Lane Strathearn; Buyun Liu; Wei Bao, Year: 2018, "Prevalence of Autism Spectrum Disorder Among Us Children and Adolescents, 2014–2016," Jama, Vol. 319, No. 1, pp. 81–82.
3. Zhi-An Huang; Zexuan Zhu; Chuen Heung Yau; Kay Chen Tan, Year: 2021, "Identifying Autism Spectrum Disorder From Resting-State Fmri Using Deep Belief Network", IEEE Transactions on neural networks and learning systems, Vol: 32, No: 7, pp. 2847-2861.
4. Shahadat Uddin; Arif Khan; Md Ekramul Hossain; Mohammad Ali Moni, Year: 2019, "Comparing different supervised machine learning algorithms for disease prediction", BMC medical informatics and decision making, Vol. 19, pp. 281.
5. Rahul Katarya; Sunit Kumar Meena, Year: 2021, "Machine Learning Techniques for Heart Disease Prediction: A Comparative Study and Analysis", Health and Technology, Vol: 11, pp. 87–97.

6. A. Sazzadur Rahman; F. J. M. Shamrat; Zarrin Tasnim; Joy Roy; Syed Akhter Hossain, Year: 2019, "A comparative study on liver disease prediction using supervised machine learning algorithms", *International Journal of Scientific & Technology Research*, Vol: 8, No: 11, pp. 419–22.
7. F. Javed Mehedi Shamrat; Md Asaduzzaman; A. Sazzadur Rahman; Raja Tariqul Hasan Tusher; Zarrin Tasnim, Year: 2019, "A comparative analysis of parkinson disease prediction using machine learning approaches", *Int J Sci Technol Res.*; Vol: 8, No: 11, pp. 2576–80.
8. Parul Sinha; Poonam Sinha, Year: 2015, "Comparative study of chronic kidney disease prediction using KNN and SVM", *International Journal of Engineering Research and Technology*, Vol: 4, No: 12, pp. 608–12.
9. Shahadat Uddin; Ibtisham Haque; Haohui Lu; Mohammad Ali Moni; Ergun Gide, Year: 2022, "Comparative performance analysis of K-nearest neighbour (KNN) algorithm and its different variants for disease prediction", *Scientific Reports*, Vol: 12, No: 1, pp. 1–11.
10. Mohamed Alloghani; Dhiya Al-Jumeily; Jamila Mustafina; Abir Hussain; Ahmed J. Aljaaf, Year: 2020, "A systematic review on supervised and unsupervised machine learning algorithms for data science", In: *Supervised and unsupervised learning for data science*. Springer. pp. 3–21.
11. Pierre Geurts; Damien Ernst; Louis Wehenkel, Year: 2006, "Extremely randomized trees", *Machine learning*, Vol: 63, No: 1, pp. 3–42.
12. Shasha Zhang; Dan Chen; Yunbo Tang; Lei Zhang, Year: 2021, "Children ASD evaluation through joint analysis of EEG and eye-tracking recordings with graph convolution network", *Frontiers Human Neurosci.*, Vol: 15, pp. 651349.
13. Arnaud Delorme; Scott Makeig, Year: 2004, "EEGLAB: An open source toolbox for analysis of single-trial EEG dynamics including independent component analysis", *Journal of neuroscience methods*, Vol: 134, No: 1, pp. 9–21.
14. Arthur Asuncion; David Newman, Year: 2007, "UCI machine learning repository", Irvine. USA: CA.
15. Yoshua Bengio; Salem Lahlou; Tristan Deleu; Edward J. Hu; Mo Tiwari; Emmanuel Bengio, Year: 2023, "Journal of Machine Learning Research", Vol: 24, pp. 1–55.
16. Bogdan Alexandru Cociu; Saptarshi Das; Lucia Billeci; Wasifa Jamal; Koushik Maharatna; Sara Calderoni; Antonio Narzisi; Filippo Muratori, Year: 2017, "Multimodal functional and structural brain connectivity analysis in autism: A preliminary integrated approach with EEG, fMRI, and DTI," *IEEE Transactions on Cognitive and Developmental Systems*, No: 99, pp. 1–1.
17. Lisa E. Mash; Brandon Keehn; Annika C. Linke; Thomas T. Liu; Jonathan L. Helm, Frank Haist; Jeanne Townsend; Ralph-Axel Müller, Year: 2020, "Atypical relationships between spontaneous EEG and fMRI activity in autism", *Brain Connectivity*, Vol: 10, No: 1, pp. 18–28.
18. Juan Camilo Vásquez-Correa; Tomas Arias-Vergara; Juan Rafael Orozco-Arroyave; Björn Eskofier; Jochen Klucken; Elmar Nöth, Year: 2019, "Multimodal assessment of Parkinson's disease: A deep learning approach", *IEEE journal of biomedical and health informatics*, Vol: 23, No: 4, pp. 1618–1630.
19. Jun Shi; Xiao Zheng; Yan Li; Qi Zhang; Shihui Ying, Year: 2018, "Multimodal neuroimaging feature learning with multimodal stacked deep polynomial networks for diagnosis of

- Alzheimer's disease", IEEE journal of biomedical and health informatics, Vol: 22, No: 1, pp. 173–183.
20. Hamdi Dibeklioglu; Zakia Hammal; Jeffrey F. Cohn, Year: 2018, "Dynamic multimodal measurement of depression severity using deep autoencoding", IEEE journal of biomedical and health informatics, Vol: 22, No: 2, pp. 525–536.
 21. Arnaud Delorme; Scott Makeig, Year: 2004, "EEGLAB: An open source toolbox for analysis of single-trial EEG dynamics including independent component analysis", Journal of neuroscience methods, Vol: 134, No: 1, pp. 9–21.
 22. Andrej Miklosik; Nina Evans, Year: 2020, "Impact of big data and machine learning on digital transformation in marketing: A literature review", IEEE Access, Vol. 8, pp. 101284–92.
 23. Rahul Katarya; Sunit Kumar Meena, Year: 2021, "Machine learning techniques for heart disease prediction: a comparative study and analysis", Heal Technol, Vol: 11, pp. 87–97.
 24. Kanika Pahwa; Ravinder Kumar, Year: 2017, "Prediction of heart disease using hybrid technique for selecting features", 2017 4th IEEE Uttar Pradesh Section International Conference on Electrical, Computer and Electronics (UPCON), Mathura, pp. 500–504.
 25. Ridhi Saini; Namita Bindal; Puneet Bansal, Year: 2015, "Classification of heart diseases from ECG signals using wavelet transform and kNN classifier", International Conference on Computing, Communication & Automation, Noida, pp. 1208–1215.