

Comparative Analysis of Machine Learning Techniques for Automated Autism Spectrum Disorder Detection

Naren Ramakrishnan

Chirec International School, 1-55/12, Botanical Garden Rd, Sri Ram Nagar, Kondapur, Hyderabad, Telangana 500084, India;
narenramakrishnan5@gmail.com

ABSTRACT: Autism spectrum disorder (ASD) is a developmental disorder linked to intellectual and neurobehavioral disorders, impacting an individual's comprehension and communication skills. Early detection and diagnosis of ASD are crucial for various reasons, including educational planning, implementing effective treatments, providing necessary family support, and timely medical intervention, making automated ASD diagnosis methods essential. People diagnosed with ASD share similar traits. We will be doing a comparative analysis of different machine learning models to identify which models perform the best in determining whether an individual has autism or does not. The six machine learning models used in this study are Artificial Neural Network, Decision Tree Classifier, Ensemble Techniques, K Nearest Neighbors, Logistic Regression, and Support Vector Machine. These models will be used to predict whether an individual does or does not have ASD. We use publicly available datasets- which consist of information about those with ASD and those without- to train our model by trying to identify these common traits. Based on our study, we recommend using a random forest classifier, which had the highest accuracy, recall, and f1 score for automated ASD detection.

KEYWORDS: Robotics and Intelligent Machines; Machine Learning; Autism spectrum disorder; Comparative analysis; Autism detection; Machine Learning models.

■ Introduction

Affecting one out of every fifty-four children in the United States,¹ autism spectrum disorder, or ASD, is a major problem in today's world. Stemming from genetic factors, this disorder poses several challenges to those with it, including difficulties with social interaction, challenges in academic settings,¹ and heightened sensory sensitivities, thus affecting their everyday lives in one way or another. The significance of autism as a problem is underscored by its ability to manifest along a spectrum, where individuals with autism grapple with a unique amalgamation of strengths and challenges, requiring tailored support and intervention.

It is not possible to prevent an individual from developing ASD, and neither is there a cure for ASD; however, early detection of this disorder is important as this gives time for doctors and family to find an appropriate method of treatment to help the individual.

Unfortunately, there is no genetic test that can diagnose someone with autism. Traditionally, specialists diagnose children by observing the child's communication, social interaction, play, and restricted and repetitive behavior. Others do so with a structured interview. A neuropsychological evaluation can also be carried out, yielding important information regarding the child's intellectual ability. Autism diagnosis may also involve carrying out a simple auditory and visual examination. Unfortunately, there are long waiting times for these lengthy tests, and the procedures are not cost-effective. The models presented in this paper would provide a method for family members to assess whether an individual has ASD, potentially providing a more accessible alternative.

Machine learning can help simplify detection as machine learning models can identify features or common trends in the data that are not easily observed by the naked eye. These models can also evolve, improving accuracy over time by learning from new data. This would help in the earlier detection of ASD as the model may be able to identify features present in an individual with ASD that may not be observable by the naked eye. Machine learning models also provide a cheaper alternative to the more expensive existing detection methods.

We will conduct this comparative study by training different models with a dataset and evaluating their performance across multiple metrics against a test dataset. These metrics are accuracy score, precision score, recall score, and f1 score. Each model will be given a score per metric depending on how the model predicts the samples in the test dataset. The model with the highest score for a metric will be the best-performing model for that metric. Our studies suggest that the random forest classifier is best suited for detecting ASD. The random forest classifier model had the best accuracy, recall, and f1 score, while the K Nearest Neighbors classifier had the highest precision score.

Before discussing the results that lead to this conclusion in the results and discussion section, we first cover the materials and methodology used in the methods section. This section will describe the dataset, the machine learning methods evaluated, and the metrics used. Lastly, we will give our final recommendation in the conclusion section.

■ Methods

Dataset:

In this study, we have used three existing datasets available on Kaggle. These datasets are part of UC Irvine's Machine Learning Repository and make use of the Q-Chat-10. The Q-CHAT was developed as a checklist filled out by the child's parents based on behavioral observations. The Q-Chat-10 is an improvement on its predecessors: it is more accessible and is shortened to only ten questions. It comprises ten questions with five possible answers: Always, Usually, Sometimes, Rarely, and Never.

By combining the three different datasets, we aim to build a model that generalizes to a broader range for ASD detection.² The first dataset is focused on adults and has 704 samples, with the ages ranging from 17 years old to 64 years old.³ The second dataset is focused on children/adolescents. It has 1985 samples, ranging from one year to eighteen years old.⁴ The third dataset is focused on toddlers/children and has 1054 samples ranging from 12 months to 36 months old. This brings the number of samples in the new dataset to 3743. These datasets had initially comprised 18 attributes each; however, as some were not shared between the datasets, some had to be excluded. We note that none of the excluded attributes were part of the questionnaire filled in by the participants but were part of the metadata collected for each participant instead.

The new dataset encompasses diverse ethnicities, including Asian, Black, Hispanic, Latino, Middle Eastern, Mixed, Native Indian, Others, Pasifika, South Asian, Turkish, and White European. This diversity indicates that the dataset is not confined to specific geographical locations, and our development models should be generalized.

These datasets were developed by Dr Fadi Faye Thabtah using a mobile app called ASDTests. Since this data was collected through a mobile app, some input data may need to be more accurate. As a result, the model may be trained with misleading information, leading to a reduction in the model's accuracy. A complete list of the questions asked can be found in Table 1.

To merge the three datasets, we initially loaded them onto Pandas DataFrames and concatenated them into one. We then identified the similar questions asked in each of the DataFrames (the questions in Table 1) and filtered the new DataFrame to contain only the inputs to these questions.

Table 1: Questions/Data used.²

Q1	Does your child look at you when you call his/her name?
Q2	How easy is it for you to make eye contact with your child?
Q3	Does your child point to indicate that he/she wants something? (e.g., a toy that is out of reach)
Q4	Does your child point to share interests with you? (e.g., pointing at an interesting sight)
Q5	Does your child pretend? (e.g., care for dolls, talk on a toy phone)
Q6	Does your child follow where you are looking?
Q7	If you or someone else in the family is visibly upset, does your child show signs of wanting to comfort them? (e.g., stroking hair) hugging them)
Q8	Would you describe your child's first words as:
Q9	Does your child use simple gestures? (e.g., wave goodbye)
Q10	Does your child stare at nothing with no apparent purpose?

Models:

In this work, we will be comparing numerous models against each other to identify which one works the best. Each model is designed to identify patterns in the dataset and make a prediction about the newly entered data based on the pattern it has identified. The models we will work with are Logistic regression, Decision Tree classifier, Ensemble techniques, Support Vector Machine, K Nearest Neighbors, and Artificial Neural Network.

A few critical parameters, known as hyperparameters, affect each model's performance. We will experimentally select the optimal value for each model for these hyperparameters. This evaluation is presented in the results and discussion section.

All the models are implemented using the TensorFlow framework.⁵ As such, only the possible hyperparameter values implemented by this framework are considered.

Logistic regression:

Logistic regression is an extension of linear regression.⁶ Linear regression is a machine learning technique that predicts a quantitative response Y based on a single predictor X.⁷ Logistic regression was developed by statistician David Cox in 1959 and is based on the principle of probability. Logistic regression helps us estimate the probability of Y based on one or more predictor variables. It determines how much the probability of an event changes based on numerous other events. Logistic regression can only output binary values and is thus suitable for use here as our final output is one of two options (autistic or non-autistic).

The main parameters that affect how the model works are the regularization parameters, types of functions used, and the penalty term. The different penalty terms are l1, l2, and elasticnet.

The logistic regression model has been used for the identification of other psychological disorders:⁸ mental health problems among children achieving an accuracy score of 70%;⁹ Depression detection achieving an accuracy score of 48%.

Decision Tree classifier:

The decision tree classifier algorithm is similar to that of a dichotomous key in biology. A dichotomous key works by asking numerous yes or no questions until, finally, after answering all the questions, you land on one output.¹⁰ A decision tree classifier is identical in the way it works to a dichotomous key wherein the root is the main node or main question, the stem connecting each question to the next, the branches being follow-up questions based on the previous questions (the path followed) and finally, the leaves are the output. A decision tree is suitable for this scenario. The questions mentioned in the dataset above can be made into a tree-like structure split into yes or no branches, with each yes or no branch leading to a different path and a different output.

The main parameter that affects how the model works is the criterion used to split. The different criteria are gini, entropy, and log loss.¹¹ The decision tree classifier model has also been used to detect PTSD, achieving an accuracy score of 77%.

Ensemble techniques:

Ensemble techniques combine several individual techniques to arrive at a final output. This algorithm is inspired by how

humans work together to solve a problem. We combine the opinions of several experts to make crucial decisions.¹² Ensemble techniques similarly induce several individual functions into a bigger algorithm and produce an output.

Specifically, we consider an ensemble of Decision Tree classifiers, also known as a random forest. The most critical parameter for this ensemble technique is the number of estimators, which can be any number.

The random forest model has been used for identification of other psychological disorders:⁸ mental health problems among children were detected, achieving an accuracy score of 73.9%;¹³ PTSD cases were detected, achieving an accuracy score of 89.1%; and¹⁴ Schizophrenia cases achieved an accuracy score of 73.7%.

Support Vector Machine:

Support vector machine (SVM) algorithms work by dividing lines that partition the data samples into different regions.¹⁵ These dividing lines are defined using additional artificial data points called support vectors, hence the name. The model's prediction is consecutively determined by which side of dividing the new input falls on.

In its most basic form, SVMs are only usable for datasets with few variables but can be extended by processing the data using a predefined kernel. The type of kernel used is, therefore, an important hyperparameter. The kernels are linear, poly, radial basis function, sigmoid, and precomputed.

The SVM model has been used for the identification of other psychological disorders:¹⁶ PTSD detection achieving an accuracy score of 83.59%;¹⁷ bipolar disorder detection achieving an accuracy score of 96.36%.

K-Nearest Neighbors:

The k-nearest neighbors (KNN) algorithm works based on distance.¹⁸ Important parameters affect this model's performance: the function to calculate the distance of the new input from its counterparts and the number of neighbors (data points) already existing. There can be any number of neighbors.⁹ The KNN model has also been used for depression detection, achieving an accuracy score of 49%.

Artificial neural network:

The artificial neural network algorithm is inspired by the human brain, in the sense that our brain has several neurons that finally connect to the brain and a body part.¹⁹ The artificial network similarly has several nodes that connect layer by layer. There is an input layer where the data is inputted, then the hidden layers where a function is carried out on the values output from the previous layers to modify these values, and finally, the output layer. Nonlinearity is introduced into the model through activation functions.

We consider the neural network's architecture as the main hyperparameter for our evaluation. As such, we consider the number of layers in the model and the number of nodes used in each layer.⁸ The artificial neural network has also been used to detect mental health problems in children, achieving an accuracy score of 70.5%.

Metrics:

In this subsection, we will discuss how we compare the models. The metrics we have chosen to compare the models against each other are accuracy, precision, recall, and F1 score.

Accuracy:

Accuracy is a measure of the correctness of the model. The formula to calculate accuracy is

$$\text{Accuracy} = \frac{\text{correct predictions}}{\text{total number of predictions}}$$

Correct predictions are the count of instances where the model predicted the correct class, and the total number of predictions is the total number of instances the model made predictions for.

Precision:

Precision is a measure that checks the quality of the model's positive predictions. The formula to calculate precision is

$$\text{Precision} = \frac{\text{True positives}}{\text{True positives} + \text{False positives}}$$

True positives refer to the number of people the model correctly predicts as having autism. False positives refer to the number of people that the model predicts as having autism, but in reality, they do not.

Recall:

Recall is similar to precision in the way that it works. This metric also checks how many of the true cases the model finds. The formula to calculate recall is

$$\text{Recall} = \frac{\text{True positives}}{\text{True positives} + \text{False negatives}}$$

False negatives are the frequency of instances where the model incorrectly predicted an individual as not having autism when, in reality, they do.

F1 score:

There is usually a trade-off between precision and recall: If we wanted a highly precise model, the model would be designed to predict instances it is highly confident about. This would lead the model to miss out on positive instances, leading to a lower recall score. However, if we want the model to have a high recall score, it will predict instances with lower confidence as positive, resulting in false positives and leading to a lower precision score. Which of the two approaches is preferred depends on the context of the problem. We will touch on this aspect in the results and discussion section.

To bypass the trade-off between recall and precision, we also use the F1 score. The F1 score is a combination of the recall plus the precision. The formula to calculate the F1 score is

$$F1 = \frac{2 \times \text{Precision} \times \text{Recall}}{\text{Precision} + \text{Recall}}$$

■ Result and Discussion

In this section, we present our study's empirical evaluation. First, we discuss the impact of the different hyperparameters for each model. Then, using the best scoring model for each approach, we compare all approaches.

Hyperparameters affecting each model:

We will focus only on the accuracy score to perform hyperparameter tuning and identify which hyperparameters lead to each model's best performance.

Upon trying out the different penalty terms for the logistic regression model, we realized that regardless of the penalty term, the model performed the same (each model having an accuracy of 0.798). As such, we used the TensorFlow framework's default penalty term for the next step of our evaluation.

For the decision tree classifier model, we found that the best criterion for this model is the default criterion (having an accuracy of 0.856).

Finding the best number of estimators for the ensemble techniques (random forest) model poses a challenge. Ten estimators (accuracy of 0.864) perform better than 1000 estimators (accuracy of 0.862) but worse than 100 estimators (accuracy of 0.865). Thus, for this model, we will be using 100 estimators.

After testing the different kernels for the support vector machine model, we found that the best kernel was the radial basis function kernel (accuracy of 0.856).

After testing out the different numbers of neighbors for the KNN model, we realized that the model with one neighbor (accuracy of 0.854) had outperformed both the model with ten neighbors (accuracy of 0.834) and the model with 100 neighbors (accuracy of 0.846).

Performance of the models:

With the hyperparameters selected, we compare the different machine learning approaches using all 4 (accuracy, precision, recall, and f1-score) metrics introduced in the methods section. Additionally, we consider the running time of every approach as an additional metric.

Table 2: Accuracy, Precision, Recall, and F1 model scores. Random Forest Classifier has the highest accuracy, recall, and f1 score, while the KNeighbors Classifier has the highest precision score.

	accuracy	precision	recall	f1
Logistic Regression	0.798398	0.830189	0.777778	0.803129
Decision Tree Classifier	0.855808	0.828767	0.916667	0.870504
Random Forest Classifier	0.865154	0.835991	0.926768	0.879042
SVM Classifier	0.855808	0.827273	0.919192	0.870813
KNeighbors Classifier	0.854473	0.886792	0.830808	0.857888
Neural Network	0.849132	0.883469	0.823232	0.852288

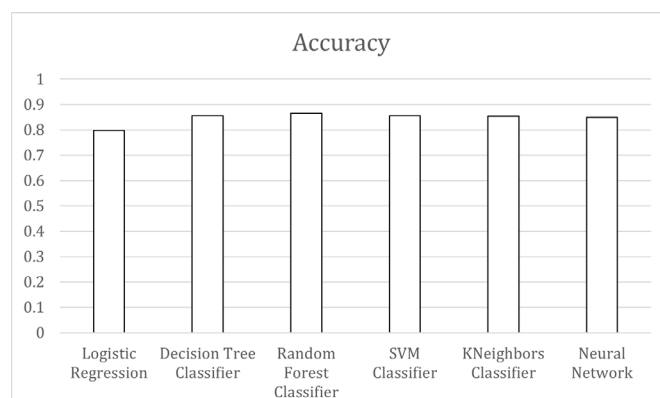


Figure 1: Accuracy vs model name chart. The Random Forest Classifier model has the highest accuracy score, while the Logistic Regression model has the lowest.

Accuracy:

As we can see from Table 2, the model with the highest accuracy score is the Random Forest Classifier Model, which has an accuracy of 0.865 and is used as an ensemble technique. This would mean that the Random Forest Classifier is the best model at making its predictions, with 0.865 being correct. To compare this against other models, the Logistic Regression model has the lowest accuracy with an accuracy of 0.798; the neural network has the second lowest accuracy with an accuracy of 0.849; the KNN Classifier has an accuracy of 0.854; the SVM Classifier and Decision Tree Classifier both have the second highest accuracy with an accuracy of 0.856. Figure 1 allows for easier visualization of this data.

Table 3 consists of models from five papers used for ASD detection, with their accuracy scores, for comparison.

Table 3: Models from existing literature for ASD detection with their accuracy scores.

Paper Name	Models Used	Best Accuracy Achieved
20 The Periodic Risk Evaluation: A new tool to link Medicaid-enrolled autistic adults to services and support	Logistic Regression, Random Forest, XGBoost	75.5%
21 Bag-of-features model for ASD fMRI classification using SVM	SVM	81%
22 A machine learning approach for diagnosing neurological disorders using longitudinal resting-state fmri	SVM	80.76%
23 Combining anatomical and functional networks for neuropathology identification: A case study on autism spectrum disorder	Decision Tree	75%
24 A robust DWT-CNN-based CAD system for early diagnosis of autism using task-based fMRI	Random Forest	72%

Recall vs. Precision:

In Figure 2, we visualize the recall and precision metrics for the different models considered. The trade-off between these two metrics can be seen as the models with the two highest precisions have the 2nd and 3rd lowest recall. The models with the best precision score are the KNN classifier, with a precision score of 0.887, and the Neural Network, with a precision score of 0.883. This means these two models would have predicted fewer positive cases than the other models, resulting in a higher precision score, as seen in Table 2. This would mean that the model only predicts cases it is highly confident about as positive, leading it to miss out on a few true positive cases, resulting in a lower recall score than other models. The KNN classifier has a recall score of 0.831, while the Neural Network has a score of 0.823.

The models with the two highest recall scores are the Random Forest Classifier Model, with a recall score of 0.927, and the Support Vector Machine (SVM) model, with a recall score of 0.919. This means that these two models have predicted a higher number of people with autism who have autism compared to the rest of the models. This is important when it comes to the diagnosis of autism, as it is better to classify one who does not have autism as having autism when compared to diagnosing someone who has autism as not having ASD. Once again, the trade-off is visible as the Random Forest Classifier Model has a precision score of 0.836, making it the model with the third highest precision score, and the Support Vector Machine model has a precision score of 0.827, making it the algorithm with the lowest precision score.

As Table 2 shows, some models have a higher precision score than the others, and others have a higher recall score, making it difficult to determine which model performs best. Thus, we use the F1 score. The model with the highest F1 score is the Random Forest Classifier Model, which suggests it has the best trade-off between precision and recall.

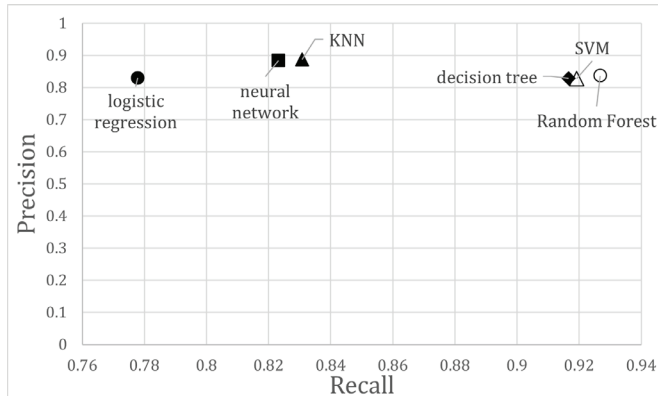


Figure 2: Precision vs Recall Chart. The Random Forest model had the highest recall score, the KNN had the highest precision score, and logistic regression had the lowest precision and recall score.

Running time:

As an additional metric, we consider the running time of the various methods. Running time is a useful proxy for determining an approach's usability. For ASD detection, the method must also be applicable in low-resource settings where only limited access to computing is available. A long-running time might make using the method in these settings impossible.

We found that all approaches, except for the artificial neural network, could produce an output in less than a second. In contrast, the artificial neural network took 20 seconds to run, which is considerably longer than the other models.

Conclusion

This work presents a comparative study of 6 machine-learning methods for automated ASD detection. We compared all six methods based on accuracy, recall, precision, and f1-score. We found different approaches to have the best performance depending on the metric. The Random Forest Classifier was found to have the highest accuracy, recall, and f1-score (in all cases, closely followed by the decision tree classifier and the SVM); in contrast, the KNN classifier was found to have the best precision (closely followed by the artificial neural network). In terms of running time, we found all models to be equally quick, except the neural network, which was orders of magnitude slower.

Based on our presented results, we recommend using a random forest classifier for automated ASD detection. However, some researchers may choose other models that prioritize the precision score over the recall score—which would mean that the model would not be the random forest classifier model—as a follow-up test to check if one has ASD is expensive. Thus, researchers may choose different models for ASD detection based on whether they prefer to minimize false positives or false negatives.

Acknowledgments

Lars Holdijk (Oxford)

References

1. National Institute on Deafness and Other Communication Disorders. *Autism Spectrum Disorder: Communication Problems in Children*. NIDCD. <https://www.nidcd.nih.gov/health/autism-spectrum-disorder-communication-problems-children>. (accessed June 7, 2023)
2. Thabtah, F. *Autism Screening Adult*. UCI Machine Learning Repository. <https://doi.org/10.24432/C5F019> (accessed Oct 10, 2023).
3. Thabtah, F. *Autistic Spectrum Disorder Screening Data for Children*. UCI Machine Learning Repository. <https://doi.org/10.24432/C5659W> (accessed Oct 10, 2023).
4. Thabtah, F. *Autistic Spectrum Disorder Screening Data for Adolescent*. UCI Machine Learning Repository. <https://doi.org/10.24432/C5V89T> (accessed Oct 10, 2023).
5. Goldsborough, P. A Tour of TensorFlow. *arXiv:1610.01178* [cs] 2016.
6. James, G.; Witten, D.; Hastie, T.; Tibshirani, R.; Taylor, J. Linear regression. In *An Introduction to Statistical Learning: With Applications in Python*; Springer: New York, 2023; pp. 69–134
7. LaValley, M. P. Logistic Regression. *Circulation* 2008, 117 (18), 2395–2399. <https://doi.org/10.1161/circulationaha.106.682658>.
8. Tate, A. E.; McCabe, R. C.; Larsson, H.; Lundström, S.; Lichtenstein, P.; Kuja-Halkola, R. Predicting Mental Health Problems in Adolescence Using Machine Learning Techniques. *PLOS ONE* 2020, 15 (4), e0230389. <https://doi.org/10.1371/journal.pone.0230389>.
9. Gao, Y.; Flannery, M. A.; Assan, J. X.; Wu, Y.; Resom, Adonay Zenebe. *Machine Learning for Mental Health Detection*. Digital WPI. https://digital.wpi.edu/concern/student_works/9306t094r?locale=en (accessed 2024-04-27).
10. Patel, H. H.; Prajapati, P. Study and Analysis of Decision Tree Based Classification Algorithms. *International Journal of Computer Sciences and Engineering* 2018, 6 (10), 74–78. <https://doi.org/10.26438/ijcse/v6i10.7478>.
11. Vergyri, D.; Knoth, B.; Shriberg, E.; Mitra, V.; McLaren, M.; Ferrer, L.; Garcia, P.; Marmar, C. Speech-based assessment of PTSD in a military population using diverse feature classes. *www.isca-archive.org*. <https://doi.org/10.21437/Interspeech.2015-740>.
12. Przemyslaw, K.; Lughofer, E.; Bogdan, T. Hybrid and Ensemble Methods in Machine Learning. *Journal of Universal Computer Science* 2013, 19, 457–461.
13. Reece, A. G.; Reagan, A. J.; Lix, K. L. M.; Dodds, P. S.; Danforth, C. M.; Langer, E. J. Forecasting the Onset and Course of Mental Illness with Twitter Data. *Scientific Reports* 2017, 7 (1). <https://doi.org/10.1038/s41598-017-12961-9>.
14. Greenstein, D.; Malley, J. D.; Weisinger, B.; Clasen, L.; Gogtay, N. Using Multivariate Machine Learning Methods and Structural MRI to Classify Childhood Onset Schizophrenia and Healthy Controls. *Frontiers in Psychiatry* 2012, 3. <https://doi.org/10.3389/fpsy.2012.00053>.
15. Cervantes, J.; Garcia-Lamont, F.; Rodríguez-Mazahua, L.; Lopez, A. A Comprehensive Survey on Support Vector Machine Classification: Applications, Challenges and Trends. *Neurocomputing* 2020, 408(1), 189–215. <https://doi.org/10.1016/j.neucom.2019.10.118>.
16. Rangaprakash, D.; Deshpande, G.; Daniel, T. A.; Goodman, A. M.; Robinson, J. L.; Salibi, N.; Katz, J. S.; Denney, T. S.; Dretsch, M. N. Compromised Hippocampus-Striatum Pathway as a Potential Imaging Biomarker of Mild-Traumatic Brain Injury and

- Posttraumatic Stress Disorder. *Human Brain Mapping* 2017, 38 (6), 2843–2864. <https://doi.org/10.1002/hbm.23551>.
17. Gokay Akinci; Polat, E.; Orhan Murat Kocak. A Video Based Eye Detection System for Bipolar Disorder Diagnosis. 2012. <https://doi.org/10.1109/siu.2012.6204617>.
 18. Zhang, Z. Introduction to Machine Learning: K-Nearest Neighbors. *Annals of Translational Medicine* 2016, 4 (11), 218–218. <https://doi.org/10.21037/atm.2016.03.37>.
 19. Madhiarasan, M.; Louzazni, M. Analysis of Artificial Neural Network: Architecture, Types, and Forecasting Applications. *Journal of Electrical and Computer Engineering* 2022, 2022, 1–23. <https://doi.org/10.1155/2022/5416722>.
 20. Shea, L.; Miller, K. K.; Nonnemacher, S.; Becker, A.; Treadway, P.; Alford, A.; Newschaffer, C.; Lee, B. K. The Periodic Risk Evaluation: A New Tool to Link Medicaid-Enrolled Autistic Adults to Services and Support. *Research in Autism Spectrum Disorders* 2022, 98, 102037. <https://doi.org/10.1016/j.rasd.2022.102037>.
 21. Shale Ahammed; Niu, S.; Ahmed, R.; Dong, J.; Gao, X.; Chen, Y. Bag-of-Features Model for ASD fMRI Classification Using SVM. 2021. <https://doi.org/10.1109/acctcs52002.2021.00019>.
 22. Devika K; Murthy, R. A Machine Learning Approach for Diagnosing Neurological Disorders Using Longitudinal Resting-State fMRI. 2021. <https://doi.org/10.1109/confluence51648.2021.9377173>.
 23. Itani, S.; Thanou, D. Combining Anatomical and Functional Networks for Neuropathology Identification: A Case Study on Autism Spectrum Disorder. *Medical Image Analysis* 2021, 69, 101986. <https://doi.org/10.1016/j.media.2021.101986>.
 24. Haweel, R.; Shalaby, A.; Mahmoud, A.; Seada, N.; Ghoniemy, S.; Ghazal, M.; Casanova, M. F.; Barnes, G. N.; El-Baz, A. A Robust DWT–CNN-Based CAD System for Early Diagnosis of Autism Using Task-Based fMRI. *Medical Physics* 2021, 48 (5), 2315–2326. <https://doi.org/10.1002/mp.14692>.

■ Author

Naren is in his junior year at Chirec International School and has been mesmerized by the wonders of technology since childhood. He hopes to build on his knowledge by majoring in computer science, specifically AI/ML, in college.