



UNIVERSIDADE FEDERAL DO OESTE DA BAHIA
Centro Multidisciplinar de Bom Jesus da Lapa
Colegiado de Engenharia Elétrica

MATHEUS PEREIRA ALVES

**TÉCNICAS DE PROCESSAMENTO DE SINAIS E
MACHINE LEARNING APLICADO NA
IDENTIFICAÇÃO DE DEFEITOS EM ESTRUTURAS
INDUSTRIAIS**

Bom Jesus da Lapa–BA
Abril de 2026

UNIVERSIDADE FEDERAL DO OESTE DA BAHIA

Centro Multidisciplinar de Bom Jesus da Lapa

Colegiado de Engenharia Elétrica

Matheus Pereira Alves

**TÉCNICAS DE PROCESSAMENTO DE SINAIS E MACHINE
LEARNING APLICADO NA IDENTIFICAÇÃO DE DEFEITOS EM
ESTRUTURAS INDUSTRIAIS**

Trabalho de conclusão de curso apresentado ao Colegiado de Engenharia Elétrica do Centro Multidisciplinar de Bom Jesus da Lapa da Universidade Federal do Oeste da Bahia, como requisito parcial para obtenção do grau de Bacharel em Engenharia Elétrica.

Orientador: Prof. Dr. Manoel Messias Silva Junior

Bom Jesus da Lapa–BA

Abril de 2026

FICHA CATALOGRÁFICA

A474

Alves, Matheus Pereira

Técnicas de processamento de sinais *E Machine Learning* aplicado na identificação de defeitos em estruturas industriais. / Matheus Pereira Alves. – 2026.

63f.: il.

Orientador: Prof. Dr. Manoel Messias Silva Junior

TCC - Graduação em Engenharia Elétrica, Universidade Federal do Oeste da Bahia. Centro Multidisciplinar de Bom Jesus da Lapa - BA, 2026.

1. Engenharia Elétrica. 2. Aprendizado de Máquina. 3. *E Machine Learning*. 4. Instrumentação Industrial. I. Silva Junior, Manoel Messias. II. Universidade Federal do Oeste da Bahia – Centro Multidisciplinar de Bom Jesus da Lapa - BA. III. Título.

CDD 621.3

Biblioteca Universitária de Bom Jesus da Lapa – UFOB

FOLHA DE APROVAÇÃO

MATHEUS PEREIRA ALVES

TÉCNICAS DE PROCESSAMENTO DE SINAIS E MACHINE LEARNING APLICADO NA IDENTIFICAÇÃO DE DEFEITOS EM ESTRUTURAS INDUSTRIAIS

Esta monografia foi aprovada como requisito parcial para obtenção do título de Bacharel em Engenharia Elétrica, aprovada em sua forma final pelo Colegiado de Engenharia Elétrica do Centro Multidisciplinar de Bom Jesus da Lapa da Universidade Federal do Oeste da Bahia.

Bom Jesus da Lapa, 23 de Abril de 2026

Prof. Dr. Manoel Messias Silva Junior (UFOB)
(Orientador)

Profa. Dra. Stefania de Oliveira Silva (UFOB)

Mestranda Mariana Farias da Silva (UFBA)

Prof. Dr. Kleymilson do Nascimento Souza (UFOB)

*Esta monografia é dedicada aos meus pais, pilares da
minha formação como ser humano.*

Agradecimentos

Início agradecendo a Deus, que me conduziu com sabedoria nos momentos fáceis e foi meu amparo constante quando os obstáculos surgiram. Sem a Sua força, esta jornada não teria sido possível.

Agradeço ao meu orientador, Prof. Dr. Manoel Messias Silva Junior, que foi muito mais que um professor: um mentor e amigo. Seus ensinamentos, sugestões, paciência e companherismo foram essenciais para a conclusão deste trabalho.

Agradeço a meus pais, Manoel da Cruz Alves e Maria de Lourdes Pereira de Oliveira. Obrigado por me ensinarem que o conhecimento é enriquecedor, algo que ilumina a alma e a mente. Agradeço também por todo suporte — moral, afetivo, financeiro — que nunca me faltou. Espero que uma fração desta vitória seja de vocês também. Extendendo os agradecimentos também a minha família como um todo.

Agradeço à minha namorada, Slene Cardoso, por todo o apoio incondicional e por dividir comigo cada preocupação e felicidade ao longo desta jornada. Sem você, esta conquista certamente não teria o mesmo significado. Obrigado por tudo, meu amor.

À Universidade Federal do Oeste da Bahia, em especial ao seu corpo docente e técnico-administrativo, principalmente ao colegiado de Engenharia Elétrica, meu sincero reconhecimento pelo empenho em construir um ambiente acadêmico inclusivo e propício ao desenvolvimento intelectual e humano.

Aos meus amigos, principalmente Jadson, Valdeilson, Mariana, Rafael e João Pedro que tanto me apoiaram e me encorajaram ao longo dessa jornada.

Agradeço também aos colegas de classe e de curso que sempre estiveram presentes e auxiliando no meio acadêmico.

*“Não é o conhecimento, mas o ato de aprender,
não a posse mas o ato de chegar lá, que concede a
maior satisfação.”*

(Carl Friedrich Gauss)

*“O navio que navega mais rápido não é
necessariamente aquele construído com a ciência
mais rigorosa.”*

(Oliver Heaviside)

Resumo

A integridade estrutural de equipamentos industriais é crítica para a segurança e eficiência operacional, demandando técnicas robustas de inspeção, como os Ensaios Não Destrutivos por Correntes Parasitas Pulsadas. Contudo, a interpretação manual dos sinais oriundos dessas inspeções é suscetível à subjetividade humana e à complexidade dos dados. Este trabalho propõe o desenvolvimento de um sistema inteligente para a classificação automática de defeitos em tubulações de aço carbono revestidas, distinguindo entre condições de integridade, corrosão interna e corrosão externa. A metodologia integrou um sistema de aquisição microprocessado baseado em sensores de Magnetorresistência Gigante e um pipeline computacional. Técnicas de engenharia de atributos e aumento de dados foram aplicadas para mitigar a escassez de amostras, seguidas por uma avaliação comparativa entre a Análise de Componentes Principais e a Análise Discriminante Linear para a redução de dimensionalidade. Cinco algoritmos de aprendizado de máquina foram implementados e otimizados via Otimização Bayesiana e validação cruzada aninhada: Gated Recurrent Unit, K-Nearest Neighbors, Máquina de Vetores de Suporte, Random Forest e CatBoost. Os resultados demonstraram a superioridade da Análise Discriminante Linear na separabilidade das classes, permitindo que a rede neural Gated Recurrent Unit alcançasse 94% de acurácia e o CatBoost 92% em ambiente desktop. A validação estatística confirmou a equivalência de desempenho entre esses modelos de topo. Para atestar a aplicabilidade em cenários reais, o modelo CatBoost associado à Análise Discriminante Linear foi implementado de forma embarcada em um dispositivo de borda Raspberry Pi Zero 2W. A avaliação final em hardware registrou uma acurácia de 90,56%, tempo médio de inferência de 1,80 ms por amostra e baixo consumo de memória operacional, com pico de 8,54 KB. Conclui-se que o sistema desenvolvido apresenta viabilidade prática, oferecendo uma solução preditiva robusta, eficiente e de baixo custo computacional capaz de operar em tempo real no setor industrial.

Palavras-chave: Correntes Parasitas Pulsadas. Aprendizado de Máquina. Redução de Dimensionalidade. GRU. CatBoost. Instrumentação Industrial.

Abstract

The structural integrity of industrial equipment is critical for operational safety and efficiency, demanding robust inspection techniques, such as Pulsed Eddy Current non-destructive testing. However, the manual interpretation of signals from these inspections is susceptible to human subjectivity and data complexity. This work proposes the development of an intelligent system for the automatic classification of defects in coated carbon steel pipes, distinguishing between integrity conditions, internal corrosion, and external corrosion. The methodology integrated a microprocessor-based acquisition system using Giant Magnetoresistance sensors and a computational pipeline. Feature engineering and data augmentation techniques were applied to mitigate sample scarcity, followed by a comparative evaluation between Principal Component Analysis and Linear Discriminant Analysis for dimensionality reduction. Five machine learning algorithms were implemented and optimized via Bayesian Optimization and nested cross-validation: Gated Recurrent Unit, K-Nearest Neighbors, Support Vector Machine, Random Forest, and CatBoost. The results demonstrated the superiority of Linear Discriminant Analysis in class separability, enabling the Gated Recurrent Unit neural network to achieve 94% accuracy and CatBoost 92% in a desktop environment. Statistical validation confirmed the performance equivalence between these top models. To verify applicability in real-world scenarios, the CatBoost model combined with Linear Discriminant Analysis was embedded in a Raspberry Pi Zero 2W edge device. The final hardware evaluation recorded an accuracy of 90.56%, an average inference time of 1.80 ms per sample, and a low operational memory consumption with a peak of 8.54 KB. It is concluded that the developed system presents practical viability, offering a robust, efficient, and low computational cost predictive solution capable of operating in real-time within the industrial sector.

Keywords: Pulsed Eddy Current. Machine Learning. Dimensionality Reduction. GRU. CatBoost. Industrial Instrumentation..

Lista de Figuras

1	Diagrama básico de um sistema EC na configuração diferencial	26
2	Descrição de um Perceptron	30
3	Representação de uma célula GRU	31
4	Margens de decisão com diferentes kernels para SVM: a) e b) Linear, c) Função de base radial e d) Polinomial	32
5	Comparativo entre <i>Bagging</i> e <i>Boosting</i>	33
6	Dimensões do corpo de prova utilizado	36
7	Diagrama Experimental Montado	37
8	Assinaturas PEC no domínio do tempo: medido x referência (a) e Diferencial (b)	39
9	Fluxograma simplificado do processo de classificação de sinais PEC.	41
10	Fluxograma do pipeline de processamento, otimização e avaliação.	42
11	Comparação entre o sinal diferencial original (esquerda) e o sinal aumentado sinteticamente (direita) com aplicação de ruído, <i>jitter</i> e escala.	44
12	Loops Interno e Externos da Validação Cruzada Aninhada	45
13	Distribuição dos valores das Componentes Principais (PCAs) utilizadas nos modelos de ML.	51
14	Comparativo das Matrizes de Confusão: GRU+LDA (esquerda) x CatBoost+PCA (direita).	53
15	Comparativo das Curvas ROC (esquerda) e Ranking de Performance AUC (direita).	53
16	Comparação de Modelos com Intervalos de Confiança.	54
17	Heatmap de Comparação Par-a-Par (ANOVA + Tukey HSD).	55

Lista de Tabelas

1	Parâmetros da Sonda PEC	38
2	Métricas de Desempenho Detalhadas	52
3	Métricas do CatBoost+LDA no Raspberry Pi Zero 2W	56

Lista de quadros

1	Comparativo entre técnicas de ENDs	24
2	Especificações do Hardware	40
3	Ferramentas de Software e Bibliotecas	40
4	Espaço de busca dos Hiperparâmetros	46
5	Hiperparâmetros Otimizados	47

Lista de abreviaturas, acrônimos e siglas

ADC	Analog-to-Digital Converter
ANOVA	Analysis of Variance
AUC	Area Under the Curve
CV	Cross-Validation
DE	Defeito Externo
DFT	Discrete Fourier Transform
DI	Defeito Interno
DMA	Direct Memory Access
DSP	Digital Signal Processing
EC	Eddy Current
END	Ensaaios Não Destrutivos
FPR	False Positive Rate
FPU	Floating Point Unit
GMR	Giant Magnetoresistance
GPU	Graphics Processing Unit
GRU	Gated Recurrent Unit
HSD	Honest Significant Difference
K-NN	K-Nearest Neighbors
LDA	Linear Discriminant Analysis
LSTM	Long Short-Term Memory
ML	Machine Learning
PCA	Principal Components Analysis
PEC	Pulsed Eddy Current

ReLU	Rectified Linear Unit
RF	Random Forest
RMS	Root Mean Square
RNA	Rede Neural Artificial
RNR	Redes Neurais Recorrentes
ROC	Receiver Operating Characteristic
SD	Sem Defeito
SVM	Support Vector Machine
TanH	Hyperbolic Tangent
TPE	Tree-structured Parzen Estimator
TPR	True Positive Rate

Lista de Símbolos

f	Frequência
H_0	Hipótese nula
h_t	Estado oculto da GRU no instante t
J	Score de Fisher
k	Índice de frequência discreta ou número de folds
N	Número total de amostras
n	Índice de tempo discreto
p	Valor de probabilidade (p-valor)
r_t	Porta de redefinição (<i>reset gate</i>)
S_b	Matriz de dispersão entre classes
S_w	Matriz de dispersão intraclasse
$S[k]$	Sinal no domínio da frequência
$s[n]$	Sinal no domínio do tempo
t	Tempo
W_N	Fator de rotação
x_t	Entrada da rede neural no instante t
z_t	Porta de atualização (<i>update gate</i>)
Ω	Ohm

Sumário

1	INTRODUÇÃO	16
1.1	Objetivos	17
1.1.1	Objetivos Específicos	18
1.2	Metodologia da Pesquisa	18
1.3	Organização do Texto	19
1.4	Produções Científicas Associadas ao Trabalho	19
2	REVISÃO DA LITERATURA	21
2.1	Trabalhos Correlatos	21
3	FUNDAMENTAÇÃO TEÓRICA	23
3.1	Técnicas de Ensaios Não-Destrutivos	23
3.1.1	Ensaio por Correntes Parasitas	25
3.2	Técnicas de Pré-Processamento e <i>Feature Engineering</i>	26
3.2.1	Redução de dimensionalidade	27
3.3	Algoritmos de Aprendizado de Máquina	29
3.3.1	Redes Neurais Artificiais	29
3.3.2	Algoritmos Baseados em Instância e Estatística	31
3.3.3	Árvores de Decisão	32
4	DESENVOLVIMENTO	35
4.1	Sistema de Aquisição de Dados	35
4.1.1	Caracterização do Objeto de Ensaio	36
4.1.2	Sistema Microprocessado de Aquisição	37
4.2	Ferramentas de <i>software</i> e <i>hardware</i>	39
4.3	Sistema Proposto	41
4.3.1	Técnicas de Aumentação de Dados	43
4.3.2	Otimização Bayesiana e Validação Cruzada	45

4.4	Metodologia de Avaliação dos Resultados	46
4.4.1	Matriz de Confusão e Métricas Derivadas	47
4.4.2	Curva ROC e Área sob a Curva (AUC)	48
4.4.3	Validação Estatística: ANOVA e Teste de Tukey HSD	48
5	RESULTADOS	50
5.1	Extração de Características e Redução de Dimensionalidade	50
5.2	Desempenho dos Classificadores	51
5.3	Validação Estatística	54
5.4	Validação Embarcada na Plataforma Raspberry Pi Zero 2W	55
6	CONSIDERAÇÕES	57
6.1	Sugestões Para Trabalhos Futuros	58
	REFERÊNCIAS	61

INTRODUÇÃO

As técnicas de monitoramento estrutural não destrutivas estão cada vez mais sendo empregadas no meio industrial. Elas permitem identificar problemas em ferramentas e materiais sem comprometer a integridade e a disponibilidade do sistema inspecionado (HELLIER, 2012). A identificação precoce e precisa de defeitos contribui para a segurança operacional, reduzindo custos de manutenção e aumentando a confiabilidade nos sistemas produtivos (MELGAR; CHAMBI; TORRES, 2025).

Os ensaios por correntes parasitas pulsadas (Pulsed Eddy Current - PEC) são frequentemente empregados para avaliar estruturas metálicas, fazendo parte das técnicas de Ensaio Não Destrutivo (END). Esses testes utilizam campos magnéticos pulsados para identificar defeitos e corrosão em materiais, principalmente em aplicações nas quais o acesso à superfície é limitado, como em materiais com proteção ou revestimento térmico (CHEN et al., 2023).

No contexto atual, os equipamentos utilizados em inspeções por correntes parasitas geralmente se baseiam na análise da amplitude ou do ângulo de fase dos sinais em um gráfico do plano de impedância, no qual a resistência da bobina é representada no eixo X e a reatância indutiva no eixo Y (SOPHIAN; TIAN; FAN, 2017). Em outras situações, a detecção de defeitos é realizada a partir de parâmetros como amplitude e tempo de pico durante o ciclo de excitação da bobina. Nesses procedimentos, é comum que o operador analise o comportamento do sinal de inspeção e o compare com um padrão de referência, previamente estabelecido a partir de uma amostra com as mesmas características do material avaliado (SOPHIAN; TIAN; FAN, 2017). Contudo, a interpretação dessas medições depende intrinsecamente da experiência e da habilidade técnica do operador.

A utilização de sistemas automáticos para identificação de defeitos tem se mos-

trado uma alternativa promissora para apoiar os operadores durante os testes realizados com correntes parasitas. Nos últimos anos, técnicas de processamento de sinais associadas a métodos de aprendizado de máquina (*Machine Learning* - ML) têm sido aplicadas com sucesso no tratamento das assinaturas obtidas em ensaios de Pulsed Eddy Currents (PEC) (SILVA et al., 2021)(FU et al., 2019). Um aspecto central nesse processo é a construção de um pipeline de dados robusto, que geralmente se inicia com a extração de atributos (*feature engineering*). Esta etapa converte os sinais brutos em um vetor de características de dimensão significativamente menor, capturando informações cruciais sobre a morfologia do sinal. Em seguida, técnicas de redução de dimensionalidade, como a Análise de Componentes Principais (PCA) ou a Análise de Discriminante Linear (LDA), são frequentemente aplicadas sobre esses atributos para eliminar redundâncias e maximizar a separabilidade entre as classes. Esse processamento em múltiplos estágios permite que os sistemas de classificação operem de forma mais eficiente (MUNIR et al., 2024).

Além disso, ao identificar falhas de forma preditiva de componentes e maquinários é possível reduzir custos operacionais, manter em funcionamento o fluxo de produção e garantir um nível elevado de segurança industrial. Dessa forma, a combinação de técnicas avançadas de processamento de sinais, aprendizado de máquina e sistemas inteligentes representa um avanço significativo para a modernização e otimização das práticas de inspeção em ambientes industriais (MELGAR; CHAMBI; TORRES, 2025).

Diante do exposto, esta pesquisa propõe a avaliação de diferentes combinações de técnicas de processamento de sinais, como PCA e LDA, engenharia de atributos e de *Machine Learning* para o estudo de caso de um tubo de aço carbono com revestimento compósito. Serão avaliadas três classes distintas: uma sem defeito (SD) e outras duas com defeitos de corrosão, sendo interna (DI) e externa (DE). As avaliações propostas neste trabalho buscam contribuir para o entendimento e implementações práticas de sistemas inteligentes de monitoramento de falhas, auxiliando no avanço da manutenção preditiva industrial.

1.1 Objetivos

Realizar um estudo comparativo de técnicas de processamento de sinais e algoritmos de aprendizado de máquina aplicados à identificação de defeitos em estruturas industriais.

1.1.1 Objetivos Específicos

- Revisar métodos de processamento de sinais comumente empregados na análise de dados provenientes de ensaios não destrutivos;
- Implementar diferentes algoritmos de aprendizado de máquina (Support Vector Machines (SVM), Redes Neurais Recorrentes, Árvores de Decisão, K-Nearest Neighbors (K-NN), entre outros) voltados à classificação de defeitos.
- Comparar o desempenho das técnicas de processamento e dos algoritmos de classificação por meio de métricas quantitativas e gráficas (acurácia, precisão, F1-score, Curva Característica de Operação do Receptor (ROC), entre outros).
- Identificar quais combinações entre técnicas de pré-processamento e modelos de aprendizado apresentam maior robustez e aplicabilidade para ambientes industriais.
- Propor recomendações para a integração de tais métodos em sistemas de monitoramento inteligente e manutenção preditiva.

1.2 Metodologia da Pesquisa

Para o desenvolvimento deste projeto, foram seguidas as seguintes etapas, que delimitam desde a análise exploratória de dados até a avaliação e comparação dos modelos de *Machine Learning*.

1. Rotulação, normalização e separação dos dados, como também, as devidas importações dos dados e bibliotecas ao Ambiente de Desenvolvimento Integrado;
2. Realização da engenharia de atributos, com a extração de *features* (do inglês, atributos) dos sinais da PEC, seguida pela aplicação de técnicas de Aumento de Dados e, por fim, a divisão do conjunto de dados nas partições de treino, validação e teste;
3. Utilização de PCA ou LDA para redução de dimensionalidade ou separação de classes das *features*;
4. Implementação das *features* para otimização de hiperparâmetros via Busca Bayesiana e validação cruzada *k-fold*;
5. Utilização dos dados pré-processados acrescidos dos melhores resultados da otimização de hiperparâmetros para treinamento e classificação dos modelos de aprendizado de máquina;

6. O desempenho dos classificadores foi mensurado por meio de métricas como acurácia, precisão, recall, F1-score e pela análise da matriz de confusão. Foram também geradas curvas ROC para análise visual da capacidade de discriminação de cada modelo. Por fim, houve a comparação e discussão entre os modelos.

1.3 Organização do Texto

- **Capítulo 02** - Realiza-se a revisão da literatura, em que são apresentados e analisados os principais trabalhos correlatos que fundamentam a pesquisa, com foco em estudos que aplicam Aprendizado de Máquina em Ensaios Não Destrutivos, processamento de sinais e detecção de defeitos.
- **Capítulo 03** - Apresenta-se a fundamentação teórica relacionada ao tema do trabalho, focando nos Ensaios Não Destrutivos (END), principalmente PEC, técnicas de processamento de atributos e sinais e por último, algoritmos de *Machine Learning*;
- **Capítulo 04** - É Apresentado o fluxograma do sistema proposto da combinação das técnicas de processamento e de algoritmos de ML para classificação. Além disso, descreve-se a sistematização do procedimento experimental, parâmetros do sistema, *Hardware* e *Software* utilizados, metodologia dos resultados;
- **Capítulo 05** - Expõe os resultados das combinações entre diferentes técnicas de processamento, engenharia de atributos e algoritmos de *Machine Learning*. Nesse capítulo também ocorre as discussões acerca das melhores combinações, as suas diferenças e possível implementação física do sistema proposto na área industrial;
- **Capítulo 06** - Reservado para as devidas conclusões sobre o trabalho realizado, como também, sugestões para trabalhos futuros.

1.4 Produções Científicas Associadas ao Trabalho

1. CARVALHO, G. O.; SILVA, M. M.; ALVES, M, P.. Banana Leaf Disease Classification System Using Convolutional Neural Networks. In: Conferência Brasileira de Inteligência Computacional, 2025, Belo Horizonte. Anais do XVII CBIC. v. 1. p. 1-8. This article presents the development of a disease classification system for banana plantations using Convolutional Neural Networks (CNN). The main objective is to assist producers in the early identification of diseases such as Sigatoka, Cordana, and Pestalotiopsis, which compromise production and cause financial losses. The methodology involves processing images of banana leaves to enable the

classification of crop conditions, using a dataset extracted from the Banana Leaf Spot Diseases (BananaLSD) dataset. The proposed model was trained and validated using various hyperparameters to optimize disease detection performance. Performance evaluation used confusion matrices alongside precision, recall, and F1 score metrics to assess the performance of the designed classification system. The best result was obtained by combining CNN with the application of the SVM post-processing technique, achieving an overall accuracy of 89%, considering the classes cordana, healthy, pestalotiopsis, and sigatoka.

REVISÃO DA LITERATURA

Este capítulo apresenta uma revisão da literatura fundamental para esta pesquisa, com foco em trabalhos que aplicam Aprendizado de Máquina (Machine Learning - ML) em Ensaios Não Destrutivos (END). A análise de estudos correlatos, abordando desde o processamento de sinais com Análise de Componentes Principais (PCA) e Análise Discriminante Linear (LDA) até a classificação de defeitos com redes neurais.

2.1 Trabalhos Correlatos

Sophian et al. (2021) propuseram um método baseado em Regressão por Processo Gaussiano (GPR) para prever a espessura de aço (simulando corrosão) a partir de sinais PEC. O objetivo foi eliminar a necessidade da média de múltiplos sinais, acelerando assim o tempo de inspeção. O modelo GPR demonstrou ser mais preciso que a abordagem tradicional, alcançando um Erro Médio Absoluto (MAE) de 0,21 mm, em comparação com 0,36 mm, mostrando-se promissor para ensaios não destrutivos que exigem varredura rápida.

No contexto nacional, o estudo de Silva et al. (2021) é de particular importância para o presente trabalho, pois serve como base metodológica para a detecção de defeitos em tubos de aço carbono. Os autores desenvolveram um sistema embarcado inteligente para aquisição, processamento e classificação de sinais de Correntes Parasitas Pulsadas (PEC) na detecção de defeitos em tubos de aço carbono. O sistema integra o uso de redes neurais artificiais (RNAs) e PCA. Os autores alcançaram um alto nível de eficiência na detecção de defeitos e compararam o desempenho entre os modelos MLP (Multilayer Perceptron) e ELM (Extreme Learning Machine) em problemas de classificação.

Munir et al. (2024) realizaram uma extensa revisão da literatura sobre a aplicação de Aprendizado de Máquina (ML) em Ensaios por Correntes Parasitas (ECT). O estudo teve como objetivo avaliar o cenário atual, identificar lacunas de pesquisa e apontar direções futuras. Como principais resultados, os autores destacaram a necessidade de desenvolver datasets públicos para garantir a reprodutibilidade dos experimentos, a carência de métodos para definir o tamanho ideal de amostras, a pouca exploração de técnicas de aumento de dados (data augmentation) e a ausência de modelos de ML que quantifiquem suas próprias incertezas.

Finalmente, Wang et al. (2025) apresentaram um método de aprendizado profundo para a segmentação de trincas em túneis com fundos complexos. A solução combina um algoritmo de rotulagem automática com uma rede U-Net otimizada com estruturas residuais (Res-UNet), visando reduzir o trabalho manual na criação de datasets. A abordagem atingiu uma precisão de segmentação de 98,78% e uma Interseção sobre União (IoU) de 58,50% após otimização, demonstrando a eficácia do modelo e do método de rotulagem automática para melhorar a eficiência do processo de anotação.

O seguinte trabalho se diferencia por utilizar e testar diferentes configurações de algoritmos e técnicas de processamento para avaliar os melhores resultados e sua possível aplicação embarcada.

FUNDAMENTAÇÃO TEÓRICA

Este capítulo apresenta os conceitos importantes que sustentam o desenvolvimento do presente trabalho. A abordagem se inicia com a contextualização dos Ensaaios Não Destrutivos (ENDs) com foco na técnica de Correntes Parasitas (EC). Em seguida, são detalhadas as técnicas de processamento de sinais utilizadas para extrair informações pertinentes, por exemplo a Transformada Discreta de Fourier (DFT). Por fim, são explorados os princípios do Aprendizado de Máquina, incluindo o estudo de diversas arquiteturas como Redes Neurais Artificiais (RNA), Redes Neurais Recorrentes (RNR), Gated Recurrent Unit (GRU), K-NN, CatBoost, SVM e Random Forest bem como as técnicas para otimização e validação desses modelos.

3.1 Técnicas de Ensaaios Não-Destrutivos

Técnicas de ENDs podem ser definidas como examinações ou testes de materiais sem alterá-los em nenhuma forma ou propriedade (HELLIER, 2012). Tais técnicas podem ser utilizadas para as mais diversas finalidades como controle de qualidade, avaliação de produto em etapas de processo e manutenção preditiva em produtos em uso (HASSAN; PANDURU; WALSH, 2025). No contexto da integridade de materiais, é relevante distinguir entre uma imperfeição e um defeito. Imperfeições são descontinuidades físicas, por vezes em escala atômica, que estão presentes no material original ou que se formam durante a fabricação ou no tempo de serviço. Uma imperfeição evolui para um defeito quando compromete a integridade e/ou a função do componente. Isso ocorre por meio de qualquer alteração em seu tamanho, forma ou propriedades, tornando-o inadequado para a sua finalidade e podendo, conseqüentemente, causar a ruptura do componente (SILVA

et al., 2023).

Diversas técnicas de ENDs são aplicadas no contexto industrial, cada uma com suas particularidades. As vantagens e limitações dos métodos mais comuns, incluindo aqueles relevantes para este trabalho como o ensaio por Correntes Parasitas (Eddy Currents - EC), são apresentadas de forma comparativa no Quadro 1.

Quadro 1 – Comparativo das principais características de técnicas selecionadas de Ensaios Não Destrutivos (ENDs).

Método	Vantagens	Limitações
Ensaio Ultrassônico	• Portátil e não perigoso. • Alta profundidade de penetração.	• Requer operador experiente. • Necessita de acoplamento com a superfície.
Partículas Magnéticas	• Método simples e de rápida execução. • Alta sensibilidade para defeitos superficiais.	• Aplicável apenas a materiais ferromagnéticos. • Requer operador experiente para a interpretação.
Correntes Parasitas	• Não exige contato direto com a peça. • Inspecciona através de tintas e revestimentos.	• Limitado a materiais condutores. • Sensibilidade reduzida para defeitos profundos.
Ensaio Radiográfico	• Capacidade de detectar defeitos internos com maior precisão. • Gera um registro permanente (filme).	• Exige procedimentos de segurança rigorosos. • Custo elevado e menor portabilidade.

Fonte: Autoria própria

O método de teste por ultrassom é uma técnica que emprega ondas mecânicas com frequências acima da audição humana (>20 kHz). No método convencional, um transdutor em contato com o componente gera ondas acústicas que se propagam através

do material. Quando essas ondas interagem com descontinuidades, elas retornam à superfície, onde o transdutor as converte em sinais elétricos que podem ser monitorados e analisados (SHALOO et al., 2022). Este método convencional é eficaz para detectar defeitos profundos em metais e plásticos. No entanto, devido a sua necessidade de contato e baixa resistência a altas temperaturas, se viu necessário o desenvolvimento de técnicas mais avançadas como o ultrassom a laser, que opera sem contato direto com o corpo de prova (SHALOO et al., 2022).

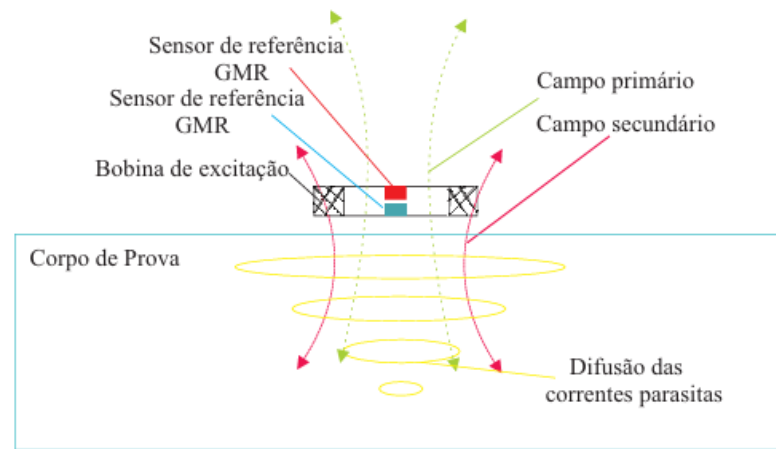
O ensaio por partículas magnéticas se destaca por ser um método rápido, confiável e simples, amplamente utilizado na detecção de defeitos em superfícies metálicas (SHRI-FAN; AKBAR; ISA, 2019). Durante o teste, a descontinuidade no corpo de prova ocasiona uma fuga do fluxo magnético, o que leva a uma concentração das partículas magnéticas no local da falha.

A inspeção por radiografia é um método não invasivo utilizado para examinar defeitos internos em componentes. O princípio se baseia na penetração de radiação ionizante por meio do corpo de prova, que devido a diferença de composição, formato ou propriedade no componente, a radiação é absorvida de maneira desigual, evidenciando descontinuidades (JIN et al., 2021). Essa variação pode ser detectada via filme ou detectores eletrônicos que são utilizados para indicar os defeitos.

3.1.1 Ensaio por Correntes Parasitas

No ensaio não destrutivo por Correntes Parasitas (EC), uma bobina de excitação alimentada por corrente alternada induz correntes elétricas (chamadas de correntes parasitas) na amostra, as quais, por sua vez, criam um campo magnético secundário de mesma direção e sentido oposto (SILVA, 2021). A detecção de defeito torna-se possível porque este campo magnético induzido irá variar se houver uma falha que impeça o fluxo das correntes parasitas. O campo também pode variar devido a mudanças na condutividade elétrica, na permeabilidade magnética ou na espessura da amostra (SOPHIAN; TIAN; FAN, 2017). A alteração no campo é então captada por um dispositivo de detecção, que normalmente é uma bobina ou um sensor magnético, como os sensores de efeito *Giant Magnetoresistance* (GMR).

Sistema de medição por EC geralmente utilizam diferentes topologias de aquisição dos sinais (direta, planar, matricial), porém, com foco maior para a forma diferencial na detecção de defeitos (SOPHIAN; TIAN; FAN, 2017). Nesta configuração, utiliza-se um sensor de medição e um sensor de referência, que podem ou não compartilhar o mesmo corpo de sonda, conforme ilustrado na Figura 1. Essa técnica permite que o sistema minimize os efeitos de variações indesejadas, como o *lift-off*, realçando as alterações no sinal causadas exclusivamente pelas descontinuidades no material.

Figura 1 – Diagrama básico de um sistema EC na configuração diferencial

Fonte: Silva (2021)

A técnica de Correntes Parasitas Pulsadas (PEC) é uma variação do método convencional (EC) que se diferencia pelo sinal de excitação, ao invés de uma onda senoidal de frequência única, no qual o PEC utiliza pulsos de onda quadrada. A principal vantagem dessa abordagem é que um pulso de onda quadrada é composto por uma soma de senoides com múltiplas frequências e amplitudes (SILVA et al., 2021). Isso permite que a inspeção alcance simultaneamente diferentes profundidades de penetração no material, simplificando o processo de avaliação. Os sinais transitórios resultantes, no entanto, são suscetíveis a ruídos, por isso, frequentemente exigem o uso de técnicas de pré-processamento, como a Transformada de Fourier Discreta (DFT) e a Análise de Componentes Principais (PCA), para extrair características relevantes e reduzir a dimensionalidade dos dados antes da análise (SILVA et al., 2021).

3.2 Técnicas de Pré-Processamento e *Feature Engineering*

A importância da fase do processamento dos dados é justificada pela natureza do sinal bruto, que frequentemente contém um alto grau de ruído e redundância (WANG et al., 2025). Uma etapa de pré-processamento é, portanto, interessante para a atenuação de ruídos (SILVA et al., 2023) e para a redução da dimensionalidade dos dados (GAO; HU; CHEN, 2024). Além de facilitar a extração de características eficazes para a classificação, essa abordagem melhora a eficiência computacional, resultando na aceleração do treinamento, da inferência e na redução do consumo de recursos computacionais, como a memória GPU (GAO; HU; CHEN, 2024).

A etapa de engenharia de características ou *feature engineering*, é frequentemente aplicada no processamento de sinais PEC (SILVA et al., 2021), cujo sinal PEC bruto pode estar contaminado por ruído e apresentar informações redundantes (SOPHIAN; TIAN; FAN, 2017). O objetivo dessa etapa é extrair informações úteis do sinal e reduzir sua dimensionalidade. Algumas características, também chamadas de atributos, extraídas do domínio do tempo incluem o valor de pico, o tempo até o pico e o ponto de cruzamento por zero (SOPHIAN; TIAN; FAN, 2017).

Transformada Discreta de Fourier

A Transformada de Fourier Discreta (DFT) é uma técnica utilizada para representar sinais de tempo discreto no domínio da frequência. Essa transformação permite identificar as componentes espectrais mais relevantes do sinal, explicitando características que não são diretamente observáveis na representação original no domínio do tempo (SILVA et al., 2021). No contexto da metodologia proposta para ENDS, a DFT é aplicada como um dos algoritmos para a extração de características relevantes da informação contida no domínio do tempo (SILVA et al., 2021). A DFT pode ser definida, conforme a Equação (1):

$$S[k] = \sum_{n=0}^{N-1} s[n]W_N^{nk}, \quad k = 0, 1, \dots, N-1 \quad (1)$$

em que,

$$W_N = e^{-j\frac{2\pi}{N}} \quad (2)$$

Dessa forma, $S[k]$ representa o sinal no domínio da frequência, $s[n]$ representa o sinal no domínio do tempo, N é o número total de amostras do sinal $s[n]$ e W_N^{nk} é o fator de rotação, que implementa a base de exponenciais complexas.

3.2.1 Redução de dimensionalidade

A redução de dimensionalidade é um processo de pré-processamento que consiste em transformar dados de um espaço de alta dimensão para um espaço de dimensão inferior (SILVA et al., 2023), mantendo suas características mais relevantes. Esta abordagem visa mitigar a alta dimensionalidade dos dados (SILVA et al., 2021), o que pode reduzir o custo computacional (CHEN et al., 2023) e acelerar o treinamento e a inferência dos modelos (GAO; HU; CHEN, 2024).

Análise de Componentes Principais

A Análise de Componentes Principais (PCA) é uma ferramenta clássica de análise estatística não supervisionada. A ideia central da técnica é transformar dados de alta dimensão em um novo espaço de baixa dimensão (ZHANG; ZHONG; DIAN, 2023). Isso é alcançado por meio da identificação de componentes principais (PCs), que são um conjunto de variáveis mutuamente ortogonais que representam as direções de variância máxima nos dados (GAO; HU; CHEN, 2024). O primeiro componente principal (PC1) captura a maior variância possível, e cada componente subsequente captura a maior variância restante (ZHANG; ZHONG; DIAN, 2023).

No contexto de análise de sinais, o PCA atua como uma ferramenta eficiente para a redução de ruído, pois os componentes principais de ordem superior (com menor variância) frequentemente capturam ruído e variações sutis, enquanto os primeiros componentes retêm a informação que melhor descrevem o conjunto original (GAO; HU; CHEN, 2024). A técnica demonstra preservar as características estatísticas críticas dos dados originais, como média, soma e valores de pico. A aplicação do PCA como um método de pré-processamento antes da modelagem pode, portanto, acelerar o treinamento e a inferência dos modelos de *machine learning* e reduzir o consumo de recursos computacionais como a memória da GPU (GAO; HU; CHEN, 2024).

Análise Discriminante Linear

A Análise Discriminante Linear (LDA) é um método estatístico supervisionado frequentemente utilizado, tanto para a separação linear de dados, quanto para a redução de dimensionalidade. O foco desta técnica é encontrar uma combinação linear de características (*features*) que maximize a razão da variância entre as classes pela variância interna de cada classe (SZOSTAK; WALKOWIAK; WŁODARCZYK, 2020).

Na prática, este método reduz a dimensionalidade do espaço de características, ao mesmo tempo que preserva a informação discriminatória da classe (SZOSTAK; WALKOWIAK; WŁODARCZYK, 2020). O número de dimensões resultantes é determinado pela Equação (3):

$$\text{Número de dimensões} = C - 1 \quad (3)$$

no qual C é a quantidade de classes.

3.3 Algoritmos de Aprendizado de Máquina

Machine Learning é um campo da ciência da computação focado no desenvolvimento de algoritmos que permitem a um sistema computacional aprender com dados (BISHOP, 2006). O propósito é dado pela descoberta automática de regularidades ou padrões nos dados, permitindo que o sistema tome ações, como fazer previsões. Esta capacidade de generalização, ou seja, de aplicar o conhecimento aprendido em um conjunto de dados de treinamento para fazer previsões precisas em dados novos e nunca vistos, é um dos principais objetivos da área (BISHOP, 2006).

Em problemas de aprendizado supervisionado, por exemplo, o foco recai sobre a estimativa de uma função que mapeia variáveis de entrada a variáveis de saída, buscando minimizar o erro entre a resposta predita e a real (HASTIE; TIBSHIRANI; FRIEDMAN, 2009).

3.3.1 Redes Neurais Artificiais

As Redes Neurais Artificiais (RNAs) são modelos de processamento de dados que simulam o funcionamento de um cérebro biológico por meio de elementos computacionais simples conhecidos como neurônios (SHRIFAN; AKBAR; ISA, 2019). Embora inspiradas na biologia, essas redes são estruturadas matematicamente como uma série de transformações funcionais projetadas para aprender padrões a partir de exemplos, ajustando seus parâmetros internos para realizar tarefas como regressão ou classificação (BISHOP, 2006). O neurônio artificial atua como a unidade fundamental de processamento da rede, cujo funcionamento consiste em receber sinais de entrada (x_j), que são multiplicados por pesos sinápticos (w_j) e somados a um termo de viés (b_j) (SILVA et al., 2021). A Equação (4) descreve matematicamente v_k , um potencial de ativação:

$$v_k = \sum_{j=1}^N w_{kj} x_j + b_k \quad (4)$$

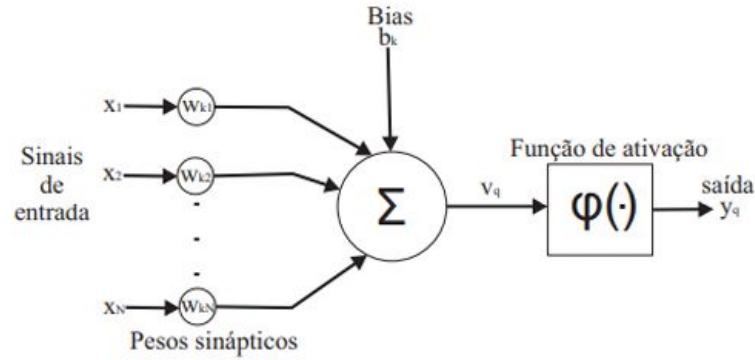
Assim, o resultado é passado por uma função de ativação não linear (φ) para produzir a saída do neurônio. Essa função de ativação não-linear pode possuir diversos formatos e comportamentos diferentes (ReLU, TanH, Sigmoid, Leaky ReLU, etc).

$$y_k = \varphi(v_k) \quad (5)$$

O resultado dessa combinação linear é então submetido a uma função de ativação, muitas vezes não linear, que determina a saída do neurônio. É essa não-linearidade que

permite à rede modelar relações complexas e extrair características latentes dos dados de entrada (FU et al., 2019). A Figura 2 representa a arquitetura de um Perceptron, um neurônio artificial de uma camada.

Figura 2 – Descrição de um Perceptron

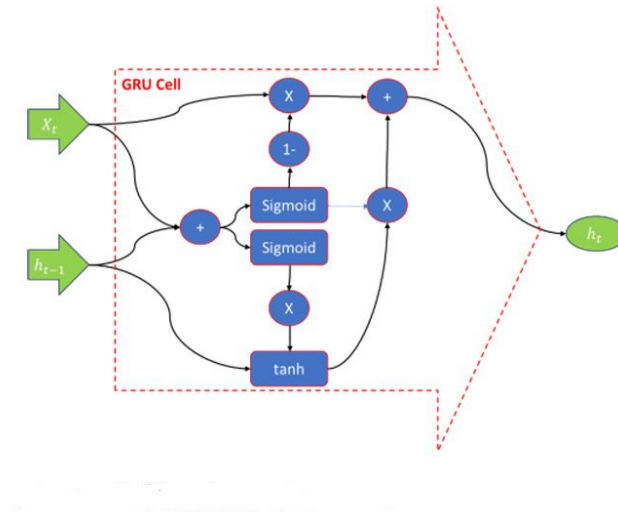


Fonte: Silva (2021)

Gated Recurrent Unit

As Redes Neurais Recorrentes (RNRs) caracterizam-se por possuírem conexões de *feedback*, o que lhes permite manter um estado interno de memória e processar dados sequenciais de forma dinâmica (BISHOP, 2006). O treinamento dessas estruturas é geralmente realizado por meio do desdobramento da rede ao longo do tempo, aplicando-se o algoritmo de *backpropagation* para o cálculo das derivadas. Contudo, uma limitação crítica dessas redes é a dificuldade prática em aprender dependências de longo prazo. Desta forma, à medida que o erro é propagado para camadas anteriores no tempo, os gradientes tendem a desaparecer ou explodir, impedindo que a rede retenha informações do início da sequência (BISHOP, 2006).

A Gated Recurrent Unit (GRU) é uma variação otimizada das RNRs, projetada para mitigar o problema do desvanecimento do gradiente em sequências longas (AMIRI et al., 2024). A arquitetura distingue-se pela introdução do mecanismo de controle de fluxo de informação denominado porta de atualização (*update gate*) e porta de redefinição (*reset gate*). Estas portas determinam quais informações passadas devem ser retidas ou descartadas, permitindo que a rede capture dependências temporais relevantes com maior eficiência computacional que as LSTMs (Long Short-Term Memory), devido à sua estrutura simplificada (AMIRI et al., 2024). Em tarefas de análise de séries temporais, a GRU processa a série histórica de entrada para extrair características latentes em seu estado oculto, que são subsequentemente transformadas para gerar as previsões do modelo (GAO; HU; CHEN, 2024), desta forma a Figura 3 representa uma célula GRU.

Figura 3 – Representação de uma célula GRU

Fonte: Amiri et al. (2024)

- x_t é a entrada do estado atual;
- h_{t-1} é o estado oculto anterior, traz consigo informações aprendidas dos passos anteriores;
- r_t é o *reset gate*, responsável por identificar o quão importante h_{t-1} é para a informação do passo atual; e por último
- z_t representa o *update gate*, ele que define se o passo h_t deve manter a informação de h_{t-1} ou ser atualizado por h_{t+1} .

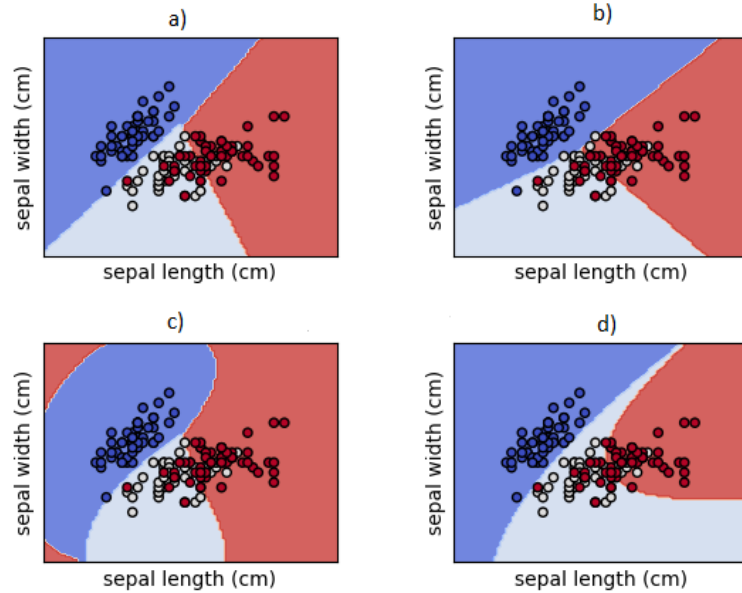
3.3.2 Algoritmos Baseados em Instância e Estatística

A Máquina de Vetores de Suporte (SVM) é uma técnica de aprendizado supervisionado fundamentada na teoria da aprendizagem estatística, muito utilizada para tarefas de classificação e regressão (BISHOP, 2006). O foco do algoritmo é determinar um limite de decisão ótimo, denominado hiperplano, que maximize a margem de separação entre as diferentes classes de dados (HASTIE; TIBSHIRANI; FRIEDMAN, 2009).

Uma propriedade distinta do SVM é a esparsidade da sua solução em que o hiperplano de decisão é definido inteiramente por um subconjunto de dados de treinamento, conhecido como vetores de suporte, que são os pontos situados mais próximos da fronteira de decisão (BISHOP, 2006). Para cenários em que os dados não são linearmente separáveis no espaço original, o SVM emprega funções de *kernel* (vide Figura 4). Esta técnica, conhecida como *kernel trick*, mapeia implicitamente os dados de entrada para um espaço de características de dimensão superior, no qual a separação linear se torna

possível, permitindo que o algoritmo modele relações não lineares complexas (SHRIFAN; AKBAR; ISA, 2019).

Figura 4 – Margens de decisão com diferentes kernels para SVM: a) e b) Linear, c) Função de base radial e d) Polinomial



Fonte: Adaptado de Scikit-learn (2025)

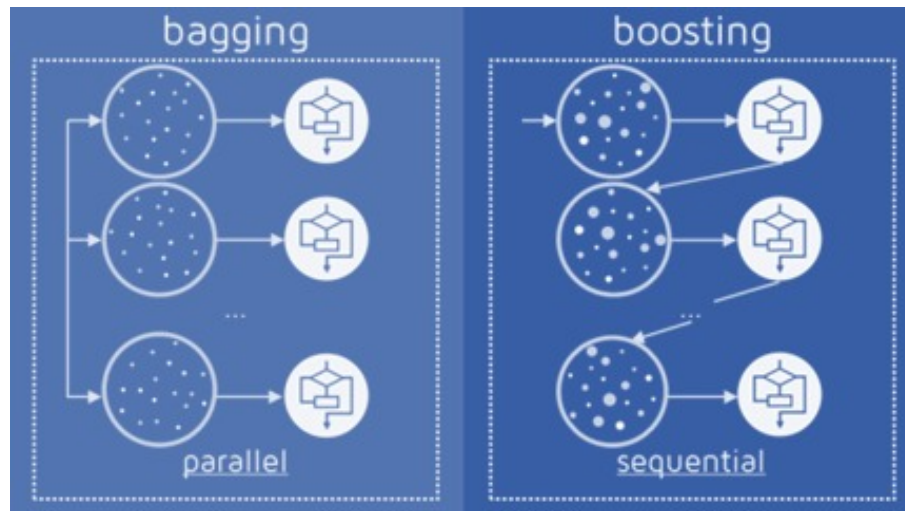
Já para o algoritmo K -Nearest Neighbors (KNN), tem-se um método de classificação não-paramétrico e baseado em memória (BISHOP, 2006), diferente de outros algoritmos que constroem um modelo interno explícito durante o treinamento, o KNN simplesmente armazena os dados de treinamento. A classificação de uma nova amostra é realizada identificando-se os K exemplos mais próximos no espaço de características e atribuindo a classe mais frequente entre eles, por meio de uma votação majoritária (HASTIE; TIBSHIRANI; FRIEDMAN, 2009). Embora conceitualmente simples, o KNN pode ser altamente eficaz, mas seu desempenho depende crucialmente da escolha adequada da métrica de distância (como a Euclidiana) e do valor do parâmetro K , que controla o equilíbrio entre o ruído e a suavização da fronteira de decisão (BISHOP, 2006).

3.3.3 Árvores de Decisão

As Árvores de Decisão particionam o espaço de características por meio de uma estrutura hierárquica de decisões sequenciais, mas tendem a apresentar instabilidade e alta variância (HASTIE; TIBSHIRANI; FRIEDMAN, 2009). Para mitigar essas limitações e aprimorar a capacidade preditiva, aplicam-se técnicas de *Ensemble Learning* como o *Bagging*, que reduz a variância ao combinar árvores treinadas em subconjuntos aleatórios

de dados (HASTIE; TIBSHIRANI; FRIEDMAN, 2009), e o *Boosting*, que constrói o modelo sequencialmente para corrigir os erros das iterações anteriores (BISHOP, 2006). A Figura 5 descreve as diferenças entre *Bagging* e *Boosting*.

Figura 5 – Comparativo entre *Bagging* e *Boosting*



Fonte: Adaptado de Terko et al. (2019)

CatBoost

O CatBoost é um algoritmo avançado que se fundamenta na técnica de *Gradient Boosting* aplicada a árvores de decisão (HASTIE; TIBSHIRANI; FRIEDMAN, 2009). A premissa central do *Boosting* é a construção de um comitê de modelos preditivos, em que múltiplos modelos simples — denominados aprendizes fracos — são combinados para formar um estimador robusto com alta precisão (BISHOP, 2006). Diferentemente de abordagens paralelas, o *Gradient Boosting* constrói o modelo de forma sequencial, cada nova árvore é adicionada especificamente para corrigir os erros residuais cometidos pelas árvores anteriores, ajustando-se na direção do gradiente negativo da função de perda (HASTIE; TIBSHIRANI; FRIEDMAN, 2009).

Random Forest

O algoritmo Random Forest é uma modificação substancial e aprimorada da técnica de *Bagging*, projetada para construir uma coleção de árvores de decisão decorrelacionadas (HASTIE; TIBSHIRANI; FRIEDMAN, 2009). Enquanto o *Bagging* padrão, constrói árvores usando todas as características disponíveis, o que pode levar a árvores muito semelhantes e correlacionadas, se houver preditores dominantes. O Random Forest introduz uma aleatoriedade adicional no processo de crescimento da árvore, acarretando em melhores resultados de predição (HASTIE; TIBSHIRANI; FRIEDMAN, 2009).

Para garantir a descorrelação, o algoritmo seleciona aleatoriamente um subconjunto de características para cada divisão, forçando diferenças estruturais entre as árvores. A agregação final das predições reduz a variância e aumenta a robustez do modelo em comparação com árvores individuais ou *Bagging* simples (HASTIE; TIBSHIRANI; FRIEDMAN, 2009).

DESENVOLVIMENTO

Este capítulo apresenta as metodologias, ferramentas e especificações de *Software* e *Hardware* aplicadas no desenvolvimento deste trabalho. Adicionalmente, são evidenciados os fluxos que envolvem coleta e processamento dos dados, implementação e otimização dos algoritmos de aprendizado e, por fim, avaliação e análise dos resultados e métricas obtidas.

4.1 Sistema de Aquisição de Dados

O banco de dados (*dataset*) utilizado para o treinamento, validação e testes dos modelos de *Machine Learning* consiste em arquivos estruturados (formato .csv), segmentados em duas categorias de sinal (medido e referência) e três classes de integridade estrutural: Sem Defeito (SD), Defeito Interno (DI) e Defeito Externo (DE). O conjunto totaliza 6 arquivos, no qual cada classe possui 100 amostras, sendo cada amostra composta por uma série temporal de 200 pontos (*steps*). Esse banco de dados foi disponibilizado pelo Doutor Manoel Messias Silva Junior - UFOB, podendo ser visualizado no repositório *Github*: <https://github.com/Beroradin/Dados-Sinais-PEC>

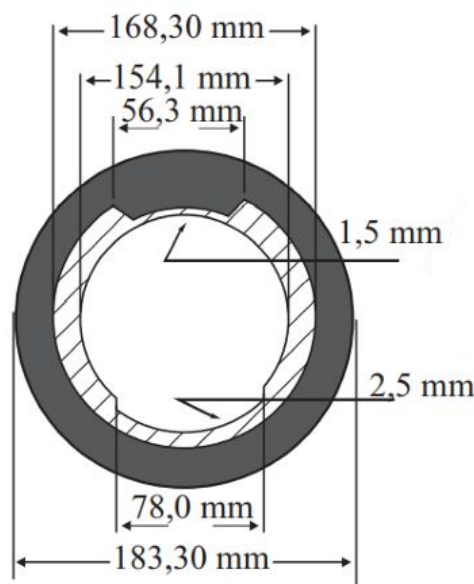
Os dados brutos foram obtidos experimentalmente por meio de um sistema embarcado de instrumentação dedicado, desenvolvido para inspeção não destrutiva por correntes parasitas pulsadas (PEC).

4.1.1 Caracterização do Objeto de Ensaio

Os ensaios experimentais foram conduzidos em um corpo de prova padronizado: um tubo de aço carbono revestido com material compósito, representando as condições de isolamento térmico típicas de tubulações industriais e seus eventuais defeitos por condições atípicas como corrosão. As dimensões geométricas do objeto de ensaio (SILVA, 2021), ilustradas na Figura 6, são:

- **Diâmetro Interno do Aço:** 56,3 mm;
- **Diâmetro Externo do Aço:** 78,0 mm;
- **Diâmetro Total (com revestimento):** 168,30 mm.

Figura 6 – Dimensões do corpo de prova utilizado



Fonte: Adaptado de Silva (2021)

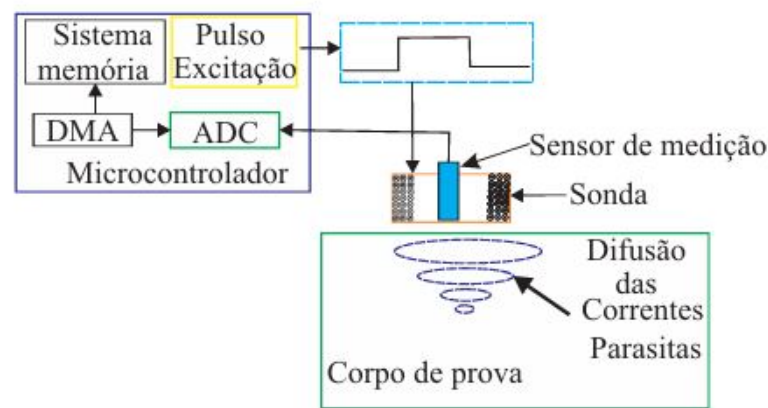
O conjunto de dados contempla três classes distintas para avaliação da integridade do material:

- **Sem Defeito (SD):** Região do tubo em condições íntegras, utilizada como linha de base;
- **Defeito Interno (DI):** Perda de massa simulada na parede interna do tubo, com profundidade aproximada de 1,5 mm;
- **Defeito Externo (DE):** Perda de massa simulada na parede externa do tubo, com profundidade aproximada de 2,5 mm.

4.1.2 Sistema Microprocessado de Aquisição

O sistema de aquisição de sinais PEC, cuja arquitetura conceitual é apresentada na Figura 7, baseia-se no kit STM32F4-Discovery, equipado com o microcontrolador STM32F429ZI (ARM Cortex-M4, 180 MHz). Este dispositivo utiliza FPU e instruções DSP para processamento eficiente, além de empregar DMA para transferir dados do ADC de 12 bits diretamente para a memória RAM, otimizando a aquisição contínua (SILVA, 2021).

Figura 7 – Diagrama Experimental Montado



Fonte: Adaptado de Silva (2021)

O sistema é composto por módulos analógicos e digitais integrados:

- **Sonda e Sensores:** A detecção do campo magnético utiliza sensores baseados no efeito de Magnetorresistência Gigante (GMR), modelo AAL002-02. A configuração da sonda é diferencial, empregando um sensor de referência (para captar o campo primário) e um sensor de medição (para captar a resposta do material induzido).
- **Módulo de Excitação:** A bobina da sonda foi excitada por um sinal de onda quadrada (pulso PEC) com amplitude de 4 V, frequência de 1 kHz e *duty cycle* de 40%. Para prover a corrente necessária, utilizou-se um circuito de potência tipo Ponte H (L298N). As especificações técnicas e parâmetros construtivos da sonda estão detalhados na Tabela 1.
- **Condicionamento de Sinais:** Antes da digitalização, os sinais provenientes dos sensores GMR foram amplificados por um amplificador de instrumentação **INA125** e filtrados por um filtro analógico passa-baixas (topologia Sallen-Key, Butterworth de 4ª ordem) com frequência de corte em 10 kHz, visando mitigar ruídos de alta frequência e evitar *aliasing* do sinal (SILVA, 2021).

- **Conversão Analógico-Digital:** A digitalização dos sinais condicionados foi realizada pelo conversor A/D (ADC) interno do microcontrolador STM32F4, configurado com resolução de 12 bits. A taxa de amostragem foi definida em 200 kHz, resultando em um intervalo entre amostras de 5 μ s.

Tabela 1 – Características geométricas e elétricas da sonda de inspeção.

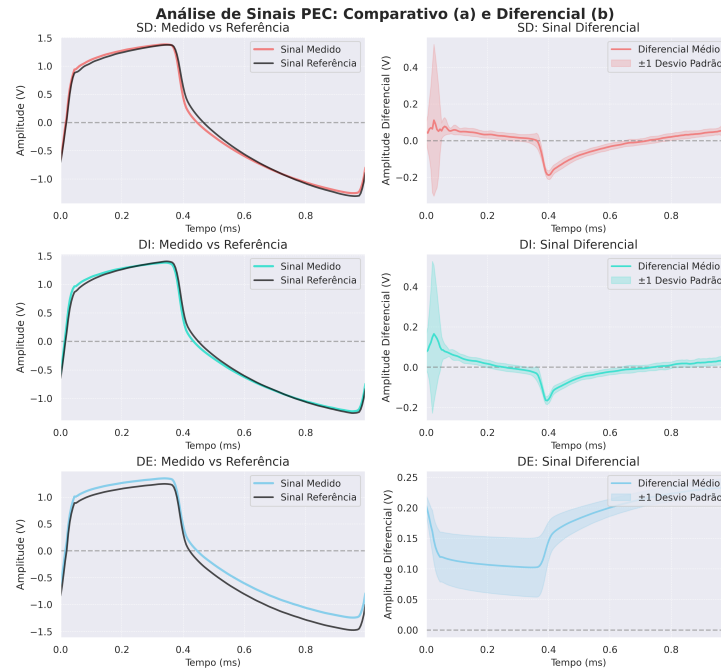
Diâmetro Interno (mm)	Diâmetro Externo (mm)	Altura da Bobina (mm)	Diâmetro do Fio (mm)	Resistência Elétrica (Ω)	Indutância (mH)
10	20	15	0,27	11	3,53

Fonte: Adaptado de Silva (2021)

Os dados brutos adquiridos consistem em séries temporais que representam a variação da amplitude da tensão induzida nos sensores GMR. A Figura 8 ilustra o comportamento desses sinais para as classes Sem Defeito (SD), Defeito Interno (DI) e Defeito Externo (DE).

A coluna da esquerda apresenta a comparação direta entre o sinal medido pelo sensor ativo e o sinal de referência (sensor passivo). A coluna da direita destaca o sinal diferencial médio, obtido pela subtração ($V_{dif} = V_{medido} - V_{ref}$), juntamente com a variabilidade estatística (± 1 desvio padrão) observada no conjunto de amostras. Esta técnica de subtração é fundamental para atenuar componentes comuns do sinal e evidenciar as perturbações específicas causadas por cada tipo de defeito.

Figura 8 – Assinaturas PEC no domínio do tempo: medido x referência (a) e Diferencial (b)



Fonte: Autoria Própria

4.2 Ferramentas de *software* e *hardware*

O desenvolvimento e a validação dos modelos computacionais propostos neste trabalho demandaram recursos de *hardware* capazes de lidar com o treinamento de algoritmos de aprendizado de máquina e o processamento de grandes volumes de dados. As especificações técnicas do computador pessoal utilizado durante o desenvolvimento desse trabalho estão detalhadas no Quadro 2.

A Unidade de Processamento Gráfico (GPU) dedicada foi fundamental para a aceleração dos cálculos matriciais, aproveitando sua arquitetura paralela otimizada para o processamento simultâneo de grandes volumes de dados, essencial no treinamento de redes neurais e outros algoritmos (LI et al., 2025). Para gerenciar essa unidade, utilizou-se a plataforma CUDA (Compute Unified Device Architecture) para computação paralela de propósito geral, aliada à biblioteca cuDNN (CUDA Deep Neural Network), que otimiza primitivas específicas de deep learning no *hardware*.

No que tange ao ambiente de *software*, todo o código foi desenvolvido na linguagem de programação Python, utilizando o Visual Studio Code (VSCode) como Ambiente de Desenvolvimento Integrado (IDE). A escolha da linguagem Python justifica-se por sua posição consolidada como a primeira preferência quando se trata de computação científica e aprendizado de máquina, impulsionada por uma vasta comunidade e um ecossistema robusto de ferramentas e bibliotecas (RASCHKA; PATTERSON; NOLET, 2020).

Quadro 2 – Especificações técnicas do *hardware* utilizado no desenvolvimento.

Componente	Especificação
Processador (CPU)	Intel Core i5-13420H
Memória RAM	16 GB
Placa de Vídeo (GPU)	NVIDIA GeForce RTX 3050 (6 GB VRAM)

Fonte: Autoria própria

As principais bibliotecas e ferramentas empregadas, identificadas a partir do código-fonte do projeto, estão categorizadas e descritas no Quadro 3.

Quadro 3 – Resumo das ferramentas de *software* e bibliotecas utilizadas.

Categoria	Ferramenta	Aplicação no Projeto
Linguagem e IDE	Python 3.11.5, VS Code	Desenvolvimento dos algoritmos e scripts de automação.
Processamento Numérico e Dados	NumPy 2.1.2, Pandas 2.3.3, SciPy 1.16.2	Manipulação de matrizes, estruturação de <i>datasets</i> (.csv), transformadas (FFT) e cálculos estatísticos.
Visualização de Dados	Matplotlib 3.10.7, Seaborn 0.13.1, Plotly 6.3.1	Geração de gráficos estáticos e interativos (curvas ROC, matrizes de confusão, gráficos 3D).
Machine Learning e Deep Learning	Scikit-learn 1.7.2, PyTorch 2.5.1, CatBoost 1.2.8	Implementação dos modelos classificadores (SVM, KNN, RF), redes neurais, pré-processamento (PCA, LDA) e métricas de avaliação.
Otimização	Optuna 4.5.0	Otimização Bayesiana automática de hiperparâmetros.
Aceleração de Hardware	CUDA 11.8.89, cuDNN	Bibliotecas da NVIDIA para paralelização de processos na GPU.
Utilitários	OS, Sys	Manipulação de arquivos do sistema.

Fonte: Autoria própria

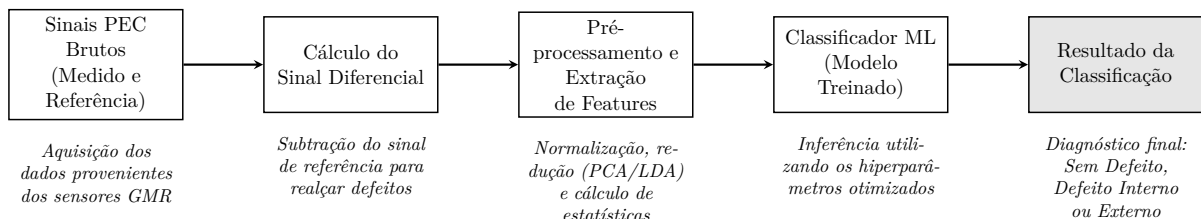
4.3 Sistema Proposto

O sistema proposto consiste em uma arquitetura completa para a classificação automática de defeitos em tubulações industriais, integrando a aquisição de sinais de Correntes Parasitas Pulsadas (PEC), técnicas de processamento e algoritmos de aprendizado de máquina. O estudo avalia o desempenho de modelos baseados em GRU, KNN, SVM, Catboost e Random Forest (RF), aplicados tanto isoladamente quanto em combinação com as técnicas de redução de dimensionalidade PCA e LDA¹. O fluxo de processamento foi desenhado para maximizar a extração de informações relevantes e garantir a robustez do diagnóstico, mitigando problemas comuns como ruído, desbalanceamento de dados e variância/viés dentro do *dataset*.

A Figura 9 ilustra o fluxo simplificado de operação do sistema, desde a captura do sinal bruto até a classificação final. Enquanto a Figura 10 demonstra o fluxograma de todo o *pipeline* do sistema. O processo inicia-se com a aquisição dos sinais diferenciais, obtidos pela subtração da resposta do sensor de referência em relação ao sensor de medição. Assim, os dados passam por etapas de pré-processamento e extração de características antes de alimentar o classificador inteligente, que foi otimizado para distinguir entre condições de integridade estrutural (Sem Defeito, Defeito Interno e Defeito Externo).

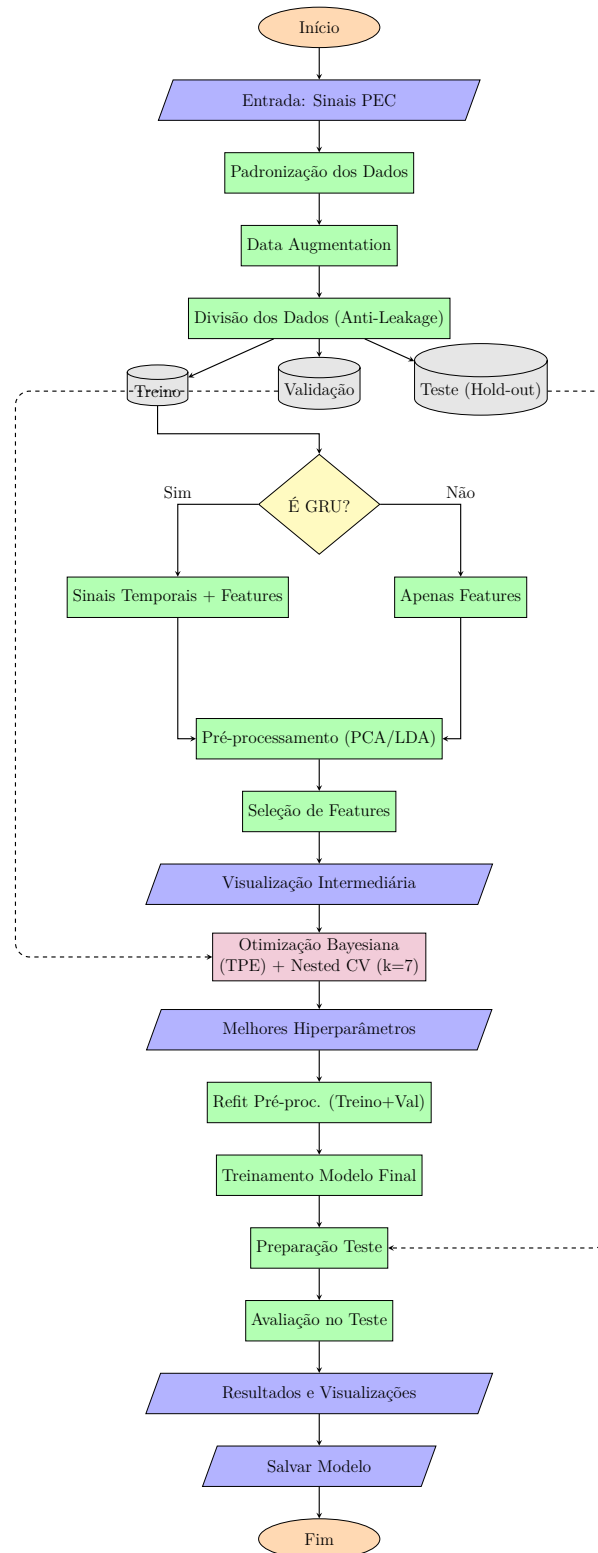
Para a validação estatística do método, adotou-se uma estratégia de particionamento dos dados na proporção aproximada de 70% para treinamento, 15% para validação e 15% para teste. Ressalta-se que o conjunto de validação foi gerenciado dinamicamente via validação cruzada (*cross-validation*) para evitar o vazamento de dados (*data leakage*) durante o ajuste fino dos hiperparâmetros, enquanto o conjunto de teste (*hold-out*) permaneceu isolado para a avaliação definitiva do modelo.

Figura 9 – Fluxograma simplificado do processo de classificação de sinais PEC.



Fonte: Autoria Própria

¹ Combinações utilizadas: GRU, GRU+PCA, GRU+LDA, KNN+PCA, KNN+LDA, SVM+PCA, SVM+LDA, Catboost+PCA, Catboost+LDA, RF+PCA e RF+LDA.

Figura 10 – Fluxograma do pipeline de processamento, otimização e avaliação.

Fonte: Autoria Própria

4.3.1 Técnicas de Aumentação de Dados

Devido à limitação inerente à obtenção de grandes volumes de dados experimentais, tanto em termos de quantidade quanto de qualidade, em ensaios destrutivos e não destrutivos, essenciais para treinamentos de modelos ML (AIS Field, 2024), este trabalho implementa uma estratégia de aumento de dados (*Data Augmentation*). Essa abordagem visa enriquecer o conjunto de treinamento e mitigar o *overfitting*, fenômeno no qual os classificadores aprendem padrões irrelevantes ou ruídos, prejudicando sua capacidade de generalização para dados não vistos. O objetivo é gerar variações sintéticas realistas que simulem condições adversas do ambiente industrial, preservando a natureza fundamental do sinal PEC.

Para cada sinal original adquirido, foram geradas novas amostras aplicando-se três tipos de transformações:

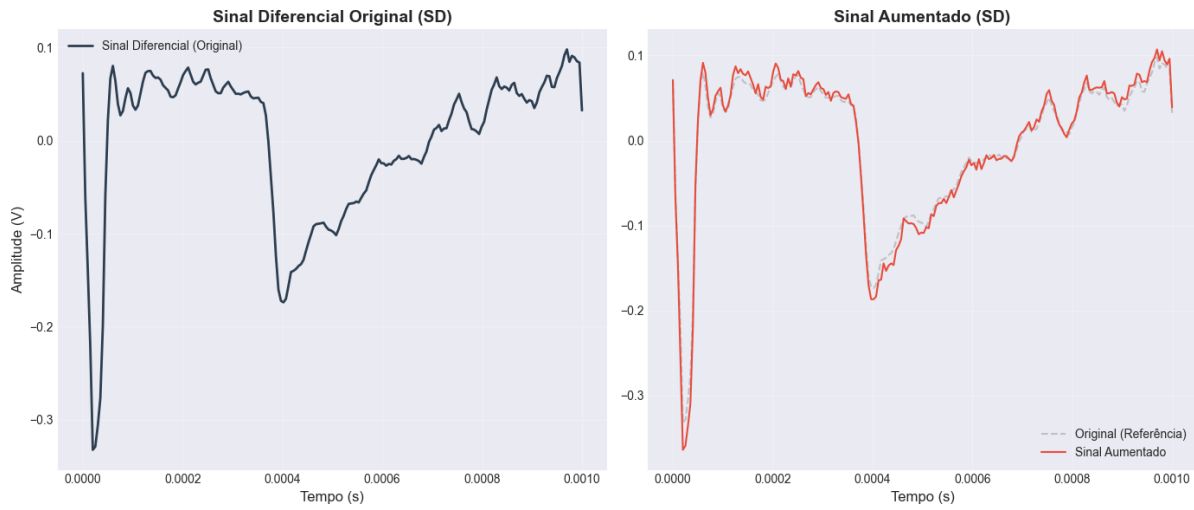
- **Ruído Gaussiano Aditivo:** Simula o ruído térmico e eletrônico inerente aos sensores GMR e aos circuitos de condicionamento, valores de ruído variaram de 1% a 3% do sinal PEC.
- **Jitter Temporal (Deslocamento):** Introduz pequenos deslocamentos circulares no vetor de tempo, simulando imprecisões de sincronização ou atrasos no disparo da excitação. O deslocamento foi aplicado aleatoriamente dentro de um intervalo de $\pm 0,5\%$ em relação ao comprimento total do vetor de amostras.
- **Variação de Amplitude (Escala):** Aplica fatores de ganho multiplicativos para simular variações na distância de *lift-off* ou na impedância do material. Os fatores de escala foram amostrados uniformemente no intervalo de $[0,95; 1,05]$, representando uma variação de amplitude de $\pm 5\%$ em relação ao sinal original.

Essa abordagem quadruplicou o tamanho efetivo do dataset disponível para treinamento, aumentando de 300 amostras para 1200 amostras, elevando a capacidade de generalização dos classificadores.

Visualização e Validação da Aumentação de Dados

A eficácia da estratégia de *Data Augmentation* é evidenciada visualmente na Figura 11, que confronta o sinal diferencial original de um defeito (SD) com sua contraparte sintética. Nota-se que a topologia fundamental do sinal permanece inalterada, preservando a assinatura física do defeito, ao passo que a introdução controlada de ruído e distorções temporais simula a variabilidade típica de ambientes industriais. Essa característica é

Figura 11 – Comparação entre o sinal diferencial original (esquerda) e o sinal aumentado sinteticamente (direita) com aplicação de ruído, *jitter* e escala.



Fonte: Autoria Própria

fundamental para evitar a memorização de formas de onda ideais, forçando o modelo a desenvolver uma capacidade de generalização robusta frente às incertezas de medição.

A análise estrutural do conjunto de dados expandido para 1200 amostras valida estatisticamente o método de geração. Observou-se uma redução drástica no Número de Condicionamento (Kappa) da matriz de correlação, que decaiu de $1,42 \times 10^{12}$ nos dados originais para $4,40 \times 10^6$ no conjunto aumentado. Essa alteração de seis ordens de grandeza indica que a inclusão de dados sintéticos mitigou a multicolinearidade excessiva e a redundância quase perfeita das amostras originais. O novo patamar de condicionamento sugere que o *dataset* adquiriu uma variância estatística adequada, necessária para a estabilidade numérica dos algoritmos, sem descaracterizar as correlações físicas inerentes ao fenômeno de correntes parasitas.

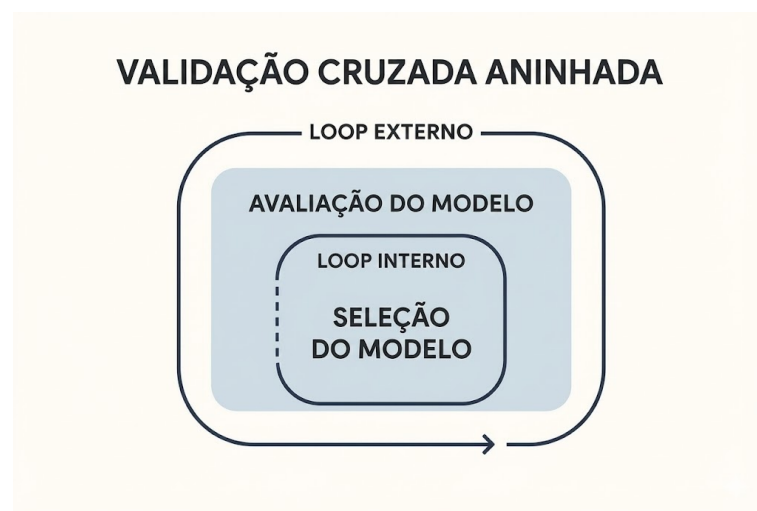
De forma análoga para a LDA, a avaliação das matrizes de dispersão demonstra o impacto realista da aumentação na estrutura das classes. Como reflexo direto da adição de ruído e variações de escala, a dispersão intraclasse (S_w) aumentou aproximadamente quatro vezes ($1,24 \times 10^4$ para $4,97 \times 10^4$), proporcional ao fator de expansão do banco de dados. Esse aumento na dispersão interna resultou em uma redução natural do Score de Fisher (J) de aproximadamente 1179 para 284,72. Essa queda no índice J não representa uma perda de qualidade significativa, mas sim um aumento de complexidade na separação das classes, no qual aumenta a robustez dos modelos. Mesmo com essa redução, o valor final permanece suficientemente elevado, com J muito maior que 1, para garantir que as fronteiras de decisão entre os tipos de defeitos se mantenham bem definidas e linearmente separáveis.

4.3.2 Otimização Bayesiana e Validação Cruzada

Para garantir a seleção dos melhores hiperparâmetros para cada modelo de aprendizado de máquina, adotou-se uma metodologia rigorosa de Otimização Bayesiana utilizando o estimador TPE (Tree-structured Parzen Estimator). Diferente da busca em grade (Grid Search), que é computacionalmente custosa e ineficiente, ou da Random Search, que nem sempre pode alcançar o *fine tuning* mais adequado de um hiperparâmetro, a otimização bayesiana constrói um modelo probabilístico da função objetivo, permitindo explorar o espaço de hiperparâmetros de forma inteligente e focar nas regiões mais promissoras, podando (em inglês, *pruning*) caminhos de exploração que não alcancem bons resultados iniciais. A avaliação de desempenho e a robustez dos modelos foram asseguradas por meio da técnica de Validação Cruzada Aninhada (Nested Cross-Validation), com $k = 7$ folds. Este método, detalhado na Figura 12, consiste em dois laços de repetição:

- Loop Interno: Utilizado exclusivamente para a otimização dos hiperparâmetros, prevenindo o vazamento de dados (*data leakage*) do conjunto de teste para a etapa de seleção do modelo.
- Loop Externo: Responsável pela avaliação imparcial da performance do modelo otimizado em dados nunca vistos, posteriormente utilizado para testes estatísticos.

Figura 12 – Loops Interno e Externos da Validação Cruzada Aninhada



Fonte: Autoria Própria

Espaço de Busca e Configuração dos Hiperparâmetros

A otimização dos modelos via TPE explorou espaços de busca definidos especificamente para as características intrínsecas de cada algoritmo. O objetivo foi cobrir

uma ampla gama de configurações, desde a arquitetura de redes neurais até penalidades de regularização em métodos de *ensemble*, garantindo que o potencial máximo de cada classificador fosse avaliado.

Para a rede neural GRU, a busca priorizou a definição topológica (número de camadas e neurônios) e a estabilidade do treinamento (taxa de aprendizado e otimizadores). No KNN e SVM, o foco recaiu sobre as métricas de distância e a definição das fronteiras de decisão (kernels e regularização). Já para os modelos baseados em árvores, Random Forest e CatBoost, os hiperparâmetros foram selecionados para controlar a complexidade das árvores e mitigar o *overfitting*, ajustando profundidade, número de estimadores e taxas de regularização.

O Quadro 4 detalha o intervalo completo de busca, enquanto o Quadro 5 apresenta a configuração final vencedora para cada combinação de modelo e técnica de redução de dimensionalidade, resultante do processo de otimização bayesiana.

Quadro 4 – Intervalos de busca e configurações dos hiperparâmetros para otimização bayesiana.

Modelo	Estrutura & Arquitetura	Otimização & Aprendizado	Regularização & Controle	Outros Parâmetros
GRU	Camada Oculta: 32–256 Camadas: 1–3	LR: 10^{-4} a 10^{-2} Opt: AdamW, SGD	Dropout: 0,1–0,4 Weight Decay: 10^{-5}	Batch: 16, 32, 64
KNN	Vizinhos (k): 1 a 40	Métrica: Euclidiana, Manhattan, Cosseno	—	Pesos: Uniforme, Distância
SVM	Kernel: RBF, Linear, Poly, Sigmoid	—	C (Penalidade): 10^{-3} a 10^3	Pesos de Classe: None, Balanced
Random Forest	Estimadores: 50–500 Profundidade: 5–50	Critério: Gini, Entropia	Min. Split: 2–20 Min. Leaf: 1–10	Max Features: Sqrt, Log2
CatBoost	Iterações: 100–600 Profundidade: 4–8	LR: 0,005 a 0,3 Boosting: Ordered/Plain	L2 Leaf Reg: 0,01 a 20,0	Rand. Strength: 0–5 Border Count: 64–255

Fonte: Autoria Própria

4.4 Metodologia de Avaliação dos Resultados

Para garantir a robustez e a confiabilidade dos modelos de aprendizado de máquina desenvolvidos, a fase de avaliação de desempenho é essencial. Conforme detalhado nas seções anteriores, os dados foram divididos em três subconjuntos: treinamento (70%), validação (15%, gerenciado via *nested cross-validation* para otimização) e teste (15%, mantido isolado para a avaliação final). A performance real dos classificadores é determinada nesta fase de teste, utilizando dados nunca vistos pelos modelos. Para mensurar quantitativamente a eficácia das Redes Neurais e dos demais algoritmos, foram utilizadas métricas derivadas da Matriz de Confusão, além da análise de curvas características de operação e testes de hipótese estatística para comparação entre modelos.

Quadro 5 – Configurações finais dos hiperparâmetros após a otimização bayesiana (TPE).

Modelo	Redução	Configuração Otimizada (Principais Parâmetros)
GRU	Puro	<i>Hidden</i> : 128, <i>Layers</i> : 2, <i>Drop</i> : 0.15, LR: $9,6 \times 10^{-4}$, <i>Opt</i> : AdamW
	+ PCA	<i>Hidden</i> : 32, <i>Layers</i> : 1, <i>Drop</i> : 0.1, LR: $1,7 \times 10^{-3}$, <i>Batch</i> : 32
	+ LDA	<i>Hidden</i> : 64, <i>Layers</i> : 2, <i>Drop</i> : 0.2, LR: $1,7 \times 10^{-3}$, <i>Batch</i> : 16
KNN	+ PCA	$k = 4$, Pesos: <i>Distance</i> , Métrica: Manhattan
	+ LDA	$k = 15$, Pesos: <i>Uniform</i> , Métrica: Manhattan
SVM	+ PCA	$C \approx 0,41$, <i>Kernel</i> : Poly, <i>Gamma</i> : 0.16, <i>Coef0</i> : 0.98
	+ LDA	$C \approx 0,44$, <i>Kernel</i> : RBF, <i>Gamma</i> : 0.54, <i>Class Weight</i> : Balanced
CatBoost	+ PCA	<i>Iter</i> : 443, <i>Depth</i> : 4, LR: 0.23, <i>Type</i> : Plain, L2: 6.1
	+ LDA	<i>Iter</i> : 370, <i>Depth</i> : 7, LR: 0.068, <i>Type</i> : Ordered, L2: 16.6
Random Forest	+ PCA	<i>Est</i> : 179, <i>Depth</i> : 32, <i>Crit</i> : Gini, <i>Feat</i> : Sqrt
	+ LDA	<i>Est</i> : 445, <i>Depth</i> : 29, <i>Crit</i> : Entropy, <i>Feat</i> : 0.7

Fonte: Autoria Própria

4.4.1 Matriz de Confusão e Métricas Derivadas

A matriz de confusão é uma ferramenta fundamental para a avaliação de modelos de classificação supervisionada. Ela consiste em uma tabela quadrada de dimensão $N \times N$, no qual N é o número de classes (neste trabalho, $N = 3$). As linhas representam as classes reais (rótulos verdadeiros) e as colunas representam as classes preditas pelo modelo. A partir desta disposição, definem-se quatro conceitos básicos para cada classe:

- **Verdadeiro Positivo (TP):** Para uma classe de interesse, corresponde às amostras corretamente classificadas como pertencentes a essa classe, situadas na diagonal principal da matriz de confusão quando analisadas sob a abordagem *one-vs-rest*.
- **Verdadeiro Negativo (TN):** Amostras corretamente identificadas como não pertencentes à classe de interesse.
- **Falso Positivo (FP):** Também conhecido como Erro Tipo I, ocorre quando o modelo classifica incorretamente uma amostra como pertencente à classe de interesse, quando na realidade ela pertence a outra classe.
- **Falso Negativo (FN):** Conhecido como Erro Tipo II, ocorre quando o modelo falha em identificar uma amostra pertencente à classe de interesse, classificando-a como outra classe.

Com base nesses valores, foram calculadas as seguintes métricas de desempenho:

1. **Acurácia (Accuracy):** Representa a proporção global de acertos do modelo em relação ao total de amostras. É a métrica mais intuitiva, calculada pela razão entre

a soma da diagonal principal e o número total de amostras:

$$Acuracia = \frac{TP + TN}{TP + TN + FP + FN} \quad (6)$$

2. **Precisão (Precision):** Indica a qualidade das predições positivas. Quantifica a proporção de amostras classificadas como positivas que realmente são positivas. Uma precisão de 1,0 indica que não houve falsos positivos.

$$Precisao = \frac{TP}{TP + FP} \quad (7)$$

3. **Recall (Sensibilidade):** Mede a capacidade do modelo em detectar todas as amostras positivas existentes. É crucial em cenários em que a não detecção de um defeito (Falso Negativo) é crítica.

$$Recall = \frac{TP}{TP + FN} \quad (8)$$

4. **F1-Score:** É a média harmônica entre a precisão e o recall. Esta métrica é particularmente útil quando se busca um equilíbrio entre a qualidade da predição e a capacidade de detecção, penalizando modelos que privilegiam uma métrica em detrimento da outra.

$$F1-Score = 2 \cdot \frac{Precisao \cdot Recall}{Precisao + Recall} \quad (9)$$

4.4.2 Curva ROC e Área sob a Curva (AUC)

Para uma avaliação independente do limiar de decisão, utilizou-se a curva ROC (*Receiver Operating Characteristic*). Este gráfico plota a Taxa de Verdadeiros Positivos (TPR ou Recall) em função da Taxa de Falsos Positivos (FPR) para diversos limiares de classificação. Uma curva ROC ideal aproxima-se do canto superior esquerdo do gráfico (TPR=1, FPR=0). Para quantificar esse desempenho em um único valor escalar, calcula-se a Área sob a Curva (AUC - *Area Under the Curve*). O valor da AUC varia de 0,5 (classificação aleatória) a 1,0 (classificador perfeito). No contexto multiclasse deste trabalho, a AUC foi calculada utilizando a estratégia *One-vs-Rest*, calculando a média das áreas para cada classe.

4.4.3 Validação Estatística: ANOVA e Teste de Tukey HSD

A simples comparação direta das médias de acurácia pode não ser suficiente para afirmar, com confiança estatística, que um modelo é superior a outro. Para mitigar a

incerteza e validar as diferenças de desempenho observadas entre os diversos algoritmos testados (GRU, CatBoost, etc.), aplicou-se a Análise de Variância (ANOVA) *one-way*. A ANOVA testa a hipótese nula H_0 de que as médias de desempenho de todos os modelos são estatisticamente iguais. Caso a ANOVA indique uma diferença significativa ($p < 0,05$), rejeita-se H_0 . No entanto, a ANOVA não indica quais modelos diferem entre si, sendo necessária uma análise posterior para identificar os subgrupos distintos (AGBANGBA et al., 2024). Para identificar essas diferenças par-a-par, utilizou-se o teste *post-hoc* de Tukey HSD (*Honest Significant Difference*). Este teste compara todas as combinações possíveis de modelos e ajusta os valores-p para múltiplas comparações, permitindo afirmar se a superioridade de um modelo sobre outro é estatisticamente significativa ou fruto da aleatoriedade dos dados.

RESULTADOS

Este capítulo apresenta a análise do desempenho do sistema proposto. Inicialmente, discute-se a relevância das características extraídas e o impacto das técnicas de redução de dimensionalidade na separabilidade das classes. Em seguida, são comparadas as métricas de performance dos classificadores otimizados (GRU, KNN, SVM, Random Forest e CatBoost) e, por fim, apresenta-se a validação estatística dos resultados para atestar a robustez do método.

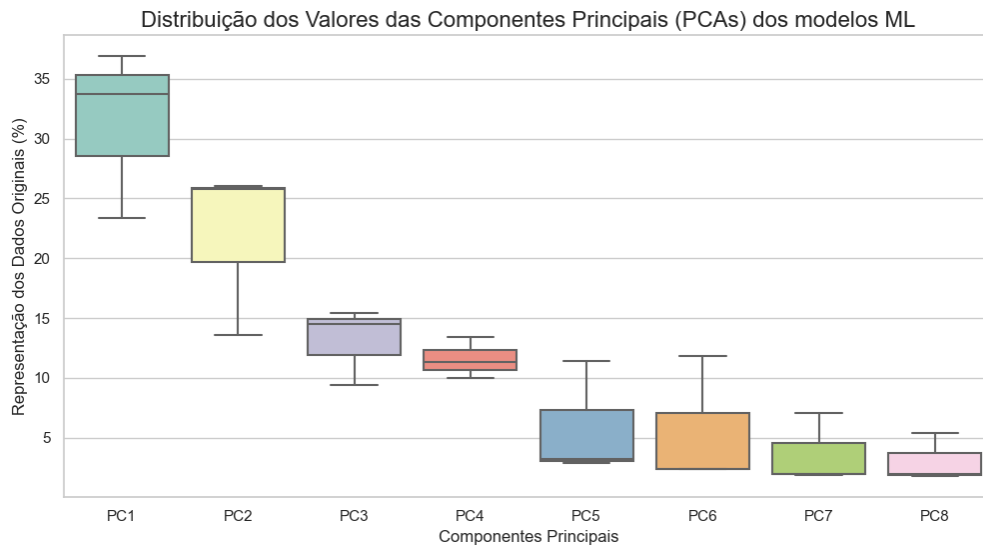
5.1 Extração de Características e Redução de Dimensionalidade

Após a etapa de pré-processamento, foram extraídas características estatísticas e temporais dos sinais para compor o vetor de entrada dos classificadores. A análise de importância das *features* ao longo dos testes revelou que os atributos mais determinantes para a distinção dos defeitos foram, em ordem decrescente de relevância: o valor mínimo do sinal, o tempo de decaimento, o centróide espectral, o valor RMS (*Root Mean Square*) e o valor de pico. Essa hierarquia faz sentido físico, uma vez que o valor mínimo e o decaimento estão diretamente associados à profundidade do defeito e a constante de tempo das correntes parasitas, respectivamente.

Para lidar com a dimensionalidade e otimizar o custo computacional, as técnicas PCA e LDA foram avaliadas comparativamente. A Figura 13 ilustra a distribuição dos valores das oito primeiras Componentes Principais (PCs) extraídas do conjunto de dados. Observa-se que a primeira componente (PC1) concentra a maior parte da variância, apresentando a maior amplitude e dispersão de valores (mediana próxima a 34), seguida

por um decaimento progressivo nas componentes subsequentes. As componentes a partir da PC5 apresentam valores significativamente menores e distribuições mais compactas, indicando uma menor contribuição para a variância total do sistema.

Figura 13 – Distribuição dos valores das Componentes Principais (PCAs) utilizadas nos modelos de ML.



Fonte: Autoria Própria

A análise estrutural demonstrou a superioridade do LDA para a organização deste problema específico. Em contraste com a ampla dispersão observada nas componentes principais, as componentes discriminantes (LD1 e LD2) apresentaram desvios padrão significativamente reduzidos ($\sim \pm 0,002\%$) em comparação ao PCA, bem como um *Silhouette Score* superior a 2 vezes o obtido pela técnica não supervisionada (0,76 contra 0,31), o que evidencia uma maior compactação e separabilidade dos dados dentro de cada classe. Além disso, a primeira componente discriminante (LD1) concentrou em média, sozinha, 97,1% de todo o poder de separação do modelo. Isso comprova que a abordagem supervisionada do LDA, ao maximizar a razão entre a variância entre classes e a variância intraclasse (S_b/S_w), é mais eficiente para a distinção dos defeitos PEC do que a estratégia de maximização de variância global proposta pelo PCA, fato que se confirma posteriormente nos resultados de classificação e significância estatística.

5.2 Desempenho dos Classificadores

A Tabela 2 resume o desempenho dos modelos avaliados no conjunto de teste, utilizando as métricas de Acurácia e F1-Score. Observa-se uma tendência geral de melhoria no desempenho quando os classificadores são associados à técnica LDA, corroborando a

análise estrutural prévia.

O modelo baseado em redes neurais recorrentes (GRU) apresentou o desempenho mais consistente entre todas as configurações, atingindo 94% de acurácia quando combinado com LDA. Destaca-se também a recuperação de performance do algoritmo CatBoost: enquanto sua versão com PCA sofreu uma degradação severa (77% de acurácia), a combinação CatBoost+LDA alcançou 92%, posicionando-se entre os melhores avaliados.

Tabela 2 – Métricas de desempenho dos modelos no conjunto de teste.

Modelo Avaliado	Acurácia (Acc)	Precisão (Pond.)	Recall (Pond.)	F1-Score (Média)
GRU	0,90	0,90	0,90	0,89
GRU + PCA	0,93	0,93	0,92	0,93
GRU + LDA	0,94	0,94	0,94	0,94
KNN + PCA	0,82	0,82	0,81	0,81
KNN + LDA	0,84	0,84	0,84	0,84
SVM + PCA	0,81	0,82	0,81	0,81
SVM + LDA	0,83	0,83	0,83	0,83
CatBoost + PCA	0,77	0,78	0,77	0,76
CatBoost + LDA	0,92	0,92	0,92	0,92
RF + PCA	0,77	0,78	0,77	0,77
RF + LDA	0,83	0,83	0,82	0,82

Fonte: Autoria Própria

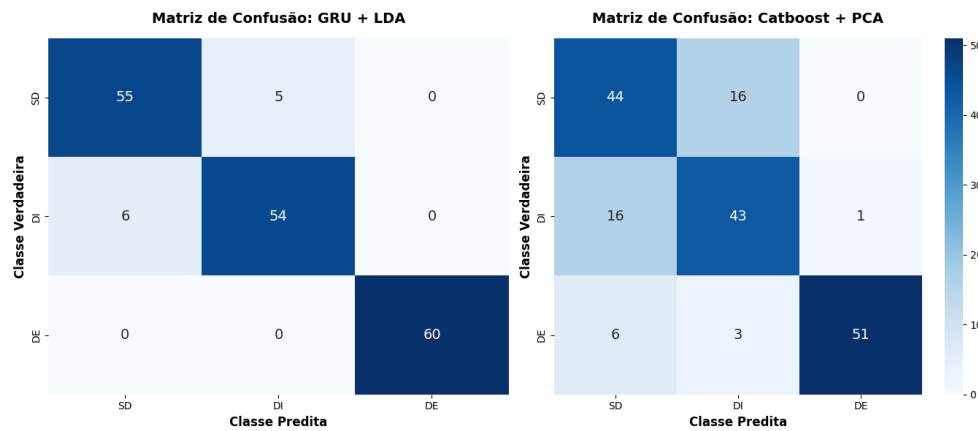
Análise de Erros via Matriz de Confusão

Para compreender detalhadamente o comportamento dos classificadores e identificar quais classes apresentam maior dificuldade de distinção, a Figura 14 compara as matrizes de confusão do modelo com melhor desempenho geral (GRU+LDA) e de um modelo com desempenho inferior representativo (CatBoost+PCA).

A matriz à esquerda (GRU+LDA) evidencia a robustez da rede neural recorrente quando auxiliada pela análise discriminante: o modelo obteve uma taxa de acerto perfeita (100%) na identificação de Defeitos Externos (DE) e cometeu erros mínimos na distinção entre Sem Defeito (SD) e Defeito Interno (DI), com apenas 11 amostras classificadas incorretamente no total.

Em contrapartida, a matriz à direita (CatBoost+PCA) ilustra as limitações de um classificador baseado em árvores quando operando sobre dados ruidosos sem a projeção discriminante adequada. Nota-se uma confusão significativa entre as classes SD e DI (32 erros cruzados), além de uma dificuldade moderada em identificar corretamente o Defeito

Figura 14 – Comparativo das Matrizes de Confusão: GRU+LDA (esquerda) x CatBoost+PCA (direita).

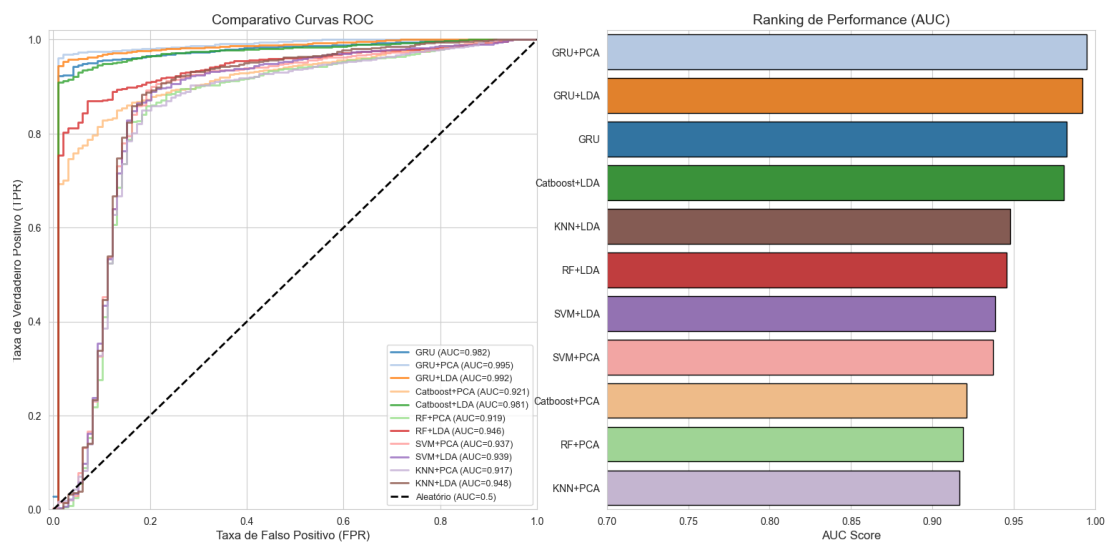


Fonte: Autoria Própria

Externo, com 9 amostras sendo confundidas com outras classes. Essa comparação visual reforça que o LDA não apenas melhora a acurácia global, mas especificamente "limpa" as fronteiras de decisão entre defeitos morfologicamente similares, o que é importante para a confiabilidade da inspeção industrial.

Para uma avaliação independente do limiar de decisão, a Figura 15 apresenta as curvas ROC e os respectivos valores de Área Sob a Curva (AUC). Os modelos GRU+PCA e GRU+LDA apresentaram curvas quase ideais, com AUC de 0,995 e 0,992, respectivamente, indicando uma excelente capacidade de separação entre as classes de defeito e integridade. O gráfico de ranking à direita reforça a dominância da família GRU e a eficácia da combinação com LDA para os modelos de *boosting*.

Figura 15 – Comparativo das Curvas ROC (esquerda) e Ranking de Performance AUC (direita).

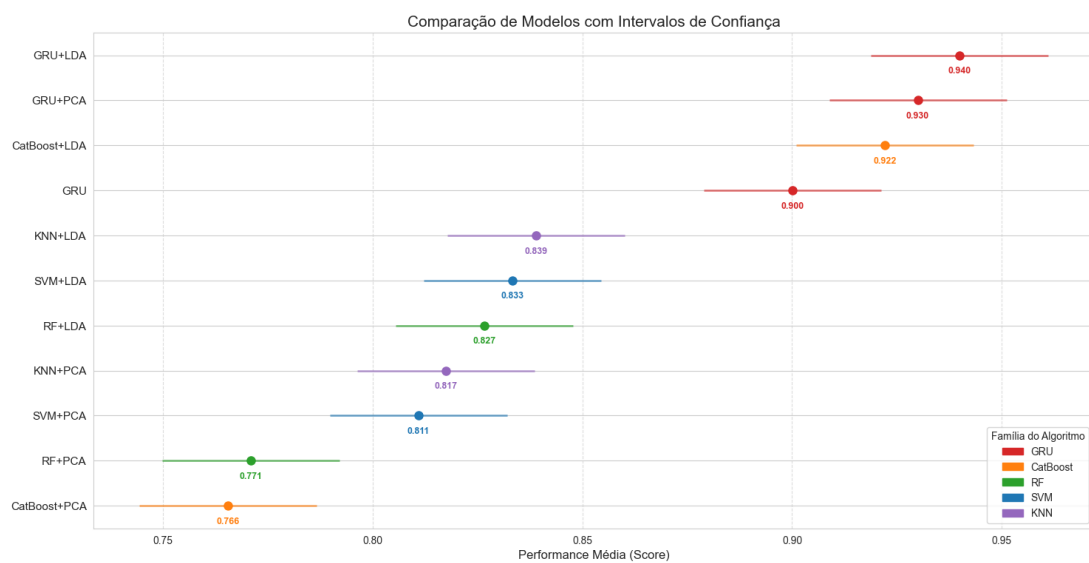


Fonte: Autoria Própria

5.3 Validação Estatística

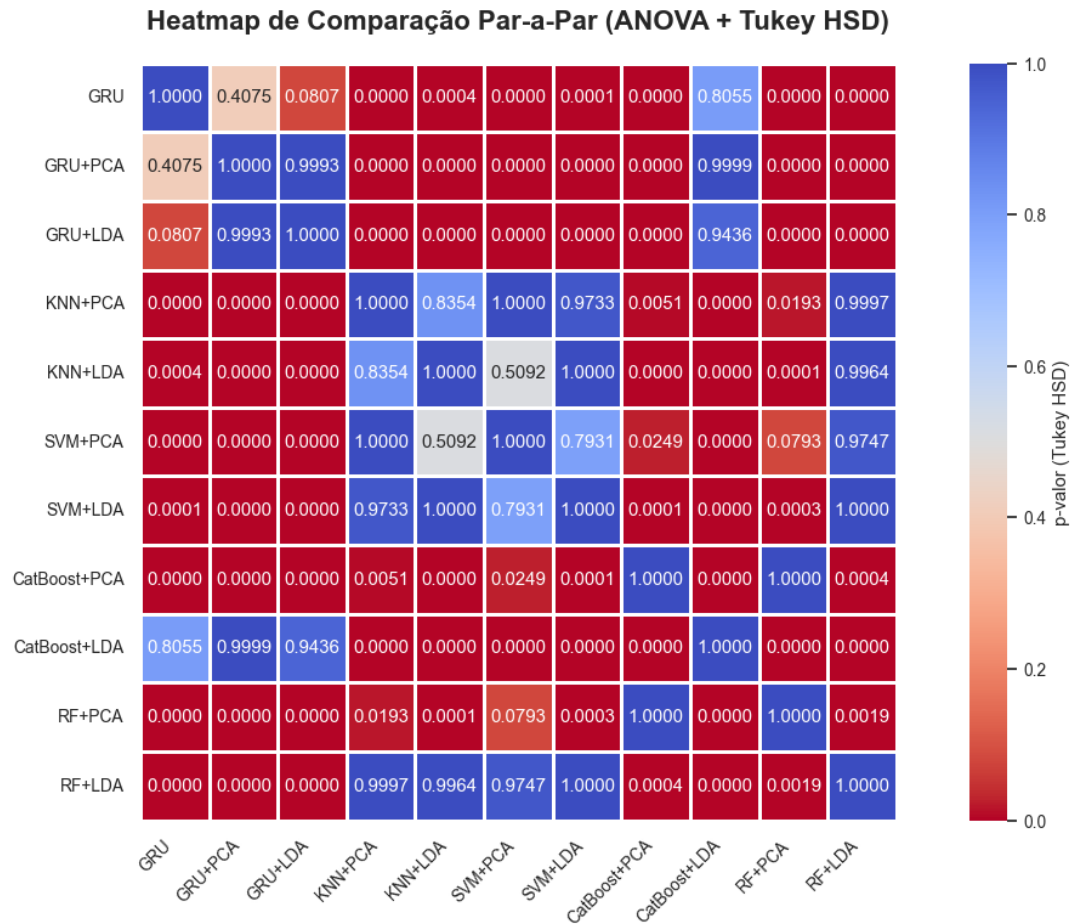
Embora a comparação direta de médias indique a superioridade numérica da GRU+LDA, a validação estatística é necessária para confirmar se essas diferenças são significativas. A Figura 16 apresenta os intervalos de confiança para a performance média dos modelos. Nota-se uma sobreposição entre os intervalos dos modelos de topo (GRU+LDA, GRU+PCA e CatBoost+LDA), sugerindo que eles podem ser estatisticamente equivalentes.

Figura 16 – Comparação de Modelos com Intervalos de Confiança.



Fonte: Autoria Própria

Para verificar as diferenças par-a-par, aplicou-se o teste *post-hoc* de Tukey HSD após a rejeição da hipótese nula na ANOVA. O mapa de calor (*heatmap*) dos p-valores é apresentado na Figura 17.

Figura 17 – Heatmap de Comparação Par-a-Par (ANOVA + Tukey HSD).

5.4 Validação Embarcada na Plataforma Raspberry Pi Zero 2W

Conforme discutido nas Considerações do Capítulo 6, a equivalência estatística entre os modelos GRU+LDA e CatBoost+LDA motivou a seleção do CatBoost+LDA para a implementação embarcada, dado seu menor custo computacional. Para validar essa decisão, o modelo otimizado foi portado para a plataforma Raspberry Pi Zero 2W, um dispositivo de borda com processador ARM Cortex-A53 quad-core de 1 GHz e 512 MB de memória RAM. O experimento de inferência foi conduzido com o conjunto completo de 180 amostras de teste, mantendo o mesmo particionamento utilizado na avaliação em ambiente *desktop*. A Tabela 3 apresenta as métricas de desempenho obtidas na plataforma embarcada.

Tabela 3 – Métricas de desempenho do modelo CatBoost+LDA no Raspberry Pi Zero 2W.

Métrica	Valor
Acurácia	0,9056 (90,56%)
Precisão (Ponderada)	0,9089
Recall (Ponderado)	0,9056
F1-Score (Média)	0,9052
Tempo Total de Inferência (180 amostras)	0,3233 s
Tempo Médio por Amostra	1,80 ms
Frequência de Classificação	556,7 Hz
Pico de Alocação de Memória (Operação)	8,54 KB

Fonte: Autoria Própria

Os resultados demonstram que o modelo CatBoost+LDA manteve um desempenho de classificação elevado na plataforma embarcada, com acurácia de 90,56%. A classe Sem Defeito (SD) apresentou 55 acertos em 60 amostras, com 5 classificações incorretas como Defeito Interno (DI). A maior fonte de erro concentrou-se na classe Defeito Interno (DI), com 12 amostras erroneamente classificadas como SD, resultando em 48 acertos de 60. Essa confusão entre SD e DI é consistente com o comportamento observado na avaliação em ambiente desktop e pode ser atribuída à similaridade morfológica entre os sinais diferenciais dessas duas classes, uma vez que o defeito interno, por estar localizado na parede interna do tubo e sob o revestimento compósito, produz perturbações de menor amplitude no campo magnético secundário.

Do ponto de vista de engenharia, os indicadores de desempenho temporal são particularmente relevantes para a viabilidade da aplicação industrial. O tempo médio de inferência de 1,80 ms por amostra corresponde a uma frequência de classificação de 556,7 Hz, o que é significativamente superior à taxa de aquisição dos sinais PEC (1 kHz de frequência de excitação com 200 pontos por ciclo). Isso indica que o sistema embarcado é capaz de processar e classificar os sinais em tempo real, sem introduzir gargalos no fluxo de inspeção. Adicionalmente, o pico de alocação de memória de apenas 8,54 KB durante a operação de inferência confirma o baixo impacto computacional do modelo baseado em árvores de decisão, validando a escolha do CatBoost em detrimento da GRU para cenários de computação de borda.

CONSIDERAÇÕES

Este trabalho propôs e avaliou um sistema inteligente para a classificação automática de defeitos em tubulações de aço carbono revestidas, integrando técnicas de processamento de sinais de Correntes Parasitas Pulsadas (PEC), engenharia de atributos, redução de dimensionalidade e algoritmos de aprendizado de máquina. A análise comparativa entre as técnicas de redução de dimensionalidade demonstrou a superioridade da Análise Discriminante Linear (LDA) sobre a Análise de Componentes Principais (PCA) para o problema em questão, evidenciada por um Silhouette Score mais de duas vezes superior (0,76 contra 0,31). Essa superioridade refletiu-se diretamente no desempenho dos classificadores, com ganhos consistentes de acurácia ao se substituir o PCA pelo LDA em todos os modelos avaliados. Entre os cinco algoritmos implementados e otimizados via Otimização Bayesiana com validação cruzada aninhada, a rede neural GRU+LDA alcançou a maior acurácia nominal de 94%, seguida pelo CatBoost+LDA com 92%. A validação estatística por meio da Análise de Variância (ANOVA) e do teste *post-hoc* de Tukey-HSD confirmou a equivalência estatística de desempenho entre esses dois modelos de topo, permitindo que a decisão de implementação fosse guiada por critérios de viabilidade computacional.

A validação embarcada no Raspberry Pi Zero 2W confirmou a viabilidade prática do sistema proposto. O modelo CatBoost+LDA manteve uma acurácia de 90,56% na plataforma de borda, com tempo médio de inferência de 1,80 ms por amostra e pico de alocação de memória de apenas 8,54 KB. A leve redução de acurácia em relação ao ambiente desktop (de 92% para 90,56%) é atribuída às diferenças de precisão numérica entre as plataformas, sendo um compromisso aceitável frente aos ganhos de portabilidade e autonomia operacional.

Conclui-se que a combinação de CatBoost e Análise Discriminante Linear constitui uma solução robusta e eficiente para a classificação automática de defeitos por cor-

rentes parasitas pulsadas em sistemas embarcados, equilibrando precisão diagnóstica e baixo custo computacional. O sistema proposto contribui para o avanço da manutenção preditiva industrial, oferecendo uma ferramenta de apoio ao diagnóstico que reduz a dependência da subjetividade do operador.

6.1 Sugestões Para Trabalhos Futuros

Como desdobramentos desta pesquisa, sugere-se:

- A ampliação do conjunto de dados experimentais com maior variedade de geometrias de defeito e diferentes materiais de revestimento, visando aumentar a robustez e a generalização dos classificadores;
- A investigação de técnicas de quantização e compressão de modelos para viabilizar a implementação em microcontroladores de recursos ainda mais restritos;
- O desenvolvimento de uma interface de comunicação sem fio integrada ao sistema embarcado, permitindo o envio dos diagnósticos em tempo real para plataformas de supervisão e controle (SCADA);
- A avaliação da adaptabilidade do sistema a outros tipos de ensaios não destrutivos, como ultrassom e termografia, ampliando o escopo de aplicação da metodologia proposta.

Referências Bibliográficas

AGBANGBA, C. E. et al. On the use of post-hoc tests in environmental and biological sciences: A critical review. *Heliyon*, v. 10, 2024. Disponível em: <<https://api.semanticscholar.org/CorpusID:267274860>>.

AIS Field. *Non-Destructive Testing (NDT) & Machine Learning*. 2024. <<https://aisfield.com/machine-learning-in-non-destructive-testing/>>. Acessado em: 09 dez. 2025.

AMIRI, A. F. et al. Fault detection and diagnosis of a photovoltaic system based on deep learning using the combination of a convolutional neural network (cnn) and bidirectional gated recurrent unit (bi-gru). *Sustainability*, MDPI, v. 16, n. 3, p. 1234, 2024.

BISHOP, C. M. *Pattern Recognition and Machine Learning*. New York: Springer, 2006.

CHEN, H. et al. Study on the effect of metal mesh on pulsed eddy-current testing of corrosion under insulation using an early-phase signal feature. *Materials*, v. 16, n. 4, 2023. ISSN 1996-1944. Disponível em: <<https://www.mdpi.com/1996-1944/16/4/1451>>.

FU, X. et al. Towards end-to-end pulsed eddy current classification and regression with cnn. *arXiv preprint arXiv:1902.08553*, feb 2019.

GAO, J.; HU, W.; CHEN, Y. Revisiting pca for time series reduction in temporal dimension. 2024. Disponível em: <<https://arxiv.org/abs/2412.19423>>.

HASSAN, I. ul; PANDURU, K.; WALSH, P. J. Non-destructive testing methods for condition monitoring: A review of techniques and tools. *Procedia Computer Science*, v. 257, p. 420–427, 2025. ISSN 1877-0509. The 16th International Conference on Ambient Systems, Networks and Technologies Networks (ANT)/ the 8th International Conference on Emerging Data and Industry 4.0 (EDI40). Disponível em: <<https://www.sciencedirect.com/science/article/pii/S1877050925007914>>.

HASTIE, T.; TIBSHIRANI, R.; FRIEDMAN, J. *The Elements of Statistical Learning: Data Mining, Inference, and Prediction*. 2nd. ed. New York: Springer, 2009.

HELLIER, C. *Handbook of Nondestructive Evaluation*. 2nd. ed. New York: McGraw-Hill Education, 2012. ISBN 978-0071777114.

JIN, G. soo et al. Extracting weld bead shapes from radiographic testing images with u-net. *Applied Sciences*, 2021. Disponível em: <<https://api.semanticscholar.org/CorpusID:245366137>>.

LI, M. et al. *Deep Learning and Machine Learning with GPGPU and CUDA: Unlocking the Power of Parallel Computing*. 2025. Disponível em: <<https://arxiv.org/abs/2410.05686>>.

MELGAR, P. C.; CHAMBI, L. H.; TORRES, R. S. Predictive maintenance based on deep learning: Early identification of failures in heavy machinery components. *International Journal of Advanced Computer Science and Applications*, 2025. Disponível em: <<https://api.semanticscholar.org/CorpusID:279088029>>.

MUNIR, N. et al. Machine learning based eddy current testing: A review. *Results in Engineering*, v. 25, p. 103724, 12 2024.

RASCHKA, S.; PATTERSON, J.; NOLET, C. Machine learning in python: Main developments and technology trends in data science, machine learning, and artificial intelligence. *Information*, v. 11, p. 193, 04 2020.

Scikit-learn. *Support Vector Machines*. 2025. Documentação da biblioteca Scikit-learn. Disponível em: <<https://scikit-learn.org/stable/modules/svm.html>>. Acesso em: 24 nov. 2025.

SHALOO, M. et al. A review of non-destructive testing (ndt) techniques for defect detection: Application to fusion welding and future wire arc additive manufacturing processes. *Materials*, v. 15, n. 10, 2022. ISSN 1996-1944. Disponível em: <<https://www.mdpi.com/1996-1944/15/10/3697>>.

SHRIFAN, N. H. M. M.; AKBAR, M. F.; ISA, N. A. M. Prospect of using artificial intelligence for microwave nondestructive testing technique: A review. *IEEE Access*, v. 7, p. 110628–110650, 2019.

SILVA, M. et al. Intelligent embedded system for decision support in pulsed eddy current corrosion detection using extreme learning machine. *Measurement*, v. 185, 2021. Disponível em: <<https://api.semanticscholar.org/CorpusID:239636754>>.

SILVA, M. et al. Review of conventional and advanced non-destructive testing techniques for detection and characterization of small-scale defects. *Progress in Materials Science*, v. 138, p. 101155, 06 2023.

SILVA, M. M. *Sistema Embarcado Inteligente para Auxílio ao Diagnóstico em Inspeções por Correntes Parasitas Pulsadas*. Tese (Doutorado) — Universidade Federal da Bahia, Salvador, BA, 2021.

SOPHIAN, A. et al. Machine-learning-based evaluation of corrosion under insulation in ferromagnetic structures. *IIUM Engineering Journal*, 2021. Disponível em: <<https://api.semanticscholar.org/CorpusID:238213013>>.

SOPHIAN, A.; TIAN, G.; FAN, M. Pulsed eddy current non-destructive testing and evaluation: A review. *Chinese Journal of Mechanical Engineering*, v. 30, p. 500 – 514, 2017. Disponível em: <<https://api.semanticscholar.org/CorpusID:7009742>>.

SZOSTAK, D.; WALKOWIAK, K.; WŁODARCZYK, A. Short-term traffic forecasting in optical network using linear discriminant analysis machine learning classifier. In: *2020 22nd International Conference on Transparent Optical Networks (ICTON)*. [S.l.: s.n.], 2020. p. 1–4.

TERKO, A. et al. Credit scoring model implementation in a microfinance context. p. 1–6, 10 2019.

WANG, H. et al. A deep learning-based watershed feature fusion approach for tunnel crack segmentation in complex backgrounds. *Materials*, v. 18, 2025. Disponível em: <https://api.semanticscholar.org/CorpusID:275223602>.

ZHANG, N.; ZHONG, Y.; DIAN, S. Rethinking unsupervised texture defect detection using pca. *Optics and Lasers in Engineering*, v. 163, p. 107470, 2023. ISSN 0143-8166. Disponível em: <https://www.sciencedirect.com/science/article/pii/S014381662200519X>.

CorujaT_EX



Este volume foi tipografado em L^AT_EX na classe CorujaT_EX como uma demanda do Colegiado do curso de Engenharia Elétrica da **UFOB!** (**UFOB!**)
(https://github.com/ademariocarvalho/CCEE_UFOB_CorujaTEX).