

UTILIZAÇÃO DE COMITÊS DE CLASSIFICADORES PARA MELHORIA DA PRECISÃO EM DIAGNÓSTICOS DE FALHAS EM COMPONENTES ELETRÔNICOS

Autores

Julyana Nazario Alves ¹

Vinicius de Oliveira Querencia ²

Antonio Cesar Galhardi ³

Resumo

A crescente complexidade dos processos de serviços de reparo de equipamentos eletrônicos impõe desafios à padronização, à redução de custos e à confiabilidade operacional. Este estudo propõe a utilização de algoritmos de aprendizado de máquina organizados em um comitê de classificadores como ferramenta de apoio à tomada de decisão em uma empresa multinacional do setor de eletrônicos. Foram aplicados os modelos *k-Nearest Neighbors (KNN)*, *Random Forest* e *Multi-Layer Perceptron (MLP)* sobre um conjunto de dados reais de reparo de placas de notebooks e desktops, combinados por meio de votação majoritária. Os resultados demonstraram que o comitê de classificadores atingiu acurácia de 98,6%, superando os modelos individuais, e apresentou robustez em métricas de precisão, *recall* e *F1-score*. O projeto também contemplou o desenvolvimento de uma interface gráfica interativa, construída em *Python*, que permite a operadores da linha de reparo acessarem o modelo preditivo sem a necessidade de conhecimento avançado em *machine learning*. A ferramenta facilita a inserção dos dados de entrada, pré-processamento automático e exibe os resultados de forma clara, agilizando a tomada de decisão e promovendo a padronização dos diagnósticos. A aplicação prática da ferramenta ocorre em um contexto *B2B*, contribuindo para a redução do tempo do processo, diminuição do retrabalho e aumento da confiabilidade na entrega dos serviços, aspectos cruciais para a competitividade da organização. Do ponto de vista gerencial, a aplicação contribui para a padronização do diagnóstico de falhas, otimização do tempo de reparo, redução de retrabalho e aumento da confiabilidade na entrega ao cliente. O estudo reforça o potencial de aplicação da Inteligência Artificial não apenas no setor industrial, mas também em empresas de serviços intensivos em conhecimento, com possibilidades de expansão para manutenção preditiva, suporte técnico e setores correlatos.

Palavras-chave: *Ensemble Learning*. Predição de Falhas. Equipamentos Eletrônicos. Interface Gráfica. *Machine Learning*.

USING CLASSIFICATION COMMITTEES TO IMPROVE ACCURACY IN DIAGNOSING FAULTS IN ELECTRONIC COMPONENTS

Abstract

The increasing complexity of electronic equipment repair service processes imposes challenges on standardization, cost reduction, and operational reliability. This study proposes the use of machine learning algorithms organized into an ensemble of classifiers as a decision support tool within a multinational electronics company. The k-Nearest Neighbors (KNN), Random Forest, and Multi-Layer Perceptron (MLP) models were applied to a real dataset of notebook and desktop motherboard repairs,

¹ ORCID 0009-0002-6534-6005

² ORCID 0009-0006-1373-4582

³ ORCID 0000-0002-8838-6870

combined through majority voting. The results demonstrated that the ensemble of classifiers achieved an accuracy of 98.6%, surpassing individual models, and exhibited robustness in precision, recall, and F1-score metrics. The project also included the development of an interactive graphical user interface, built in Python, which allows repair line operators to access the predictive model without the need for advanced machine learning knowledge. The tool facilitates input data entry, automates preprocessing, and clearly displays results, thereby streamlining decision-making and promoting diagnostic standardization. The practical application of this tool occurs within a B2B context, contributing to reduced process time, decreased rework, and enhanced reliability in service delivery—critical aspects for organizational competitiveness. From a managerial perspective, the application aids in standardizing failure diagnostics, optimizing repair time, reducing rework, and increasing reliability in customer delivery. The study underscores the potential for Artificial Intelligence applications not only in the industrial sector but also in knowledge-intensive service companies, with possibilities for expansion into predictive maintenance, technical support, and related sectors.

Keywords: Ensemble Learning. Failure Prediction. Electronic Equipment. Graphical User Interface (GUI). Machine Learning.

INTRODUÇÃO

A crescente complexidade dos processos produtivos e de serviços tem imposto novos desafios às organizações contemporâneas, em especial às que atuam em setores intensivos em tecnologia. O setor de serviços, responsável por mais de dois terços da atividade econômica global, representa atualmente o eixo central de geração de valor em economias desenvolvidas e emergentes (Fitzsimmons & Fitzsimmons,

2014). No Brasil, segundo dados da Agência Brasileira de Desenvolvimento Industrial (ABDI, 2024), mais de 52% das empresas em operação estão direcionadas à prestação de serviços, consolidando o setor terciário como o principal motor da atividade econômica. Dentro desse contexto, empresas de reparo de equipamentos eletrônicos assumem papel estratégico, uma vez que oferecem soluções que prolongam o ciclo de vida de produtos, reduzem custos de substituição e mitigam impactos ambientais associados ao descarte de resíduos tecnológicos (Slack *et al.*, 2019).

Apesar da relevância, as operações de reparo enfrentam desafios significativos relacionados à variabilidade dos processos, à heterogeneidade dos componentes e à influência direta do cliente sobre os critérios de qualidade e tempo de resposta (Johnston & Clark, 2008). Essa variabilidade compromete a padronização de diagnósticos e dificulta a consolidação de controles operacionais. Em linhas de reparo de placas-mãe de notebooks e desktops, por exemplo, a diversidade de modelos, fornecedores e tipos de falhas acarreta cenários de incerteza que impactam diretamente o planejamento da capacidade, a alocação de mão de obra e o cumprimento de prazos contratuais (SLA – *Service Level Agreement*). A gestão eficiente dessas

operações exige, portanto, ferramentas capazes de apoiar a tomada de decisão em ambientes de alta incerteza e complexidade.

Nesse cenário, a Inteligência Artificial (IA) emerge como um recurso estratégico para o setor de serviços. Avanços recentes em aprendizado de máquina e mineração de dados possibilitam a análise de grandes volumes de informações históricas, a identificação de padrões ocultos e a predição de falhas antes mesmo da intervenção humana (Russell & Norvig, 2021; Jordan & Mitchell, 2015). Em particular, a utilização de classificadores supervisionados e de técnicas de *ensemble learning* tem se mostrado eficaz em problemas de diagnóstico e manutenção preditiva (Zhang & Ma, 2012; Lee *et al.*, 2014). Ao reduzir a subjetividade do julgamento humano e fornecer previsões baseadas em evidências, tais ferramentas contribuem para maior confiabilidade no diagnóstico, otimização do tempo de reparo e redução do retrabalho.

Diante desse panorama, este estudo justifica-se pela necessidade de aplicar e avaliar modelos preditivos supervisionados no contexto de serviços de reparo de equipamentos eletrônicos, buscando apoiar a gestão operacional. A proposta é investigar como um comitê de classificadores – composto por algoritmos de diferentes naturezas – pode atuar como guia na previsão da reparabilidade de placas de *notebooks* e *desktops*. Espera-se que os resultados ofereçam subsídios para decisões estratégicas relacionadas à alocação de mão de obra, ao cumprimento de prazos e ao planejamento de recursos, contribuindo para a redução de incertezas e o aumento da eficiência operacional em empresas de serviços intensivos em conhecimento.

1.1 Objetivos

O presente estudo tem como objetivo principal desenvolver e avaliar uma ferramenta de apoio à tomada de decisão, baseada em algoritmos de aprendizado de máquina supervisionados, voltada à otimização de operações de reparo de placas-mãe de *notebooks* e *desktops* em uma empresa multinacional do setor eletrônico. Ao propor a utilização de um comitê de classificadores, busca-se demonstrar como a combinação de diferentes modelos preditivos pode aumentar a acurácia dos diagnósticos de falhas, contribuindo para a padronização de processos e a redução de incertezas em cenários de alta variabilidade operacional.

De forma mais específica, busca-se:

- Aplicar e comparar o desempenho de três algoritmos supervisionados amplamente utilizados – *k-Nearest Neighbors* (k-NN), *Random Forest* e

- *MultiLayer Perceptron* (MLP) – em um conjunto de dados reais provenientes de processos de reparo de *notebooks* e *desktops*;
- Implementar um comitê de classificadores, baseado em votação majoritária, para verificar se a combinação de diferentes modelos aumenta a precisão preditiva em relação aos classificadores individuais;
- Analisar os resultados obtidos à luz dos impactos gerenciais, destacando a contribuição da inteligência artificial para a padronização de diagnósticos, a redução de retrabalho e a melhoria da qualidade operacional em serviços de reparo;
- Implementar uma ferramenta para apoio a tomada de decisão na operação de reparo de placas-mãe de *notebooks* e *desktops*;
- Propor uma abordagem preditiva baseada em inteligência artificial que seja replicável em outros contextos de serviços intensivos em conhecimento, alinhada à busca por eficiência, inovação tecnológica e vantagem competitiva.

2. REFERENCIAL TEÓRICO

A gestão de serviços distingue-se da manufatura pela maior intensidade da participação do cliente e pela presença de variabilidade nos processos (Fitzsimmons & Fitzsimmons, 2014). Em linhas de reparo de equipamentos eletrônicos, essa variabilidade é ainda mais acentuada, uma vez que cada componente pode apresentar falhas de naturezas distintas, em diferentes modelos e gerações de dispositivos, o que dificulta a padronização dos diagnósticos e a previsibilidade do tempo de execução. Slack, Brandon-Jones e Burgess (2019) ressaltam que a variabilidade do fluxo de operações gera impactos diretos sobre custos, produtividade e qualidade percebida, tornando essencial o desenvolvimento de mecanismos de controle e previsão.

Em empresas de reparo, especialmente em modelos *business-to-business* (B2B), a previsibilidade do processo é estratégica não apenas para o cumprimento de prazos contratuais (SLA), mas também para a consolidação de vantagem competitiva baseada em confiabilidade e agilidade (Johnston & Clark, 2008). Dessa forma, a capacidade de prever a reparabilidade de placas eletrônicas antes mesmo da intervenção completa reduz desperdícios, possibilita melhor alocação da mão de obra e favorece a utilização racional dos recursos disponíveis.

Diante de um cenário marcado pela incerteza e pela complexidade, a inteligência artificial (IA) surge como uma abordagem promissora para apoiar a gestão operacional desses serviços. A IA, em especial o aprendizado de máquina, permite o processamento e a análise de

grandes volumes de dados, possibilitando a identificação de padrões complexos e a realização de previsões em contextos de alta variabilidade (Russell & Norvig, 2021). Segundo Van Klompenburg *et al.* (2022), a aplicação de técnicas de IA em serviços incrementa a precisão das decisões e reduz a dependência de análises subjetivas, contribuindo para a padronização dos processos e a otimização do uso dos recursos. A inteligência artificial facilita a construção de modelos que antecipam resultados operacionais, essencial em ambientes onde a variabilidade e a dinâmica do mercado impõem desafios contínuos.

Por sua vez, a aprendizagem de máquina é um recurso da inteligência artificial que utiliza dados e suas características para gerar informações que auxiliam pesquisadores a resolver e compreender relações de alta complexidade (Kaul *et al.*, 2005; Van Klompenburg *et al.*, 2022; Çetin e Sağlam, 2022). Nesse contexto, os classificadores supervisionados *k-NN*, *Random Forest* e *MLP* podem ser utilizados para identificar categorias ou realizar previsões a partir de dados de entrada.

O classificador *k-NN* (*k-Nearest Neighbour*) baseia-se na formação de clusters que calculam a distância dos objetos vizinhos para distinguir entre características de imagem, permitindo a classificação de amostras semelhantes. Trata-se de um classificador de aprendizagem não paramétrica, que utiliza grandes conjuntos de dados e apresenta maior eficácia com dados ruidosos, podendo facilitar a implementação (Acharya *et al.*, 2020). De acordo com o trabalho de Shahbeik *et al.* (2023), o classificador *k-NN* pode ser empregado na reconstituição de dados ausentes por meio da detecção de pontos similares com valores nulos, a fim de preenchê-los e estimar lacunas em amostras de interesse. Além disso, o *k-NN* pode ser utilizado para prever rendimentos em processos (İrik *et al.*, 2023).

O classificador *Random Forest* (RN) tem como base árvores de decisão que realizam previsões a partir dos dados de entrada por meio de divisões sucessivas em seus nós, baseadas nos atributos do conjunto de dados. No final, o classificador analisa as previsões de todas as árvores e define a saída com base na votação da maioria (comitê) (Oshiro, Perez e Baranauskas, 2012). O modelo apresenta grande potencial para prever rendimento em processos (İrik *et al.*, 2023). Além disso, o RN evita *overfitting* devido à grande variabilidade de árvores decisórias, impedindo que o classificador perca a capacidade de predição de casos já conhecidos, mesmo com a entrada de novos dados (Liu *et al.*, 2020).

O classificador *Multilayer Perceptron* (MLP) é um tipo de rede neural artificial composta por, no mínimo, uma camada de entrada, uma ou mais camadas ocultas e uma camada de saída. Cada neurônio de uma camada é totalmente conectado aos neurônios da

camada seguinte. Esse modelo é capaz de aprender padrões complexos; no entanto, modelos com ruídos e *outliers* podem interferir nas previsões do classificador (Rezaie-Balf, M. *et al.*, 2020).

Os classificadores (Tabela 1) aplicados para o treinamento dos modelos são supervisionados, ou seja, precisam de dados rotulados para o processo de aprendizagem. Os rótulos apresentam as saídas esperadas, tornando possível a aprendizagem do modelo durante o treinamento.

Tabela 1 - Classificadores de aprendizagem supervisionada.

Classificador	Tipo	Aprendizado	Descrição
Random Forest (RF)	Árvores de decisão	Supervisionado	Modelo baseado na agregação das decisões de várias árvores para classificar os dados.
k-NN (k-Nearest Neighbors)	Baseado em instâncias	Supervisionado	Classifica com base nos vizinhos mais próximos do ponto
MLP (Multilayer Perceptron)	Rede neural	Supervisionado	Modelo com camadas de neurônios que aprende padrões complexos

Fonte: Elaborado pelos autores (2025)

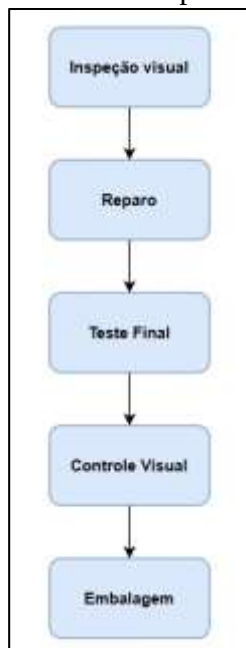
Desta forma, no contexto de serviços de reparo, a aplicação de um comitê de classificadores possibilita maior confiabilidade na previsão da reparabilidade dos equipamentos, fornecendo subsídios para decisões estratégicas de gestão de recursos e planejamento da produção. Assim, a utilização de técnicas de ensemble learning conecta-se diretamente às necessidades práticas das empresas de serviços em ambientes de alta incerteza e variabilidade.

3. MÉTODO

A pesquisa foi realizada em uma empresa prestadora de serviço de equipamentos eletrônicos, situada em São Paulo, no município de Barueri. A empresa possui diversificação de equipamentos e processos de reparo; para o presente estudo, foi escolhida a linha de reparo de *notebooks* e *desktops* de um cliente global, com grande relevância na organização.

O processo analisado inclui a linha de reparo (Figura 1), com foco na reparabilidade de placas. A linha é constituída por mão de obra especializada e não especializada. A mão de obra não especializada é responsável por análises de baixa complexidade, como inspeção visual e validações gerais de acordo com especificações requeridas. A mão de obra especializada é responsável pelo diagnóstico e reparo de placas de diferentes graus de complexidade. Essa segmentação operacional reflete a necessidade de apoio na previsão da reparabilidade dos componentes, de modo a otimizar a alocação de recursos e o cumprimento dos prazos de atendimento.

Figura 1 - Fluxo reparo



Fonte: Elaborado pelos autores (2025)

As informações obtidas desta linha de reparo foram reunidas em um conjunto de dados (*dataset*) em formato de planilha eletrônica com os metadados apresentados no Quadro 1:

Quadro 1 - Metadados do dataset

Coluna	Descrição	Tipo da variável
<i>Product</i>	Código do produto	entrada (<i>feature</i>)
PN	<i>Partnumber</i> do produto	entrada (<i>feature</i>)
<i>Model</i>	Modelo do produto	entrada (<i>feature</i>)
<i>Category</i>	Tipo do produto	entrada (<i>feature</i>)
<i>Part Type</i>	Equipamento com defeito	entrada (<i>feature</i>)
<i>Result</i>	Resultado dos testes após o reparo	saída (<i>target</i>)

Fonte: Elaborado pelos autores (2025)

Este *dataset* contém informações relevantes para a tarefa de classificação, cujo objetivo é prever quais equipamentos puderam ser reparados com sucesso e quais precisaram ser devolvidos ao cliente sem que o reparo tenha sido bem-sucedido.

Para realizar a predição, as colunas de entrada (*features*) selecionadas para o treinamento dos modelos foram: 'PN', 'Category', 'Model' e 'Part Type'. O préprocessamento dos dados seguiu as seguintes etapas:

- Carregamento dos dados:** O arquivo de planilha eletrônica foi lido utilizando a biblioteca *pandas*.
- Codificação de variáveis categóricas:** As variáveis categóricas presentes nas *features* ('PN', 'Category', 'Model', 'Part Type') foram convertidas em representações numéricas

utilizando LabelEncoder da biblioteca *sklearn.preprocessing*. Este passo foi necessário, pois a maioria dos algoritmos de aprendizagem de máquina requer entradas numéricas.

- c) **Separação de *features* e *target*:** As colunas de entrada (*features*) foram separadas da coluna de saída (*target*), que é a variável a ser prevista ('*Result*').
- d) **Normalização dos dados:** As *features* numéricas (após a codificação das categóricas) foram normalizadas utilizando *StandardScaler* da biblioteca *sklearn.preprocessing*. A normalização foi realizada para garantir que todas as *features* contribuam igualmente para o treinamento do modelo, evitando que *features* com escalas maiores dominem o processo de aprendizagem.
- e) **Divisão em conjuntos de treino e teste:** O conjunto de dados pré-processado foi dividido em conjuntos de treino e teste, com 70% dos dados destinados ao treinamento dos modelos e 30% para a avaliação de seu desempenho. A divisão foi realizada de forma aleatória, garantindo a reprodutibilidade dos resultados por meio da definição de uma semente.

Para composição do comitê de classificadores, foram selecionados três algoritmos supervisionados que apresentaram complementares características de aprendizagem:

1. ***K-Nearest Neighbors (k-NN)*:** Um classificador não paramétrico que classifica uma instância com base na maioria das classes de seus *k* vizinhos mais próximos no espaço de características. Para este estudo, foi configurado com 4 vizinhos.
2. ***MultiLayer Perceptron (MLP)*:** Uma rede neural artificial *feedforward*, treinada com o algoritmo de retropropagação. Para esta aplicação foi configurada com duas camadas ocultas compostas por 10 e 5 neurônios, função de ativação ReLU, otimizador Adam e máximo de 10.000 iterações.
3. ***Random Forest (Random Forest)*:** Um método de *ensemble* baseado em árvores de decisão. Ele constrói múltiplas árvores de decisão durante o treinamento e produz a classe que é a moda das classes (classificação) ou a previsão média
4. (regressão) das árvores individuais. Para este trabalho, foi utilizado com 500 árvores na floresta. Cada um desses modelos foi treinado independentemente utilizando o conjunto de dados de treino e, posteriormente, utilizado para fazer previsões no conjunto de teste.

Os classificadores *k-NN*, *MLP* e *Random Forest* foram treinados e avaliados individualmente. Em seguida, um comitê de classificadores foi implementado para combinar suas previsões. O comitê de classificadores foi implementado utilizando uma estratégia de

votação majoritária. Nesta abordagem, as previsões de cada um dos classificadores individuais (*k*-NN, MLP e *Random Forest*) são coletadas para cada instância do conjunto de teste. A classe final prevista pelo comitê é aquela que recebe o maior número de votos entre os modelos. Em caso de empate na votação, uma regra de desempate foi estabelecida, onde a previsão do modelo *Random Forest* é utilizada como decisão final.

Essa função recebe as previsões de cada modelo para o conjunto de teste e retorna um *array* com as previsões do comitê. A avaliação do desempenho do comitê foi realizada utilizando a acurácia e o relatório de classificação, que inclui precisão, *recall* e *F1-score* para cada classe, além do suporte. Essa metodologia permite uma análise do desempenho do comitê em comparação com os modelos individuais.

Para ampliar a acessibilidade e facilitar a implementação dos modelos na operação diária da linha de reparo, desenvolveu-se uma interface gráfica interativa utilizando a biblioteca *Tkinter* em linguagem de programação *Python*. Essa ferramenta foi projetada para oferecer aos operadores, independentemente do grau de familiaridade com aprendizado de máquina, uma forma intuitiva e eficiente de obter as predições do modelo. A interface apresenta campos de entrada com caixas de seleção pré-populadas com os valores únicos extraídos do *dataset* de treino, minimizando erros de digitação e garantindo a consistência dos dados inseridos. Após o preenchimento dos campos (PN, *Category*, *Model* e *Part Type*), o operador aciona o comando para realizar a predição, momento em que os dados são automaticamente pré-processados — com codificação e normalização — e enviados aos três modelos para obtenção das respectivas previsões. O resultado do comitê é então exibido, acompanhado de uma indicação textual e visual clara (“APROVADO” ou “RETORNAR AO CLIENTE”), facilitando a tomada de decisão imediata.

A ferramenta também contempla funcionalidades de validação automática dos dados e opção para limpeza dos campos, favorecendo múltiplas utilizações em sequência sem necessidade de reinício da aplicação. Essa solução tecnológica representa uma ponte entre a complexidade dos modelos de *machine learning* e a operacionalização prática, contribuindo para a redução do tempo de diagnóstico, diminuição da variabilidade associada ao fator humano e aumento da confiabilidade do processo produtivo na linha de reparo.

4. RESULTADOS E DISCUSSÃO

Após o treinamento dos classificadores individuais (*k*-Nearest Neighbors, *MultiLayer Perceptron* e *Random Forest*) e a implementação do comitê de classificadores, os modelos

foram avaliados no conjunto de teste, que correspondeu a 30% da base de dados total. Os resultados revelaram desempenhos consistentes dos algoritmos individuais, com destaque para o *Random Forest* e o *MultiLayer Perceptron*, que atingiram acurácia de 98,36%. O classificador k-NN apresentou desempenho ligeiramente inferior, refletindo sua maior sensibilidade a ruídos e à distribuição dos dados. A Figura 2 apresenta um resumo da acurácia alcançada por cada classificador individual e pelo comitê.

Figura 2 - Acurácia dos Classificadores Individuais e do Comitê

---	Desempenho Individual dos Classificadores	---
Acurácia KNN:	0.9269	
Acurácia MLP:	0.9836	
Acurácia Random Forest:	0.9836	
Acurácia do Comitê:	0.9858	

Fonte: Elaborado pelos autores (2025)

A análise dos resultados revela o desempenho dos classificadores individuais e do comitê. O *Random Forest* e o *MLP* foram os modelos individuais de melhor desempenho, alcançando as maiores acurácias. O desempenho superior do *Random Forest* pode ser atribuído à sua natureza de *ensemble*, que combina múltiplas árvores de decisão para reduzir o *overfitting* e melhorar a generalização. A robustez do *Random Forest* em lidar com diferentes tipos de dados e sua capacidade de capturar relações complexas contribuem para seu desempenho superior. O *MLP*, por sua vez, demonstrou sua capacidade de aprender padrões complexos, alcançando boa acurácia mesmo que, em algumas execuções, tenha gerado um aviso de convergência (*ConvergenceWarning*), indicando que o otimizador estocástico atingiu o número máximo de iterações (10.000) e a otimização ainda não havia convergido.

A agregação dos modelos por meio da estratégia de votação majoritária resultou em um ganho marginal, mas estatisticamente relevante, com o comitê de classificadores alcançando acurácia de 98,58%. Embora o incremento em relação aos modelos de melhor desempenho individual seja pequeno, esse resultado confirma o princípio fundamental do *ensemble learning*, segundo o qual a combinação de modelos heterogêneos tende a aumentar a robustez e a confiabilidade preditiva (Zhang & Ma, 2012). Esse comportamento foi evidenciado na análise das métricas de precisão, recall e F1-score, que demonstraram equilíbrio entre as classes “Aprovado” e “Retornar ao Cliente”, reduzindo assim o risco de vieses na classificação.

Para uma análise mais detalhada do desempenho do comitê, o relatório de classificação é apresentado a seguir, detalhando métricas como precisão, *recall* e *F1score* para cada classe (P – *Pass* e T – *Return to Customer*).

Figura 3 – Relatório de Classificação do Comitê

	precision	recall	f1-score	support
P - Pass	0.98	1.00	0.99	691
T - Return to Customer	1.00	0.94	0.97	225
accuracy			0.99	916
macro avg	0.99	0.97	0.98	916
weighted avg	0.99	0.99	0.99	916

Fonte: Elaborado pelos autores (2025)

A análise detalhada por classe mostrou que, para equipamentos reparáveis (“Aprovado”), o comitê alcançou precisão de 0,98 e *recall* de 1,00, resultando em *F1score* de 0,99. Já para os equipamentos não reparáveis (“Retornar ao Cliente”), obteve-se precisão de 1,00 e *recall* de 0,94, com *F1-score* de 0,97. Esses resultados indicam que o sistema foi eficaz em evitar falsos positivos, ou seja, em não classificar equivocadamente um equipamento irrecuperável como reparável. Tal característica é crítica em contextos de reparo, pois diagnósticos incorretos implicam custos adicionais de retrabalho, atrasos no cumprimento de SLA e perda de confiabilidade junto aos clientes.

Do ponto de vista gerencial, os resultados demonstram que a adoção de modelos preditivos baseados em inteligência artificial pode contribuir para a padronização de diagnósticos e a otimização da operação de reparo. Ao reduzir a

variabilidade associada ao julgamento humano, a ferramenta auxilia na alocação mais eficiente de recursos, uma vez que permite direcionar casos mais complexos à mão de obra especializada e casos rotineiros a técnicos de menor senioridade. Essa prática contribui tanto para a redução do tempo médio de reparo quanto para o aumento da taxa de aproveitamento da mão de obra, o que está diretamente relacionado à produtividade da linha.

Para tornar o modelo de comitê de classificadores acessível e aplicável em um ambiente de linha de reparo, foi desenvolvida uma ferramenta interativa em *Python*. Esta ferramenta oferece uma interface gráfica que permite a operadores, mesmo sem conhecimento aprofundado em *machine learning*, inserir dados de equipamentos e obter instantaneamente o resultado da predição de falhas. A ferramenta abstrai a complexidade dos modelos de *aprendizagem de máquina*, fornecendo uma solução prática e eficiente para o diagnóstico. Outro impacto observado refere-se à redução de retrabalho, pois a utilização do comitê de classificadores como apoio à decisão pode minimizar diagnósticos equivocados, que são uma das principais fontes de desperdício em linhas de serviço de alta variabilidade. Esse resultado converge com os apontamentos de Slack, Brandon-Jones e Burgess (2019), segundo os quais

a previsibilidade e a consistência do fluxo de operações estão entre os fatores determinantes da eficiência em serviços.

Por sua vez, a implementação da interface gráfica em *Python* pode trazer benefícios ao tornar a solução acessível a operadores com diferentes níveis de conhecimento técnico. Essa interface não apenas traduz a complexidade matemática dos algoritmos em um resultado de fácil interpretação, mas também viabiliza a integração da solução ao fluxo de trabalho cotidiano da linha de reparo. A interface gráfica foi desenvolvida utilizando a biblioteca *Tkinter* em *Python*, garantindo compatibilidade com diversos sistemas operacionais.

Ao executar a aplicação a janela principal é exibida conforme demonstrado na Figura 4. Neste momento a ferramenta verifica automaticamente a existência dos modelos treinados. Caso não os encontre, um aviso é exibido, e o usuário é instruído a executar previamente o script de treinamento.

Figura 4 – Interface gráfica para predição de falhas

Fonte: Elaborado pelos autores (2025)

A interface apresenta quatro campos de entrada essenciais para a predição: o **PN (Part Number)**: código de identificação único do produto.

- **Category**: categoria geral do equipamento (ex.: *MAINBOARD*, *CPU/PROCESSOR*).
- **Model**: modelo específico do equipamento.
- **Part Type**: tipo de componente defeituoso identificado.

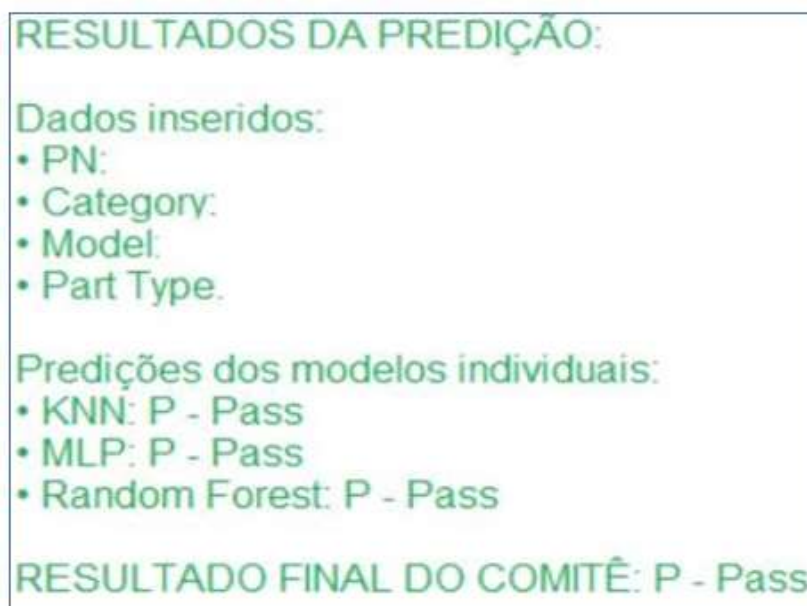
Para facilitar a entrada e garantir a consistência dos dados, os campos são prépopulados com todos os valores únicos existentes no *dataset* de treinamento. Isso minimiza erros de digitação e garante que apenas valores válidos, com os quais os modelos foram treinados, sejam inseridos. Se um valor não presente no *dataset* for inserido manualmente, o sistema emitirá um aviso e utilizará um valor padrão para a predição, embora isso possa impactar a precisão.

Após o preenchimento dos campos, o operador clica no botão "Fazer Predição". A ferramenta então coleta os valores inseridos e pré-processa esses valores da mesma forma que os dados de treinamento. Então a mesma alimenta os três modelos individuais (*k-NN*, *MLP*, *Random Forest*) com os dados pré-processados para obter suas respectivas previsões e aplica a lógica de votação majoritária do comitê para determinar o resultado.

Os resultados são apresentados na interface gráfica da ferramenta, onde são exibidos os dados de entrada fornecidos pelo operador, as previsões individuais geradas pelos modelos *k-Nearest Neighbors (k-NN)*, *Multi-Layer Perceptron (MLP)* e *Random Forest*, bem como o resultado final do comitê de classificadores, que constitui a predição mais relevante para o operador. Acompanhando a predição final, uma indicação visual e textual é fornecida para facilitar a interpretação imediata: "APROVADO" para a classe *P – Pass* ou "RETORNAR AO CLIENTE" para a classe *T – Return to Customer*.

A Figura 5 ilustra um exemplo de saída da predição para resultados positivos, onde o equipamento é classificado como "APROVADO".

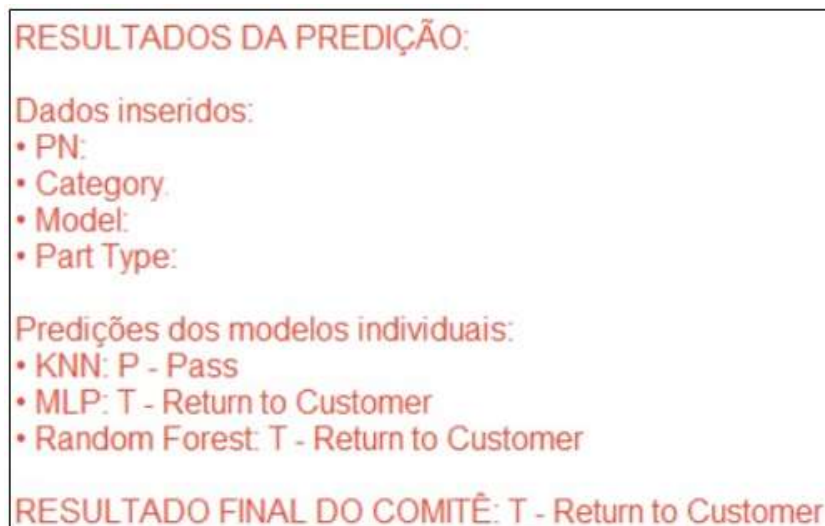
Figura 5 - Exemplo de saída da predição para resultados positivos "APROVADO"



Fonte: Elaborado pelos autores (2025)

De maneira complementar, a Figura 6 apresenta um exemplo de saída da predição para resultados negativos, classificando o equipamento como "RETORNAR AO CLIENTE".

Figura 6 - Exemplo de saída da predição para resultados negativos "RETORNAR AO CLIENTE"



Fonte: Elaborado pelos autores (2025)

A implementação da ferramenta gráfica representa um avanço na aplicabilidade prática dos modelos desenvolvidos. A interface permite a validação automática de dados, garantindo que o sistema possa ser utilizado efetivamente por operadores com diferentes níveis de experiência técnica. A pré-população dos campos com valores válidos do *dataset* minimiza erros de entrada e maximiza a confiabilidade das predições. Do ponto de vista operacional, a ferramenta oferece benefícios tangíveis que se traduzem em melhorias mensuráveis de eficiência. A redução do tempo de diagnóstico representa uma possibilidade de otimização do fluxo de trabalho da linha de reparo. Além disso, a consistência das predições elimina a variabilidade associada ao fator humano, possibilitando maior confiabilidade do processo.

CONSIDERAÇÕES FINAIS

O presente estudo explorou a eficácia da utilização de comitês de classificadores compostos pelos algoritmos *k-Nearest Neighbors (k-NN)*, *Multi-Layer Perceptron (MLP)* e *Random Forest* para o aprimoramento da precisão em diagnósticos de falhas em componentes eletrônicos, especificamente na linha de reparo de placas-mãe de notebooks e desktops. A superação das acurácias individuais pelos resultados do ensemble confirma a premissa de que a combinação de modelos baseados em diferentes princípios de funcionamento pode gerar previsões mais robustas e confiáveis, mitigando as limitações inerentes a cada classificador isoladamente.

Desta forma, o presente trabalho evidencia a importância da interface entre aprendizado de máquina e gestão de serviços industriais, demonstrando a viabilidade e o impacto das técnicas de *ensemble learning* aplicadas a dados reais de operações complexas e suscetíveis à variabilidade. O estudo reforça a importância de abordagens multidisciplinares que fundamentam a tomada de decisão em análises quantitativas robustas, estimulando futuras investigações na aplicação de técnicas híbridas, otimizações de hiper parâmetros e ampliação do escopo para áreas correlatas como manutenção preditiva e suporte técnico.

Por sua vez, a gestão operacional pode se beneficiar do desenvolvimento de uma ferramenta que facilita a padronização dos diagnósticos e reduz a variabilidade associada a avaliações subjetivas e ao fator humano. A implementação de uma interface gráfica intuitiva, desenvolvida em *Python*, amplia a aplicabilidade prática do modelo na linha de reparo, democratizando o acesso às previsões e promovendo maior agilidade e confiabilidade no processo decisório. Isso pode implicar em redução de custos operacionais, diminuição de retrabalho e melhoria na satisfação do cliente, fortalecendo a competitividade da empresa no mercado *B2B*.

As perspectivas de pesquisa futura incluem a exploração de novas técnicas de combinação de classificadores, bem como o ajuste fino dos hiper parâmetros dos modelos empregados para potencializar ainda mais a acurácia e a robustez das previsões. Adicionalmente, recomenda-se a incorporação de conjuntos de dados mais variados e detalhados, que possam captar outras dimensões do processo de reparo, como variáveis temporais e contextuais. A investigação da aplicabilidade da metodologia em outros setores de serviços intensivos em conhecimento constitui outra oportunidade importante, possibilitando a expansão do uso da inteligência artificial para promover ganhos expressivos em eficiência operacional em diversos contextos industriais e comerciais.

Referências

- ACHARYA, S. et al. A long short term memory deep learning network for the classification of negative emotions using EEG signals. In: **INTERNATIONAL JOINT CONFERENCE ON NEURAL NETWORKS**, 2020. Anais. 2020. DOI: <https://doi.org/10.1109/IJCNN48605.2020.9207280>.
- AGÊNCIA BRASILEIRA DE DESENVOLVIMENTO INDUSTRIAL (ABDI). **PIB do 2º Trimestre de 2024: indústria e serviços garantem recuperação, 2024**. Disponível em: <https://www.abdi.com.br>. Acesso em: 23 ago. 2025.
- ÇETIN, N.; SAĞLAM, C. Rapid detection of total phenolics, antioxidant activity and ascorbic acid of dried apples by chemometric algorithms. *Food Bioscience*, v. 47, p. 101670, 2022. DOI: <https://doi.org/10.1016/j.fbio.2022.101670>.

FITZSIMMONS, J. A.; FITZSIMMONS, M. J. **Administração de Serviços**: operações, estratégia e tecnologia da informação. 7. ed. Porto Alegre: McGraw-Hill, 2014.

JOHNSTON, R.; CLARK, G. **Service Operations Management: Improving Service Delivery**. 3. ed. Harlow: Pearson, 2008.

JORDAN, M. I.; MITCHELL, T. M. Machine learning: Trends, perspectives, and prospects. **Science**, v. 349, n. 6245, p. 255-260, 2015.

LEE, J.; BAGHERI, B.; KAO, H. A. A cyber-physical systems architecture for industry 4.0-based manufacturing systems. **Manufacturing Letters**, v. 3, p. 18-23, 2014.

LIU, D. *et al.* Random Forest regression evaluation model of regional flood disaster resilience based on the whale optimization algorithm. **Journal of Cleaner Production**, v. 250, art. 119468, 2020.

MINISTÉRIO DA INDÚSTRIA, COMÉRCIO E SERVIÇOS (Brasil). **Mapa de Empresas: novembro de 2024**. 2024. Disponível em: <https://www.gov.br/mapa-de-empresas>. Acesso em: 22 ago. 2025.

OSHIRO, T. M.; PEREZ, P. S.; BARANAUSKAS, J. A. How many trees in a random forest? In: **INTERNATIONAL WORKSHOP ON MACHINE LEARNING AND DATA MINING IN PATTERN RECOGNITION**, 2012, Berlin. *Proceedings...* Berlin: Springer, 2012. p. 154-168.

REZAIE-BALF, M. *et al.* Physicochemical parameters data assimilation for efficient improvement of water quality index prediction: comparative assessment of a noise suppression hybridization approach. **Journal of Cleaner Production**, v. 271, art. 122576, 2020.

RUSSELL, S.; NORVIG, P. **Artificial Intelligence: A Modern Approach**. 4. ed. Harlow: Pearson, 2021.

SLACK, N.; BRANDON-JONES, A.; BURGESS, N. **Operations Management**. 8. ed. Harlow: Pearson, 2019.

TASAN, S. *et al.* Estimation of eggplant yield with machine learning methods using spectral vegetation indices. **Computers and Electronics in Agriculture**, v. 202, p. 107367, 2022. DOI: <https://doi.org/10.1016/j.compag.2022.107367>.

VAN KLOMPENBURG, T. *et al.* Artificial Intelligence in Service Operations: A Review and Research Agenda. **International Journal of Operations & Production Management**, v. 42, n. 5, p. 631-653, 2022.

Zhang, C.; Ma, Y. **Ensemble Machine Learning: Methods and Applications**. Springer, 2012.