



CLASSIFICAÇÃO BINÁRIA DO ESTADO OPERACIONAL DE TRANSFORMADORES DE POTÊNCIA UTILIZANDO ÍNDICE DE DESEMPENHO E MODELOS DE APRENDIZADO DE MÁQUINA

Binary classification of operating state of power transformers using Performance Index and Machine Learning Models

Clasificación binaria del estado de operación de transformadores de potencia utilizando Índice de Rendimiento y Modelos de Aprendizaje Automático

Vinícius Faria Costa Mendanha¹, André Pereira Marques², & Cacilda de Jesus Ribeiro^{3*}

^{1,3} Universidade Federal de Goiás, Escola de Engenharia Elétrica, Mecânica e de Computação

² Instituto Federal de Goiás, Engenharia Elétrica

¹vinicius.fcfariacosta@gmail.com ²ap.marques@ifg.edu.br ^{3*}cacilda@ufg.br

ARTIGO INFO.

Recebido: 07.07.2025

Aprovado: 19.02.2026

Disponibilizado: 11.03.2026

PALAVRAS-CHAVE: Aprendizado de máquinas; classificador; índice de desempenho; técnicas preditivas; transformador de potência.

KEYWORDS: Classifier; machine learning; performance index; power transformer; predictive technique.

PALABRAS CLAVE: Aprendizaje automático; clasificador; índice de desempeño; técnicas predictivas; transformador de potencia.

*Autor Correspondente: Ribeiro, C. de J.

RESUMO

Transformadores de potência são ativos valiosos e estratégicos em sistemas elétricos, uma vez que falhas inesperadas deles podem resultar em prejuízos operacionais e financeiros expressivos às empresas do setor elétrico e aos consumidores. Embora haja avanços no monitoramento de suas condições operativas, algumas metodologias ainda requerem interpretação especializada, além de serem pouco padronizadas ou adotarem modelos cuja complexidade pode dificultar a integração com práticas operacionais usuais dos profissionais de manutenção. Nesse sentido, o objetivo deste trabalho é desenvolver classificadores binários relacionados a algoritmos de aprendizado de máquina para predição rápida e eficiente do estado operacional de transformadores de potência, rotuladas como *Satisfatório* ou *Insatisfatório*, com base em dados oriundos de ensaios físico-químicos, Análise de gases dissolvidos (AGD), e índices de desempenho, a partir de exemplos reais de equipamentos. Na metodologia, desenvolvem-se modelos supervisionados de aprendizado de máquinas, como *Random Forest*, *HistGradientBoosting*, Regressão Logística Balanceada e *XGBoost*, implementados com validação cruzada estratificada. Os resultados indicam que os classificadores são capazes de identificar satisfatoriamente os transformadores em condição crítica, mesmo em um cenário com relevante dispersão nos dados. Portanto, a abordagem desenvolvida representa uma ferramenta promissora à decisão técnica em estratégias de manutenção preventiva, aliando confiabilidade, escalabilidade e facilidade de aplicação em ambientes de campo.

ABSTRACT

Power transformers are strategic and valuable assets in electrical systems, as their unexpected failures can lead to significant operational and financial losses for power sector companies and consumers. Although there have been advancements in monitoring their operational conditions, some methodologies still require specialized interpretation, lack standardization, or adopt models whose complexity

can hinder integration with the usual operational practices of maintenance professionals. In this context, the objective of this work is to develop binary classifiers based on machine learning algorithms for fast and efficient prediction of the operational state of power transformers, labeled as *Satisfactory* or *Unsatisfactory*, using data derived from physicochemical tests, Dissolved Gas Analysis (DGA), and performance indices, based on real equipment samples. The methodology involves the development of supervised machine learning models, such as *Random Forest*, *HistGradientBoosting*, *Balanced Logistic Regression*, and *XGBoost*, implemented with stratified cross-validation. The results indicate that the classifiers can satisfactorily identify transformers in critical condition, even in a scenario with considerable data dispersion. Therefore, the proposed approach represents a promising tool for technical decision-making in preventive maintenance strategies, combining reliability, scalability, and ease of application in field environments.

RESUMEN

Los transformadores de potencia son activos valiosos y estratégicos en los sistemas eléctricos, ya que sus fallas inesperadas pueden generar pérdidas operativas y financieras significativas para las empresas del sector eléctrico y los consumidores. Aunque se han logrado avances en el monitoreo de sus condiciones operativas, algunas metodologías aún requieren interpretación especializada, carecen de estandarización o adoptan modelos cuya complejidad puede dificultar su integración con las prácticas operativas habituales de los profesionales de mantenimiento. El objetivo de este trabajo es desarrollar clasificadores binarios basados en algoritmos de aprendizaje automático para la predicción rápida y eficiente del estado operativo de transformadores de potencia, etiquetados como *Satisfactorio* o *Insatisfactorio*, utilizando datos provenientes de ensayos fisicoquímicos, análisis de gases disueltos (AGD) e índices de desempeño, a partir de muestras reales de equipos. La metodología incluye el desarrollo de modelos supervisados de aprendizaje automático, como *Random Forest*, *HistGradientBoosting*, *Regresión Logística Balanceada* y *XGBoost*, implementados con validación cruzada estratificada. Los resultados indican que los clasificadores son capaces de identificar satisfactoriamente los transformadores en condición crítica, incluso en un escenario con una dispersión significativa de los datos. Por lo tanto, el enfoque propuesto representa una herramienta prometedora para la toma de decisiones técnicas en estrategias de mantenimiento preventivo, combinando confiabilidad, escalabilidad y facilidad de aplicación en entornos de campo.

INTRODUÇÃO

Os ensaios físico-químico e a Análise de Gases Dissolvidos (AGD) são técnicas preditivas utilizadas em transformadores de potência para auxiliar na determinação do índice de desempenho deles (Campi, 2014; Dias, 2023; Marques et al., 2017; Marques, 2018; Moura, 2018). Embora algumas metodologias tradicionais desenvolvam expressões analíticas e ponderadas para relacionar os parâmetros que compõem cada uma dessas técnicas (Campi, 2014; Dias, 2023; Marques et al., 2017; Marques, 2018; Moura, 2018), a diversidade de variáveis envolvidas e a necessidade de etapas adicionais para a verificação de discrepâncias nos valores amostrados podem tornar oneroso o processo de classificação da condição operativa dos equipamentos. Além disso, por consequência de uma insatisfatória etapa anterior, a indicação das ações corretivas a serem tomadas podem ser prejudicadas, levando o equipamento a, em casos extremos, falhas catastróficas, com danos direto ao patrimônio das empresas do setor elétrico e aos consumidores de energia elétrica.

Primeiramente, quando se analisa o óleo mineral isolante de transformadores de potência, são determinadas suas propriedades físico-químicas com o intuito de verificar se ainda preservam as características necessárias à função primordial de isolamento. Entre os principais parâmetros avaliados, destacam-se: fator de potência, tensão interfacial, teor de água, rigidez dielétrica, índice de neutralização e índice de cor (Campi, 2014; Dias, 2023; Marques et al., 2017; Marques, 2018). A deterioração progressiva desses parâmetros indica processos de envelhecimento do óleo, contaminação por agentes externos ou degradação térmica. Portanto, a análise dessas variáveis é fundamental para a tomada de decisões relativas à manutenção corretiva e preventiva, incluindo a filtragem, a substituição do fluido isolante e, em última instância, do papel celulósico.

Outrossim, em decorrência de fenômenos elétricos e térmicos, por exemplo, o óleo mineral isolante produz quantidades de gases (Bechara, 2010; Marques, Azevedo et al., 2017; Marques et al., 2017). Assim, uma das formas de diagnosticar anormalidade na operação dos transformadores de potência, de forma não intrusiva e sem desligamentos, é por meio da técnica preditiva de Análises de Gases Dissolvidos (AGD), cujo diagnóstico por ser feito, dentre outros, pelo método de gás chave (Duval, 2002; Rogers, 1978; Senoussaoui et al., 2018), além da interpretação da correlação entre gases e sua evolução na série histórica (relação temporal). A Análise de Gases Dissolvidos (AGD) contempla os seguintes compostos: hidrogênio (H_2), metano (CH_4), etano (C_2H_6), etileno (C_2H_4), acetileno (C_2H_2), monóxido de carbono (CO) e dióxido de carbono (CO_2), além do nitrogênio (N_2) e oxigênio (O_2). A partir deles, é possível identificar eventos distintos em transformadores (Bechara, 2010; IEEE, 2019), como: maus contatos, superaquecimento do óleo, superaquecimento da celulose do papel *kraft* isolante, descargas parciais e arco elétrico. Todavia, não deve ser feita análise de AGD de maneira isolada (Bechara, 2010) mas sim correlacionando os resultados desta com outros ensaios, inclusive os elétricos, de forma a ser obter uma visão holística do que ocorre.

Embora os métodos tradicionais utilizados em AGD, como o de gás chave e o de Rogers, dentre outros, desempenhem papel fundamental na interpretação diagnóstica e possuam boa aceitação na prática técnica, os avanços proporcionados por técnicas de Inteligência Artificial (IA), como as redes neurais, têm se mostrado promissores no aprimoramento da acurácia e na melhoria do processo decisório. Por exemplo, foram obtidos resultados mais acurados na classificação de equipamentos em relação à aplicação da metodologia de gás chave e Rogers,

respectivamente, com acurácias de 45,19% e 40,00%, em relação à 72,24% obtidos de uma rede *Multilayer Perceptron* (MLP) (Barbosa et al., 2012). Assim, a aplicação de Inteligência Artificial (IA) pode funcionar como ferramenta complementar ao suporte decisório em práticas de manutenção preventiva de transformadores.

Nesse sentido, o objetivo é desenvolver classificadores binários relacionados a algoritmos de aprendizado de máquina para predição rápida e eficiente do estado operacional de transformadores de potência, rotuladas como *Satisfatório* ou *Insatisfatório*, com base em dados, oriundos de ensaios físico-químicos e AGD, e índices de desempenho, a partir de exemplares reais de equipamentos de vários países e em condições operacionais diversas. Embora nenhuma técnica preditiva deva ser utilizada isoladamente para o diagnóstico do equipamento, pretende-se estudar se a quantidade de parâmetros usada descreve, com satisfatória abrangência, o desempenho geral do transformador de potência.

Para tanto, são desenvolvidas estratégias específicas para lidar com o desbalanceamento de dados entre classificações e para validar, de maneira rigorosa, o desempenho dos modelos por meio de técnicas de validação cruzada. Como principal contribuição, este estudo resulta no desenvolvimento de uma ferramenta inédita, capaz de auxiliar profissionais de engenharia de manutenção à tomada de decisão quanto ao estado operacional de transformadores de potência, promovendo agilidade e confiabilidade no diagnóstico com base em dados acessíveis de rotina.

METODOLOGIA

Os dados dos parâmetros preditivos do óleo mineral isolante utilizados neste trabalho, formados por resultados das técnicas preditivas de ensaios físico-químico (FQ) e Análise de Gases Dissolvidos (AGD), estão disponíveis em (Arias, 2020). Eles foram, originalmente, publicados (Velásquez et al., 2019; Velásquez & Lara, 2018) e complementados por amostras de óleo coletadas de transformadores de potência, totalizando 470 registros. Especificamente, foram extraídos pontos de (Castro & Miranda, 2005; Duraisamy et al., 2007; Duval & dePabla, 2001; Naresh et al., 2008; Richardson et al., 2008; Velásquez et al., 2019; Velásquez & Mejía Lara, 2018; Yadaiah & Ravi, 2011). Ao todo, 360 amostras de óleo foram obtidas a partir de transformadores localizados em diferentes regiões do Peru, Colômbia e Chile, abrangendo equipamentos com distintas idades, especificações técnicas e condições operacionais (Velásquez & Lara, 2020).

Inicialmente, o arquivo de entrada é carregado e a variável-alvo *Health Index (HI)* é lida. Cada exemplo é classificado em A (Excelente), B (Bom), C (Marginal), D (Ruim) ou E (Muito Ruim), conforme (1), com base no valor calculado do *HI* (Dias, 2023; Marques, Azevedo, et al., 2017; Naderian et al., 2008; Wang et al., 2017; Wattakapaiboon & Pattanadech, 2016; Zeinoddini-Meymand & Vahidi, 2016). Para cada classificação, há ações recomendadas correspondentes, sendo (Dias, 2023; Marques, 2018): A (continuar a operar o equipamento normalmente), B (continuar a operar o equipamento atento à evolução dos resultados nos próximos registros), C (investigar e realizar outros testes a curto prazo para confirmar resultados e tendências), D (planejar a retirada de operação do equipamento para inspeção interna, localização e correção de defeitos) e E (remover o equipamento de operação imediatamente para inspeção interna, localização e correção de defeitos).

$$\text{Classificação} = \begin{cases} A, & \text{se } HI \geq 80 \\ B, & \text{se } 65 \leq HI < 80 \\ C, & \text{se } 50 \leq HI < 65 \\ D, & \text{se } 35 \leq HI < 50 \\ E, & \text{se } HI < 35 \end{cases} \quad (1)$$

Após esta etapa, cada atributo x passa por um processo de normalização z -score, conforme (2), em que x'_i é o valor normalizado, e $\bar{x}$ e σ_x são, nessa ordem, a média e o desvio padrão de x . Assim, todas as variáveis preditoras têm média zero e variância unitária, evitando que atributos com magnitudes maiores dominem o processo de aprendizado dos algoritmos de classificação. Em seguida, a partir do coeficiente de Pearson para os dados normalizados, r_{xy} , conforme (3), sendo n o número de tuplas, x_i e y_i os respectivos valores de x e y na i -ésima tupla, $\bar{x}$ e $\bar{y}$ os respectivos valores médios de x e y , e σ_x e σ_y os respectivos desvios-padrão de x e y , é possível construir a matriz de correlação entre as variáveis preditoras. Além disso, o coeficiente $r_{xy} \in [-1,1]$. Se r_{xy} for maior que 0, então as variáveis são positivamente correlacionadas, isto é, os valores de uma aumentam à medida que os valores de outra também. Quanto maior o valor, mais forte é a correlação (ou seja, mais cada atributo implica no outro) (Bishop, 2006; Han et al., 2011).

$$x'_i = \frac{x_i - \bar{x}}{\sigma_x} \quad (2) \quad r_{xy} = \frac{\sum_{i=1}^n (x_i y_i) - n \bar{x} \bar{y}}{n \sigma_x \sigma_y} \quad (3)$$

Neste trabalho, os algoritmos classificadores são treinados e testados para reconhecer duas classificações distintas de transformadores, sendo elas: a *Satisfatório*, composta por exemplares oriundos das classificações A, B e C; e a *Insatisfatória*, formada por exemplos das classificações D e E. O agrupamento adotado fundamenta-se na interpretação técnica dos níveis de desempenho dos transformadores, conforme preconizado (Dias, 2023; Marques, Azevedo et al., 2017; Naderian et al., 2008; Wang et al., 2017; Wattakapaiboon & Pattanadech, 2016; Zeinoddini-Meymand & Vahidi, 2016). Desse modo, as classificações A, B e C correspondem a equipamentos com condição operacional satisfatória, variando entre Excelente e Marginal. Assim, a atuação corretiva indicada pode ser postergada ou, várias vezes, avaliada com cautela em um curto espaço de tempo. Todavia, as classificações D e E são compostas por equipamentos com severa degradação e que exigem atuação de intervenção em curto prazo ou até imediata e priorização em rotinas de manutenção corretiva, bem como substituição para inspeção interna, localização e correção de defeitos.

Além disso, a estratégia discutida visa facilitar a aplicação dos modelos preditivos no contexto real de gestão de ativos, permitindo a identificação binária (Satisfatório ou Insatisfatório) de transformadores de potência que necessitam de intervenção imediata e daqueles que podem seguir operando com boa confiabilidade, com base na evolução histórica dos parâmetros avaliados. Outrossim, esse procedimento reduz a complexidade da tarefa de classificação, ao mesmo tempo que que preserva sua utilidade prática para a tomada de decisão por equipes de manutenção nas empresas de energia elétrica. Para tanto, aplicam-se quatro algoritmos de aprendizado supervisionado, a saber:

- *Random Forest* (RF): algoritmo baseado em conjuntos (*ensemble*). Ele constrói múltiplas árvores de decisão e árvores de classificação e regressão (CART) utilizando subconjuntos aleatórios dos dados e das variáveis preditoras, sendo que a decisão final é obtida por meio da votação majoritária entre as árvores, dada em (4) (Li et al., 2011; Sekulić et al., 2020), sendo $h_t(x)$ a predição da t -ésima árvore para a instância x e $mode(.)$ a classificação mais votada entre as árvores. Essa abordagem reduz o risco de sobreajuste (*overfitting*) e melhora a generalização do modelo (Li et al., 2011; Sekulić et al., 2020).

Na estrutura, as regras de divisão são representadas por nós, as decisões por ramos e as previsões finais por folhas e cada árvore construída maximizando o ganho de impureza, dado pelo índice de Gini, em (5) (Li et al., 2011; Sekulić et al., 2020), em que S é o conjunto de amostras em um nó da árvore de decisão, K é o número total de classificações possíveis e p_k é a proporção de amostras da classificação k no conjunto S , com $p_k = \frac{n_k}{n}$, sendo: n_k o número de amostras da classificação k no nó e n o número total de amostras no nó.

$$\hat{y} = \text{mode}(h_t(x)_{t=1}^T) \quad (4) \quad \text{Gini}(S) = 1 - \sum_{k=1}^K p_k^2 \quad (5)$$

- **Regressão Logística Balanceada (LogReg)**: estima a probabilidade de uma instância pertencer à classificação positiva usando a função logística, em (6) (Dempster et al., 2025; Wang et al., 2025; Zhang et al., 2022), sendo $x \in R^n$ o vetor de atributos, $w \in R^n$ o vetor de pesos, b o viés e $\sigma(\cdot)$ a função sigmoide.

$$P(y = 1|x) = \sigma(w^T x + b) = \frac{1}{1 + e^{-(w^T x + b)}} \quad (6)$$

A otimização é feita minimizando a função de perda logística regularizada, $L(w, b)$, dada em (7) (Dempster et al., 2025; Wang et al., 2025; Zhang et al., 2022), em que N representa o número total de amostras no conjunto de dados, $y_i \in \{0,1\}$ é o rótulo verdadeiro da i -ésima amostra, $p_i = \sigma(w^T x_i + b)$ corresponde à probabilidade prevista pelo modelo para a classificação e λ é o hiperparâmetro de regularização. O primeiro termo da equação corresponde à entropia cruzada, usada para medir a discrepância entre os rótulos verdadeiros e as previsões do modelo, enquanto o segundo, $\lambda|w|^2$, corresponde à regularização do tipo L_2 , penalizando coeficientes elevados de w com o objetivo de evitar *overfitting*.

$$L(w, b) = \frac{-1}{N} \sum_{i=1}^N [y_i \log p_i + (1 - y_i) \log(1 - p_i)] + \lambda |w|^2 \quad (7)$$

- **HistGradientBoosting (HistGB)**: algoritmo otimizado de *Gradient Boosting* (Miller et al., 2022). A predição é uma soma de modelos fracos (árvores), dado em (8) (Karbasi, 2022; Y. Shi et al., 2023), sendo f_t a árvore adicionada na iteração t , η a taxa de aprendizado e $\hat{y}_i^{(t)}$ a predição acumulada após t iterações. Ademais, o modelo f_t é treinado para minimizar o gradiente da função de perda L com relação às previsões anteriores, conforme (9) (Y. Shi et al., 2023; Velarde et al., 2024), em que $\hat{y}_i^{(t)}$ é a predição do modelo para a i -ésima amostra após t iterações, $\hat{y}_i^{(t-1)}$ é a predição acumulada até a iteração anterior, η é a taxa de aprendizado, $f_t(x_i)$ é a predição fornecida pela árvore f_t para a entrada x_i , $L(y_i, \hat{y}_i^{(t-1)})$ é a função de perda entre o valor real y_i e a predição anterior e $g_i^{(t)}$ é o gradiente da função de perda em relação à predição anterior para a i -ésima amostra.

$$\hat{y}_i^{(t)} = \hat{y}_i^{(t-1)} + \eta \cdot f_t(x_i) \quad (8) \quad g_i^{(t)} = \frac{\partial L(y_i, \hat{y}_i^{(t-1)})}{\partial \hat{y}_i^{(t-1)}} \quad (9)$$

- **Extreme Gradient Boosting (XGBoost)**: baseia-se no algoritmo de *Gradient Boosting*, com regularização explícita (Sheridan et al., 2016). A predição é modelada como uma soma de funções pertencentes ao espaço das árvores de decisão F , conforme (10), em que f_k representa a k -ésima árvore de decisão e K é o total de árvores. Além disso, a função objetivo a ser minimizada, dada em (11), combina a perda associada à predição com um termo de regularização, sendo $l(y_i, \hat{y}_i)$ a função de perda para classificação binária, $\Omega(f_k)$ a complexidade da árvore f_k , com T_k número de folhas, ω_j a predição na j -ésima folha, e γ e λ hiperparâmetros de regularização (Cherif & Kortebi, 2019; H. Shi et al., 2019).

$$\hat{y}_i = \phi(x_i) = \sum_{k=1}^K f_k(x_i), f_k \in F \quad (10)$$

$$L(\phi) = \sum_{i=1}^n l(y_i, \hat{y}_i) + \sum_{k=1}^K \Omega(f_k), \Omega(f_k) = \gamma T_k + \frac{1}{2} \lambda \sum_{j=1}^{T_k} \omega_j^2 \quad (11)$$

Na classificação binária, a função de perda adotada é l , dada em (12), com $\sigma(\hat{y}_i) = \frac{1}{1+e^{-\hat{y}_i}}$ sendo a função sigmoide que transforma a predição em probabilidade.

$$l(y_i, \hat{y}_i) = -[y_i \log(\sigma(\hat{y}_i)) + (1 - y_i) \log(1 - \sigma(\hat{y}_i))] \quad (12)$$

Para diminuir o desbalanceamento entre as classificações originais A, B, C, D e E, aplica-se o algoritmo SMOTE (*Synthetic Minority Over-sampling Technique*) com a estratégia *sampling_strategy = 'not majority'*. Essa abordagem preserva a quantidade de amostras da classificação majoritária e gera instâncias sintéticas apenas para as demais classificações, por meio da interpolação entre uma amostra da classificação minoritária e um de seus k vizinhos mais próximos. O processo é definido em (13) (Chawla et al., 2002), em que x_i representa uma instância da classificação minoritária, x_{nn} é um dos seus vizinhos mais próximos e $\delta \sim U(0,1)$ é uma variável aleatória com distribuição uniforme.

$$x_{novo} = x_i + \delta \cdot (x_{nn} - x_i) \quad (13)$$

Assim, as amostras sintéticas são utilizadas, apenas, no conjunto de treinamento, enquanto o conjunto de teste permanece composto apenas pelas originais. Essa estratégia evita o sobre ajuste associado à replicação simples de instâncias e deve contribuir para um aprendizado mais robusto dos classificadores. Também, para todos os modelos desenvolvidos, aplica-se validação cruzada estratificada (*Stratified K-Fold*). Ela consiste em dividir o conjunto de dados em $k_f = 5$ subconjuntos (*folds*) de forma que cada um mantenha, aproximadamente, a proporção original das classificações. O treinamento é realizado k_f vezes, sendo que, em cada uma, utilizam-se $k_f - 1$ subconjuntos para treinar o modelo e o subconjunto restante para teste. As métricas de desempenho são calculadas, então, como os valores médios a cada iteração.

Posteriormente, para o modelo com melhor desempenho, utiliza-se a validação cruzada *Leave-One-Out* (LOOCV), a qual representa um caso limite da validação cruzada com $k_f = n$, sendo n o número total de observações. Nesse caso, cada amostra é usada uma única vez como conjunto de teste, enquanto as demais compõem o conjunto de treino. Essa estratégia é adotada, pois permite uma avaliação menos enviesada do modelo final, o que é essencial no conjunto de dados avaliado, que possui grande desbalanceamento e número de amostras limitado para a classificação *Satisfatório*. As métricas de desempenho avaliadas neste trabalho são: Acurácia, dada em (14), Precisão, dada em (15), *Recall*, dada em (16), *F1-score*, dada em (17), e *F1-score weighted*, dada em (18), sendo TP_k o número de verdadeiros positivos da classificação k , FP_k o número de falsos positivos, FN_k o número de falsos negativos, n_k o número de amostras reais da classificação k , n o número total de amostras e K é o número total de classificações.

$$\text{Acurácia} = \frac{TP_k + TN_k}{TP_k + TN_k + FP_k + FN_k} \quad (14)$$

$$\text{Precisão} = \frac{TP_k}{TP_k + FP_k} \quad (15)$$

$$\text{Recall} = \frac{TP_k}{TP_k + FN_k} \quad (16)$$

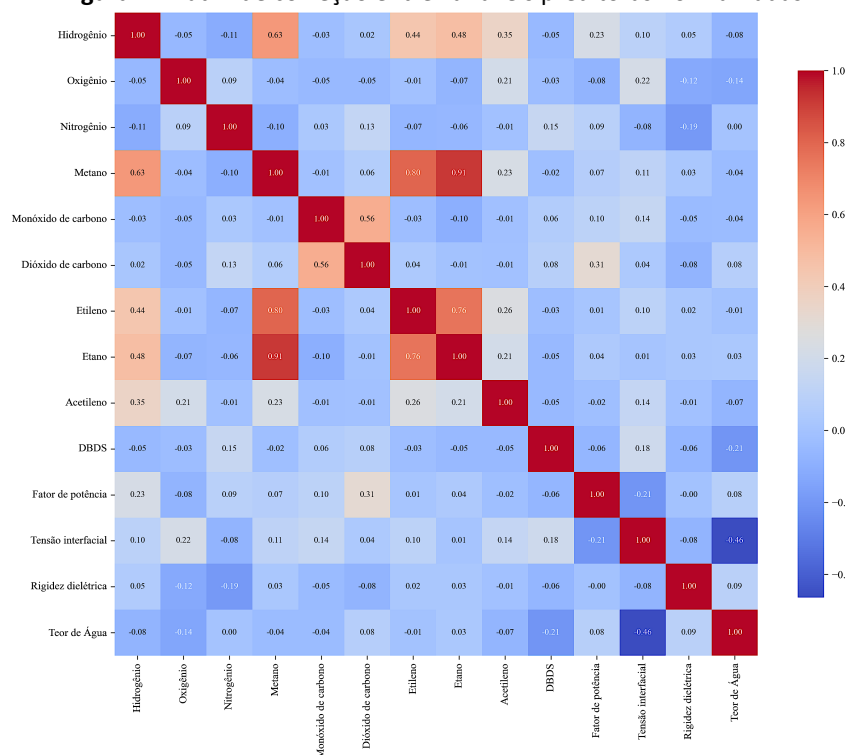
$$\text{F1-score} = 2 \cdot \frac{\text{Precisão} \cdot \text{Recall}}{\text{Precisão} + \text{Recall}} \quad (17)$$

$$\text{F1-score weighted} = \sum_{k=1}^K \frac{n_k}{n} \cdot \text{F1}_k \quad (18)$$

RESULTADOS E DISCUSSÃO

Utiliza-se o *StandardScaler* do *sklearn*, em *Python*, para normalizar as variáveis preditoras numéricas de interesse, sendo elas: hidrogênio, em ppm; metano, em ppm; etileno, em ppm; acetileno, em ppm; fator de potência, adimensional; tensão interfacial, em mN/m; monóxido de carbono, em ppm; dióxido de carbono, em ppm; oxigênio, em ppm; nitrogênio, em ppm; teor de água, em mg/kg; rigidez dielétrica, em kV; e DBDS (Dibenzil Dissulfeto), em mg/kg (Figura 1).

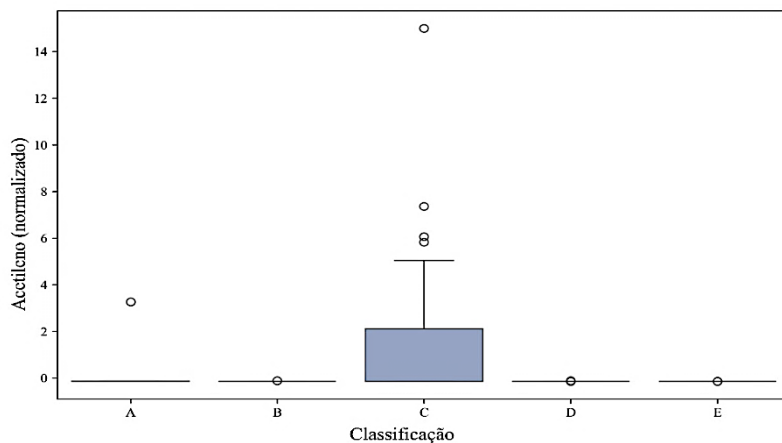
Figura 1. Matriz de correção entre variáveis preditoras normalizadas



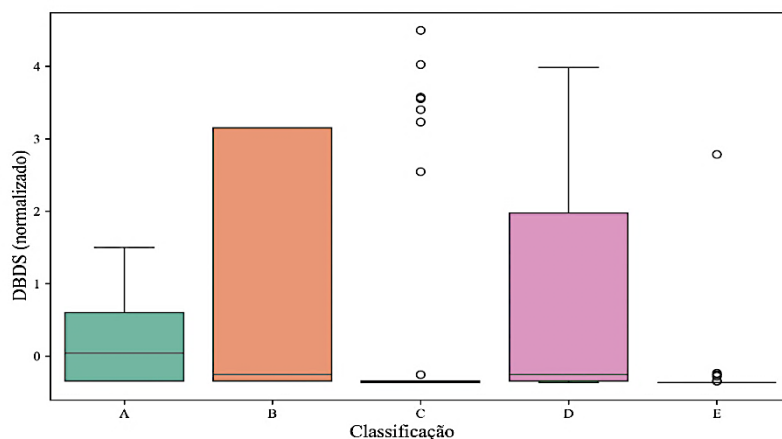
Fonte: Autores (2025).

Inicialmente, a distribuição por classificações constituiu: A (5 equipamentos), B (5 equipamentos), C (41 equipamentos), D (132 equipamentos) e E (287 equipamentos), totalizando 470 exemplares. Evidencia-se (Figura 1) que o par de gases metano e etileno apresenta correlação igual a 0,91, expondo que são gerados ou evoluem simultaneamente, fenômeno típico do processo de pirólise térmica do óleo mineral isolante. Além disso, há alta correlação entre etileno e etano de 0,76, reforçando a hipótese de que os compostos derivam de mecanismos análogos de degradação térmica. Já os gases monóxido de carbono e dióxido de carbono exibem correlação moderada de 0,56, coerente com a origem comum na decomposição da celulose presente no papel isolante. Por outro lado, o composto DBDS não apresenta correlação significativa com os gases dissolvidos. Já o teor de água expôs correlação negativa com a tensão interfacial, igual a $-0,46$, e com o fator de potência, a 0,08, sugerindo que a presença de umidade contribui para a degradação da qualidade dielétrica do óleo, o que pode comprometer a eficiência e confiabilidade do sistema de isolamento.

Posteriormente, plotam-se os *boxplots* referentes a cada variável por classificação (Figuras 2 e 3). Neles, observa-se o comportamento estatístico das variáveis ao longo das diferentes categorias de desempenho, por meio do destaque à mediana, quartis e presença de possíveis valores discrepantes (*outliers*). Dados discrepantes não foram removidos e entende-se que eles compõem o conjunto de observações, representando comportamentos atípicos da operação dos transformadores de potência, todavia, verdadeiros.

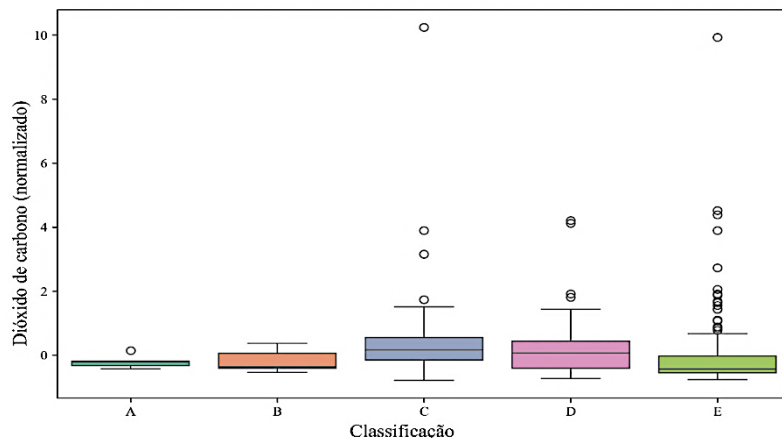
Figura 2. *Boxplots* da variável acetileno normalizada por classificação (A até E)

Fonte: Autores (2025).

Figura 3. *Boxplots* da variável DBDS normalizada por classificação (A até E)

Fonte: Autores (2025).

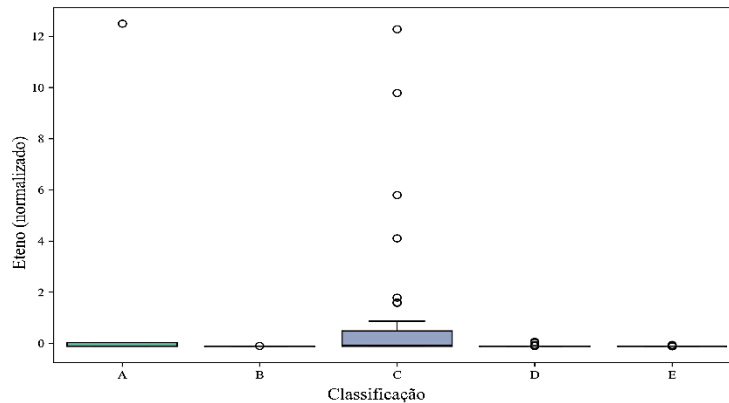
O acetileno (Figura 2), gás associado a falhas elétricas por arco elétrico, é encontrado em concentração elevada, principalmente, na classificação C, com valores residuais ou nulos nas demais. Isso indica que esse tipo de falha é pontual e menos comum nos extremos do espectro de desempenho (A, B, D e E). Por sua vez, o composto DBDS (Figura 3) apresenta maior concentração na classificação B, seguida pela D, sugerindo que a presença de enxofre corrosivo está associada a estágios iniciais de degradação. Doravante, no dióxido de carbono, subproduto da degradação térmica da celulose, evidencia-se distribuição mais homogênea ao longo das classificações, porém, com aumento progressivo da dispersão nas classificações C, D e E, reforçando sua relação com o envelhecimento da isolação sólida (Figura 4).

Figura 4. *Boxplots* da variável dióxido de carbono normalizada por classificação (A até E)

Fonte: Autores (2025).

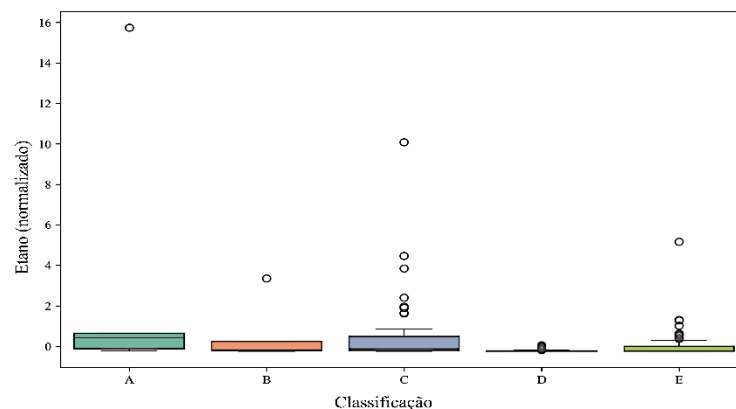
As variáveis, etileno (Figura 5), etano (Figura 6) e metano (Figura 7), relacionadas à pirólise do óleo mineral isolante, exibem máximos na classificação C, indicando que, ainda em transformadores classificados com satisfatório desempenho, podem ocorrer eventos térmicos localizados. Esse comportamento realça a hipótese de que a presença dos gases não é restrita a casos de desempenhos classificados como ruins ou muito ruins, podendo ocorrer, até, em equipamentos com índices de desempenho altos e desejáveis. Assim, exige-se uma abordagem atenta na interpretação dos resultados.

Figura 5. Boxplots da variável etileno normalizada por classificação (A até E)



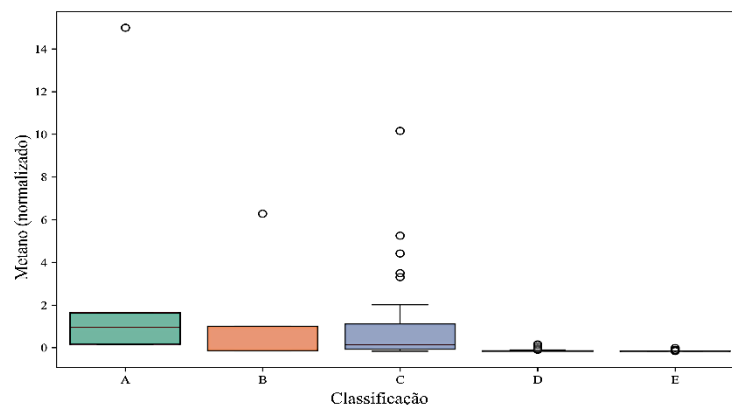
Fonte: Autores (2025).

Figura 6. Boxplots da variável etano normalizada por classificação (A até E)



Fonte: Autores (2025).

Figura 7. Boxplots da variável metano normalizada por classificação (A até E)



Fonte: Autores (2025).

É válido ressaltar a presença de *outliers* pronunciados em diversas variáveis, como etileno (Figura 5), etano (Figura 6) e metano (Figura 7), sobretudo, nas classificações C, D e E. Sugere-se que há casos isolados, com comportamento atípico ou intensificação repentina de processos de falha, que podem inibir a generalização de modelos classificadores baseados em Inteligência Artificial.

A presença de registros discrepantes, foi detectada na própria visualização dos *boxplots*, utilizando o intervalo interquartil (*IQR*). A partir dessa abordagem, consideram-se como valores extremos aqueles que satisfazem (19), em que Q_1 e Q_3 são, nessa ordem, o primeiro e o terceiro quartis da distribuição da variável analisada, sendo o intervalo interquartil definido como $IQR = Q_3 - Q_1$. Assim, valores abaixo de $Q_1 - 1,5 \times IQR$ ou acima de $Q_3 + 1,5 \times IQR$ são considerados *outliers*.

$$x_{\text{outlier}} \in (-\infty, Q_1 - 1,5 \times IQR] \cup [Q_3 + 1,5 \times IQR, +\infty) \quad (19)$$

A variável fator de potência, (Figura 8), por exemplo, destaca-se com elevado número de *outliers* nas classificações A e E, o que pode estar associado tanto à variabilidade do processo de medição quanto à natureza não normal da distribuição. Desse modo, a análise de valores atípicos constitui uma etapa fundamental no diagnóstico da qualidade dos dados, permitindo a identificação de amostras que destoam do comportamento geral da população.

Importante destacar, todavia, que esses dados não foram descartados na análise. A decisão de mantê-los fundamenta-se no entendimento de que, em muitos casos, os valores atípicos não correspondem a erros de medição ou registros incorretos, mas podem refletir condições extremas de operação, como altos carregamentos ou situações de estresse no sistema. Tais ocorrências carregam informações relevantes sobre o fenômeno estudado e, portanto, merecem ser preservadas. Além disso, a exclusão desses pontos poderia introduzir vieses, comprometendo a representatividade da amostra e reduzindo a capacidade do modelo de capturar toda a variabilidade do sistema. Dessa forma, optou-se por considerar esses dados nas análises subsequentes, assegurando uma interpretação mais fidedigna e abrangente dos resultados.

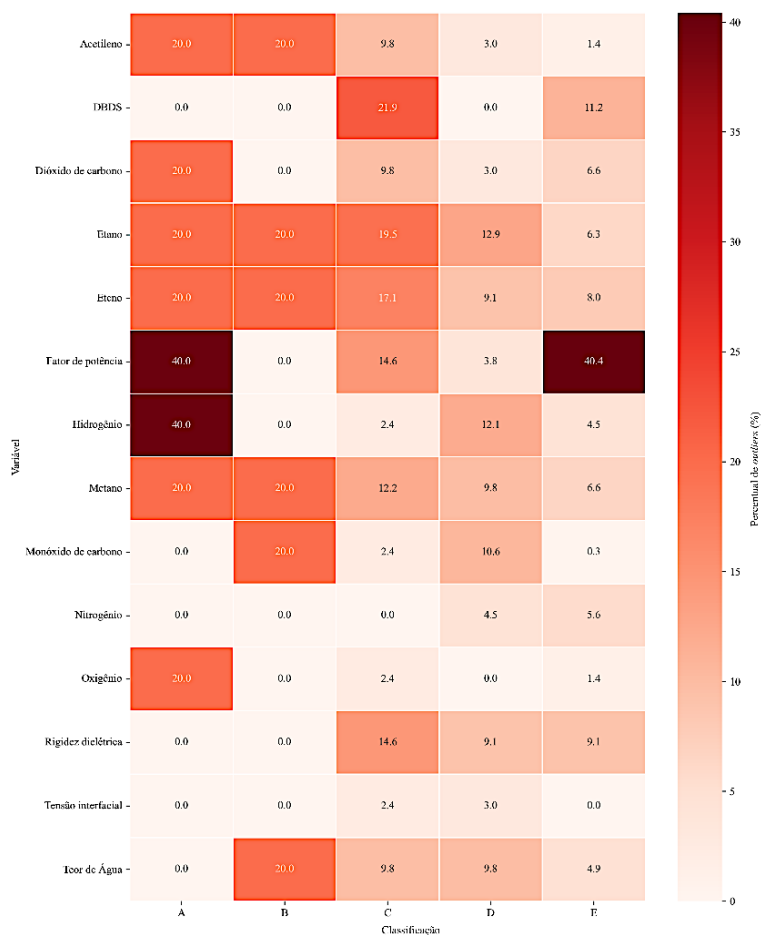
Assim, a Tabela 1 explicita os hiperparâmetros usados em cada classificador binário e em seguida, para avaliar o desempenho dos algoritmos de aprendizado de máquina, na atividade de classificação dos transformadores, entre *Satisfatório* e *Insatisfatório*, foram consideradas três métricas: acurácia, *F1-score* e *F1-score weighted*.

Tabela 1. Valores de hiperparâmetros em cada modelo de classificador binário utilizado

Modelo	Hiperparâmetros	Valores
Random Forest	<i>random_state</i>	42
	Demais parâmetros	Padrão (<i>n_estimators=100, max_depth=None</i>)
	<i>class_weight</i>	<i>balanced</i>
LogReg	<i>max_iter</i>	1000
	Demais parâmetros	Padrão (<i>penalty='l2', solver='lbfgs'</i>)
	<i>random_state</i>	42
HistGB	Demais parâmetros	Padrão (<i>learning_rate=0,1, max_iter=100</i>)
	<i>use_label_encoder</i>	<i>False</i>
	<i>eval_metric</i>	<i>logloss</i>
XGBoost (K-Fold)	<i>random_state</i>	42
	Demais parâmetros	Padrão (<i>n_estimators=100, learning_rate=0,3</i>)
	<i>use_label_encoder</i>	<i>False</i>
XGBoost (LOOCV)	<i>eval_metric</i>	<i>logloss</i>
	<i>random_state</i>	42
	<i>scale_pos_weight</i>	10
	Demais parâmetros	Padrão

Fonte: Autores (2025).

Figura 8. Mapa de calor da frequência relativa de *outliers* identificados por IQR por variável preditora e classificação de desempenho dos transformadores de potência



Fonte: Autores (2025).

A Tabela 2, por sua vez, apresenta os resultados obtidos para cada um dos modelos, permitindo uma análise comparativa quanto à capacidade de generalização e ao equilíbrio na classificação entre as classificações.

Tabela 2. Métricas de desempenho dos algoritmos de classificação binária utilizando aprendizado de máquinas e validação *Stratified K-Fold* com cinco dobras

Algoritmo	Acurácia	F1-score	F1-score weighted
Random Forest	96,38%	90,73%	96,40%
LogReg	92,98%	83,14%	93,23%
HistGB	96,60%	91,20%	96,60%
XGBoost	96,60%	91,63%	96,98%

Fonte: Autores (2025).

Os modelos baseados em *ensemble*, que são o Random Forest, o HistGB e o XGBoost, apresentaram desempenhos bastante similares, com acurácias superiores a 96% e F1-scores acima de 90%. O XGBoost se destacou levemente nas três métricas, consolidando-se como a melhor técnica com desempenho global no conjunto de dados avaliado. A regressão logística, embora tenha apresentado resultados promissores, com métricas acima de 83%, é superada pelas demais abordagens, especialmente, no que se refere ao F1-score, indicando menor equilíbrio na classificação entre as classificações.

Já a Tabela 3 exibe as métricas de desempenho de precisão, *recall* e F1-score para o XGBoost, algoritmo com o melhor desempenho, utilizando a validação por *Leave-One-Out Cross-Validation* (LOOCV). Esses resultados podem ser visualizados, também, na matriz de confusão, (Figura 9), em que são exibidas as classificações feitas pelo XGBoost e o estado real do transformador de potência.

Tabela 3. Métricas de desempenho do algoritmo de classificação binária XGBoost para cada classificação utilizando validação por *Leave-One-Out Cross-Validation* (LOOCV)

Classificação	Precisão	Recall	F1-score
<i>Satisfatório</i>	75,41%	90,20%	82,14%
<i>Insatisfatório</i>	98,78%	96,42%	97,58%

Fonte: Autores (2025).

Figura 9. Matriz de confusão do algoritmo de classificação binária XGBoost para cada classificação utilizando validação por *Leave-One-Out Cross-Validation* (LOOCV)



Fonte: Autores (2025).

Considerando a aplicação da validação *Leave-One-Out Cross-Validation* (LOOCV), na Tabela 3 e na Figura 9, consta-se que o XGBoost obteve desempenho excelente, com acurácia final de 96% e destaque para a classificação *Insatisfatório*, com precisão de 98,78%. Ainda que a classificação *Satisfatório* tenha apresentado F1-score inferior à outra, observa-se uma taxa de *recall* de 90,20%, indicando que poucas amostras foram erroneamente classificadas. A acurácia final foi de cerca de 96%.

CONSIDERAÇÕES FINAIS

Este estudo demonstra que dados de ensaios físico-químicos e de AGD, quando tratados estatisticamente e processados por algoritmos de aprendizado supervisionado, podem ser utilizados de forma eficaz para classificação binária do estado operacional de transformadores de potência por meio do índice de desempenho. Ademais, a adoção da estratégia para tornar menos significativo o desbalanceamento original de exemplares entre classificações, a partir do algoritmo SMOTE e da aplicação rigorosa de validação cruzada, permitiu avaliar com o desempenho dos modelos preditivos em cenários realistas.

Os resultados obtidos indicam que XGBoost e *HistGradientBoosting* são capazes de alcançar altos índices de F1-score, respectivamente, 91,63% e 91,20%, ainda que em conjuntos de dados heterogêneos e com alta variabilidade. Isso atesta a aplicabilidade da metodologia desenvolvida em ambientes de campo. Além disso, como principal contribuição, há a construção de uma ferramenta de classificação que, embora baseada em dados de rotina, fornece subsídios técnicos confiáveis para a priorização de ações de manutenção, reduzindo, por consequência, o risco de falhas incipientes e catastróficas e os custos operacionais associados, com acurácia máxima de 96,60%.

Como limitações, tem-se a dependência de bases históricas consistentes, bem como a ausência de séries temporais que possibilitem análises preditivas mais refinadas do desempenho ao longo do tempo para cada equipamento. Nesse sentido, pesquisas futuras podem integrar dados adicionais, tais como medições *on-line* e variáveis ambientais (quando disponíveis).

Logo, a ferramenta projetada possui amplas possibilidades de uso prático. Além de robusta e de fácil replicação nos ambientes das empresas do setor elétrico, pode ser incorporada a sistemas de apoio à decisão rápida nas empresas, sem a necessidade de aquisição de equipamentos adicionais, como sensores, ou de técnicas invasivas de diagnóstico, uma vez que os ensaios físico-químicos e da AGD não exigem o desligamento dos equipamentos analisados.

AGRADECIMENTOS

Escola de Engenharia Elétrica, Mecânica e de Computação (EMC) Universidade Federal de Goiás (UFG), Laboratório de Pesquisa em Engenharia de Alta Tensão (LAPEAT-UFG) e ao Centro de Excelência em Redes Inteligentes sem fio e Serviços Avançados (CERISE) pelos apoios institucionais, e à Fundação de Amparo à Pesquisa do Estado de Goiás (FAPEG) pelo financiamento.

REFERÊNCIAS

- Arias, R. (2020). Data for: Root cause analysis improved with machine learning for failure analysis in power transformers [Dataset]. *Mendeley*. <https://doi.org/10.17632/RZ75W3FKXY.1>
- Barbosa, F. R., Almeida, O. M., Braga, A. P. S., Amora, M. A. B., & Cartaxo, S. J. M. (2012). Application of an artificial neural network in the use of physicochemical properties as a low-cost proxy of power transformers DGA data. *IEEE Transactions on Dielectrics and Electrical Insulation*, 19(1), 239-246. <https://doi.org/10.1109/TDEI.2012.6148524>
- Bechara, R. (2010). Análise de falhas de transformadores de potência [Dissertação (Mestrado)]. *Universidade de São Paulo*.
- Bishop, C. M. (2006). Pattern recognition and machine learning. *Springer*.
- Campi, R. L. (2014). Modelagem fuzzy da concentração dos gases dissolvidos em óleo mineral isolante de transformadores baseada em resultados de ensaios físico-químicos [Dissertação (Mestrado)]. *Universidade de São Paulo*.
- Castro, A. R. G. & Miranda, V. (2005). Knowledge discovery in neural networks with application to transformer failure diagnosis. *IEEE Transactions on Power Systems*, 20(2), 717-724. <https://doi.org/10.1109/TPWRS.2005.846074>
- Chawla, N. V., Bowyer, K. W., Hall, L. O., & Kegelmeyer, W. P. (2002). SMOTE: Synthetic Minority Over-sampling Technique. *Journal of Artificial Intelligence Research*, 16, 321-357. <https://doi.org/10.1613/jair.953>
- Cherif, I. L., & Kortebi, A. (2019). On using extreme gradient boosting (XGBoost) machine learning algorithm for home network traffic classification. *Wireless Days (WD)*, 1-6. <https://doi.org/10.1109/WD.2019.8734193>
- Dempster, A., Webb, G. I., & Schmidt, D. F. (2025). Prevalidated ridge regression is a highly efficient drop-in replacement for logistic regression for high-dimensional data (No. arXiv:2401.15610). *arXiv*. <https://doi.org/10.48550/arXiv.2401.15610>
- Dias, Y. A. (2023). Method for predicting the performance indices of power transformers immersed in mineral insulating oil and medium and high-voltage circuit breakers [Tese (Doutorado)]. *Universidade Federal de Goiás*.
- Duraisamy, V., Devarajan, N., Somasundareswari, D., Vasanth, A. A. M., & Sivanandam, S. N. (2007). Neuro fuzzy schemes for fault detection in power transformer. *Applied Soft Computing*, 7(2), 534-539. <https://doi.org/10.1016/j.asoc.2006.10.001>
- Duval, M. (2002). A review of faults detectable by gas-in-oil analysis in transformers. *IEEE Electrical Insulation Magazine*, 18(3), 8-17. <https://doi.org/10.1109/MEI.2002.1014963>
- Duval, M. & dePabla, A. (2001). Interpretation of gas-in-oil analysis using new IEC publication 60599 and IEC TC 10 databases. *IEEE Electrical Insulation Magazine*, 17(2), 31-41. <https://doi.org/10.1109/57.917529>
- Han, J., Kamber, M., & Pei, J. (2011). Data Mining: Concepts and Techniques (3o ed). Elsevier.
- IEEE. (2019). IEEE Guide for the interpretation of gases generated in mineral oil-immersed transformers. *IEEE*. <https://doi.org/10.1109/IEEESTD.2019.8890040>
- Shen, Z., Hassani, H., Kale, S., & Karbasi, A. (2022). Federated functional gradient boosting. *Proceedings of the 25th International Conference on Artificial Intelligence and Statistics*, 151, 7814-7840. Recuperado de <https://proceedings.mlr.press/v151/shen22a.html>
- Li, J., Heap, A. D., Potter, A., & Daniell, J. J. (2011). Application of machine learning methods to spatial interpolation of environmental variables. *Environmental Modelling & Software*, 26(12), 1647-1659. <https://doi.org/10.1016/j.envsoft.2011.07.004>
- Marques, A. P. (2018). Optimized diagnosis of power transformers through the integration of predictive techniques [Tese (Doutorado)]. *Universidade Federal de Goiás*.
- Marques, A. P., Azevedo, C. H. B., Santos, J. A. L., Machado, S. G., Ribeiro, C. J. R., Moura, N. K., Dias, Y. A., & Brito, L. C. (2017). Metodologia para Reenergização de Transformadores de Potência após Interrupções não Programadas no Sistema Elétrico. *XXIV Seminário Nacional de Produção e Transmissão de Energia Elétrica (SNPTEE)*.
- Marques, A. P., Azevedo, C. H. B., Santos, J. A. L. dos., Sousa, F. R. de C., Moura, N. K., Ribeiro, C. J., Dias, Y. A., Rodrigues, A., Rocha, A. S., & da C. Brito, L. da C. (2017). Insulation resistance of power transformers - method for optimized analysis. *2017 IEEE 19th International Conference on Dielectric Liquids (ICDL)*, Manchester, UK, pp. 1-4, <https://doi.org/10.1109/ICDL.2017.8124629>
- Marques, A. P., Moura, N. K. de., Dias, Y. A., Jesus Ribeiro, C. de., Rocha, A. S., Brito, L. da C., Azevedo, C. H. B., & Santos, J. A. L. dos. (2017). Method for the evaluation and classification of power transformer

insulating oil based on physicochemical analyses. *IEEE Electrical Insulation Magazine*, 33(1), 39-49. <https://doi.org/10.1109/MEI.2017.7804315>

Miller, C., Hao, L., & Fu, C. (2022). Gradient boosting machines and careful pre-processing work best: ASHRAE Great Energy Predictor III lessons learned. <https://doi.org/10.48550/ARXIV.2202.02898>

Moura, N. K. de. (2018). Otimização Computacional da Avaliação de Resultados de Ensaio Físico-químicos em Transformadores de Potência [Dissertação (Mestrado)]. *Universidade Federal de Goiás*.

Naderian, A., Cress, S., Piercy, R., Wang, F., & Service, J. (2008). An approach to determine the health index of power transformers. *Conference Record of the 2008 IEEE International Symposium on Electrical Insulation*, 192-196. <https://doi.org/10.1109/ELINSL.2008.4570308>

Naresh, R., Sharma, V., & Vashisth, M. (2008). An integrated neural fuzzy approach for fault diagnosis of transformers. *IEEE Transactions on Power Delivery*, 23(4), 2017-2024. <https://doi.org/10.1109/TPWRD.2008.2002652>

Richardson, Z. J., Fitch, J., Tang, W. H., Goulermas, J. Y., & Wu, Q. H. (2008). A Probabilistic classifier for transformer dissolved gas analysis with a particle swarm optimizer. *IEEE Transactions on Power Delivery*, 23(2), 751-759. <https://doi.org/10.1109/TPWRD.2008.915812>

Rogers, R. (1978). IEEE and IEC Codes to Interpret incipient faults in transformers, using gas in oil analysis. *IEEE Transactions on Electrical Insulation*, EI-13(5), 349-354. <https://doi.org/10.1109/TEI.1978.298141>

Sekulić, A., Kilibarda, M., Heuvelink, G. B. M., Nikolić, M., & Bajat, B. (2020). Random forest spatial interpolation. *Remote Sensing*, 12(10), 1687. <https://doi.org/10.3390/rs12101687>

Senoussaoui, M. E. A., Brahami, M., & Fofana, I. (2018). Combining and comparing various machine-learning algorithms to improve dissolved gas analysis interpretation. *IET Generation, Transmission & Distribution*, 12(15), 3673-3679. <https://doi.org/10.1049/iet-gtd.2018.0059>

Sheridan, R. P., Wang, W. M., Liaw, A., Ma, J., & Gifford, E. M. (2016). Extreme gradient boosting as a method for quantitative structure-activity relationships. *Journal of Chemical Information and Modeling*, 56(12), 2353-2360. <https://doi.org/10.1021/acs.jcim.6b00591>

Shi, H., Wang, H., Huang, Y., Zhao, L., Qin, C., & Liu, C. (2019). A hierarchical method based on weighted extreme gradient boosting in ECG heartbeat classification. *Computer Methods and Programs in Biomedicine*, 171, 1-10. <https://doi.org/10.1016/j.cmpb.2019.02.005>

Mendanha, V. F. C., Marques, A. P., & Ribeiro, C. de J.

Shi, Y., Ke, G., Chen, Z., Zheng, S., & Liu, T.-Y. (2023). Quantized training of gradient boosting decision trees, (No. arXiv:2207.09682). *arXiv*. <https://doi.org/10.48550/arXiv.2207.09682>

Velarde, G., Weichert, M., Deshmunkh, A., Deshmane, S., Sudhir, A., Sharma, K., & Joshi, V. (2024). Tree boosting methods for balanced and imbalanced classification and their robustness over time in risk assessment. *Intelligent Systems with Applications*, 22, 200354. <https://doi.org/10.1016/j.iswa.2024.200354>

Velásquez, R. M. A., Lara, J. V. M., & Melgar, A. (2019). Converting data into knowledge for preventing failures in power transformers. *Engineering Failure Analysis*, 101, 215-229. <https://doi.org/10.1016/j.engfailanal.2019.03.027>

Velásquez, R. M. A. & Lara, J. V. M. (2020). Root cause analysis improved with machine learning for failure analysis in power transformers. *Engineering Failure Analysis*, 115, 104684. <https://doi.org/10.1016/j.engfailanal.2020.104684>

Velásquez, R. M. A. & Mejía Lara, J. V. (2018). Corrosive Sulphur effect in power and distribution transformers failures and treatments. *Engineering Failure Analysis*, 92, 240-267. <https://doi.org/10.1016/j.engfailanal.2018.05.018>

Wang, J., Wu, K., Zhu, W., & Gu, C. (2017). Condition assessment for power transformer using Health Index. *IOP Conference Series: Materials Science and Engineering*, 199, 012046. <https://doi.org/10.1088/1757-899X/199/1/012046>

Wang, J., Xie, X., Wang, P., Sun, J., Liu, Y., & Zhang, L. (2025). Incorporating symmetric smooth regularizations into sparse logistic regression for classification and feature extraction. *Symmetry*, 17(2), 151. <https://doi.org/10.3390/sym17020151>

Wattakapaiboon, W. & Pattanadech, N. (2016). The new developed Health Index for transformer condition assessment. *International Conference on Condition Monitoring and Diagnosis (CMD)*, 32-35. <https://doi.org/10.1109/CMD.2016.7757760>

Yadaiah, N. & Ravi, N. (2011). Internal fault detection techniques for power transformers. *Applied Soft Computing*, 11(8), 5259-5269. <https://doi.org/10.1016/j.asoc.2011.05.034>

Zeinoddini-Meymand, H. & Vahidi, B. (2016). Health index calculation for power transformers using technical and economical parameters. *IET Science, Measurement & Technology*, 10(7), 823-830. <https://doi.org/10.1049/iet-smt.2016.0184>

Zhang, Y., Wei, C., & Liu, X. (2022). Group logistic regression models with lp,q regularization. *Mathematics*, 10(13), 2227. <https://doi.org/10.3390/math10132227>