

# Uma prova de conceito para o diagnóstico diferencial do transtorno do espectro autista usando aprendizado de máquina

## *A proof of concept for the differential diagnosis of autism spectrum disorder using machine learning*

Flávio Secco Fonseca<sup>1</sup> 

Maria Vitória Soares Muniz<sup>2</sup> 

Catarina Victória Nascimento de Oliveira<sup>2</sup> 

Thailson Caetano Valdeci da Silva<sup>2</sup> 

Wellington Pinheiro Dos Santos<sup>1</sup> 

<sup>1</sup>Universidade de Pernambuco (UPE), Programa de Pós-graduação em Engenharia da computação (PPGEC), Recife, PE, Brasil.

<sup>2</sup>Universidade Federal de Pernambuco (UFPE), Departamento de Engenharia Biomédica (DEBM), Recife, PE, Brasil.

## Resumo

**Objetivo:** O Transtorno do Espectro do Autismo (TEA) é um transtorno do neurodesenvolvimento caracterizado por déficits na comunicação social e na interação social. Atualmente, o diagnóstico do TEA é clínico. Sendo assim, este experimento busca apresentar uma nova proposta de ferramenta para diagnóstico diferencial do TEA, dando suporte à decisão do especialista, baseada na utilização de aprendizagem de máquina e sinais de EEG.

**Método/abordagem:** Utilizou-se o dataset 1 da base de dados Sheffield após processamento. Foram extraídos e selecionados atributos para avaliação utilizando 2 métodos: otimização por enxame de partículas e busca evolucionária. Os dois conjuntos foram divididos em treino (80%) e teste (20%), aplicou-se validação cruzada com 10 folds e, em seguida, avaliou-se 12 modelos de classificação distintos. Os experimentos foram repetidos 30 vezes.

**Contribuições teóricas/práticas/sociais:** O modelo com melhores resultados foi o SVM Rbf 0.5, com bons valores de acurácia (99,31%), índice Kappa (0,986), sensibilidade (0,994), especificidade (0,992) e área ROC (0,993). **Originalidade/relevância:** Os resultados sugerem que a aprendizagem de máquina é uma ferramenta eficaz no diagnóstico do TEA.

**Palavras-chave:** Aprendizagem de Máquina; Autismo; Diagnóstico; Eletroencefalograma.

## Abstract

**Purpose:** Autism Spectrum Disorder (ASD) is a neurodevelopmental disorder characterized by deficits in social communication and social interaction. Currently, the diagnosis of ASD is clinical. In this context, this work aims to validate a differential diagnosis proposal, based on the use of machine learning, to support the diagnosis of ASD based on electroencephalography analysis.

**Design/methodology/approach:** Dataset 1 from the Sheffield database was used after data processing. Furthermore, attributes were extracted and selected for evaluation using 2 different methods: particle examination optimization and evolutionary search. The dataset was divided into training (80%) and testing (20%). For each of the databases, 12 classification configurations were evaluated. 10 fold cross-validation was applied. The experiments were repeated 30 times.

**Research, Practical & Social implications:** The most successful model was the SVM Rbf 0.5, reaching good results for accuracy values (99.31%), kappa index (0.986), sensitivity (0.994), specificity (0.992) and AUC (0.993).

**Originality/value:** The results suggest that machine learning is an effective tool in ASD diagnosis..

**Keywords:** Machine Learning; Autism; Diagnosis; Electroencephalogram;

## Introdução

O Transtorno do Espectro do Autismo (TEA) é um transtorno do neurodesenvolvimento caracterizado por déficits na comunicação social e na interação social, além de padrões de comportamento restritos e repetitivos (American Psychiatric Association, 2013). Até o ano de 2013, o Transtorno do Espectro do Autismo era dividido em cinco grupos: Síndrome de Asperger, Síndrome de Rett, Transtorno Desintegrativo da Infância, Transtorno Autista ou Autismo Clássico, e Transtorno Invasivo do Desenvolvimento. Atualmente, é adotado o Diagnostic and Statistical Manual of Mental Disorders (DSM-5) e os subtipos no diagnóstico do TEA foram substituídos por 3 divisões dado o nível de gravidade: (a) nível 1: paciente necessita de suporte; (b) nível 2: paciente necessita de suporte substancial; e (c) nível 3: o paciente precisa de suporte muito substancial (Kulage et al., 2014; Volkmar e Reichow, 2013; Kulage et al., 2020).

No ano de 2022, entrou em vigor a Classificação Estatística Internacional de Doenças e Problemas Relacionados à Saúde, CID-11, vinculada à Organização Mundial de Saúde (OMS) e de uso obrigatório no Sistema Único de Saúde (SUS) para fins de diagnósticos. A CID-11 adota o termo “Transtorno do Espectro Autista (TEA)” seguindo os padrões da DSM-5. Anteriormente, na CID-10, era utilizado o termo “autismo infantil” como parte dos Transtornos Invasivos do

Desenvolvimento (TID) (Dias et al., 2022).

As manifestações clínicas do TEA são muito precoces, sendo evidentes até mesmo antes dos dois anos de idade. No entanto, o quadro clínico inicial mimetiza um atraso psicomotor global (Oliveira, 2009), o que dificulta o diagnóstico ainda nos primeiros anos de vida. Atualmente, o diagnóstico do TEA é essencialmente clínico: equipes de psiquiatras, psicólogos clínicos e neuropsicólogos fazem uso de observações e questionários para identificação do TEA (Heinsfeld et al., 2018; Khodatars et al., 2021). Todavia, o diagnóstico é, muitas vezes, pouco assertivo e tardio. Além disso, não existem tratamentos específicos para o TEA, mas diversas abordagens terapêuticas vêm sendo desenvolvidas para minimizar os sintomas e melhorar as capacidades cognitivas e comunicativas, as habilidades sociais e a qualidade de vida do paciente. Nessa perspectiva, é evidente a necessidade do desenvolvimento de tecnologias que apoiem e contribuam para o diagnóstico diferencial do TEA, tendo por objetivo principal o diagnóstico precoce, preciso e democrático.

Por esse motivo, diversas pesquisas têm sido feitas para gerar diagnósticos diferenciais baseados na utilização de sinais eletroencefalográficos (EEG), avaliando a atividade cerebral. Ademais, técnicas de aprendizado de máquina vêm sendo investigadas para construir

ferramentas de apoio ao diagnóstico (Heinsfeld et al., 2018; Khodatars et al., 2021). Neste contexto, este trabalho busca validar uma nova proposta de ferramenta para diagnóstico diferencial, de suporte à decisão do especialista, do Transtorno do Espectro Autista (TEA), baseada na utilização de aprendizagem de máquina e aprendizado estatístico, a partir da análise de sinais de eletroencefalografia (EEG) rotulados previamente por diagnóstico clínico.

## Fundamentação Teórica

Nesta seção foram apresentados alguns trabalhos de destaque envolvendo tanto o diagnóstico do autismo, quanto a utilização do EEG e do aprendizado de máquinas.

### Trabalhos Relacionados

O trabalho precursor “EEG complexity as a biomarker for autism spectrum disorder risk” mostra que mesmo antes de aparecerem distúrbios comportamentais, mudanças muito sutis na atividade cerebral podem ser percebidas através de sinais de EEG. Para desvendar a arquitetura dos cérebros em questão, os pesquisadores propõem a utilização da entropia multiescala modificada (mMSE) como biomarcador na diferenciação de cérebros com um desenvolvimento normal e de pessoas que estão no espectro do autismo. Os resultados foram bastante promissores, atingindo valores superiores a 80% de acurácia (Bols et al., 2011).

Partindo do pressuposto de que sinais de EEG e aprendizado de máquinas já são utilizados no auxílio ao diagnóstico de epilepsia, doença de Alzheimer dentre outros distúrbios neurológicos, Mohi-ud Din e Jayanthi (2021) apresentam em seu trabalho uma forma de manipular este mesmo sinal para a detecção do TEA. Com o auxílio da aprendizagem por transferência, os autores utilizaram as redes profundas pré-treinadas GoogLeNet e SqueezeNet na classificação de indivíduos com e sem TEA. Os algoritmos utilizaram como atributos de entradas representações visuais da transformada de wavelet do sinal EEG e atingiram 75% e 82% de acurácia com a GoogLeNet e SqueezeNet, respectivamente.

No artigo de Grossi, et al. (2021), os pesquisadores projetam uma otimização das análises de EEG para o diagnóstico do TEA. Com a utilização de apenas 2 canais na coleta do EEG, os pesquisadores conseguiram apresentar um método robusto o suficiente para um bom desempenho do modelo. O método passou inicialmente por uma extração de 1588 atributos quantitativos, para depois serem selecionados os mais significativos pelo algoritmo evolucionário chamado Gen-D aliado a uma rede neural. Após esse pré-processamento, a rede neural para classificação atingiu 100% de acurácia para os casos de autismo.

Com o objetivo de auxiliar médicos no diagnóstico do transtorno do espectro autista, o estudo de Ari, et

al. (2022) desenvolveu um novo método automatizado para detectar o TEA através de gravações de EEG. O método utiliza o algoritmo Douglas-Peucker (DP) para reduzir o número de amostras do sinal EEG, seguido pela extração de atributos usando a transformada de wavelet e codificação esparsa. Redes neurais convolucionais profundas (CNNs) são usadas para classificar os sinais, com Autoencoders (AE) baseados em Extreme Learning Machines (ELM) para aumentar os dados. Os resultados mostraram alta precisão (98,88%), sensibilidade (100%) e especificidade (96,4%) na detecção automática de TEA, tornando o modelo promissor para aplicações clínicas e sujeito a mais testes para validação.

Abdolzadegan, et al. (2020) apresenta um método para diagnóstico precoce do TEA utilizando sinais de EEG. O estudo envolveu 45 voluntários. O estudo utiliza métodos como Power Spectrum, Wavelet Transform, FFT, Fractal Dimension, Correlation Dimension, Lyapunov Exponent, Entropy, Detrended Fluctuation Analysis e Synchronization Likelihood. O uso de clustering baseado em densidade e seleção de características melhora a robustez e remove artefatos. Os modelos KNN e SVM foram utilizados para a classificação, alcançando os resultados: precisão de 90,57% (SVM) e 72,77% (KNN); e sensibilidade de 99,91% (SVM) e 91,96% (KNN).

O artigo de Chawla, et al. (2023) propõe uma nova abordagem para o

diagnóstico automatizado do TEA através de sinais de EEG usando a transformada de wavelet analítica flexível (FAWT). O método utiliza a rede neural convolucional (CNN) para obter precisão de 99,19%, sensibilidade de 99,34%, especificidade de 99,21% e área sob a curva de 0,9997 na identificação de TEA em segmentos de EEG de 10 segundos, tornando-se uma ferramenta promissora para auxiliar neurologistas no diagnóstico do autismo.

Alves, Caroline L., et al. (2023) propõe um método de diagnóstico automático preciso para o transtorno do espectro do autismo usando imagens cerebrais funcionais e análise de rede. Ele supera resultados anteriores, identificando diferenças na conectividade cerebral, especialmente na região do córtex cingulado posterior ventral esquerdo, e oferece interpretabilidade médica. O método é aplicável a conjuntos de dados menores e pode fornecer insights valiosos sobre as redes cerebrais em pacientes com autismo.

Com uma base de dados envolvendo diversas faixas etárias, o trabalho de Farooq, Muhammad Shoaib, et al. (2023) utiliza Aprendizagem Federada (FL) para diagnóstico de Transtorno do Espectro do Autismo (TEA), treinando dois classificadores de Machine Learning. Os resultados são enviados a um servidor central, onde um meta classificador determina a abordagem mais precisa. Utilizando quatro

conjuntos de dados, o modelo prevê TEA com 98% de precisão em crianças e 81% em adultos, destacando a utilidade da FL no diagnóstico precoce do TEA.

Por fim, o trabalho de Mellema, Cooper J., et al destaca a prevalência e a escassez de especialistas para diagnóstico preciso do TEA. Além do mais, foram avaliados 12 modelos de aprendizado de máquina com recursos de ressonância magnética funcional e estrutural (fMRI e sMRI). Os modelos de aprendizagem profunda demonstraram a maior precisão diagnóstica e identificaram biomarcadores cerebrais associados ao TEA, incluindo novos no cerebelo. Esses resultados sugerem o potencial de desenvolver biomarcadores generalizáveis para uma compreensão mais profunda das alterações neuronais no TEA.

Esta seção deixa claro o estado atual do aprendizado de máquinas para o apoio ao diagnóstico do TEA. Mesmo em fases de validação, são abordagens que trazem resultados promissores e comprovam a acurácia dos sistemas. Um ponto importante a se salientar é que a maior parte dos trabalhos não foca no diagnóstico precoce, tendo em vista que a maioria das bases de dados envolveram pessoas adultas. Além do mais, foi observado que apenas 2 dos trabalhos abordam a questão da simplificação do sistema de coleta de

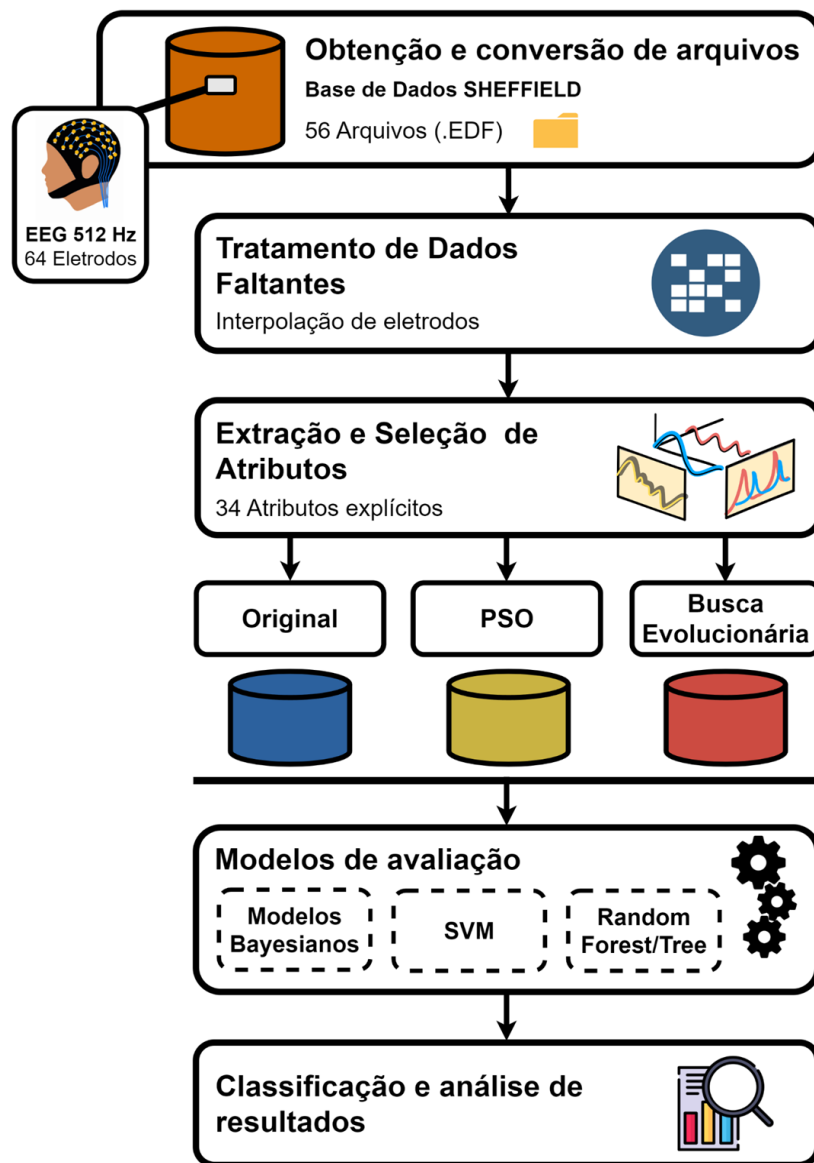
dados, seja por EEG ou por Imagem de Ressonância Magnética. Tendo em mente estas lacunas, este trabalho servirá como abordagem exploratória de um projeto maior, que tem como foco a criação de uma base de dados própria com crianças, e o desenvolvimento de um sistema de eletrodos de baixa densidade para apoio ao diagnóstico do TEA.

## Metodologia

Nesta seção, serão abordados os materiais e métodos empregados no desenvolvimento deste trabalho, incluindo as particularidades da base de dados selecionada, o tratamento adotado para lidar com dados faltantes e os classificadores testados, bem como as métricas escolhidas para avaliá-los. A Figura 1 ilustra o fluxo metodológico aqui seguido.

## Base de dados

A priori, foi feito um levantamento de bases de dados públicas disponíveis. Foi escolhido o dataset 1 da base de dados Sheffield, que contém sinais eletroencefalográficos (EEG) de 56 indivíduos. Sabe-se que os sinais de EEG foram adquiridos através do sistema Biosemi Active Two EEG e que a idade dos participantes é de 18 a 68 anos. O paradigma que gerou os dados



**Figura 1.** Diagrama do fluxo metodológico aplicado no presente trabalho  
Fonte: Autoria própria (2023).

foi um período de 150 segundos de olhos fechados em repouso e de momentos de estímulo visual. Os dados foram filtrados com o filtro passa banda de 0.01 - 140 Hz e o canal de referência utilizado foi o Cz. Posteriormente, os dados foram disponibilizados com frequência de amostragem reduzida de 512 Hz (Dickinson et al., 2022).

## Tratamento dos dados

Os arquivos adquiridos foram convertidos para edf e csv, formatos mais apropriados para o tratamento. Em seguida, foi feito um levantamento dos canais faltantes e foi criado um *script* usando a linguagem de programação Python 3.11.1 para o tratamento desses dados segundo a

sequência: 1) Identificaram-se os canais faltantes; 2) Adicionaram-se as colunas dos canais faltantes e do eletrodo de referência; 3) Todas as colunas foram ordenadas de acordo com a sequência pré-estabelecida; 4) Leu-se uma tabela pré-montada que contém os vizinhos mais próximos de cada canal; 5) Para cada canal faltante, foi identificado se seus vizinhos estavam presentes; 6) Quantificaram-se os vizinhos presentes; 7) O canal que possuía pelo menos 2 vizinhos presentes foi interpolado.

A interpolação dos canais ausentes foi realizada mediante o cálculo da média aritmética dos valores dos sinais dos vizinhos presentes. Os passos V, VI e VII foram repetidos 4 vezes, dado que a cada interpolação surgem mais vizinhos presentes para os canais que não puderam ser interpolados à princípio. Finalmente, obteve-se a primeira versão da base de dados tratada.

Com auxílio do software GNU Octave 8.2.0, o sinal de eletroencefalografia foi janelado e foram extraídos atributos relevantes a partir dos dados conhecidos. De cada janela, foram extraídos 34 atributos explícitos ao sinal, tais como valor médio, amplitude máxima e frequência média. As janelas utilizadas foram de 2s com sobreposição de 0,5s e frequência de amostragem 512Hz. Posteriormente, foi feita a seleção dos atributos mais relevantes utilizando os métodos de otimização por enxame de partículas (PSO) e busca evolucionária. Dessa

forma, foram gerados três arquivos finais: base de dados original, base de dados otimizada com PSO e base de dados otimizada com busca evolucionária, todas tratadas.

## Classificação e avaliações

Após a obtenção das 3 bases de dados distintas a partir de cada um dos métodos detalhados anteriormente, foram treinados, validados e testados diferentes classificadores: Bayes Net, Naive Bayes, Random Tree, Random Forest (10, 100 e 500 árvores), Máquina de Vetor de Suporte (Kernel Linear, Polinomial de Grau 2, Polinomial de Grau 3, Rbf de Gama 0,1, Gama 0,2 e Gama 0,3).

As Redes Bayesianas (Bayesian Networks - BN), também conhecidas como redes de crença, pertencem à família de modelos gráficos probabilísticos (Ben-gal, 2008). Esse tipo de rede pode ser definida como um grafo direcionado acíclico, no qual os nós correspondem às variáveis do domínio e as arestas correspondem às dependências probabilísticas diretas entre elas (Friedman et al., 2003). A força da correlação entre essas variáveis é definida nas Tabelas de Probabilidade Condicional (CPTs) anexadas a cada nó, as quais especificam o grau de crença do nó em um estado específico, dado os estados dos nós pais (Chen, 2012).

O algoritmo Naive Bayes é derivado do teorema de Bayes (Vijay et al., 2023). Ele é a forma mais simples de rede bayesiana, na qual todos os atributos são independentes, dado o

valor da variável de classe, caracterizando o que é conhecido como independência condicional (Zhang, 2004). O Classificador Naive Bayes tem a vantagem de exigir uma pequena quantidade de dados de treinamento para estimar os parâmetros (média e variância das variáveis) necessários para a classificação (Aditya et al., 2022).

O Random Tree é um modelo de árvore de decisão construído a partir de um subconjunto aleatório de atributos (Niranjan et al., 2018). Nesse contexto, uma árvore é selecionada aleatoriamente a partir de um conjunto de árvores possíveis, garantindo que cada árvore tenha uma probabilidade equivalente de ser escolhida para a análise (Shubho et al., 2019). Em outras palavras, a Random Tree é um algoritmo de aprendizado de máquina supervisionado que combina dois algoritmos: uma árvore de modelo único e o conceito de Random Forest (Hermawan et al., 2021).

Por outro lado, Random Forest é um algoritmo de aprendizado supervisionado que se baseia na geração de múltiplas árvores de decisão, selecionando amostras e características aleatórias (Zhou et al., 2020). O propósito é utilizar essas árvores como uma unidade para alcançar previsões mais precisas. Cada árvore vota em uma classe para cada observação, sendo a classe mais votada considerada como a previsão do modelo (Briouza et al., 2022).

As Máquinas de Vetores de Suporte (SVMs) foram desenvolvidas

por Vapnik como uma implementação do SRM (Burbidge et al., 2001). Essa abordagem pode ser descrita como uma ferramenta de predição que visa encontrar um hiperplano específico, uma linha ou limite de decisão, capaz de separar eficientemente conjuntos de dados ou classes, minimizando o problema de sobreajuste de dados (Somvanshi et al., 2016). As SVMs têm demonstrado ampla eficácia em aplicações de aprendizado de máquina envolvendo conjuntos de dados volumosos e de alta dimensionalidade (Ruping, 2001).

Utilizando a biblioteca Weka 3.9.6, cada conjunto de dados foi dividido em treino (80%) e teste (20%). Além disso, os modelos e configurações supracitados foram aplicados. Foram feitas 30 repetições por experimento e validação cruzada de 10 folds para uma maior segurança estatística. Como resultado desse processo, foram obtidas a acurácia, o índice Kappa, a sensibilidade, a especificidade e a área sob a curva ROC para análise do desempenho de cada classificador.

A acurácia é uma das métricas mais simples e comuns usadas para avaliar a performance de modelos de classificação. Ela fornece uma medida geral da taxa de acertos do modelo, ou seja, a proporção de previsões corretas em relação ao total de previsões feitas pelo modelo (Srivastava, 2014). A acurácia é especialmente útil quando as classes estão balanceadas, ou seja, têm aproximadamente o mesmo número de amostras em cada classe. Nesses casos,

a acurácia fornece uma boa indicação do desempenho geral do modelo. Para calcular a acurácia se usa a seguinte fórmula:

$$A = \frac{\text{VerdadeirosPositivos} + \text{VerdadeirosNegativos}}{\text{TotaldeAmostras}}$$

Índice Kappa, também conhecido como Kappa de Cohen, é uma métrica de avaliação de concordância usada para medir o grau de acordo entre as previsões de um modelo de classificação e as classes reais dos dados (Merlini e Rossini, 2022). Diferentemente da acurácia, o Kappa leva em consideração o acordo que pode ser esperado ao acaso, o que o torna especialmente útil quando se lida com classes desbalanceadas. O índice Kappa é calculado comparando as frequências observadas de concordância entre o modelo e os dados reais com as frequências que seriam esperadas ao acaso. Ele fornece um valor que varia de -1 a 1. A fórmula para calcular esse índice é:

$$Kappa = \frac{(Po - Pe)}{(1 - Pe)}$$

Onde Po é a proporção de concordância observada entre o modelo e os dados reais. É calculada como a soma dos verdadeiros positivos e verdadeiros negativos dividida pelo número total de amostras. Pe é a proporção de concordância esperada ao acaso. É calculada como a multiplicação das proporções marginais correspondentes das previsões do modelo e das classes reais, somadas e

divididas pelo número total de amostras ao quadrado.

A sensibilidade, também conhecida como recall ou taxa de verdadeiros positivos, é uma métrica de avaliação de modelos de classificação que mede a capacidade do modelo de identificar corretamente as amostras positivas (classe positiva) em relação ao total de amostras reais que pertencem à classe positiva (Mun e Jumadi, 2020). A fórmula para calcular a sensibilidade é a seguinte:

$$Sen = \frac{\text{VerdadeirosPositivos}}{\text{VerdadeirosPositivos} + \text{FalsosNegativos}}$$

Especificidade é outra métrica de avaliação usada em modelos de classificação. Ao contrário da sensibilidade, que se concentra nas amostras positivas, a especificidade avalia a capacidade do modelo de identificar corretamente as amostras negativas (classe negativa) em relação ao total de amostras reais que pertencem à classe negativa (Abuhaija et al., 2023). A fórmula para calcular a especificidade é a seguinte:

$$Esp = \frac{\text{VerdadeirosNegativos}}{\text{VerdadeirosNegativos} + \text{FalsosPositivos}}$$

A área sob a curva ROC (Receiver Operating Characteristic) ou AUC-ROC é uma métrica de avaliação usada em modelos de classificação binária para medir a qualidade geral do modelo e sua capacidade de distinguir entre as duas classes (positiva e negativa) (Smelser, 2001). A Curva ROC é uma representação gráfica que mostra a taxa de verdadeiros positivos (sensibilidade) em função da taxa de

falsos positivos (1 - especificidade) para diferentes pontos de corte de decisão do modelo. Cada ponto na curva ROC corresponde a um limiar diferente usado para classificar as amostras como positivas ou negativas. O ponto (0,0) representa o ponto de corte em que todas as amostras são classificadas como negativas, enquanto o ponto (1,1) representa o ponto de corte em que todas as amostras são classificadas como positivas. A AUC-ROC mede a área sob essa curva ROC e varia de 0 a 1. Quanto maior o valor da AUC-ROC, melhor é o desempenho do modelo.

Por fim, as métricas supracitadas foram avaliadas e cada modelo pôde ser comparado para escolha do melhor classificador e configuração.

## Resultados e Discussão

Dando prosseguimento à análise, são apresentados os resultados obtidos por cada uma das três abordagens adotadas. Inicialmente, foram utilizados os dados da base original e completa para realizar os experimentos. Em seguida, aplicamos as técnicas de seleção de atributos por PSO e por Busca Evolucionária, buscando otimizar ainda mais o desempenho dos modelos e reduzir a complexidade do problema. Ao final

desta seção, as melhores configurações encontradas são comparadas durante a fase de testes.

### Base de dados original

Os resultados dos experimentos das melhores configurações para cada grupo de classificadores, considerando a utilização da base completa sem seleção de atributos, são apresentados na Tabela 1. Observa-se que o desempenho dos classificadores variou significativamente de acordo com a técnica adotada. O modelo Naive Bayes registrou o pior desempenho dentre todos os classificadores avaliados, seguido pelo *Bayes Net* e *Random Tree*. Tanto as taxas de acurácia quanto das demais métricas, durante o treinamento e validação, ficaram abaixo dos demais modelos.

Por outro lado, os modelos Random Forest com 10, 100 e 500 árvores se destacaram positivamente. Durante o processo de treinamento e validação, esses modelos demonstraram acurácia superior a 95%, reforçando sua robustez e precisão na classificação dos dados. Além da acurácia, outras métricas de avaliação também seguiram o mesmo padrão de desempenho para os modelos baseados na floresta aleatória.

| Classificador |         | Acurácia     | Kappa         | Sensibilidade | Especificidade | Área ROC      |
|---------------|---------|--------------|---------------|---------------|----------------|---------------|
| Bayes Net     |         | 72,83 ± 2,23 | 0,457 ± 0,044 | 0,765 ± 0,030 | 0,692 ± 0,030  | 0,785 ± 0,022 |
| Naive Bayes   |         | 61,1 ± 1,82  | 0,213 ± 0,036 | 0,291 ± 0,027 | 0,919 ± 0,020  | 0,671 ± 0,032 |
| Random Tree   |         | 85,95 ± 1,85 | 0,719 ± 0,037 | 0,858 ± 0,026 | 0,860 ± 0,025  | 0,859 ± 0,018 |
| Random Forest | 10      | 95,54 ± 0,94 | 0,910 ± 0,018 | 0,970 ± 0,011 | 0,937 ± 0,015  | 0,991 ± 0,003 |
|               | 100     | 98,52 ± 0,59 | 0,970 ± 0,011 | 0,988 ± 0,001 | 0,982 ± 0,009  | 0,998 ± 0,001 |
|               | 500     | 98,68 ± 0,53 | 0,973 ± 0,010 | 0,989 ± 0,007 | 0,983 ± 0,008  | 0,998 ± 0,001 |
| SVM           | Linear  | 98,09 ± 0,62 | 0,961 ± 0,012 | 0,981 ± 0,008 | 0,980 ± 0,009  | 0,980 ± 0,006 |
|               | Poli 2  | 98,98 ± 0,45 | 0,979 ± 0,009 | 0,989 ± 0,006 | 0,990 ± 0,006  | 0,989 ± 0,004 |
|               | Poli 3  | 99,05 ± 0,47 | 0,981 ± 0,009 | 0,989 ± 0,007 | 0,991 ± 0,006  | 0,990 ± 0,004 |
| SVM           | RBF 0,5 | 96,22 ± 0,98 | 0,924 ± 0,019 | 0,996 ± 0,003 | 0,928 ± 0,018  | 0,962 ± 0,009 |
|               | RBF 0,2 | 98,77 ± 0,59 | 0,975 ± 0,011 | 0,995 ± 0,995 | 0,980 ± 0,010  | 0,987 ± 0,005 |
|               | RBF 0,1 | 99,30 ± 0,40 | 0,986 ± 0,008 | 0,994 ± 0,004 | 0,991 ± 0,006  | 0,993 ± 0,003 |

**Tabela 1.** Resultados de treinamento e validação da base de dados completa.

Fonte: Dados da pesquisa, 2023.

As demais métricas, como índice Kappa, sensibilidade, especificidade e área sob a curva ROC, obtiveram valores superiores a 91%, o que fortalece a confiança na eficácia desses modelos em aplicações práticas.

Observa-se que com exceção da Árvore de Decisão Aleatória, os demais classificadores apresentaram resultados bastante similares. Nas avaliações das Florestas Aleatórias, a configuração com 500 árvores se destacou, exibindo um desempenho ligeiramente superior em relação às outras configurações. Essa diferença sutil pode indicar que um maior número de árvores contribui para um melhor poder de generalização do modelo. Da mesma forma, o SVM polinomial de grau 2 e o SVM com Kernel Rbf 0,1 sobressaíram-se em seus respectivos grupos de classificadores. Essa superioridade sugere que as características particulares desses

métodos são bem adequadas ao conjunto de dados em questão, possibilitando um aprendizado mais eficiente e representativo dos padrões presentes nos dados.

## Base de dados com PSO

Após a aplicação do PSO, houve uma notável redução na base original, passando de 2176 atributos (34 atributos explícitos x 64 canais) para apenas 494 atributos. Essa significativa diminuição, no entanto, não impactou o desempenho dos classificadores, uma vez que eles mantiveram uma performance semelhante à obtida anteriormente. A Tabela 2 apresenta um resumo das melhores configurações encontradas dentre os modelos mais simples (Bayes Net, Naive Bayes e Random Tree), bem como árvores de decisão aleatórias, SVM polinomiais e de Kernel Rbf.

| Classificador |                | Acurácia             | Kappa                | Sensibilidade        | Especificidade       | Área ROC             |
|---------------|----------------|----------------------|----------------------|----------------------|----------------------|----------------------|
| Bayes Net     |                | 74,50 ± 1,88%        | 0,490 ± 0,037        | 0,786 ± 0,024        | 0,705 ± 0,031        | 0,831 ± 0,019        |
| Naive Bayes   |                | 62,40 ± 1,72%        | 0,239 ± 0,035        | 0,309 ± 0,029        | 0,927 ± 0,016        | 0,773 ± 0,021        |
| Random Tree   |                | 86,33 ± 1,78%        | 0,726 ± 0,036        | 0,862 ± 0,024        | 0,864 ± 0,026        | 0,863 ± 0,018        |
| Random Forest | 10             | 95,52 ± 1,01%        | 0,910 ± 0,020        | 0,974 ± 0,010        | 0,936 ± 0,017        | 0,991 ± 0,003        |
|               | 100            | 98,29 ± 0,62%        | 0,966 ± 0,012        | 0,987 ± 0,008        | 0,979 ± 0,009        | 0,998 ± 0,001        |
|               | 500            | 98,46 ± 0,59%        | 0,969 ± 0,012        | 0,988 ± 0,007        | 0,981 ± 0,009        | 0,998 ± 0,001        |
| SVM           | Linear         | 94,75 ± 1,05%        | 0,895 ± 0,021        | 0,940 ± 0,015        | 0,954 ± 0,014        | 0,947 ± 0,010        |
|               | Poli 2         | 98,83 ± 0,48%        | 0,976 ± 0,010        | 0,986 ± 0,007        | 0,990 ± 0,006        | 0,988 ± 0,005        |
|               | Poli 3         | 98,90 ± 0,49%        | 0,978 ± 0,010        | 0,987 ± 0,007        | 0,991 ± 0,006        | 0,989 ± 0,005        |
| SVM           | <b>RBF 0,5</b> | <b>99,19 ± 0,40%</b> | <b>0,984 ± 0,008</b> | <b>0,993 ± 0,005</b> | <b>0,990 ± 0,006</b> | <b>0,992 ± 0,004</b> |
|               | RBF 0,2        | 99,05 ± 0,44%        | 0,980 ± 0,009        | 0,990 ± 0,007        | 0,991 ± 0,006        | 0,990 ± 0,004        |
|               | RBF 0,1        | 98,51 ± 0,54%        | 0,970 ± 0,011        | 0,984 ± 0,008        | 0,986 ± 0,008        | 0,985 ± 0,005        |

**Tabela 2.** Resultados de treinamento e validação da base após o PSO.

Fonte: Dados da pesquisa, 2023.

É possível observar que a configuração SVM com Kernel Rbf 0,5 obteve melhor desempenho tanto nos níveis de acurácia (99,19 % e desvio padrão 0,4%), quanto no índice Kappa (98,4% e desvio padrão 0,8%), demonstrando a alta confiabilidade do modelo.

A performance semelhante dos classificadores, mesmo após a redução dos atributos, reforça a eficácia do PSO como uma técnica de seleção de atributos. Essa abordagem permite que o modelo se concentre apenas nas características mais relevantes para a tarefa de classificação, eliminando a redundância e o ruído desnecessário que poderiam comprometer a precisão do modelo.

## Base de dados com busca evolucionária

Na última abordagem, a base de dados original foi submetida a uma seleção de atributos por busca evolucionária, que identificou um total de 66 atributos, dentre os 2176, como sendo os mais significativos para a classificação do problema. A Tabela 3 traz os resultados dos melhores modelos encontrados nesta abordagem.

Analisando o comportamento geral das mesmas configurações, observa-se uma forte semelhança aos resultados após a seleção por PSO. Entretanto, desta vez os modelos SVM Polinomial 3 e SVM Rbf 0,5 tiveram uma perda mais significativa de acurácia e Kappa. O classificador Random Forest com 500 árvores teve o melhor desempenho com uma acurácia de 97,76%.

| Classificador |         | Acurácia             | Kappa                | Sensibilidade        | Especificidade       | Área ROC             |
|---------------|---------|----------------------|----------------------|----------------------|----------------------|----------------------|
| Bayes Net     |         | 71,24 ± 2,10%        | 0,426 ± 0,041        | 0,769 ± 0,029        | 0,657 ± 0,030        | 0,796 ± 0,022        |
| Naive Bayes   |         | 58,02 ± 1,79%        | 0,150 ± 0,036        | 0,235 ± 0,028        | 0,912 ± 0,020        | 0,734 ± 0,024        |
| Random Tree   |         | 85,13 ± 2,03%        | 0,702 ± 0,040        | 0,847 ± 0,028        | 0,855 ± 0,026        | 0,851 ± 0,020        |
| Random Forest | 10      | 94,07 ± 1,23%        | 0,881 ± 0,024        | 0,962 ± 0,013        | 0,919 ± 0,019        | 0,986 ± 0,004        |
|               | 100     | 97,34 ± 0,75%        | 0,948 ± 0,015        | 0,975 ± 0,010        | 0,972 ± 0,010        | 0,996 ± 0,001        |
|               | 500     | <b>97,76 ± 0,69%</b> | <b>0,955 ± 0,013</b> | <b>0,977 ± 0,009</b> | <b>0,977 ± 0,009</b> | <b>0,997 ± 0,001</b> |
| SVM           | Linear  | 77,73 ± 1,82%        | 0,553 ± 0,036        | 0,724 ± 0,029        | 0,828 ± 0,023        | 0,776 ± 0,018        |
|               | Poli 2  | 92,26 ± 1,18%        | 0,845 ± 0,023        | 0,913 ± 0,019        | 0,931 ± 0,015        | 0,922 ± 0,011        |
|               | Poli 3  | 96,28 ± 0,84%        | 0,925 ± 0,016        | 0,958 ± 0,012        | 0,967 ± 0,011        | 0,962 ± 0,008        |
| SVM           | RBF 0,5 | 91,24 ± 1,26%        | 0,824 ± 0,025        | 0,895 ± 0,021        | 0,928 ± 0,015        | 0,912 ± 0,012        |
|               | RBF 0,2 | 82,57 ± 1,69%        | 0,650 ± 0,033        | 0,776 ± 0,028        | 0,873 ± 0,023        | 0,824 ± 0,016        |
|               | RBF 0,1 | 77,02 ± 1,79%        | 0,539 ± 0,035        | 0,705 ± 0,028        | 0,832 ± 0,025        | 0,769 ± 0,017        |

**Tabela 3.** Resultados de treinamento e validação após busca evolucionária.

Fonte: Dados da pesquisa, 2023.

## Fase de testes

A fase de testes foi executada para cada modelo uma única vez com 20% dos dados, separados da base original. A Tabela 4 traz o teste das melhores configurações encontradas durante o treinamento em cada uma das três abordagens.

Nesta etapa, é possível observar a perda significativa de performance do modelo SVM de Kernel Rbf 0,1 com relação a etapa de treinamento, com cerca de 9% a menos de acurácia e 17%

para o índice Kappa. Dentre os três, destaca-se a configuração SVM com Kernel Rbf 0,5, utilizado em conjunto com a seleção de atributos por PSO. Além de obter maior acurácia, essa abordagem obteve as melhores taxas de sensibilidade (verdadeiro positivo), ou seja, a classificação correta de indivíduos no espectro autista, de 99,4%. E especificidade (verdadeiros negativos), classificação correta de indivíduos de controle com 99,2% de acerto.

| Métricas       | SVM Rbf 0,1 | SVM Rbf 0,5  | Rdm Forest 500 |
|----------------|-------------|--------------|----------------|
| Acurácia (%)   | 90,78       | <b>99,31</b> | 93,04          |
| Kappa          | 0,815       | <b>0,986</b> | 0,860          |
| Sensibilidade  | 0,834       | <b>0,994</b> | 0,936          |
| Especificidade | 0,979       | <b>0,992</b> | 0,924          |
| Área ROC       | 0,906       | <b>0,993</b> | 0,977          |

**Tabela 4.** Resultados da fase de testes dos melhores modelos.

Fonte: Dados da pesquisa, 2023.

A matriz de confusão da Figura 2 exibe a distribuição dos registros dos

modelos testados em termos de suas classes atuais e de suas classes

previstas. Na Figura 2, a) SVM com Kernel Rbf de  $\gamma=0,1$ , na base original. b) SVM com Kernel Rbf de  $\gamma=0,5$ , após PSO. c) Random Forest com 500 árvores, após busca evolucionária.

Na melhor abordagem, com PSO e SVM de kernel Rbf 0,5, ao todo, 2846 janelas de sinais no espectro autista foram previstas corretamente, ou seja,

verdadeiros positivos. Da mesma forma, 2945 janelas de sinais do grupo de controle foram classificadas de forma correta. A baixa taxa de falsos positivos e negativos reforça ainda mais a robustez dos modelos aqui testados. As outras duas abordagens obtiveram resultados interessantes, embora ainda assim inferiores se analisados de maneira combinada.

|      |     | Previsto |      | Previsto |      | Previsto |      |
|------|-----|----------|------|----------|------|----------|------|
|      |     | TEA      | C    | TEA      | C    | TEA      | C    |
| Real | TEA | 2812     | 57   | 2846     | 23   | 2639     | 230  |
|      | C   | 504      | 2458 | 17       | 2945 | 207      | 2755 |
|      |     | a)       |      | b)       |      | c)       |      |

**Figura 2.** Matriz de confusão do melhor modelo encontrado nas três abordagens: a) SVM com Kernel Rbf de  $\gamma=0,1$ , na base original. b) SVM com Kernel Rbf de  $\gamma=0,5$ , após PSO. c) Random Forest com 500 árvores, após busca evolucionária. Fonte: Dados da pesquisa, 2023.

## Considerações Finais

Baseando-se no contexto e na problemática apresentadas, é evidente a necessidade do desenvolvimento de tecnologias que apoiam o diagnóstico do Transtorno do Espectro do Autismo (TEA). A partir dos resultados obtidos, concluímos e destacamos que a inteligência artificial é uma ferramenta com grande potencial de aplicação na área da saúde, essencialmente como ferramenta para o diagnóstico diferencial do TEA e de apoio aos profissionais da área.

Sendo assim, os resultados alcançados sugerem que é sim possível propor um novo método diagnóstico,

baseado em aprendizagem de máquina e aprendizado estatístico, para apoio ao diagnóstico do TEA a partir da análise de sinais de eletroencefalografia (EEG). Embora todos os classificadores tenham obtido boa acurácia, o modelo com melhores resultados foi o SVM Rbf 0,5, com valores promissores de acurácia (99,31%), índice kappa (0,986), sensibilidade (0,994), especificidade (0,992) e área ROC (0,993). Além do mais, com base nos resultados obtidos, sugere-se que a aplicação de técnicas de seleção de atributos, como o PSO e a Busca Evolucionária, preenchem uma lacuna importante no processo de desenvolvimento de classificadores. A

abordagem comparativa entre os melhores modelos evidencia a relevância dessas técnicas na obtenção de resultados semelhantes à base sem seleção, além de simplificar o problema, reduzindo sua dimensionalidade. Tendo em vista que muitos dos trabalhos na literatura buscam a classificação e o reconhecimento de padrões nos sinais de EEG com sistemas de alta densidade, selecionar canais e atributos explícitos mais significativos é uma forma de democratizar e facilitar a aplicação em larga escala, principalmente para sistemas com poucos recursos.

Não obstante, é válido ressaltar que este trabalho é o experimento exploratório de um projeto maior: propõe-se que, uma vez validada a hipótese inicial através destas abordagens, sejam reproduzidos os experimentos utilizando dados de pacientes locais (faixa etária de 2 a 5 anos). Mesmo sabendo-se que, a princípio, existem diferenças entre o cérebro da criança e do adulto, esse estudo serviu como norteador para os futuros modelos que serão criados com a base de dados própria. Por fim, pretende-se obter um sistema inteligente, de baixa densidade de eletrodos, capaz de apoiar tanto o diagnóstico quanto as terapias do TEA. Além disso, durante o desenvolvimento dessa pesquisa, foram encontradas algumas dificuldades. A base de dados obtida apresenta uma grande quantidade de dados faltantes, por esse motivo é ideal

que sejam exploradas outras metodologias de manipulação desses dados, minimizando a propagação de erros. Ainda, é necessário trabalhar com outros conjuntos de dados.

Por fim, destacamos que este artigo não esgota o assunto, mas fornece uma base para reflexões e passos adicionais. Os autores encorajam outros pesquisadores a explorar ainda mais ferramentas para diagnóstico diferencial do TEA, a fim de aprofundar o entendimento sobre o tema e desenvolver novas soluções. Através desse trabalho, espera-se ter contribuído para o avanço do conhecimento e da sociedade. Os resultados aqui apresentados servem como um ponto de partida para debates e ações futuras em busca de soluções para a questão.

## Referências

- Abdolzadegan, D., Moattar, M. H., & Ghoshuni, M. (2020). A robust method for early diagnosis of autism spectrum disorder from EEG signals based on feature selection and DBSCAN method. *Biocybernetics and Biomedical Engineering*, 40(1), 482-493.
- Abuhaija, B., Alloubani, A., Almatari, M., Jaradat, G. M., Hemn, B. A., Abualkishik, A. M., & Alsmadi, M. K. (2023). A comprehensive study of machine learning for predicting cardiovascular disease using Weka and SPSS tools. *International Journal of Electrical and Computer Engineering*, 13(2), 1891.

Aditya, E., Situmorang, Z., Hayadi, B. H., & Zarlis, M. (2022, October). New student prediction using algorithm naive bayes and regression analysis in universitas potensi utama. In 2022 4th International Conference on Cybernetics and Intelligent System (ICORIS) (pp. 1-6). IEEE.

Alves, C. L., Toutain, T. G. D. O., de Carvalho Aguiar, P., Pineda, A. M., Roster, K., Thielemann, C., ... & Rodrigues, F. A. (2023). Diagnosis of autism spectrum disorder based on functional brain networks and machine learning. *Scientific Reports*, 13(1), 8072.

American Psychiatric Association, D. S. M. T. F., & American Psychiatric Association. (2013). *Diagnostic and statistical manual of mental disorders: DSM-5* (Vol. 5, No. 5). Washington, DC: American psychiatric association.

Ari, B., Sobahi, N., Alçin, Ö. F., Sengur, A., & Acharya, U. R. (2022). Accurate detection of autism using Douglas-Peucker algorithm, sparse coding based feature mapping and convolutional neural network techniques with EEG signals. *Computers in Biology and Medicine*, 143, 105311.

Ben-Gal, I. (2008). Bayesian networks. *Encyclopedia of statistics in quality and reliability*.

Biau, G., & Scornet, E. (2016). A random forest guided tour. *Test*, 25, 197-227.

Bosl, W., Tierney, A., Tager-Flusberg, H., & Nelson, C. (2011). EEG complexity as a biomarker for autism

spectrum disorder risk. *BMC medicine*, 9, 1-16.

Briouza, S., Gritli, H., Khraief, N., Belghith, S., & Singh, D. (2022, March). Classification of sEMG biomedical signals for upper-limb rehabilitation using the random forest method. In 2022 5th International Conference on Advanced Systems and Emergent Technologies (IC\_ASET) (pp. 161-166). IEEE.

Burbidge, R., Trotter, M., Buxton, B., & Holden, S. (2001). Drug design by machine learning: support vector machines for pharmaceutical data analysis. *Computers & chemistry*, 26(1), 5-14.

Chawla, P., Rana, S. B., Kaur, H., & Singh, K. (2023). Computer-aided diagnosis of autism spectrum disorder from EEG signals using deep learning with FAWT and multiscale permutation entropy features. *Proceedings of the Institution of Mechanical Engineers, Part H: Journal of Engineering in Medicine*, 237(2), 282-294.

Chen, S. H., & Pollino, C. A. (2012). Good practice in Bayesian network modelling. *Environmental Modelling & Software*, 37, 134-145.

Dias, C. C. V., Maciel, S. C., Silva, J. V. C. D., & Menezes, T. D. S. B. D. (2022). Social Representations About Autism by University Students. *Psico-USF*, 26, 631-643.

Dickinson, A., Jeste, S., & Milne, E. (2022). Electrophysiological signatures

of brain aging in autism spectrum disorder. *Cortex*, 148, 139-151.

Farooq, M. S., Tehseen, R., Sabir, M., & Atal, Z. (2023). Detection of autism spectrum disorder (ASD) in children and adults using machine learning. *Scientific Reports*, 13(1), 9605.

Friedman, N., & Koller, D. (2003). Being Bayesian about network structure. A Bayesian approach to structure discovery in Bayesian networks. *Machine learning*, 50, 95-125.

Grossi, E., Valbusa, G., & Buscema, M. (2021). Detection of an autism EEG signature from only two EEG channels through features extraction and advanced machine learning analysis. *Clinical EEG and Neuroscience*, 52(5), 330-337.

Heinsfeld, A. S., Franco, A. R., Craddock, R. C., Buchweitz, A., & Meneguzzi, F. (2018). Identification of autism spectrum disorder using deep learning and the ABIDE dataset. *NeuroImage: Clinical*, 17, 16-23.

Hermawan, D. R., Fatihah, M. F. G., Kurniawati, L., & Helen, A. (2021, October). Comparative study of J48 decision tree classification algorithm, random tree, and random forest on in-vehicle CouponRecommendation data. In 2021 International conference on artificial intelligence and big data analytics (pp. 1-6). IEEE.

Khodatars, M., Shoeibi, A., Sadeghi, D., Ghaasemi, N., Jafari, M., Moridian, P., ... & Berk, M. (2021). Deep learning for neuroimaging-based diagnosis and

rehabilitation of autism spectrum disorder: a review. *Computers in biology and medicine*, 139, 104949.

Kulage, K. M., Goldberg, J., Usseglio, J., Romero, D., Bain, J. M., & Smaldone, A. M. (2020). How has DSM-5 affected autism diagnosis? A 5-year follow-up systematic literature review and meta-analysis. *Journal of autism and developmental disorders*, 50, 2102-2127.

Kulage, K. M., Smaldone, A. M., & Cohn, E. G. (2014). How will DSM-5 affect autism diagnosis? A systematic literature review and meta-analysis. *Journal of autism and developmental disorders*, 44(8), 1918-1932.

Mellema, C. J., Nguyen, K. P., Treacher, A., & Montillo, A. (2022). Reproducible neuroimaging features for diagnosis of autism spectrum disorder with machine learning. *Scientific reports*, 12(1), 3057.

Merlini, D., & Rossini, M. (2021). Text categorization with WEKA: A survey. *Machine Learning with Applications*, 4, 100033.

Mohi-ud-Din, Q., & Jayanthi, A. K. (2021, March). Detection of Autism Spectrum Disorder from EEG signals using pre-trained deep convolution neural networks. In 2021 Seventh International conference on Bio Signals, Images, and Instrumentation (ICBSII) (pp. 1-5). IEEE.

Mun, N. L., & Jumadi, N. A. (2020). Statistical evaluation on the performance of Dyslexia risk screening

system based fuzzy logic and WEKA. *Int J Adv Sci Technol*, 29(7), 638-49.

Niranjan, A., Nutan, D. H., Nitish, A., Shenoy, P. D., & Venugopal, K. R. (2018, April). ERCR TV: Ensemble of random committee and random tree for efficient anomaly classification using voting. In 2018 3rd international conference for convergence in technology (I2CT) (pp. 1-5). IEEE.

Oliveira, G. (2009). Autismo: diagnóstico e orientação. Parte I- Vigilância, rastreo e orientação nos cuidados primários de saúde. *Acta Pediatr Port*, 40(6), 278-87.

Ruping, S. (2001, November). Incremental learning with support vector machines. In *Proceedings 2001 IEEE international conference on data mining* (pp. 641-642). IEEE.

Shubho, S. A., Razib, M. R. H., Rudro, N. K., Saha, A. K., Khan, M. S. U., & Ahmed, S. (2019, December). Performance analysis of NB Tree, REP tree and random tree classifiers for credit card fraud data. In 2019 22nd International Conference on Computer and Information Technology (ICCIT) (pp. 1-6). IEEE.

Smelser, N. J., & Baltes, P. B. (Eds.). (2001). *International encyclopedia of the social & behavioral sciences* (Vol. 11). Amsterdam: Elsevier.

Somvanshi, M., Chavan, P., Tambade, S., & Shinde, S. V. (2016, August). A

review of machine learning techniques using decision tree and support vector machine. In 2016 international conference on computing communication control and automation (ICCUBEA) (pp. 1-7). IEEE.

Srivastava, S. (2014). Weka: a tool for data preprocessing, classification, ensemble, clustering and association rule mining. *International Journal of Computer Applications*, 88(10).

Uusitalo, L. (2007). Advantages and challenges of Bayesian networks in environmental modelling. *Ecological modelling*, 203(3-4), 312-318.

Vijay, V., & Verma, P. (2023, January). Variants of Naïve Bayes Algorithm for Hate Speech Detection in Text Documents. In 2023 International Conference on Artificial Intelligence and Smart Communication (AISC) (pp. 18-21). IEEE.

Volkmar, F. R., & Reichow, B. (2013). Autism in DSM-5: progress and challenges. *Molecular autism*, 4, 1-6.

Zhang, H. (2004). The optimality of naive Bayes. *Aa*, 1(2), 3.

Zhou, Q., Lan, W., Zhou, Y., & Mo, G. (2020, November). Effectiveness evaluation of anti-bird devices based on random forest algorithm. In 2020 7th International Conference on Information, Cybernetics, and Computational Social Systems (ICCSS) (pp. 743-748). IEEE.

Esse estudo foi financiado em parte pela Coordenação de Aperfeiçoamento de Pessoal de Nível Superior (CAPES). Os autores também são gratos à Fundação de Amparo à Ciência e Tecnologia do Estado de Pernambuco (FACEPE) e ao Conselho Nacional de Desenvolvimento Científico e Tecnológico (CNPq) pelo suporte financeiro. Com esforços conjuntos, podemos continuar a expandir os horizontes do saber e enfrentar os desafios que ainda se apresentam. Agradecemos também ao Departamento de Ciência e Tecnologia do Ministério da Saúde, DECIT/MS, e à Financiadora de Estudos e Projetos, FINEP ref. 2170/22, pelo financiamento parcial do projeto.

Flávio Secco Fonseca ([fsf2@ecomp.poli.br](mailto:fsf2@ecomp.poli.br))\* trabalhou no planejamento e condução dos estudos descritos pelo manuscrito, liderando a escrita do artigo.

Maria Vitória S. Muniz ([mvs3@cin.ufpe.br](mailto:mvs3@cin.ufpe.br)) trabalhou no pré-processamento dos dados do estudo, contribuindo com a escrita e revisão do artigo.

Catarina V. N. de Oliveira ([catarina.victoria@ufpe.br](mailto:catarina.victoria@ufpe.br)) trabalhou na análise dos dados do estudo descrito, contribuindo com a escrita e revisão do artigo.

Thailson C. V. da Silva ([thailson.caetano@ufpe.br](mailto:thailson.caetano@ufpe.br)) trabalhou nas revisões e na pesquisa da fundamentação teórica.

Wellington P. Dos Santos ([wellington.santos@ufpe.br](mailto:wellington.santos@ufpe.br)) trabalhou na concepção e coordenação do projeto de pesquisa, assim como na orientação e supervisão dos demais pesquisadores a cada etapa.

\*Autor-correspondente.

Data de Submissão: 11/06/2024 Data de Aprovação: 15/07/2024.

Editor-Chefe: Diogo Henrique Helal.

Editor Adjunto: Bruno Melo Moura.

Esta obra está licenciada sob uma Atribuição-Não Comercial 4.0 Internacional (CC BY NC 4.0). Esta licença permite que outros distribuam, remixem, adaptem e criem a partir do trabalho, para fins não comerciais, desde que lhe atribuam o devido crédito pela criação original. Texto da licença: [https://creativecommons.org/licenses/by-nc/4.0/deed.pt\\_BR](https://creativecommons.org/licenses/by-nc/4.0/deed.pt_BR).