

**UNIVERSIDADE FEDERAL DE PERNAMBUCO
CENTRO DE CIÊNCIAS SOCIAIS E APLICADAS
DEPARTAMENTO DE CIÊNCIAS CONTÁBEIS E ATUARIAIS
CURSO DE CIÊNCIAS CONTÁBEIS**

**O APRENDIZADO DE MÁQUINA NA CONTABILIDADE: EFETIVIDADE E
IMPLICAÇÕES ASSOCIADAS A APLICAÇÃO DA TECNOLOGIA EM CONTEXTO
CONTÁBIL**

Autor(a): Pedro Gomes Beserra Filho

Recife, 2025

**UNIVERSIDADE FEDERAL DE PERNAMBUCO
CENTRO DE CIÊNCIAS SOCIAIS E APLICADAS
DEPARTAMENTO DE CIÊNCIAS CONTÁBEIS E ATUARIAIS
CURSO DE CIÊNCIAS CONTÁBEIS**

**TÍTULO: O APRENDIZADO DE MÁQUINA NA CONTABILIDADE:
EFETIVIDADE E IMPLICAÇÕES ASSOCIADAS A APLICAÇÃO DA
TECNOLOGIA EM CONTEXTO CONTÁBIL**

**Autor(a): Pedro Gomes Beserra Filho
Orientador(a): Luiz Carlos Marques dos Anjos**

Trabalho de Conclusão de Curso apresentado ao Departamento de Ciências Contábeis e Atuariais da Universidade Federal de Pernambuco – UFPE, como requisito parcial para obtenção do grau de Bacharel em Ciências Contábeis, sob a orientação do Professor Luiz Carlos Marques dos Anjos.

Recife, 2025

Ficha de identificação da obra elaborada pelo autor,
através do programa de geração automática do SIB/UFPE

Beserra Filho, Pedro Gomes.

O aprendizado de máquina na contabilidade: Efetividade e implicações associadas a aplicação da tecnologia em contexto contábil / Pedro Gomes Beserra Filho. - Recife, 2025.

87 p. : il., tab.

Orientador(a): Luiz Carlos Marques dos Anjos

Trabalho de Conclusão de Curso (Graduação) - Universidade Federal de Pernambuco, Centro de Ciências Sociais Aplicadas, Ciências Contábeis - Bacharelado, 2025.

Inclui referências, apêndices.

1. Contabilidade. 2. Aprendizado de máquina. 3. Transformação profissional. 4. Processos organizacionais. 5. Estimativa. 6. Auditoria. I. dos Anjos, Luiz Carlos Marques. (Orientação). II. Título.

600 CDD (22.ed.)

FOLHA DE APROVAÇÃO

PEDRO GOMES BESERRA FILHO

O APRENDIZADO DE MÁQUINA NA CONTABILIDADE: EFETIVIDADE E IMPLICAÇÕES ASSOCIADAS A APLICAÇÃO DA TECNOLOGIA EM CONTEXTO CONTÁBIL

**Trabalho de Conclusão de Curso apresentado ao Curso de Ciências Contábeis da
Universidade Federal de Pernambuco – UFPE, como requisito parcial para obtenção do
grau de Bacharel em Ciências Contábeis.**

Aprovado em 14/08/2025.

BANCA EXAMINADORA

Prof. Dr. Luiz Carlos Marques dos Anjos (Orientador)
Universidade Federal de Pernambuco

Prof. Dr. Rommel de Santana Freire
Universidade Federal da Paraíba

Prof. Dr. Vinicius Gomes Martins
Universidade Federal de Pernambuco

LISTA DE FIGURAS

Figura 1 – Matriz de confusão - XGBoost - Modelo 3.....	43
Figura 2 – Matriz de confusão - XGBoost - Modelo 4.....	43
Figura 3 – Gráfico SHAP do impacto médio de cada variável.....	44
Figura 4 – Gráfico de impacto nos resultados – SHAP.....	45
Figura 5 – Matriz de confusão - XGBoost - Modelo 1 (scoring).....	46
Figura 6 – Matriz de confusão - XGBoost - Modelo 5 (ROC -AUC).....	46
Figura 7 – Matriz de confusão - RF.....	48
Figura 8 – Matriz de confusão da média dos resultados (retreino).....	48
Figura 9 – Gráfico de impacto nos resultados - Random Forest – SHAP.....	49
Figura 10 – Gráfico de impacto nos resultados - XGBoost – SHAP.....	50

RESUMO

Este trabalho investiga como o aprendizado de máquina pode ser aplicado na contabilidade, partindo da constatação de que, apesar de seu uso consolidado em setores como o financeiro, sua exploração na área contábil ainda é limitada. Os estudos analisados indicam a crescente relevância da tecnologia e seu alto potencial de transformação na área, promovendo maior capacidade preditiva, a valorização da informação contábil e a reconfiguração de estruturas organizacionais, impactando fluxos operacionais, estruturas de governança e estratégias de negócio. Também são discutidas as implicações sobre o perfil do profissional contábil, cada vez mais demandado a desenvolver competências analíticas, tecnológicas e interpretativas diante da automação de tarefas rotineiras. Na vertente aplicada, foram implementados e testados dois algoritmos supervisionados (*Random Forest* e *XGBoost*) sobre bases públicas voltadas à previsão de inadimplência e detecção de fraudes em auditoria. Os modelos apresentaram alto desempenho preditivo, com base em métricas robustas, e boa capacidade de explicação por meio de técnicas de interpretabilidade baseadas em SHAP, indicando a viabilidade prática de sua adoção em contextos contábeis. Os resultados reforçam que o aprendizado de máquina amplia não apenas a capacidade técnica da contabilidade, mas também impõe uma reconfiguração organizacional e profissional. Por fim, destaca-se a importância de pesquisas futuras que explorem dados aplicados a realidades do setor público e privado brasileiro, com ênfase na diversidade metodológica com novos algoritmos e aplicações.

Palavras-Chave: Contabilidade; Aprendizado de máquina; Inteligência artificial; Processos organizacionais; Transformação profissional; Estimativa, Auditoria.

ABSTRACT

This study investigates how machine learning can be applied in accounting, based on the finding that, despite its established use in sectors such as finance, its adoption in the accounting field remains limited. The reviewed literature indicates the growing relevance of this technology and its high transformative potential in the area, enhancing predictive capacity, increasing the analytical value of accounting information, and reshaping organizational structures by impacting operational flows, governance, and business strategies. The study also discusses implications for the accounting professional, who is increasingly expected to develop analytical, technological, and interpretive skills in response to the automation of routine tasks. On the applied front, two supervised algorithms (Random Forest and XGBoost) were implemented and tested using public datasets related to default prediction and fraud detection in auditing. The models demonstrated strong predictive performance, based on robust evaluation metrics, and good explanatory power through SHAP-based interpretability techniques, suggesting the practical feasibility of adopting such models in accounting contexts. The results reinforce that machine learning enhances not only the technical capabilities of accounting but also drives a broader professional and organizational transformation. Finally, the study highlights the importance of future research using data representative of Brazilian public and private sector realities, with emphasis on methodological diversification and the expansion of algorithmic applications.

Keywords: Accounting; Machine learning; Artificial intelligence; Organizational processes; Professional transformation; Estimates; Audit.

SUMÁRIO

1. INTRODUÇÃO	7
1.1 Contextualização Do Problema	8
1.2 Objetivos.....	9
<i>1.2.1 Geral.....</i>	<i>9</i>
<i>1.2.2 Específicos</i>	<i>9</i>
1.3 Justificativa	9
2. REFERENCIAL TEÓRICO	11
2.1 Constante Evolução Tecnológica.....	11
2.2 Machine Learning	14
<i>2.2.1 Definição e abordagens de aprendizagem</i>	<i>14</i>
<i>2.2.2 Construção de modelos ML.....</i>	<i>15</i>
<i>2.2.3 Algoritmos E Modelos</i>	<i>18</i>
<i>2.2.4 Arvore de decisão</i>	<i>18</i>
<i>2.2.5 Random Forest</i>	<i>19</i>
<i>2.2.6 Boosting e XGBoost.....</i>	<i>20</i>
2.3 Aplicação Do Aprendizado De Máquina Na Área Financeira	21
2.4 Aplicação Do Aprendizado De Máquina Na Área Contábil.....	23
<i>2.4.1 Relevância Analítica Da Informação Contábil</i>	<i>23</i>
<i>2.4.2 Estimativa Contábil</i>	<i>25</i>
<i>2.4.3 Auditoria.....</i>	<i>27</i>
<i>2.4.4 Análise Textual</i>	<i>30</i>
2.5 Implicações Da Tecnologia De IA Para As Organizações	31
2.6 Transformação Da Profissão E Profissionais Contábeis	34
3. PROCEDIMENTOS METODOLÓGICOS	37
4. ANÁLISE DOS RESULTADOS	43
4.1 Resultados – Modelos De Estimativa	43
4.2 Resultados – Modelos De Auditoria.....	49
4.3 Análise das Implicações associadas a adoção do ML	52
5. CONCLUSÕES.....	55
5.1 Limitações	56
5.2 Sugestões Para Pesquisas Futuras	57
REFERÊNCIAS	59
APÊNDICE A – INADIMPLÊNCIA - ALGORITMO BASE PRINCIPAL (RF)	74
APÊNDICE B – INADIMPLÊNCIA – RECORTE DE TREINO (XGB).....	80
APÊNDICE C – AUDITORIA – ALGORITMO PRINCIPAL (RF)	82
APÊNDICE D – AUDITORIA – RECORTE DE TREINO (XGB).....	85

1. INTRODUÇÃO

O desenvolvimento da informática ao longo do século XX transformou profundamente como as empresas processam e utilizam informações. A popularização dos computadores, da internet e de softwares específicos de contabilidade ampliou a capacidade de registrar, controlar e interpretar dados, oferecendo suporte decisivo para a gestão e para a confiabilidade das demonstrações financeiras (Dechow; Granlund; Mouritsen, 2006; Arruda; Gomes; Santos, 2013).

Em paralelo, a globalização acelerou a competição e reduziu barreiras de informação, exigindo das organizações maior agilidade e precisão na produção de relatórios, o que reposicionou o papel do contador para além da escrituração, demandando habilidades analíticas e domínio de tecnologias (Martins et al., 2012; Mohamed; Lashine, 2003).

No que tange ao progresso do desenvolvimento tecnológico, também houve o aumento de dispositivos interconectados através da internet, relacionado ao conceito de Internet das coisas e, também o crescimento da escala em que dados são gerados, ligado ao fenômeno de *big data*, que se tornou predominantemente relevante na última década (Google Ngram, 2024), referente aos desafios relacionados ao crescente volume, variedade e velocidade de dados gerados.

Com o avanço da era digital, a crescente geração de dados em larga escala — fenômeno do *big data* — e a expansão da internet elevaram os desafios relacionados ao volume, velocidade e variedade das informações. Nesse ambiente, a ciência de dados surgiu como resposta para a exploração e análise desses dados complexos, integrando estatística, computação e inteligência artificial para apoiar decisões estratégicas, assim como cresceu a relevância de profissionais aptos ao uso dessas ferramentas (Fawcett; Provost, 2013; Mathur, 2023).

Entre as soluções de maior impacto destaca-se a inteligência artificial, que possibilita análises profundas e simula processos cognitivos humanos, já consolidada em grandes corporações como Netflix, Facebook e Google (Mit Sloan, 2021; Bhattacharjee; Bose; Dey, 2022). Pertencente ao campo da IA, o aprendizado de máquina (ML) se diferencia de algumas soluções pela flexibilidade em aprender e se ajustar a partir dos dados sem depender de instruções fixas (Theobald, 2021). Embora originado na década de 1950, o ML ganhou protagonismo nas últimas décadas em áreas como saúde e serviços personalizados (Shehab, 2022; Netflix, 2023), contudo, sua inserção na contabilidade ainda não chegou à maturidade

(Bertomeu *et al.*, 2020).

Não obstante, existem diversas pesquisas do exterior que exploram o efeito de modelos de aprendizado em análises de dados estruturados ou não estruturados, diferente do observado em território nacional, onde a escassez de pesquisas sobre tema em âmbito contábil é maior.

Esse cenário evidencia um ponto central: a contabilidade, historicamente pautada por funções de registro, já avançou em direção a uma prática analítica e estratégica com a introdução de sistemas informatizados. Contudo, os atuais desafios da área e a demanda por eficiência, confiabilidade, qualidade e tempestividade das informações, requerem ferramentas robustas e de alta performance. O aprendizado de máquina se apresenta como uma dessas possíveis ferramentas, capaz de ampliar a qualidade das análises, reduzir vieses humanos e oferecer suporte decisivo em ambientes complexos.

O avanço do aprendizado de máquina na ciência contábil ainda representa possível maior aprofundamento na direção analítica da rotina do profissional, onde o mesmo alocaria maior tempo na extração de insights das análises e consultoria estratégica, em vez de passar maior tempo na preparação dos dados (IFAC, 2018). Sendo assim, a tendência e a demanda por automação e os avanços da IA andam lado a lado, e representam uma grande mudança de paradigma no campo e profissão contábil (Thomson Reuters, 2024a), o que tornam esses temas importantes nesse sentido.

Neste estudo, exploraremos a posição do *machine learning* em relação a sua utilização na contabilidade, apresentando uma visão geral sobre o tema, abordando possíveis aplicações, sua efetividade nessas, destacando oportunidades e consequências apontadas a partir da literatura sobre o tema.

1.1 Contextualização Do Problema

Os profissionais da contabilidade enfrentam atualmente um cenário cada vez mais desafiador, marcado pelo crescimento da complexidade econômica e pela necessidade de oferecer informações rápidas, confiáveis e relevantes para apoiar a tomada de decisão. As organizações, por sua vez, lidam com volumes crescentes de dados e com a pressão para transformá-los em *insights* que sustentem sua competitividade, o que torna a gestão contábil mais exigente e dependente de recursos tecnológicos.

Nesse contexto, softwares e sistemas integrados já trouxeram avanços significativos, ao automatizar tarefas operacionais e reduzir falhas humanas. Contudo, diante da escala,

velocidade e diversidade dos dados atuais, essas soluções encontram limitações. A contabilidade, para manter sua relevância estratégica, precisa ampliar sua capacidade de análise, indo além da simples automação e incorporando abordagens capazes de lidar com cenários de maior incerteza e complexidade.

Entre as tecnologias emergentes, o aprendizado de máquina (*machine learning*) tem se destacado pelo potencial de transformar dados em previsões e diagnósticos de alta precisão. Mais do que acelerar processos, esses modelos poderiam auxiliar a melhorar as informações, ajudar a lidar com questões relacionadas a gestão. Diferente de outras ferramentas já consolidadas, o ML não apenas executa, mas aprende com os dados, ampliando o alcance da análise contábil.

Entretanto, permanece a dúvida sobre como essa tecnologia poderia ser aplicada de forma efetiva dentro da profissão e quais impactos traria em termos de prática, governança e papel do profissional. Diante disso, surge a questão central deste estudo: qual a efetividade e as implicações associadas à aplicação da tecnologia em contexto contábil?

1.2 Objetivos

1.2.1 Geral

O objetivo geral deste trabalho é investigar a efetividade e as implicações da aplicação efetiva do aprendizado de máquina em âmbito contábil.

1.2.2 Específicos

Demonstrar a efetividade dos modelos supervisionados Random Forest e XGBoost em cenários de estimativa e auditoria.

Explorar o uso de ferramentas de interpretabilidade (SHAP) como suporte à análise dos resultados.

Analisar, a partir da literatura e dos resultados empíricos, as implicações e aplicabilidades do uso de ML na contabilidade, considerando aspectos organizacionais e profissionais.

1.3 Justificativa

A contabilidade desempenha papel essencial na sustentação da confiança econômica, fornecendo informações financeiras confiáveis e transparentes que embasam as decisões de

gestores, investidores, órgãos reguladores e da sociedade. No entanto, a complexidade das operações financeiras impõem múltiplos desafios, o que reforça a necessidade de ferramentas robustas analíticas capazes de lidar com essas situações, como inadimplência, fraudes, insolvências e outros riscos que impactam diretamente a sustentabilidade das organizações, estabilidade econômica, reputação etc.

Nesse contexto, o aprendizado de máquina (ML) apresenta-se como uma tecnologia promissora e de crescente relevância. Ao processar grandes volumes de dados e identificar padrões complexos, os modelos de ML superam limitações de métodos estatísticos tradicionais, oferecendo maior precisão e eficiência, conforme já demonstrado no setor financeiro. Contudo, o impacto dessa tecnologia na contabilidade permanece incipiente e pouco explorado (Bertomeu et al., 2020), abrindo espaço para pesquisas como esta que avaliem sua efetividade e adequação em tarefas relevantes da área, bem como as implicações do seu uso, fundamental que deve ser considerado no que tange sua aplicação.

Outro aspecto essencial diz respeito à interpretabilidade. A adoção de modelos complexos enfrenta resistência justamente pela dificuldade em justificar suas decisões, elemento central a contabilidade (Lundberg, 2020; Ranta; Ylinen; Järvenpää, 2022). Nesse sentido, explorar ferramentas de explicabilidade reforça a contribuição do estudo, ao se aproximar das exigências da área reforçando a confiabilidade, a segurança dos dados e a transparência analítica.

Por fim, destacam-se elementos que reforçam a originalidade do estudo. O uso do *XGBoost*, ainda pouco aplicado em pesquisas contábeis, sobretudo no Brasil (Gao et al., 2024; Zaniboni; Montini, 2019, apud Macieira, 2025), e a incorporação do SHAP, técnica recente de interpretabilidade que enfrenta a opacidade de modelos não lineares, oferecendo maior clareza na explicação dos resultados (Hu; Sun, 2021; Lundberg; Lee, 2020). Soma-se a isso o enfoque dado à inadimplência, tema de menor representatividade acadêmica em comparação à detecção de fraudes, mas de elevada relevância prática, uma vez que modelos costumam obter estimativas mais precisas e menos enviesadas do que gestores (Ding et al., 2020).

2. REFERENCIAL TEÓRICO

Nos tópicos a seguir serão apresentados o cenário de evolução tecnológica ao qual a contabilidade sempre esteve condicionada, uma visão geral do aprendizado de máquina, com a descrição de métodos e técnicas, bem como revisão da sua efetividade e consequências de sua adoção.

2.1 Constante Evolução Tecnológica

A evolução da contabilidade reflete as mudanças tecnológicas e sociais, adaptando-se às necessidades do ambiente econômico. Na antiga sociedade mesopotâmica, dada a necessidade de controle dos itens e transações comerciais, os mercadores utilizavam formas primárias de contabilidade para registrar transações com tabletas de argila.

Esta prática contábil rudimentar veio a ser aprimorada em outras sociedades ao longo dos séculos, em paralelo com a escrita, matemática, administração e, também, a economia, tendo em vista que a prática de declaração de contas só foi possível quando um padrão monetário foi adotado, permitindo que os itens de um registro pudessem ser expressos de forma equitativa (Brown, 2006).

Na renascença, em 1494, o método das partidas dobradas foi formalizado pelo matemático italiano Luca Pacioli, fornecendo uma base sólida para a elaboração de demonstrativos financeiros como temos conhecimento atualmente, e então, séculos mais tarde, com o avanço da Revolução Industrial no século XVIII, as empresas começaram a crescer em tamanho e complexidade, o que exigiu o uso de práticas contábeis mais elaboradas para garantir a eficiência operacional, o que impôs a adaptação do contador às novas circunstâncias (Martins, 2003).

Durante a Segunda Revolução Industrial, o aumento do volume de transações e a necessidade de prestação de contas aos acionistas e outras partes interessadas, criaram demanda de auditorias e registros financeiros mais detalhados, e com a aprovação do princípio de responsabilidade as empresas passaram a adotar práticas de transparência com o partilhamento de demonstrativos ao público (Carraro; Rodrigues; Xavier, 2020). Evidentemente, o ofício do profissional contábil nessa época ainda era manual e os contadores ainda não assumiam posições de influência estratégica nas organizações.

O paradigma iria mudar na segunda metade do século XX, com o avanço da informática e aumento da popularidade dos computadores pessoais, permitindo a automação de processos

contábeis tradicionais e a criação de softwares especializados que simplificaram a rotina dos profissionais da área, os sistemas garantiram não só um serviço facilitado, como menos moroso para os contadores (Chagas *et al.*, 2013). O novo contexto contribuiu para a mudança do perfil do profissional contábil, que passou de indivíduo que registra transações financeiras para quem produz informações úteis a gestão para processo de tomada de decisão.

Segundo Chapman, C. & Chua e W. F (2003), já havia a concepção da extração de conhecimento a partir dos dados e sistemas de suporte nos anos 70, objetivando o uso ao aprimoramento na tomada de decisões dos gerentes. Contudo, apenas foi possível na década seguinte, com o advento dos computadores pessoais e das tecnologias de redes locais, a partir de sistemas de apoio à decisão mais sofisticados e descentralizados.

A partir disso, o conceito de tecnologia da informação (TI) passou a ser cada vez mais difundido nas empresas, e sistemas de informação contábil e ferramentas de integração informacional e de gerenciamento, tais quais o ERP, passaram a ser adotados em maior escala.

A globalização afeta diretamente o mercado, assim como as empresas e o contador, levando a abertura de mercados e a busca das empresas em se manterem competitivas. Por consequência, levando a adoção de estratégias de redução de custos, a diversificação e a melhoria dos negócios (Agostini, 2012), o que influencia o desenvolvimento e surgimento de tecnologias, e, concomitantemente, na busca de melhorias e ampliações tecnológicas por parte das organizações, o que tem alterado as atividades dos profissionais contábeis ao longo do tempo (Chagas *et al.*, 2013).

Martins *et al.* (2012) conclui que dado o cenário globalizado e exigente em que as empresas estão atualmente inseridas, onde todas as informações são digitais, a utilização dos sistemas de informação gerencial (SIG) e tecnologias da informação (TI) é quintessencial, para que a firma se mantenha eficiente operacionalmente e tenha suporte necessário para realizar boas decisões e facilitar a gestão.

A constante evolução tecnológica consoante a da profissão contábil persiste no séc. XXI. O escalonamento dos dados gerados tem ocorrido de forma exponencial desde a década anterior, desde então, diversas empresas passaram a considerar os benefícios da capitalização desses dados para os negócios. Segundo Brownlow *et al.* (2015), principalmente dada a crescente necessidade de permanecerem competitivas, nesse sentido, o conhecimento e insight gerado possuem grande valor. Evidentemente, uma quantia considerável de investimento e planejamento teria de ser feito, visando o desenvolvimento dos recursos humanos e aspectos técnicos.

Sopesando este panorama, ferramentas comuns não oferecem solução satisfatória para o problema de escala dos dados (Theodorakopoulos; Thanas; Halkiopoulou, 2024), pois quanto maior for a fonte de pontos de dados disponíveis, maior a exigência de soluções tecnológicas para armazenamento, organização, análise e extração de informação relevante (Ibrahim; Elamer E Ezat, 2021).

Diante disto, as soluções baseadas em plataformas de nuvem se tornaram adequadas no fornecimento de infraestrutura necessária, proporcionando gerenciamento de dados de maneira eficaz e escalável, além de permitir um acesso mais rápido e flexível às informações (Hashem *et al.*, 2015).

Acerca das soluções estratégicas, vários termos entraram em evidência com o passar do tempo, como o *business intelligence*, o *data analytics* e a ciência de dados, assim como, uma miríade de instrumentos e soluções tecnológicas relacionados, advindos da necessidade de ampliar a capacidade analítica das firmas. E, por se tratar de tendências relativamente novas, permitem ao profissional contábil que possui esse tipo de conhecimento assumir uma posição de destaque no mundo dos negócios (Bhattacharjee; Bose; Dey, 2022).

Nessas circunstâncias, pertinente a ciência de dados, a inteligência artificial obteve vultosa atenção, sobretudo nos últimos anos, em geral, de grandes empresas, que veem o grande potencial da tecnologia, apesar de limitada e relativamente nova em sua aplicação (Benbya; Davenport; Pachidi, 2020). As *Big Four* representam bem este interesse, ao investir grandemente no desenvolvimento de tecnologias de automação e análise inteligente (Thomson Reuters, 2024a).

De acordo com um relatório da Deloitte (2021), numa pesquisa feita a partir de 2.875 empresas, dentre elas, 218 do Brasil, a inteligência artificial entrou na era de difusão em meados de 2020, e embora a maioria das empresas que tenham adotado em grande escala não a utilizem em seu potencial completo, isto é, mais de 25% da amostra da pesquisa, a busca pela implantação dessas soluções só tende a aumentar nos seguintes anos.

Em outro relatório, pela Thomson Reuters (2024b), sobre uso da IA generativa, viu-se que os funcionários já usam soluções de IA de propósito geral e públicas, como o ChatGPT, mas que há crescente mudança em relação a inclusão de aplicações de uso específico para indústria. No mesmo relatório, mostrou-se que o total de respondentes que não tinham planos de uso da tecnologia caiu de forma considerável entre 2023 e 2024.

O caso das IAs, conforme Theobald (2021), pode ser comparado ao que ocorreu com a Revolução Industrial, quando o trabalho manual foi substituído pela força das máquinas; no

contexto atual, essa dinâmica muda significativamente, com as máquinas assumindo tarefas de cunho intelectual. No caso da contabilidade, isso afetaria diretamente a capacidade das empresas e do profissional, ao otimizar e possibilitar a tarefas que envolvem grandes bancos de dados, auxiliar a verificação de contratos ou identificar distorções e fraudes financeiras (Ey, 2025; Ucoglu, 2020).

Apesar da semelhança, numa conceituação simples, que as IAs compartilham, trata-se uma área ampla, composta por várias subáreas diferentes. No escopo deste trabalho, iremos abordar apenas o aprendizado de máquina, uma categoria que imita a cognição humana, sendo capaz de aprender com os dados fornecidos e melhorar suas respostas a partir de tentativas anteriores (Theobald, 2021).

Dada a crescente relevância de métodos analíticos mais avançados no contexto da Quarta Revolução Industrial (Sarker, 2021), uma revisão breve da tecnologia de ML é razoável e adequada ao propósito do trabalho, antes de se seguir com a inspeção da sua utilidade em âmbito financeiro e contábil.

2.2 *Machine Learning*

O *Machine Learning* ou aprendizado de máquina se originou em 1959, quando o pioneiro Arthur Lee Samuel cunhou o processo e criou um *software* capaz de jogar damas, o sistema tinha uma função de avaliação que ajudava o programa a classificar as posições do jogo considerando a chance de vitória associada, além de aprender a jogar a partir de experiências anteriores (Samuel, 1959).

2.2.1 *Definição e abordagens de aprendizagem*

O aprendizado de máquina (ML) é um tipo específico de inteligência artificial, distinto do uso mais amplo do termo IA. Ele consiste na coleta e utilização de dados para a construção de modelos estatísticos voltados à solução de problemas práticos (Burkov, 2019). Sua principal característica é a capacidade de identificar padrões e ajustar-se por meio da experiência, realizando ações que não estão previamente definidas em código rígido (Samuel, 1959; Lecun; Bengio; Hinton, 2015; Sarker, 2021). Isso significa que, em vez de seguir instruções fixas, como operações matemáticas programadas, os modelos aprendem relações complexas entre variáveis de forma adaptativa.

Esse processo de aprendizado depende fortemente da qualidade e da quantidade de

dados utilizados, já que conjuntos pequenos ou inadequados podem comprometer os resultados (IBM, 2021). Tal dinâmica se relaciona ao raciocínio indutivo, em que generalizações são feitas a partir de amostras (Monard; Baranauskas, 2003). Os dados podem ser estruturados, de classificação relativamente simples, ou não estruturados, mais complexos e de difícil generalização, e tais características impactam a escolha do algoritmo e da técnica de aprendizagem (Sarker, 2021).

As técnicas de ML se dividem em quatro principais categorias. O aprendizado supervisionado utiliza dados rotulados, com variáveis de entrada (X) e saídas conhecidas (y), permitindo ao modelo mapear relações e realizar previsões sobre novos dados. Ele se aplica tanto em tarefas de regressão (valores numéricos), como a previsão de preços de imóveis (Prince, 2023) ou análises em contabilidade gerencial (Ranta; Ylinen; Järvenpää, 2022), quanto em classificação (atribuição de categorias), como na detecção de fraudes e distorções financeiras (Ayad; El Mezouari; Kharmoum, 2023; Bertormeu et al., 2020).

Já o aprendizado não supervisionado não utiliza rótulos, sendo voltado à descoberta da estrutura oculta nos dados. Essa abordagem é útil para identificar tendências e padrões emergentes (Cloud Google, 2024; IBM, 2021), com destaque para algoritmos de clustering, que agrupam observações por semelhança (Sarker, 2021). Suas aplicações incluem segmentação de clientes, detecção de fraudes (Cloud Google, 2024) e até previsão de insolvência e falência (Jones, 2017).

O aprendizado semissupervisionado combina as duas técnicas, com dados rotulados e não rotulados, aproveitando o primeiro conjunto para inferir rótulos sobre o segundo. Essa abordagem busca gerar modelos mais eficazes, especialmente em cenários de escassez de dados rotulados (Burkov, 2019; Theobald, 2021; IBM, 2021).

Por fim, o aprendizado de reforço envolve um processo iterativo em que o modelo aprende por tentativa e erro, tomando decisões com base em recompensas ou penalidades. Ao longo das interações, ele ajusta suas ações para maximizar o retorno esperado, sendo uma técnica poderosa, mas também complexa (Prince, 2023).

2.2.2 Construção de modelos ML

Em estruturas tabulares de dados, cada linha corresponde a uma observação e cada coluna a atributos ou variáveis, formando o conjunto que alimenta os modelos de *machine learning* (Theobald, 2021). Esses conjuntos frequentemente contêm lacunas, dados irrelevantes ou valores atípicos, que comprometem a eficácia do modelo (Chu et al., 2016). Para lidar com

isso, aplicam-se processos como *scrubbing*, voltado à limpeza dos dados (Chu et al., 2016), e *feature engineering* (engenharia dos dados), que transforma e adapta atributos conforme o problema, exigindo conhecimento de domínio do problema para ser eficaz (Ahmed; Aziz, 2010; Ridzuan; Zainon, 2019; Burkov, 2019).

Entre as técnicas disponíveis, destaca-se a *feature selection*, que identifica as variáveis mais relevantes, eliminando redundâncias e acelerando o aprendizado (Burkov, 2019). Embora possa ser feita manualmente, a complexidade crescente exige o uso de algoritmos automatizados (Jovic; Brkic; Bogunovic, 2015). Outros métodos incluem a redução de dimensionalidade pela combinação de atributos semelhantes ou compressão de linhas (Theobald, 2021), bem como a transformação de dados categóricos em variáveis binárias por meio de codificação *one-hot* (Burkov, 2019). Alternativas como *binning* também permitem converter valores contínuos em intervalos discretos (Theobald, 2021).

Quando há dados ausentes, diferentes técnicas podem ser aplicadas para preenchimento, sendo a escolha dependente do tipo de dados e do objetivo do projeto (Emmanuel et al., 2021). Em síntese, o tratamento adequado do conjunto de dados é fundamental, e a seleção das técnicas depende do contexto, das peculiaridades e do propósito analítico.

Após o tratamento dos dados, é necessário definir as variáveis que comporão as entradas (X) e a saída (y) nos modelos supervisionados, bem como estabelecer a forma de divisão do conjunto de dados. A separação entre treino e teste é essencial para avaliar a capacidade do algoritmo de generalizar, já que treinar apenas com a totalidade dos dados levaria a resultados enganosos (Burkov, 2019). No entanto, a proporção ideal dessa divisão não é fixa e depende do caso analisado, não havendo consenso sobre o padrão mais adequado (Joseph, 2022).

Embora seja comum dividir nesses dois conjuntos, é citada a importância de um conjunto intermediário de validação (Nazareth; Reddy, 2023). Um subconjunto adicional poderia gerar problemas num conjunto de dados menor (Hastie; Tibshirani; Friedman, 2009), sendo comum o uso de outros métodos para a validação do modelo. Na etapa de configuração dos dados, destaca-se o conceito de *feature scaling*, utilizado para aproximar variáveis de escalas muito diferentes, o que poderia prejudicar o resultado, dependendo do modelo.

Após o tratamento dos dados, segue o treinamento do modelo, etapa em que vários cuidados são necessários, desde a configuração a avaliação de desempenho. A performance de um modelo está diretamente ligado ao equilíbrio entre viés e variância. Um alto viés reflete baixa acurácia e pode decorrer do algoritmo utilizado, do tamanho ou da organização dos dados (Theobald, 2021), enquanto a variância está associada à sensibilidade do modelo a flutuações

nos dados (Mucci, 2024).

Esse equilíbrio explica os cenários de *overfitting* e *underfitting*: no primeiro, o modelo se ajusta excessivamente aos dados de treino, apresentando baixo viés, mas alta variância e dificuldade de generalização (Ying, 2019); no segundo, a baixa complexidade resulta em alto viés e baixa variância, prejudicando tanto o ajuste aos dados de treino quanto a capacidade de prever novos casos (Mucci, 2024).

O erro tende a ser elevado quando a complexidade é muito baixa ou muito alta, o que evidencia a necessidade de buscar um *trade-off* entre viés e variância para maximizar a performance do modelo (Mucci, 2024). Estratégias como a randomização do conjunto de dados ou o ajuste de hiperparâmetros podem contribuir para reduzir esses problemas ainda na etapa de preparação e configuração dos dados.

Os hiperparâmetros são configurações externas que controlam o processo de treinamento do modelo, influenciando sua complexidade e desempenho. Seu ajuste pode ser realizado manualmente ou por métodos automatizados, como *grid search* e *randomized search*, que otimizam o processo. Após o treinamento, torna-se essencial avaliar a performance do modelo, seja em contextos supervisionados, por meio da comparação entre previsões e rótulos reais, seja em não supervisionados, analisando a forma como os dados se agrupam.

Uma técnica amplamente utilizada é a validação cruzada (*cross-validation*), que repete ciclos de treino e validação com subconjuntos diferentes de dados, permitindo mensurar erros e desempenho (Burkov, 2019). No entanto, há risco de avaliações excessivamente otimistas, já que os mesmos dados podem ser usados tanto para ajuste quanto para validação (Cawley; Talbot, 2010). Para reduzir esse viés, pode-se aplicar a validação cruzada aninhada (*nested cross-validation*), em que os hiperparâmetros são otimizados em um primeiro ciclo e avaliados em outro, com dados independentes (Cawley; Talbot, 2010).

Caso o desempenho ainda seja insatisfatório, alternativas incluem o reajuste de hiperparâmetros, uso de técnicas de regularização para evitar *overfitting* (Jabbar; Khan, 2014; Ying, 2019), ampliação ou modificação do conjunto de dados, ou mesmo a substituição do algoritmo de aprendizado para lidar com cenários de *underfitting*.

2.2.3 Algoritmos E Modelos

A escolha do algoritmo é etapa central no aprendizado de máquina, pois está diretamente relacionada ao tipo de problema e às características dos dados, não existindo um método universalmente eficaz (Aliferis; Simon, 2024). Além disso, a escolha do algoritmo influencia

na interpretabilidade: modelos mais complexos tendem a oferecer maior poder preditivo, mas são mais difíceis de compreender, especialmente os *ensembles*, que combinam múltiplos modelos (Hoang; Wiegratz, 2023; Molnar, 2020).

Nesses modelos, a falta de transparência pode dificultar sua adoção, o que tem motivado o desenvolvimento de ferramentas de explicabilidade (Covert; Lundberg; Lee, 2020). Por outro lado, Lundberg (2020) aponta que até modelos simples podem ser difíceis de interpretar se inadequados aos dados devido a desajuste, ao passo que modelos mais sofisticados, como os baseados em árvores, podem oferecer maior clareza ao lidar com estruturas complexas.

Nesse contexto, os algoritmos podem ser classificados em dois grupos principais: os métodos rasos, que incluem técnicas tradicionais de aprendizado e redes neurais simples, e os métodos profundos (*deep learning*), representados pelas redes neurais profundas (Pasupa; Sunhem, 2016; Prince, 2023). As primeiras possuem arquitetura mais simples e custo computacional menor, enquanto as redes profundas, com múltiplas camadas ocultas, destacam-se pela capacidade de capturar relações não lineares e complexas (Lecun; Bengio; Hinton, 2015; Bengio; Lecun; Hinton, 2021).

Apesar de sua popularidade e desempenho, não há evidência de superioridade absoluta das DNNs (redes neurais profundas) sobre os modelos rasos (Athey; Imbens, 2019), já que alternativas menos complexas podem, em determinados casos, alcançar resultados equivalentes ou até melhores (Pasupa; Sunhem, 2016; Meir et al., 2023; Herrera; Ceberio; Kreinovich, 2022). Assim, a escolha do modelo deve considerar tanto a natureza do problema quanto os requisitos computacionais e interpretativos de cada aplicação.

2.2.4 Árvore de decisão

As árvores de decisão são algoritmos de aprendizado supervisionado amplamente usados para classificação, devido à sua simplicidade e interpretabilidade. Sua estrutura é hierárquica, semelhante a um fluxograma, composta por nós de decisão, ramos e folhas, de forma que os nós representam testes de atributos, os ramos indicam as opções resultantes desses testes, e as folhas mostram os resultados ou rótulos das classes (Dash; Nayak; Mishra, 2020).

Os nós avaliam características dos dados de maneira distinta, dependendo do tipo de atributo. Para atributos categóricos, cada categoria gera um nó filho diferente, enquanto para atributos contínuos, define-se um limite (*threshold*) para dividir os dados em respostas "Sim" e "Não" (Sekar, 2020).

O processo de construção da árvore começa identificando a característica que melhor

particiona os dados na raiz e continua de forma recursiva até que todas as características sejam usadas ou um limite de profundidade seja atingido (Sekar, 2020). No entanto, construir uma árvore até que todas as folhas sejam puras, isto é, quando todos os dados nela contidos pertencem à mesma classe ou categoria, geralmente resulta em *overfitting*, daí a importância de limitar a profundidade da árvore (Müller; Guido, 2016).

2.2.5 Random Forest

Random forest (florestas de decisão aleatórias) é um método que consiste no treino paralelo de diversas árvores de decisão para formar resultados. Podem ser usadas com dois objetivos principais: construir um modelo de classificação ou regressão para prever dados futuros; ou investigar a relevância das variáveis preditoras, avaliando como cada uma contribui para a previsão da variável resposta (Probst; Wright; Boulesteix, 2019).

Semelhante ao *bagging*, são criadas várias árvores de decisão com subconjuntos dos dados, porém, no RF, além de amostrar aleatoriamente os dados, também se seleciona aleatoriamente um subconjunto de características para cada árvore (*feature subsample*), contribuindo a menor correlação das árvores geradas (Hastie; Tibshirani; Friedman, 2009). Da mesma forma, o modelo combina as previsões para gerar o resultado: em casos de classificação, a previsão mais frequente será o resultado, em casos de regressão, a média dos resultados será o resultado.

O diferencial da aleatoriedade desses modelos, segundo Müller e Guido (2016), ajuda a mitigar o problema de *overfitting* que as árvores de decisão costumam apresentar, dada a criação de árvores de decisão ligeiramente diferentes e combinação dos resultados por média. Ainda, conforme os mesmos autores, as RF não são ideais para dados de alta dimensionalidade e esparsos, como dados de texto, além de consumirem maior tempo e recursos se comparado a alguns outros métodos.

Apesar dos benefícios e desempenho, um problema significativo desse modelo *ensemble* está na dificuldade de compreender as decisões resultantes do modelo, dado o tamanho e complexidade dele. Sobre a problemática, Molnar (2020, p.10, tradução nossa) discorre: “para entender como a decisão foi tomada, você teria que olhar os votos e estruturas de cada uma das centenas de árvores. Isso simplesmente não funciona, não importa o quão inteligente você seja, ou quão boa seja sua memória operacional”.

Para tratar da questão, existem estudos e abordagens referentes a medição de interações locais entre características, além de ferramentas para entender a estrutura global do modelo,

direcionadas ao aprendizado com árvores de decisão, incluindo métodos ensemble como RF ou *Boosting* (Lundberg, 2020; Covert; Lundberg; LEE, 2020).

2.2.6 *Boosting e XGBoost*

O boosting difere do bagging por adotar um processo sequencial, no qual cada novo modelo é ajustado com base nos erros do anterior, transformando classificadores fracos em um modelo mais robusto (Bühlmann, 2012; Freund; Schapire, 1996). Nesse método, observações mal classificadas recebem maior peso nas iterações seguintes, de modo que as previsões finais resultam da combinação das sucessivas correções (Mayr, 2014).

O Boosting é amplamente aplicado a árvores de decisão, e apresenta limitações, como maior suscetibilidade ao *overfitting* e dificuldades de interpretação, além de perda de desempenho quando há sobreposição significativa entre classes (VEZHNEVETS; BARINOVA, 2007).

Uma evolução dessa abordagem é o gradient boosting, que ajusta iterativamente os modelos aos resíduos gerados a cada etapa, utilizando gradiente descendente para otimizar parâmetros e reduzir a função de erro (Friedman, 2002; Mayr, 2014). Ainda, a acurácia é melhorada por meio da introdução de randomização no processo, com a seleção de amostras aleatórias em cada iteração, sem reamostragem, o que melhora a capacidade de generalização (Friedman, 2002).

Sobre essa base, destaca-se o XGBoost, uma versão otimizada e escalável do método. O algoritmo introduz uma série de inovações, como um método especializado para lidar com dados esparsos, procedimentos de aproximação para acelerar o processo de aprendizado sem comprometer significativamente sua eficácia, e outras otimizações que o tornam altamente eficiente no uso de recursos; o que possibilita seu funcionamento até mesmo em computadores comuns (Chen; Guestrin, 2016).

Entre suas principais características, estão o suporte à computação paralela e distribuída, o uso de *feature subsample* semelhante ao aplicado em *random forests*, e a possibilidade de processamento de grandes volumes de dados com armazenamento em memória externa. Além disso, sua incorporação da regularização como parte da função objetivo permite maior controle da complexidade e melhor capacidade de generalização, reduzindo o risco de *overfitting* (NIELSEN, 2016).

Dentre suas características, destacam-se a capacidade de processar grandes volumes de dados, utilizando técnicas como computação paralela e distribuída, isto é, que permitem que os

dados sejam distribuídos entre diferentes máquinas e armazenados em memória externa, além de ser possível utilizar do *feature subsample*, assim como nas *random forest* (Chen; Guestrin, 2016).

2.3 Aplicação Do Aprendizado De Máquina Na Área Financeira

O aprendizado de máquina (ML) tem sido explorado em aplicações financeiras e econômicas há várias décadas. Embora referências iniciais datem de cerca de 50 anos (Gogas; Papadimitriou, 2021), foi com o estudo de White (1988) sobre redes neurais artificiais (ANN) na predição de retornos de ações que a área ganhou destaque acadêmico, especialmente a partir das décadas de 1980 e 1990 (Nazareth; Reddy, 2023). Com o decorrer do tempo, os avanços em algoritmos e infraestrutura computacional foram fundamentais para superar as limitações iniciais de dados e processamento, permitindo maior aplicabilidade (Ghoddusi; Creamer; Rafizadeh, 2019).

No cenário atual, o ML se consolidou como uma solução inteligente para o setor financeiro, principalmente pelo seu potencial de analisar grandes volumes de dados e extrair padrões relevantes (Cao, 2017). Tanto modelos rasos, aplicáveis a diversos problemas financeiros e de negócios, quanto, principalmente, modelos profundos têm mostrado resultados promissores, especialmente em tarefas que envolvem dados não lineares e massivos (Ozbayoglu; Gudelek; Sezer, 2020). Atualmente, o ML tem se consolidado como ferramenta estratégica no setor financeiro, com aplicações em diferentes frentes.

No mercado de ações, os modelos destacam-se pela capacidade de lidar com dados não estacionários e ruidosos, incorporando informações econômicas e sentimentais (Nazareth; Reddy, 2023). Estudos indicam que o *random forest* e o *XGBoost* apresentam resultados consistentes em horizontes mais longos em relação a alguns modelos, enquanto técnicas de redes neurais aprofundam a análise temporal e de sentimentos, ampliando o poder preditivo (Basak et al., 2019; Fischer; Krauss, 2018; Li et al., 2020; Weng; Ahmed; Megahed, 2017; Yan; Ouyang, 2018). As redes neurais num geral demonstraram grande eficácia com melhores resultados na maioria dos estudos.

Na gestão de portfólios, o *random forest* mostra-se eficaz na identificação de oportunidades de investimento e, em combinação com outras abordagens, contribui para retornos mais robustos (Ma; Han; Wang, 2021) e que o modelo obtém ganhos consistentes no mercado (Tan et al., 2019). Estratégias híbridas, que integram dados financeiros e sentimentais ou diferentes abordagens de aprendizado e frameworks matemáticos, também têm se mostrado

relevantes para melhorar a construção de portfólios e prever tendências (Malandri et al., 2018; Ma et al., 2021; Picasso *et al.*, 2019; Paiva *et al.*, 2019; Chen *et al.*, 2021; Gao *et al.*, 2024).

No mercado de câmbio (Forex), marcado por elevada volatilidade, técnicas de ML são amplamente aplicadas na previsão de taxas de câmbio e desenvolvimento de estratégias de negociação. Métodos supervisionados como *random forest* e *XGBoost* figuram entre os modelos utilizados, assim como as redes profundas, aplicados na previsão direcional e no gerenciamento de riscos, ainda que a complexidade do mercado limite a precisão (Moghaddam; Momtazi, 2021; Semiromi; Lessmann; Peters, 2020; Ahmed *et al.*, 2020). Ademais, modelos híbridos profundos podem combinar a análise temporal com a extração de características para aumentar a precisão (Islam; Hossain, 2020).

A prevenção de crises financeiras também tem sido beneficiada pelo uso de ML, que supera métodos tradicionais, como regressão logística, em acurácia, utilizando indicadores como taxas de juros, PIB e preços de ativos (Chatzis *et al.*, 2018). Modelos avançados profundos podem alcançar níveis elevados de precisão ($\approx 98\%$), embora enfrentem desafios de generalização em economias emergentes (Samitas; Kampouris; Kenourgios, 2020; Ouyang; Yang; Lai, 2021).

Na previsão de falências e insolvências, iniciada com métodos estatísticos clássicos (Beaver, 1966; Altman, 1968), técnicas modernas como *random forest* apresentam desempenho superior (Barboza; Kimura; Altman, 2017; Hosaka, 2019) e são aprimoradas com variáveis adicionais, como dados de governança corporativa e análise textual de relatórios financeiros (Liang *et al.*, 2016; Matin *et al.*, 2019), atingindo acurácia acima de 75% (Petropoulos *et al.*, 2020), consolidando sua importância na gestão de riscos financeiros (Wu; Ma; Olson, 2022).

Apesar dos avanços, ainda existem desafios significativos, como o desequilíbrio dos conjuntos de dados e as dificuldades de interpretabilidade dos modelos, que podem comprometer a confiabilidade das previsões de falências e insolvências (Veganzones; Séverin, 2018; Zoričák et al., 2020).

No risco de crédito, o *credit scoring* utiliza amplamente ML para estimar a probabilidade de inadimplência em solicitação de crédito (Bao; Lianju; Yue, 2019). Métodos supervisionados como *Random forest* e *XGBoost* superam abordagens tradicionais, como a regressão logística, alcançando maior precisão preditiva (Yu, 2017; Li, 2019; Coelho et al., 2021) e trazendo melhorias significativas em redução de perdas e concessão de crédito (Paula, 2019), contribuindo para o crescimento econômico sustentável (Tsai; Hsu; Yen, 2014). Esses modelos ainda podem ser complementados por técnicas não supervisionadas, no agrupamento

dos dados dos clientes (Liu *et al.*, 2022).

Por fim, na detecção de crimes financeiros, a aplicação de ML tem se mostrado essencial para mitigar perdas significativas (Perols, 2011). *Random forest* e *XGBoost* destacam-se pela capacidade de identificar padrões ocultos e complexos em grandes bases de dados (Yao, Zhang E Wang, 2018).

Quando combinados a métodos híbridos, permitem antecipar fraudes em transações e demonstrações financeiras, reduzindo riscos e fortalecendo a segurança do sistema financeiro (Song *et al.*, 2014; Lima, 2022). A detecção de fraudes em demonstrações financeiras e transações envolve diretamente a contabilidade, já que práticas fraudulentas frequentemente surgem em registros contábeis manipulados para esconder irregularidades (Song *et al.*, 2014).

2.4 Aplicação Do Aprendizado De Máquina Na Área Contábil

2.4.1 Relevância Analítica Da Informação Contábil

A contabilidade e as finanças estão intimamente relacionadas; os dados contábeis alimentam análises financeiras, agindo como um guia a diferentes partes interessadas, servindo de base para análises de desempenho, identificação de oportunidades e mitigação de riscos financeiros. Por exemplo, demonstrações contábeis permitem que analistas avaliem a solvência de uma empresa, estimem sua lucratividade e identifiquem sua capacidade de gerar recursos internamente, influenciando decisões de concessão de crédito, alterações em políticas financeiras ou investimentos em ações (Assaf Neto, 2020).

Para que essas decisões sejam eficazes, a qualidade das informações contábeis é essencial, pois, dados contábeis confiáveis e precisos reduzem a incerteza, permitindo que decisões sejam tomadas com maior segurança. A Norma Brasileira de Contabilidade (NBC TG, 2019) define que a qualidade da informação contábil está relacionada a características fundamentais e de melhoria, são elas:

- **Relevância:** informações devem ser úteis ao ato decisório, auxiliando avaliação de eventos passados, presentes e futuros.
- **Representação fidedigna:** dados devem refletir com precisão a realidade econômica das transações, sem erros ou vieses significativos.
- **Comparabilidade:** informações devem permitir comparações entre diferentes períodos ou entidades, promovendo análises consistentes.
- **Verificabilidade:** Deve ser possível verificar a precisão dos dados, de forma que

observadores independentes cheguem às mesmas conclusões.

- **Tempestividade:** As informações precisam ser disponibilizadas em tempo hábil para que sejam úteis na tomada de decisão.
- **Compreensibilidade:** A apresentação das informações deve ser clara e acessível, facilitando a compreensão pelos usuários.

Quando essas características estão presentes, a qualidade da informação contábil é maximizada, promovendo transparência e confiança entre os *stakeholders*. Como consequência, as informações de alta qualidade podem aumentar a eficiência das decisões financeiras (Lambert; Leuz; Verrecchia, 2007).

Com o avanço tecnológico, a relevância dos dados contábeis transcendeu sua utilidade tradicional, para também servir de insumo aos algoritmos de aprendizado, o que suscita possibilidade de potencializar as informações. Como já discutido, a estrutura e a qualidade dos dados de entrada para o modelo influenciam os resultados do mesmo, o que implica que dados contábeis fidedignos produzem saídas do modelo mais próximas à realidade. Contudo, modelos de aprendizado também podem ser usados em contextos em que a qualidade dos dados seja duvidosa, ou se quer obter insight a partir desses dados, o que pode contribuir para a melhoria das informações (Hu; Sun, 2021).

De acordo com Ayad, El Mezouari e Kharmoum (2023), as ferramentas ML têm demonstrado bom potencial de análise e melhoria das informações contábeis, sendo úteis na detecção de inconsistências e identificação de padrões invisíveis a abordagens convencionais, contribuindo a previsões mais acuradas, o que representa sua utilidade e relevância para as partes, usuárias dessas informações. Ou seja, o ML pode elevar a confiabilidade dos relatórios financeiros, assegurando representações mais fiéis da realidade econômica e maior capacidade de antecipação de cenários, o que fortalece a tomada de decisão e a confiança dos investidores.

2.4.2 Estimativa Contábil

As estimativas são muito presentes no âmbito contábil (Ranta; Ylinen; Järvenpää, 2022), e se configuram como um elemento essencial no processo de reconhecimento e mensuração contábil (De Iudicibus, 2018). Conforme Assaf Neto (2020), além de critérios objetivos no processo de estimar, a intuição e experiência podem ser aplicadas para prever comportamentos financeiros, sendo relevantes, por exemplo, a estimativas relacionadas a provisão para devedores duvidosos.

Segundo a NBC TA 540 (2019), as estimativas contábeis são elaboradas pela administração com base em métodos que exigem julgamentos criteriosos e o uso de premissas e dados, frequentemente sujeitos a incertezas e limitações. Tais incertezas são inerentes ao processo, refletindo a complexidade dos sistemas econômicos e financeiros, e demandam um equilíbrio entre a aplicação de normas contábeis e a experiência prática dos gestores. Essa subjetividade, embora inevitável, exige rigor e transparência na elaboração e na comunicação das demonstrações financeiras, pois influencia diretamente a qualidade das informações fornecidas aos usuários.

O desenvolvimento de estimativas confiáveis é essencial para as demonstrações financeiras, influenciando diretamente as decisões da gestão. No entanto, ajustar essas estimativas pode se tornar desafiador com o tempo, especialmente em relação a eventos passados (CPC 23; Assaf Neto, 2020). Ademais, a subjetividade inerente ao processo aumenta o risco de distorções provenientes do julgamento da administração (NBC TA 540, 2019).

Por exemplo, em seguradoras, as reservas de perdas combinam sinistros já reportados (cujo valor final é incerto), sinistros IBNR (Sinistros ocorridos, mas não comunicados) e pagamentos que se estendem por anos, o que confere alta subjetividade e favorece vieses de superestimação para ganhos fiscais ou suavização de resultados (Gaver; Paterson, 2004). Desse modo, considerando a cotidiana presença de estimativas nas demonstrações financeiras, seus impactos, e desafios tomados pela auditoria, o ML se torna propícia pela sua capacidade preditiva baseada em padrões históricos de grandes volumes de dados (Ding *et al.*, 2020).

O estudo de Ding *et al.* (2020) apresentou que algoritmos de ML, tais como *random forest*, *gradient boosting* e ANN, demonstraram capacidade de “aprender” padrões históricos a partir de grandes volumes de dados e gerar previsões mais precisas do que as estimativas gerenciais, contudo, foi constatado que a inclusão da estimativa gerencial como variável de entrada aprimora ainda mais o desempenho do modelo, indicando também a relevância do julgamento humano e potencial da tecnologia de otimização nas projeções.

No entanto, a performance do modelo em dados contábeis, sobretudo relevante a métodos tradicionais, como de regressão, requer adoção de boas práticas específicas, tais como separar amostras de diferentes períodos para treino, validação e teste, a fim de mitigar correlações intraempresa e temporais que podem supervalorizar variáveis redundantes; realizar o *tuning* de hiperparâmetros explorando características dos algoritmos (Bertomeu, 2020).

Ademais, Bertomeu (2020) ainda discorre que a multicolinearidade entre variáveis contábeis (ativo, passivo, patrimônio) demanda transformações dos dados (log, padronização)

e uso de métricas de importância atribuídas a variáveis que possuem o problema, objetivando a interpretação de quais blocos de informação (contábil e de mercado) são realmente relevantes ao prever erros.

Na previsão de lucratividade corporativa, o uso de ML demonstrou eficácia ao empregar métricas como ROE, ROA, RNOA, CFO e FCF. Anand *et al.* (2021) mostraram que um modelo de *random forest* supera significativamente um modelo aleatório, sobretudo porque as variáveis de fluxo de caixa (CFO e FCF) tendem a fornecer sinais mais estáveis que as métricas baseadas em lucro contábil.

A inclusão de *accruals* (receitas e despesas), permitiu capturar diferenças entre resultados realizados e não realizados em caixa, elevando a acurácia das previsões de fluxos futuros, o que é implicado nos trabalhos de Dechow e Dichev (2022) e Lewellen e Resutek (2019). Além disso, esse modelo manteve-se robusto ao longo de horizontes de até cinco anos, confirmando a estabilidade temporal das previsões baseadas em ML.

Dentro desse contexto, as empresas que exibem valores extremos na relação entre *accruals* e valor de mercado, ou entre lucro contábil e valor de mercado, alcançam até 75% de precisão, pois esses extremos refletem situações contábeis ou de mercado atípicas. Altos *accruals* sinalizam dependência de resultados não convertidos em caixa, enquanto discrepâncias agudas entre lucro contábil e valor de mercado revelam desalinhamentos entre resultados e preço das ações.

Segundo Anand *et al.* (2019), ao incluir essas variáveis, o modelo capta nuances específicas de cada empresa, aprimorando suas previsões, embora, para preservar a interpretabilidade, optou-se no estudo por um conjunto limitado de *features*, o que pode ter afetado a acurácia.

Destarte, apesar das abordagens de ML nas estimativas contábeis serem promissoras e ter potencial aparente nos estudos empíricos, o tema ainda é pouco explorado segundo Bertomeu *et al.* (2020), que destacou sua crescente importância e oportunidades futuras de pesquisa, como seu uso para aprimorar controles, identificar interações relevantes ou fundamentar novas teorias contábeis. A busca por informações de alta qualidade e perspectivas impulsiona a adoção de procedimentos baseados em ML, pois dá a possibilidade de estimar, prever e extrair sinais antecipados eficientemente, princípio que pode ser usado, por exemplo, na auditoria, para identificar riscos e prever irregularidades.

2.4.3 Auditoria

De acordo com Crepaldi S. A. e Crepaldi G. S. (2016, p.28), “a auditoria das demonstrações contábeis constitui o conjunto de procedimentos técnicos que tem por objetivo a emissão de parecer sobre sua adequação”. Desse modo, a auditoria é a principal salvaguarda da confiabilidade das informações contábeis, assegurando (com nível razoável de insuspeição) que as demonstrações financeiras reflitam com fidedignidade a realidade econômica das organizações, e da ausência de maiores distorções e indícios de má conduta nas informações.

Com a evolução da economia digital e a publicação de relatórios financeiros em tempo real, como destacado por Rezaee *et al.* (2002), cresce a necessidade de auditorias contínuas e sistemas automatizados que garantam a credibilidade das informações. Isto somado a um contexto de crescente complexidade de esquemas fraudulentos, cujos efeitos a sociedade e as organizações são destrutivos, tornaram indispensáveis métodos capazes de processar grandes volumes de dados e identificar padrões sutis de manipulação (Kirkos; Spathis; Manolopoulos, 2007), visto a impraticabilidade e imprecisão decorrente de métodos tradicionais e manuais de auditoria em casos complexos (Ashtiani; Raahemi, 2021).

O aprendizado de máquina oferece esse potencial ao automatizar a análise de indicadores contábeis e de procedimentos de governança, antecipando possíveis desvios antes que gerem perdas relevantes a acionistas e ao mercado (Kirkos; Spathis; Manolopoulos, 2007).

Ademais, a capacidade de identificar irregularidades em nível transacional permite uma abordagem mais detalhada, como apontado por Bay *et al.* (2006), em que a análise de outliers e padrões incomuns nos dados contábeis foi crucial para detectar comportamentos suspeitos ou erros operacionais significativos. Da forma semelhante, Bertomeu *et al.* (2020) apontaram o ML como uma ferramenta eficaz para analisar grandes conjuntos de dados, especialmente quando não há conhecimento prévio sobre como as variáveis se relacionam, além de demonstrarem sua utilidade na identificação de distorções contábeis relevantes e interpretação dos padrões associados.

A tecnologia também beneficiaria os profissionais de auditoria, ao reduzir tarefas repetitivas e automatizar a triagem de grandes volumes de dados, tornando o trabalho mais eficiente, e permitindo que os auditores concentrem seus esforços no julgamento crítico, análise de risco e orientação estratégica (Hu; Sun, 2021; Sekar, 2022), no geral, oferecendo valor maior aos auditores (Bao *et al.*, 2020). Rezaee *et al.* (2002) complementam ao indicar que processos contínuos suportados por ML fornecem garantias mais oportunas e robustas em ambientes de negócios digitais, aumentando a confiança dos *stakeholders*.

Além da capacidade de automatizar e escalar a análise, os modelos de ML precisam se apoiar em variáveis representativas para gerar resultados confiáveis. A acurácia de um modelo está ligada à seleção e ao número de atributos escolhidos, pois conjuntos enxutos podem melhorar a interpretabilidade sem sacrificar o desempenho. Nesse sentido, um estudo de Yao, Zhang e Wang (2018) implica que selecionar variáveis-chave pode elevar o entendimento dos resultados sem alterar relevantemente a precisão.

Corroborando o ponto, Perols (2011) comprovou que seis indicadores financeiros, incluindo receita discricionária e rotatividade de auditor, são suficientes para treinar classificadores simples como regressão logística e máquinas de vetor de suporte, mantendo robustez mesmo quando fraudes são raras. Ainda, a engenharia de atributos se mostra útil no balanceamento de classes, isto é, quando há disparidade representativa muito grande entre de classes, mostrando-se eficaz na construção de variáveis representativas e no aprimoramento da performance de algoritmos (Lima, 2022).

O desempenho dos modelos é variável, embora métodos baseados em árvores de decisão tenham se destacado em vários trabalhos (Hu; Sun, 2021). O estudo de Bao *et al.* (2020) avaliou o desempenho de um algoritmo de *boosting* em comparação com métodos tradicionais, como RL e SVM, observando que o modelo *ensemble* alcançou melhor desempenho geral, apresentando AUC média de 0,725, referente a acurácia de discriminação do modelo entre fraudes e não fraudes, no período de teste entre 2003 e 2008. Apesar disso, a precisão diminui em períodos mais extensos, indicando a necessidade de retreinamento para acompanhar mudanças nos padrões de fraude.

Complementarmente, Song *et al.* (2014) mostraram, em modelos ensemble, que a votação entre múltiplos classificadores, quando enriquecida com fatores qualitativos como governança corporativa e perfil de gestão, diminui significativamente a taxa de erros, reforçando a importância da análise combinada de dados financeiros e contextuais para maior robustez dos resultados.

Este papel positivo das variáveis qualitativas também foi constatado em outros estudos. Conforme Bertomeu *et al.* (2020), embora variáveis puramente contábeis, isoladamente, não sejam muito eficazes na detecção dessas distorções, elas ganham importância quando combinadas de maneira adequada com variáveis relacionadas à auditoria e ao mercado financeiro. Ainda, Kotsiantis *et al.* (2006) destacaram não só a importância de indicadores financeiros como variáveis-chave no processo de classificação, como também sugeriram a inclusão de informações qualitativas, a fim de possivelmente aumentar a taxa de acerto na

detecção de fraudes.

No contexto de análises interpretáveis, a combinação de técnicas de otimização, como o enxame de partículas (PSO), que simula o comportamento de grupos na busca por soluções ideais, e algoritmos de árvores de decisão, como o J48, permitem identificar variáveis relevantes e criar regras claras para classificar empresas entre fraudulentas e não fraudulentas. Hooda, Bawa e Rana (2018) relataram que essa abordagem alcançou até 93 % de acurácia em auditorias práticas, sendo altamente eficaz para classificações rápidas e resultados mais acessíveis, não apenas para analistas de dados, mas também para auditores e gestores que precisam compreender as conclusões em um curto espaço de tempo.

Além da análise numérica, também é possível explorar textos gerenciais e relatórios internos por meio de técnicas ML, que são aplicáveis à extração e análise de narrativas em relatórios de gestão, atas de reunião e memorandos (Ashtiani; Raahemi, 2021). Hu e Sun (2021) destacam que métodos como análise de tópicos e detecção de outliers são eficazes na identificação de sinais precoces de manipulações em textos não rotulados, ampliando o escopo das auditorias para além dos dados numéricos tradicionais.

Esses métodos podem revelar inconsistências de linguagem, como termos vagos ou mudanças abruptas de tom, que podem indicar manipulação. Além disso, técnicas de *clustering* e detecção de outliers oferecem meios de identificar padrões atípicos em dados não rotulados, ampliando a detecção em contextos em que fraudes são mais sutis (Ashtiani; Raahemi, 2021).

2.4.4 Análise Textual

O uso combinado de ferramentas de aprendizado de máquina na análise textual, muda relevantemente a forma como informações corporativas são processadas e interpretadas, particularmente, pois permitem a extração de informações de dados em grande escala, a partir de diversos tipos de documentos que fazem parte da rotina das organizações e seus profissionais.

No contexto da detecção de fraudes financeiras, a combinação de variáveis financeiras e linguísticas tem se mostrado eficaz, conforme Hajek e Henriques (2017) demonstraram, os fatores como crescimento esperado em receita e sentimentos negativos na linguagem de relatórios são preditores importantes de fraudes em demonstrações financeiras.

Classificadores bem conhecidos como *random forest*, conhecidos por sua boa performance e robustez em diversas tarefas, ajudam a identificar padrões suspeitos, enquanto modelos mais interpretáveis, como as redes bayesianas, fornecem uma perspectiva clara para

auditorias e planejamento. Algo visto constantemente, não apenas nesta aplicação, é a abordagem mista reforça a importância de aliar dados numéricos e qualitativos para melhorar a precisão das análises.

A análise de textos corporativos também se expande para além da detecção de fraudes, abrangendo a previsão de desempenho financeiro e a identificação de padrões comportamentais. Guo, Shi e Tu (2016) destacam o uso de ferramentas de aprendizado de máquina, como Naive Bayes, para classificar sentimentos e prever retornos financeiros. Por outro lado, Li (2010) ressalta que relatórios com tom otimista nas declarações prospectivas correlacionam-se positivamente com lucros futuros e liquidez. Esses achados enfatizam como nuances textuais, como o uso de linguagem confiante, podem influenciar não apenas as percepções dos investidores, mas também as decisões estratégicas das empresas.

Ainda assim, desafios permanecem. A legibilidade de textos corporativos complexos, como os relatórios MD&A, está associada a retornos mais baixos e maior volatilidade no mercado (Guo, Shi & Tu, 2016). Li (2010) aponta que a ofuscação gerencial — uso intencional de linguagem densa para mascarar resultados negativos — compromete a transparência, afetando pequenos investidores e dificultando previsões analíticas. Esses problemas destacam a necessidade de aprimorar métodos de análise textual, desenvolvendo abordagens que combinem precisão com acessibilidade.

Além disso, o uso de aprendizado de máquina na análise de textos financeiros apresenta desafios éticos e regulatórios. A manipulação de dados textuais para influenciar o mercado ou mascarar fraudes pode aumentar o risco de litígios e comprometer a confiança dos *stakeholders*. Portanto, esforços regulatórios, como os incentivados pela SEC nos EUA, buscam melhorar a qualidade das informações prospectivas e responsabilizar as empresas pela clareza e precisão de suas comunicações (Li, 2010).

No futuro, a evolução dessas tecnologias deve focar em maior transparência e inclusão. O desenvolvimento de ferramentas acessíveis pode democratizar o uso da análise textual, capacitando empresas de pequeno e médio porte a aproveitar os benefícios dessas tecnologias. Além disso, ampliar o escopo para diferentes idiomas e contextos culturais, como sugerem Guo, Shi e Tu (2016), permitirá a expansão dessas práticas para mercados emergentes e novos setores.

Num contexto acadêmico, ML também oferece possibilidades de pesquisa na contabilidade gerencial, especialmente no que diz respeito à construção e ao refinamento de teorias. Pois, diferente das abordagens tradicionais, baseadas predominantemente em dedução

e indução clássica, os métodos de ML se destacam pela sua capacidade de favorecer processos indutivos e abduativos mais sofisticados, além de apoiar abordagens intervencionistas no campo empírico (Ranta; Ylinen; Järvenpää, 2022).

Segundo Ranta, Ylinen e Järvenpää (2022), a análise textual automatizada a partir da aplicação da ferramenta, possibilita explorar dados não estruturados de fontes externas, como redes sociais e relatórios públicos, e também, de documentos internos das organizações. Essas informações, quando processadas com técnicas adequadas, podem revelar padrões de grande valor, oferecendo meios para o avanço teórico.

Como destaca Shmueli (2010), modelos preditivos têm o potencial de revelar mecanismos causais ainda não identificados, especialmente em grandes bases de dados com padrões complexos de difícil formulação a priori. Esses modelos também permitem comparar operacionalizações distintas de construtos teóricos, avaliar sua aplicabilidade prática e, inclusive, propor novas formas de representação desses construtos. Por fim, ao confrontar os resultados dos modelos com as expectativas da teoria, é possível mensurar a distância entre teoria e prática, o que pode conduzir ao aprimoramento teórico e metodológico contínuo.

2.5 Implicações Da Tecnologia De IA Para As Organizações

A incorporação crescente de tecnologias de ciência de dados, como IA e ML, tem gerado efeitos imediatos e aponta para transformações estruturais nas organizações e no mercado em geral. No curto prazo, observa-se um aumento de produtividade e eficiência operacional, ao mesmo tempo em que surgem alguns obstáculos práticos.

No plano de desempenho organizacional, ML tem capacidade de gerar valor de negócio (*Business Value*), que se traduz em resultados tanto financeiros quanto não-financeiros. Segundo Reis *et al.* (2020), organizações maduras em plataformas de Big Data Analytics e com processos complexos conseguem acelerar o ciclo de criação de valor de ML, obtendo vantagem competitiva sustentável por meio de insights orientados por dados.

O uso da ferramenta amplia o escopo de decisão estratégica ao promover melhores informações, permitindo às empresas superarem práticas convencionais de *business-as-usual* e explorar novas oportunidades, e, a depender do desenvolvimento e atos específicos da organização, e do seu nível de expertise na ferramenta, pode afetar positivamente a performance financeira da organização (Reis *et al.*, 2020). O potencial estratégico da tecnologia é reforçado por Ayad, El Mezouari e Kharmoum (2023) e Abdi *et al.* (2021), que indica sua capacidade em

promover a qualidade e confiabilidade das informações.

Paralelamente, os ganhos de produtividade viabilizados pela automação inteligente impactam diretamente a eficiência de tarefas complexas. Wright e Schultz (2018) destacam que recentes avanços em IA, aprendizado de máquina e sensores já permitem automatizar atividades que dependiam de julgamento tácito e alto nível de cognição, elevando os níveis de produtividade de forma expressiva. Em setores como a contabilidade, essa automação tem reduzido erros, acelerado processos de auditoria e fortalecido o gerenciamento de riscos, embora também pressione profissionais a aprofundarem suas competências técnicas diante da substituição de tarefas repetitivas (Shi, 2020).

Além disso, segundo Shrestha *et al.* (2019), a incorporação da IA nas estruturas decisórias das organizações tem implicações diretas sobre como decisões são tomadas, delegadas e combinadas. Pois, a IA possibilita a delegação total de decisões algorítmicas, a tomada de decisões híbridas entre humanos e IA, ou a agregação dos dois tipos de decisão por mecanismos como votação ou consenso ponderado. Esses formatos permitem ganhos de escala, rapidez e replicabilidade, mas também expõem as organizações a riscos éticos, vieses ocultos e desafios de interpretabilidade, exigindo reavaliações nas práticas de governança e responsabilidade.

Essas tendências e reformas associadas a automação das organizações estão relacionadas ao fenômeno da Quarta Revolução Industrial, também chamada de Indústria 4.0. A indústria 4.0 representa uma fronteira de integração entre máquinas, pessoas e sistemas, mas encontra barreiras como a falta de cultura digital, capacitação e elevado investimento inicial, desafios destacados por Ślusarczyk (2018). Contudo, muitos empresários reconhecem nessa etapa uma oportunidade para fortalecer a competitividade e repensar modelos de negócio.

No entanto, a promessa dessas tecnologias trouxe consigo um receio de substituição dos profissionais, assim como ocorreu em outros períodos de grande revolução tecnológica. No que refere ao mercado de trabalho, Autor (2015) observa que, historicamente, a automação substitui e complementa a mão de obra de forma simultânea, com máquinas assumindo tarefas rotineiras, mas ampliando a demanda por habilidades de resolução de problemas, criatividade e adaptabilidade. Ou seja, quando há um avanço tecnológico que melhora a produtividade de algumas tarefas, isso tende a aumentar o valor econômico de tarefas que ainda precisam ser realizadas por humanos.

Essa dinamicidade explica porque novos setores, embora criem postos de trabalho, concentram-se em funções qualificadas, gerando polarização salarial e reduzindo oportunidades

para trabalhadores de média qualificação. A consequência é uma redistribuição de empregos, isto é, enquanto algumas ocupações ganham valor econômico e a escassez de mão de obra eleva ganhos salariais, outras veem postos e renda erodirem (Autor, 2015).

Como implicado anteriormente, as aplicações de IA também provocam reorganizações internas, pois novos papéis especializados emergem e antigas hierarquias são revistas. De Araujo e Cornacchione (2024) apontam que, apesar do potencial transformador da tecnologia, a maioria das iniciativas ainda está em fase experimental, exigindo mudanças organizacionais e requalificação de equipes para que os sistemas inteligentes deixem a fase de protótipo.

Ainda conforme Shrestha et al. (2019), o tipo de estrutura decisória adotado influencia a forma como a IA impacta as dinâmicas internas. Modelos em que a IA atua como filtro inicial (*AI-to-human*) ou realiza a decisão final (*human-to-AI*) exigem um balanceamento cuidadoso entre eficiência algorítmica e julgamento humano, sendo particularmente sensíveis à amplificação de vieses históricos, à exclusão de alternativas relevantes e à perda de transparência. Já a agregação entre decisões humanas e da IA pode mitigar esses efeitos, desde que haja regras claras de ponderação e auditabilidade dos algoritmos utilizados.

Projeções futuras ainda indicam que a governança de IA ganha centralidade. De Araujo e Cornacchione (2024) salientam a formação de centros de excelência internos e diretrizes éticas para equilibrar inovação e responsabilidade, especialmente em cenários onde grandes volumes de dados podem ser insustentáveis, e modelos mais enxutos e transparentes são defendidos por pesquisadores. Reis *et al.* (2020) reforçam que o amadurecimento das plataformas e o suporte da alta gestão são cruciais para que o valor potencial de ML se converta em resultados concretos e escaláveis.

Nesse cenário, também emergem questões de privacidade e regulação, particularmente em setores sensíveis como o financeiro e contábil. Zhang *et al.* (2020) alertam para a necessidade de fortalecer estruturas de segurança de dados e de que reguladores implementem normas eficazes contra crimes associados ao uso de IA e blockchain. O que implica que a interseção dessas tecnologias, embora prometa reduzir custos e aprimorar a precisão das análises, exige marcos regulatórios claros para resguardar informações e preservar a confiança de *stakeholders*.

Por fim, referente ao aspecto de viabilidade da tecnologia, avanços em algoritmos e infraestrutura computacional reduziram consideravelmente os custos de processamento, o que amplia a viabilidade prática dessas aplicações para organizações de diferentes portes (Reis *et al.*, 2020), embora, há de se considerar o custo x benefício dessa adoção, dado o investimento

necessário.

Isso reforça a ideia de que a adoção de modelos preditivos não é restrita a grandes corporações, mas pode ser estrategicamente incorporada por entidades menores, desde que acompanhada da definição clara de objetivos analíticos. Ademais, conforme supramencionado, modelos funcionam bem ao serem “alimentados” com conjuntos de dados maiores, com a possibilidade de uso relacionado ao Big Data, mas não quer dizer que seja necessariamente condicionado a *datasets* massivos apenas.

2.6 Transformação Da Profissão E Profissionais Contábeis

A incorporação de soluções baseadas em IA e, em particular, em ML, vem moldando de forma simultânea desafios e oportunidades no curto e no longo prazo para organizações de diferentes portes. No aspecto operacional imediato, observa-se que tarefas rotineiras de perfil cognitivo e manual, ou seja, aquelas que seguem procedimentos bem delimitados, têm sido cada vez mais codificadas em software e transferidas para máquinas (Autor, 2015).

Em pequenas e médias empresas contábeis, entretanto, essa automação não deve suprimir integralmente o trabalho dos contabilistas no médio prazo, ao contrário, permite que eles se concentrem e valorizem mais os aspectos humanos do seu trabalho, como na análise crítica, na tomada de decisão e comunicação, ampliando a qualidade das informações gerenciais (Abdi *et al.*, 2021; Zhang *et al.*, 2020).

Além disso, modelos de *deep learning* e grafos de conhecimento vêm revolucionando processos de auditoria e controle financeiro. Técnicas avançadas de análise de dados automatizam a supervisão de receitas e despesas, antecipando riscos e prevendo crises com maior precisão e agilidade do que métodos tradicionais (Hou, 2022). Esses ganhos fazem com que a função contábil migre progressivamente de tarefas reativas para papéis estratégicos, em que o profissional assume responsabilidades de interpretação de resultados e governança dos sistemas inteligentes (Shi, 2020).

Entretanto, a evolução tecnológica acentua a polarização do mercado de trabalho; conforme Frey e Osborne (2017) identificaram em um estudo sobre a susceptibilidade de ocupações dos EUA, quase metade dos empregos apresenta alto risco de ser automatizado, sobretudo aqueles de menor qualificação e rotina padronizada, enquanto funções que exigem criatividade e inteligência social tendem a manter ou até aumentar sua demanda. Essa dinâmica reforça a necessidade de investimento em educação contínua e treinamento técnico-interpessoal

para trabalhadores de média qualificação (Autor, 2015; Frey; Osborne, 2017).

No contexto organizacional, a adoção de IA também gera novos perfis de trabalho e revisita antigas hierarquias. De Araujo e Cornacchione (2024) destacam que a automação de atividades repetitivas libera contadores para tarefas analíticas de maior valor, como modelagem preditiva e monitoramento em tempo real. Ao mesmo tempo, surgem preocupações relativas a vieses algorítmicos, explicabilidade dos modelos e segurança de dados sensíveis, aspectos que exigem governança robusta e conformidade com regulamentações (Zhang *et al.*, 2020; De Araujo; Cornacchione, 2024).

Sob ótica prospectiva, Weeks, Voshaar e Zimmermann (2023) enfatizam a integração de tecnologias emergentes nos currículos universitários, o uso de projetos reais e parcerias multissetoriais, incluindo micro-certificações em análise de dados e programação, para reduzir a lacuna entre teoria e prática. Adicionalmente, propõem que o próprio ML suporte a formação de contadores, por meio de plataformas de aprendizagem personalizadas que se ajustem ao ritmo e ao nível de cada aluno, com o auxílio de dados em tempo real, promovendo maior engajamento e interesse dos alunos.

3. PROCEDIMENTOS METODOLÓGICOS

O estudo busca realizar levantamento da literatura no que tange às implicações que envolvem a aplicação da tecnologia, bem como experimento empírico da construção e aplicações de modelos, para a visão mais ampla sobre o tema na contabilidade. A primeira etapa dispõe analisar, de forma geral, os efeitos do aprendizado de máquina referentes as entidades e profissionais, enquanto a segunda demonstra sua efetividade prática. Ambas se articulam na discussão dos resultados: a comparação entre os achados permite verificar como os dados empíricos sustentam ou tensionam os impactos descritos na literatura revisada.

O levantamento feito incluiu o recolhimento de informações de diversos artigos, encontrados diretamente no sistema Google Acadêmico, a partir de termos genéricos de busca como “*machine learning*”, “contabilidade”, “qualidade” e “impacto”, em inglês, dado o maior volume de pesquisas recentes e notáveis existentes sobre o assunto, ou indiretamente, sendo encontrados nas referências de outros artigos. Os artigos usados são endereçados de diversos sites de pesquisa acadêmica, universidades, editoras e repositórios como *Springer*, *ieeexplore*, *USP* etc.

A pesquisa prática visa aplicar alguns dos conceitos explanados ao decorrer do texto, com o fim de verificar e demonstrar a performance e eficácia de modelos baseados em árvores de decisão aplicados em tarefas de classificação de fraude e estimativa contábil. Os algoritmos escolhidos foram o *random forest* e *XGBoost*, principalmente pela robustez deles ao lidar com conjuntos de dados mais complexos.

Para composição dos modelos, a linguagem de programação python foi utilizada, em conjunto com diversas bibliotecas (Quadro 1).

Quadro 1 – Bibliotecas usadas

Bibliotecas	Uso
Pandas	Manipulação do dataset
Sklearn	Construção do modelo, métricas etc.
matplotlib	Compor gráficos de resultado
xgboost	Construção do modelo xgboost
seaborn	Auxiliar para composição dos gráficos
shap	Verificação de contribuição das variáveis

Fonte: autor

A seleção de conjuntos de dados para treinamento ocorreu na plataforma *Kaggle*, compreendendo dois conjuntos. O primeiro conjunto se trata de registros de faturas emitidas

entre 2012 e 2013, fornecidos pela IBM para teste de suas ferramentas; pontua-se que não fora encontrada a origem exata e o processo de coleta desse conjunto, o que impede a verificação e confirmação da procedência das informações.

No entanto, o *dataset* apresenta características e variáveis condizentes de um cenário de gestão de recebíveis, o que possibilita, com limitações da generalização, simular um contexto de previsão de inadimplência. O conjunto conta com informações detalhadas sobre cada fatura, cliente e status da cobrança, possibilitando a construção de modelos preditivos para estimar atrasos e inadimplência com base em padrões de pagamento dos devedores.

Os dados estão estruturados em 12 features principais, com 2466 observações, cada uma relacionada a uma fatura e cliente específicos, ao todo, 1589 observações se referem a faturas com situação em dia (64%), enquanto 877 observações representam faturas em atraso (36%), o que representa um breve desbalanço do conjunto (Quadro 2).

Quadro 2 - Conjunto dos Recebíveis Original

Variáveis	Descrição
<i>countryCode</i>	Código do país
<i>customerID</i>	Identificador do cliente
<i>PaperlessDate</i>	(Variável excluída)
<i>invoiceNumber</i>	Id da fatura
<i>InvoiceDate</i>	Emissão da fatura
<i>DueDate</i>	Vencimento da fatura
<i>InvoiceAmount</i>	Valor da fatura
<i>Disputed</i>	Se clientes discordam da fatura (sim/não)
<i>SettledDate</i>	Data de pagamento
<i>PaperlessBill</i>	Fatura eletrônica (ou em papel)
<i>DaysToSettle</i>	Dias até o pagamento
<i>DaysLate</i>	Dias em atraso

Fonte: autor

Diversas variáveis (ou *features*) do conjunto de dados foram transformadas (Quadro 3) ou criadas (Quadro 4) com base em suposições ou correlações observadas entre os dados originais, com o objetivo de melhorar a capacidade do modelo de identificar padrões relevantes. As variáveis relativas à datas foram desmembradas em componentes separados — ano, mês e dia — e substituídas por essas novas variáveis (Quadro 3). A única exceção foi a variável *PaperlessDate*, que foi removida devido à inconsistência em relação às outras datas disponíveis.

A variável *countryCode* passou por codificação one-hot, gerando colunas separadas para cada código de país. Após essa transformação, a variável original foi descartada. Além disso, também foi feita uma substituição meramente representativa dos códigos por nomes de países, com fins de interpretação. Ademais, variáveis categóricas como *Disputed* e *PaperlessBill* foram

convertidas em formato binário, sendo representadas pelos valores “verdadeiro” ou “falso”.

Quadro 3 - Variáveis Tratadas/Transformadas

Variáveis	Descrição
InvoiceDay	Dia extraído a partir de InvoiceDate
InvoiceMonth	Mês extraído a partir de InvoiceDate
InvoiceYear	Ano extraído a partir de InvoiceDate
DueDateDay	Dia extraído a partir de DueDate
DueDateMonth	Mês extraído a partir de DueDate
DueDateYear	Ano extraído a partir de DueDate
Disputed_Yes	Transformação de <i>Disputed</i> ; Disputada: 1, se não, 0
PaperlessBill_Electronic	Transformação de ‘PaperlessBill’; ‘Eletrônica: 1, se Papel: 0
Argentina; Brasil; Uruguai; Paraguai; Mexico	Países extraídos de <i>countryCode</i> *nomes puramente representativos

Fonte: autor

A feature *Late* foi criada a partir de *DaysLate*, visando reformular o problema como uma tarefa de classificação, em seguida, foram criadas outras variáveis conforme Quadro 4. Variáveis como *SettledDate*, *DaysToSettle* e *DaysLate* não foram utilizadas diretamente no treinamento do modelo para evitar *data leakage* (vazamento de informações futuras), contudo, foram úteis na criação de três novas features (Quadro 4) (*).

A três novas *features*, derivadas a partir das variáveis supracitadas, foram construídas apenas com base no conjunto de treino, garantindo que informações futuras (como as do conjunto de teste) não fossem utilizadas indevidamente. Elas servem como indicadores históricos de comportamento para clientes já conhecidos. Quando aplicadas ao conjunto de teste, essas variáveis só foram utilizadas para clientes previamente observados durante o treino. No entanto, pontua-se que esse procedimento pode não ser adequado em cenários onde há uma grande quantidade de devedores novos no conjunto de aplicação.

Quadro 4 – Variáveis Criadas

Variáveis Criadas	Descrição
Late	Variável-alvo. Atrasado: 1, se não, 0
DisputedAmount	Total da fatura disputada
AvgAmount_byCustomer	Quantia média por devedor
DisputeRate_byCustomer	Proporção de faturas contestadas por devedor
TotalInvoices_byCustomer	Número de faturas por devedor
CountryZScore*	Z-score da média de DaysToSettle por país
AvgDaysToPay_byCustomer*	Média de dias até o pagamento por devedor
AvgDaysLate_byCustomer*	Média de atraso de pagamento por devedor

Fonte: autor

No Quadro 5 são exibidas as variáveis usadas para o treinamento e teste do modelo, após adições, transformações e exclusões realizadas.

Quadro 5 – Variáveis Utilizadas

<i>InvoiceDay; DueDateDay</i>
<i>InvoiceMonth; DueDateMonth</i>
<i>InvoiceYear; DueDateYear</i>
<i>Disputed Yes</i>
<i>PaperlessBill Electronic</i>
DisputedAmount
Argentina
Brasil
Uruguai
Paraguai
Mexico
AvgAmount byCustomer
DisputeRate byCustomer
TotalInvoices byCustomer
CountryZScore*
AvgDaysToPay byCustomer*
AvgDaysLate byCustomer*

Fonte: autor

Quanto ao segundo conjunto, usado no estudo de Hooda, Bawa e Rana (2018), possui origem documentada, proveniente do escritório de auditoria da Índia, e reúne informações não confidenciais de empresas no período de 2015 a 2016, destinadas ao desenvolvimento de um preditor para classificar firmas suspeitas. O conjunto contém dados financeiros, operacionais e históricos de risco de empresas públicas de diversos setores (como Agricultura, Saúde e Indústrias). Cada linha representa uma firma auditada, com variáveis numéricas e categóricas, além de uma variável alvo que indica se a empresa foi considerada fraudulenta.

Para o trabalho, foi utilizada uma versão reduzida do *dataset*, baseada nos atributos selecionados no estudo de Hooda, Bawa e Rana (2018) (onde relatam uso de PSO para seleção de features), com o único diferencial deste trabalho sendo a remoção da feature “TOTAL”, o que ocasionou um conjunto de menor dimensionalidade.

O conjunto final inclui 8 features e 776 observações, sendo 470 (61%) casos não fraudulentos e 305 (39%) fraudulentos (1). Assim como o primeiro *dataset*, o conjunto não apresenta grande desbalanceamento entre as classes. Como a seleção de features foi baseada em um estudo anterior, apenas pequenas adaptações foram necessárias para alinhar algumas

variáveis à estrutura apresentada nesse trabalho.

Para isso, foi utilizado um *dataset* complementar relacionado, que auxiliou na identificação e ajuste de informações no conjunto principal. As únicas variáveis que exigiram modificação direta foram “Loss_score” e “History_score”, que foram calculadas a partir das informações das features “PROB” e “prob”, garantindo consistência com a estrutura original utilizada no estudo de referência.

Quadro 6 - Conjunto de auditoria

Variáveis	Descrição
Sector score	Valor da pontuação de risco histórico da unidade-alvo
PARA A	Discrepâncias em gastos planejados (em milhões)
PARA B	Discrepâncias em gastos não planejados (em milhões)
numbers	Índice histórico de discrepância.
Money Value	Montante envolvido em distorções em auditorias anteriores.
District	Pontuação de risco histórico de um distrito nos últimos 10 anos.
Loss score	Quantia de prejuízo sofrido pela empresa no último ano.
History score	Média de perdas históricas sofridas pela empresa nos últimos 10 anos.

Fonte: autor

Na divisão entre treino e teste, foi adotada uma proporção de 80/20 para os modelos relacionados a recebíveis, e de 70/30 para os modelos voltados à auditoria. Para garantir que essa divisão fosse reproduzível em execuções futuras, foi utilizado o método de *shuffle* com controle de aleatoriedade (*random state*). Durante a construção dos modelos, utilizou-se a ferramenta GridSearchCV para ajustar os hiperparâmetros dos classificadores por meio de validação cruzada iterativa, permitindo encontrar combinações otimizadas.

Após esse processo, nos modelos baseados em *random forest*, os melhores parâmetros encontrados com o GridSearchCV foram reaproveitados em uma nova rodada de treinamento. Esse *retreino* foi repetido dez vezes com pequenas variações dos hiperparâmetros, com o objetivo de avaliar a estabilidade dos resultados e calcular a média das métricas de desempenho. Testaram-se também diferentes versões dos dados: uma com todas as variáveis disponíveis, outra apenas com as *features* mais relevantes, além da experimentação com diferentes *thresholds* de classificação. Para os modelos com XGBoost, como seus resultados são determinísticos, foi realizado apenas um treinamento com os parâmetros otimizados.

Na etapa de avaliação, os modelos foram testados com base nas principais métricas de desempenho: acurácia, precisão, recall, F1-score e ROC-AUC, visando verificar sua efetividade no conjunto de teste. Também foi aplicada a matriz de confusão, o que permitiu uma análise mais detalhada dos erros do modelo, especialmente os falsos positivos (erro tipo I) e falsos

negativos (erro tipo II). Por fim, para tornar os modelos mais interpretáveis, utilizou-se a API SHAP (Lundberg & Lee, 2017), que permite identificar a importância de cada variável e entender seu impacto nos resultados, com base nos valores de Shapley.

4. ANÁLISE DOS RESULTADOS

Foram treinados modelos em duas atividades contábeis comuns: auditoria, na detecção de prováveis distorções, e estimativa, na classificação de devedores inadimplentes para provisão de perdas. Além disso, os resultados são discutidos à luz de constatações de diferentes autores, considerando as implicações do uso do machine learning na contabilidade e auditoria.

4.1 Resultados – Modelos De Estimativa

Os resultados obtidos com os modelos *random forest* e XGBoost mostram desempenho consistente na tarefa de classificação binária para detecção de inadimplência. Foram testadas diversas combinações de hiperparâmetros, ajustes de *threshold* (em modelos *random forest*) e redução do número de variáveis a partir de valores SHAP.

Tabela 1 - RESULTADOS RANDOM FOREST

Modelos	Grid	Thresh,	Variáveis	Acurácia	Precisão	Recall	F1	Roc-Auc	FN	FP	Erros (%)
1	✓	0,5	Todas	0,87	0,792	0,864	0,826	0,924	24	40	12,96%
2	✗	0,5	Todas	0,868	0,786	0,867	0,824	0,924	24	42	13,16%
3	✗	0,5	SHAP (6)	0,866	0,789	0,851	0,819	0,927	26	40	13,36%
4	✗	0,48	Todas	0,87	0,785	0,876	0,828	0,924	22	42	12,96%
5	✗	0,48	SHAP (6)	0,867	0,781	0,871	0,823	0,927	23	43	13,36%

Fonte: autor

Tabela 2 - RESULTADOS XGBOOST

Modelos	Grid	Scoring	Features	Acc	Prec	Recall	F1	ROC-AUC	FN	FP	Erro (%)
1	✓	Precision	todas	0,862	0,778	0,858	0,816	0,922	25	43	13,77%
2	✗	—	SHAP (6)	0,86	0,772	0,864	0,815	0,921	24	45	13,97%
3	✗	—	todas	0,868	0,791	0,858	0,823	0,929	25	40	13,16%
4	✗	—	SHAP (6)	0,862	0,776	0,864	0,817	0,927	24	44	13,77%
5	✓	ROC-AUC	todas	0,856	0,759	0,875	0,813	0,928	22	49	14,37%
6	✗	—	SHAP (6)	0,862	0,765	0,886	0,821	0,927	20	48	13,77%
7	✗	—	todas	0,856	0,772	0,847	0,808	0,93	27	44	14,37%
8	✗	—	SHAP (6)	0,868	0,791	0,858	0,823	0,928	25	40	13,16%

Fonte: autor

O modelo Random Forest com ajuste de *threshold* (0,48) e hiperparâmetros modificados se destacou com a melhor média de combinação de recall (0,8761) e F1-score (0,8282), com um número reduzido de falsos negativos (22) e uma taxa de erro geral de 12,96%. De forma similar, foram obtidos resultados expressivos com XGBoost (4), usando hiperparâmetros ajustados manualmente a partir dos valores obtidos com o Grid (recall = 0,8580; F1 = 0,8229), com taxa de erro de predição de 13,16%.

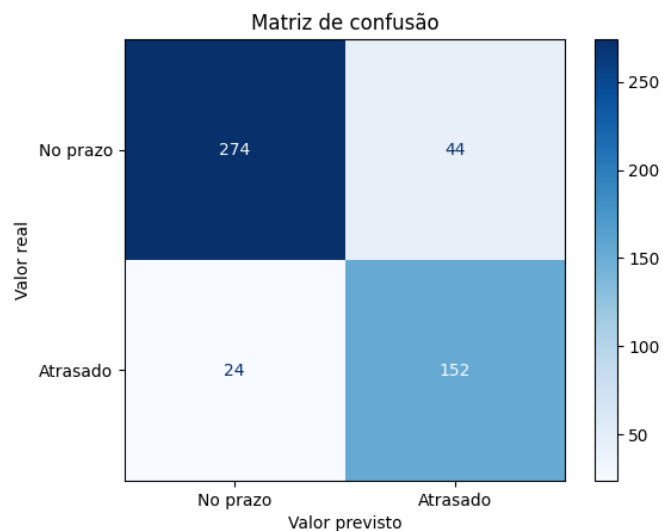
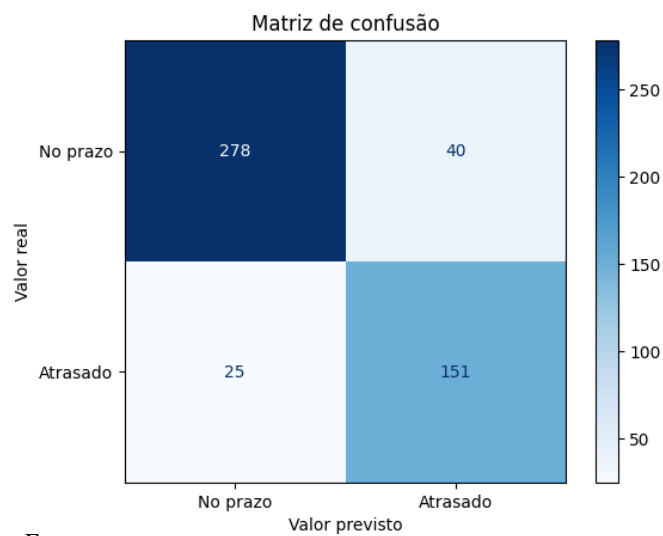
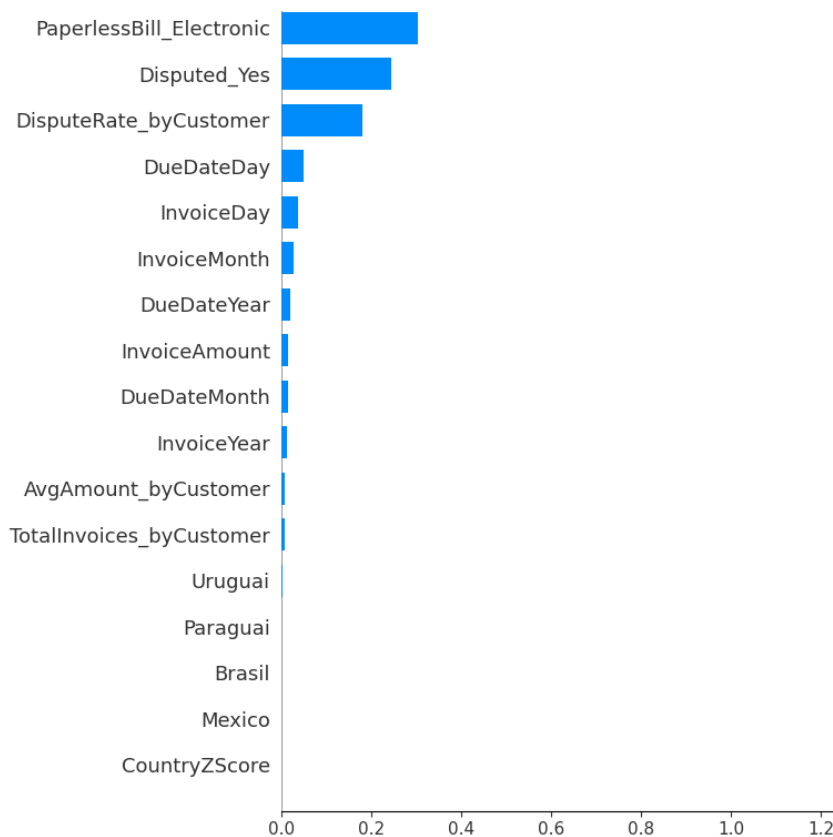


Figura 3 – Gráfico SHAP do impacto médio de cada variável



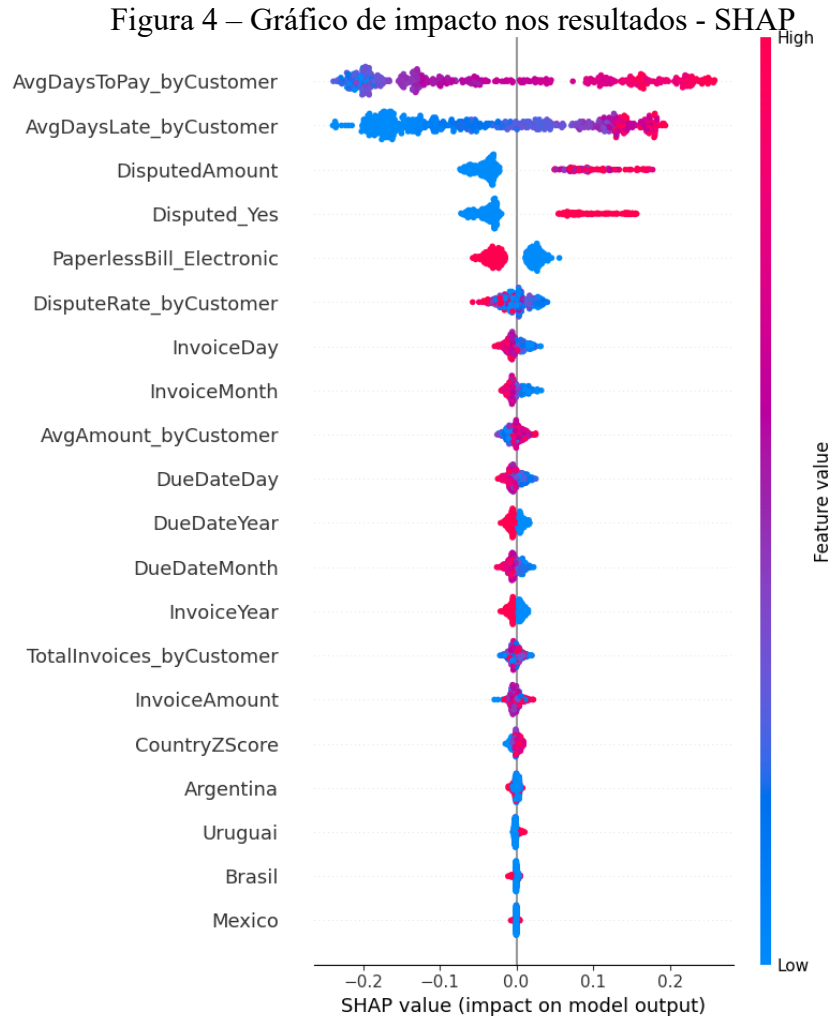
Fonte: autor

Ao comparar modelos completos com modelos que utilizam apenas as 6 variáveis mais relevantes, observa-se que a perda de desempenho é marginal na classificação de inadimplentes (Figura 1 e 2), enquanto apenas identifica alguns falsos positivos adicionais em comparação. As 6 variáveis utilizadas de maior contribuição constam na Figura 3.

Em alguns casos, como nos modelos 1 e 2 do XGBoost (Tabela 2), a performance com as 6 variáveis é praticamente idêntica à do modelo completo. A diferença máxima observada fora de 3 classificações errôneas para modelos treinados de acordo com os mesmos hiperparâmetros, considerando todos ou apenas as variáveis. O que demonstra a força explicativa desse subconjunto restrito de atributos e sua utilidade em aplicações que exigem transparência ou menor complexidade computacional.

Referente as variáveis no geral, de acordo com os valores de SHAP de um dos modelos — ressaltando que a mudança dos valores SHAP entre os modelos é pequena — os principais determinantes para a previsão de inadimplência estão associados ao comportamento histórico do cliente, com as variáveis AvgDaysToPay_byCustomer e AvgDaysLate_byCustomer, as mais influentes, indicando que atrasos frequentes e maior tempo médio de pagamento

aumentam significativamente a probabilidade de inadimplência (Figura 4).



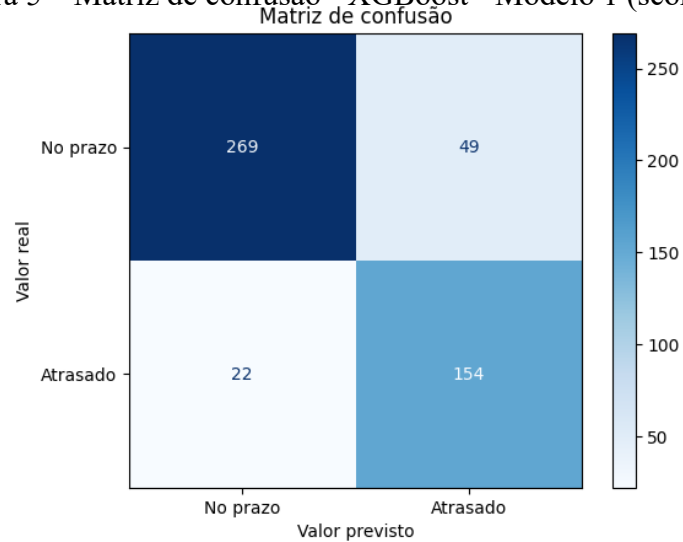
Fonte: autor

Variáveis relacionadas a disputas, como `DisputedAmount`, `Disputed_Yes` e `DisputeRate_byCustomer`, também apresentaram forte impacto, sugerindo que clientes que questionam faturas representam maior risco (Figura 4). A variável `PaperlessBill_Electronic` mostrou que clientes que recebem faturas eletrônicas tendem a ser mais adimplentes. Elementos temporais como mês e dia da fatura ou do vencimento tiveram impacto moderado, enquanto fatores como país de origem e indicadores agregados por país (como o `CountryZScore`) tiveram pouca ou nenhuma influência.

Também se observou que o uso de `GridSearchCV` com diferentes métricas de otimização (precisão ou ROC AUC) influencia o perfil do modelo: scoring por precisão tende a favorecer modelos mais conservadores (Figura 5); já o scoring por ROC AUC produziu

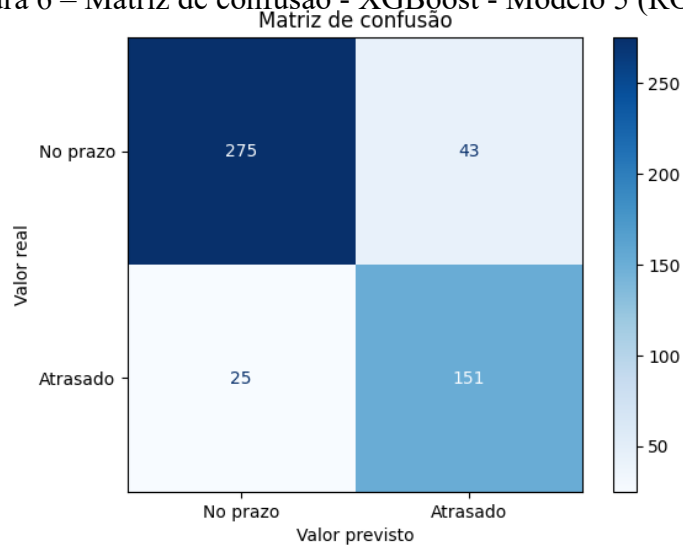
modelos com maior recall, às custas de aumento nos falsos positivos (Figura 6).

Figura 5 – Matriz de confusão - XGBoost - Modelo 1 (scoring)



Fonte: autor

Figura 6 – Matriz de confusão - XGBoost - Modelo 5 (ROC -AUC)



Fonte: autor

A análise dos resultados dos modelos aplicados à estimativa de inadimplência evidencia a robustez e a aplicabilidade dos algoritmos Random Forest e XGBoost na previsão de comportamentos financeiros de risco. Com métricas como F1-score acima de 0,82 e valores de ROC-AUC superiores a 0,92, mesmo com conjuntos reduzidos de variáveis, os modelos demonstraram capacidade consistente de distinguir devedores inadimplentes com alto grau de confiabilidade. A flexibilidade de ajuste via *threshold* permite alinhar o modelo aos objetivos

institucionais, seja priorizando a sensibilidade (ao reduzir o limite e evitar falsos negativos) ou a precisão (ao elevar o limite para conter falsos positivos).

Esse desempenho torna-se relevante ao se considerar que as estimativas são fundamentais na contabilidade e frequentemente construídas a partir de critérios objetivos e subjetivos. Conforme argumentam Gaver e Paterson (2004), estimativas como a provisão para devedores duvidosos (PDD) são particularmente vulneráveis a vieses e manipulações gerenciais, comprometendo sua fidedignidade. Nesse contexto, a adoção de modelos preditivos com respaldo estatístico e baseados em dados históricos contribui para o fortalecimento da objetividade, introduzindo maior rigor, imparcialidade e transparência ao processo.

Além disso, previsões imprecisas não apenas afetam a qualidade da informação, mas também podem comprometer decisões de gestão e investimento, como destacam Anand *et al.* (2021), ao apontarem que avaliações distorcidas podem resultar em escolhas equivocadas. O alto desempenho dos modelos na identificação de casos de inadimplência evidencia, assim, o potencial do aprendizado de máquina na mitigação de perdas financeiras e reforça sua aplicabilidade em processos como a concessão de crédito ou definição de provisões.

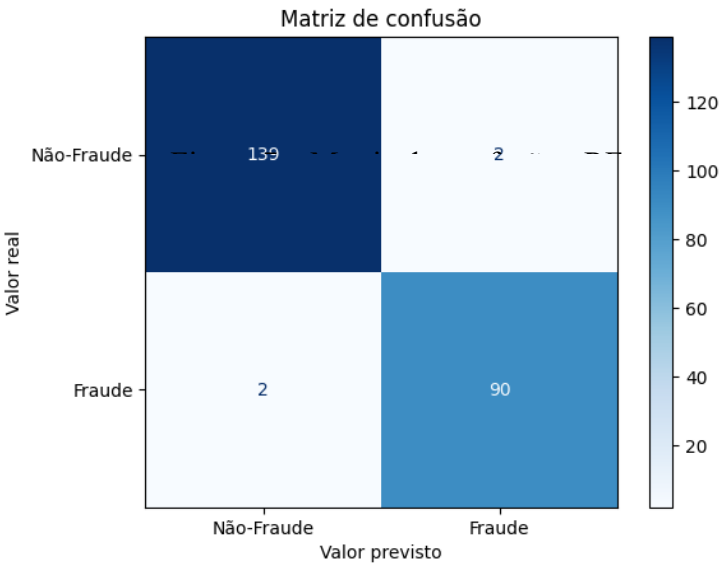
Destaca-se, ainda, que as métricas elevadas obtidas em conjunto com SHAP aumentam a explicabilidade do modelo, permitindo uma análise mais aprofundada das variáveis em determinado contexto, o que marca o potencial de aplicação na contabilidade gerencial. Os resultados também favorecem o papel do aprendizado de máquina na elevação da qualidade da informação contábil e na antecipação de cenários de risco, função já apontada por Zhang *et al.* (2020), Ayad, El Mezouari e Kharmoum (2023) e corroborada empiricamente por Ding *et al.* (2020).

Por fim, a análise permite inferir (com ressalvas do modelo) que a aplicação bem-sucedida da tecnologia nessa área pode promover ganhos de eficiência nos processos e amplia a atuação do contador. Ao dispor de instrumentos que potencializam a extração de *insights*, esse profissional passa a desempenhar funções mais analíticas e estratégicas, exigindo o domínio de competências técnicas e tecnológicas, como enfatizam De Araujo e Cornacchione (2024).

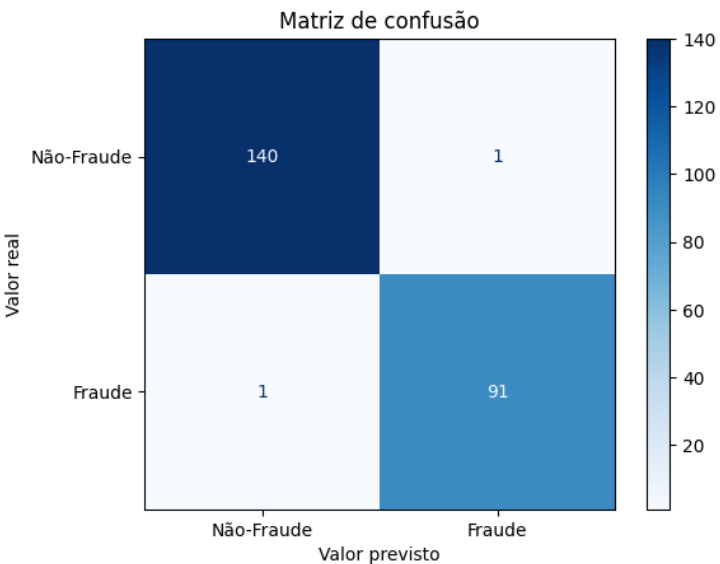
4.2 Resultados – Modelos De Auditoria

A aplicação prática dos modelos Random Forest e XGBoost no contexto de auditoria revelou desempenhos expressivos, mesmo sob condições de baixa dimensionalidade. O modelo Random Forest treinado com GridSearchCV apresentou um F1-score de 0,9783 e apenas 2

falsos positivos e falsos negativos (Figura 7). Na média do retreino em 10 execuções com os hiperparâmetros otimizados, teve-se um F1-score de 0,9864 e apenas um erro tipo I e um tipo II (Figura 8), representando uma taxa de erro de apenas 0,86%.



Fonte: autor



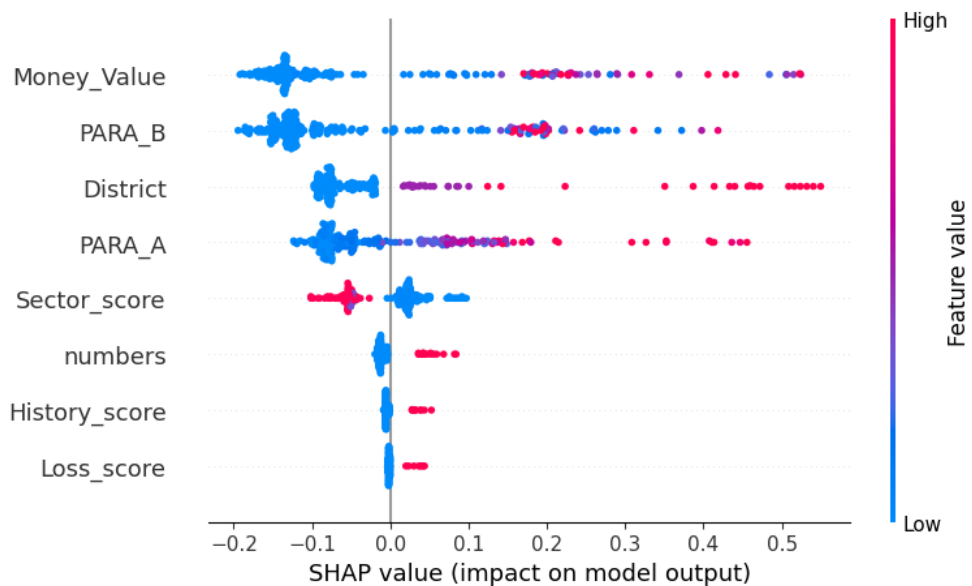
Fonte: autor

O modelo XGBoost, por sua vez, apresentou desempenho semelhante ao random Forest em termos de F1-score (0,9783), recall (0,9783), taxa de erro (1,72%) e ROC AUC (≈ 1), tanto na execução com GridSearchCV quanto na média dos retreinos. O que indica que ambos

os algoritmos têm elevada capacidade discriminativa. Esses achados estão alinhados com os resultados relatados por Hooda, Bawa e Rana (2018), que também utilizaram modelos baseados em árvores de decisão em um conjunto similar, o que reforça a robustez e eficácia desses modelos preditivos e seu potencial em aplicações em auditoria.

Quando a explicabilidade dos modelos, a análise dos valores SHAP permitiu identificar variáveis-chave que influenciaram na predição de fraude, enquanto outras variáveis exerceram impacto pequeno em modelos RF (Figura 9), ou praticamente nulo, no caso do modelo XGBoost (Figura 10), reforçando que os modelos do último tipo se concentraram fortemente em poucos atributos mais informativos. Isso não confronta diretamente o estudo indiano, pois o gráfico SHAP mostra a importância das variáveis de forma individualizada, enquanto a seleção pelo PSO considera interações sinérgicas entre elas para definir os pesos.

Figura 9 – Gráfico de impacto nos resultados - Random Forest - SHAP



Fonte: autor

Figura 10 – Gráfico de impacto nos resultados - XGBoost - SHAP



Fonte: autor

A interpretabilidade oferecida pelo SHAP permite visualizar a influência individual de cada variável sobre a predição global e localmente, promovendo maior transparência dos modelos, o que responde diretamente às exigências do campo contábil por decisões justificáveis, contribuindo para *insights* e maior compreensibilidade dos usuários (Hooda; Bawa; Rana, 2018; Sekar, 2022). Em suma, os resultados desta etapa prática não apenas confirmam a eficácia dos modelos como ferramentas auxiliares à auditoria, mas também reforçam as inferências presentes na literatura revisada.

Destaca-se, por exemplo, a possibilidade automação eficiente da análise de grandes volumes de dados, com foco dos auditores apenas em tarefas críticas (Hu; Sun, 2021), a tecnologia como alternativa frente a abordagens tradicionais menos escaláveis (Ashtiani; Raahemi, 2021), e a maior acurácia na identificação de distorções contábeis, diagnósticos mais precisos e geração de *insights* mais robustos, contribuindo para a melhoria da qualidade da informação contábil (Hu; Sun, 2021; Zhang *et al.*, 2020; Bertomeu *et al.*, 2020).

Dessa forma, o experimento reforça o papel estratégico dessas ferramentas na modernização dos processos de auditoria, o que suscita a necessidade da transformação da atuação profissional do auditor, ao mesmo tempo em que evidencia sua aplicabilidade concreta e viabilidade em cenários reais.

4.3 Análise das Implicações associadas a adoção do ML

Os resultados obtidos dos modelos construídos demonstraram a efetividade de modelos supervisionados de ML, especificamente RF e XGBoost, que alcançaram níveis satisfatórios de desempenho. Também, a utilização do método SHAP acrescentou uma camada de

interpretabilidade relevante, permitindo a identificação da contribuição das variáveis nos resultados. No entanto, a efetividade técnica desses modelos ganha relevância plena quando relacionada às implicações de sua adoção na prática contábil.

O ML tem se consolidado em múltiplas aplicações no setor financeiro, como previsão de crises, detecção de fraudes, análise de risco de crédito e até previsão de insolvências. Embora essas aplicações não pertençam exclusivamente à contabilidade, guardam proximidade com atividades de auditoria, análise financeira e mensuração de riscos, que integram o cotidiano da profissão contábil. Assim, os achados empíricos dialogam com uma tendência mais ampla de incorporação de sistemas automatizados de análise em processos críticos de tomada de decisão.

Um primeiro ponto a se destacar, derivado dessa comparação, diz respeito à eficiência na detecção de irregularidades. Estudos como o de Gao *et al.* (2024) e Chatzis *et al.* (2018) apontam que a integração de técnicas de ML permite identificar padrões sutis em grandes volumes de dados, reduzindo falhas humanas, aumentando a acurácia das análises e, até mesmo antecipando desvios, protegendo agentes do mercado (Kirkos; Spathis; Manolopoulos, 2007).

Nesse sentido, as aplicações voltadas à previsão de falências e insolvências (Petrooulos *et al.*, 2020; Wu; Ma; Olson, 2022) também demonstram que o ML fortalece mecanismos de gestão de riscos e de continuidade operacional (Ayad; El Mezouari; Kharmoum, 2023), o que também está relacionado à qualidade da informação gerada. Transposto para a contabilidade, isso amplia o papel estratégico da área, colocando-a como suporte decisivo na antecipação de crises e na proteção de credores, investidores e da própria sobrevivência empresarial.

Do ponto de vista operacional, questões de custo e disponibilidade de dados também merecem destaque. Tradicionalmente, o processamento de modelos complexos era associado a elevados custos computacionais; contudo, avanços recentes em algoritmos e infraestrutura tecnológica vêm reduzindo tais barreiras, ampliando a viabilidade prática para diferentes organizações. Além disso, embora bases extensas sejam vantajosas, observa-se que conjuntos menores, quando bem estruturados, podem produzir resultados satisfatórios (a depender do modelo utilizado), o que reforça o potencial de adoção do ML em contextos contábeis com recursos limitados.

Conforme destacado por Reis *et al.* (2020) e Wright e Schultz (2018), ganhos em produtividade, automação de decisões operacionais e geração de valor estratégico são implicações diretas da incorporação da IA nas organizações. A integração de modelos preditivos em atividades críticas também aponta para a consolidação da IA no centro da governança organizacional, e não apenas como ferramenta acessória (Shrestha *et al.*, 2019).

A adoção de ML impacta estruturas de controle interno, processos decisórios e até a cultura organizacional, ao introduzir maior dependência de sistemas tecnológicos. Isso pode gerar ganhos em eficiência e redução de custos, mas também aumenta a vulnerabilidade a falhas técnicas e de segurança.

Assim, a contabilidade, ao incorporar ML, não apenas amplia sua capacidade analítica, mas também assume responsabilidades adicionais na gestão e mitigação de novos riscos tecnológicos, garantindo responsabilidade e confiabilidade das práticas. Nesse cenário, órgãos normativos e reguladores também precisarão estabelecer parâmetros claros para o uso de algoritmos, assegurando transparência e responsabilidade.

Além das implicações organizacionais, há também a dimensão da transformação do papel do profissional contábil. A literatura implica que a incorporação dessas tecnologias pode levar à automação de funções rotineiras, como conciliações, verificações de consistência e detecção preliminar de anomalias, reduzindo a necessidade de execução manual dessas tarefas. Evidências empíricas, como as de Ding et al. (2020), mostram que modelos podem oferecer previsões mais consistentes do que gestores humanos em determinadas situações, sinalizando uma possível redistribuição de responsabilidades dentro da prática contábil.

Esse deslocamento não significa a substituição completa do profissional — embora a perda de empregos seja esperada (Abdi *et al.*, 2021) —, mas sua reconfiguração como mediador entre tecnologia e usuários da informação. O contador passa a assumir funções de validação, interpretação e comunicação dos resultados (De Araujo; Cornacchione, 2024; Shi, 2020; Zhang *et al.*, 2020), além de, no contexto de ML, monitorar riscos associados a vieses, sobreajuste ou falhas dos modelos.

Esse novo papel exige maior domínio técnico sobre ferramentas analíticas e capacidade crítica para contextualizar os *outputs* algorítmicos no ambiente de negócios. Tal cenário implica também em novas demandas formativas. Como apontam Weeks, Voshaar e Zimmermann (2023) e Frey e Osborne (2017), currículos de Ciências Contábeis precisarão integrar competências em ciência de dados, ao mesmo tempo em que programas de educação continuada deverão preparar profissionais já atuantes para lidar com essas mudanças.

A literatura também sugere que o aumento da eficiência não elimina riscos, uma vez que erros de modelagem, vieses dos dados ou sobreajuste podem comprometer resultados (Veganzones; Séverin, 2018; Zoričák *et al.*, 2020), criando um falso senso de confiabilidade. Soma-se a isso o desafio da interpretabilidade: a opacidade de modelos complexos ainda constitui barreira relevante.

Tanto nos experimentos realizados com SHAP quanto em estudos da área (Lundberg, 2020; Ranta; Ylinen; Järvenpää, 2022), evidencia-se que ferramentas de explicabilidade são fundamentais para tornar esses modelos compatíveis com as exigências da contabilidade, que demanda resultados justificáveis e interpretáveis. Contudo, essas ferramentas exigem capacitação técnica de seus usuários, bem como a adaptação dos *outputs* para diferentes públicos que utilizariam esse tipo de informação, como o corpo da gestão das organizações, *stakeholders etc.*

Portanto, a análise integrada evidencia que a implementação do aprendizado de máquina na contabilidade não deve ser entendida apenas como uma questão de performance técnica, mas como um processo de transformação estrutural. O ML tende a ampliar o alcance e a precisão das análises, possibilita avanços significativos na detecção de irregularidades e na mensuração de riscos, e aponta para uma redefinição do papel do profissional contábil. Porém, traz também implicações regulatórias, éticas e organizacionais que precisam ser cuidadosamente consideradas na sua adoção.

5. CONCLUSÕES

A literatura analisada mostrou que a adoção de algoritmos inteligentes pode contribuir substancialmente para o aprimoramento da qualidade da informação contábil, a redução de incertezas e a automatização de tarefas que antes dependiam fortemente da intervenção humana. As ferramentas de *machine learning* mostram-se particularmente promissoras em atividades como auditoria, estimativas contábeis, análise de textos e previsões financeiras, permitindo maior agilidade e acurácia nas análises.

Os resultados com os experimentos dos modelos *random forest* e *XGBoost* comprovaram a efetividade do aprendizado de máquina em dois cenários relevantes da contabilidade, oferecem desempenho consistente em ambas, em contexto de risco. Além de acurácia, a incorporação do SHAP adicionou transparência sobre a contribuição de variáveis, fortalecendo a explicabilidade exigida na prática contábil e tornando as inferências mais comunicáveis a diferentes usuários da informação.

Destaca-se, também, a robustez das soluções mesmo com a utilização de subconjuntos reduzidos de variáveis, o que indica a viabilidade da tecnologia em ambientes com limitações de dados; somam-se a isso a redução de custos computacionais e temos ampliação do acesso à tecnologia por organizações de diferentes portes.

Ademais, constatou-se que a aplicação de soluções de IA nas organizações favorece ganhos operacionais, melhora a gestão de riscos e proporciona suporte à tomada de decisão estratégica em tempo real, ainda que a implementação dessas soluções enfrente obstáculos, como a necessidade de cultura digital para uso eficaz das ferramentas, investimentos técnicos e cuidados com governança algorítmica, ética e regulação.

Do ponto de vista profissional, essas transformações alteram substancialmente o perfil e o papel do contador, em direção a funções mais estratégicas e analíticas, exigindo domínio técnico sobre dados, algoritmos e ferramentas digitais, além de capacidade crítica para interpretar resultados e orientar decisões. Ao mesmo tempo, a pesquisa corrobora que o avanço da automação tende a reduzir a demanda por funções operacionais e repetitivas, reforçando a importância de requalificação e do desenvolvimento contínuo de competências analíticas e interpessoais, como forma de garantir empregabilidade e protagonismo profissional em um cenário contábil cada vez mais orientado por dados e tecnologia.

No plano aplicado, as evidências e métodos testados no trabalho apontam caminhos concretos para a área. Em estimativas contábeis (por exemplo, provisões e perdas esperadas), modelos como Random Forest e XGBoost podem ser treinados com históricos internos e

variáveis econômicas, com ajuste explícito de limiar conforme materialidade e apetite a risco do reporte; o uso de SHAP permite justificar mudanças de estimativas, explicitar *drivers*, isto é, os fatores que influenciam os resultados, e documentar a razoabilidade para usuários internos e externos, elevando relevância e representação fidedigna da informação.

Em auditoria, classificadores podem pontuar transações, saldos ou entidades por risco, direcionando amostras e testes substantivos, apoiando rotinas de auditoria contínua e produzindo “trilhas de explicação” para os papéis de trabalho; aqui também a calibragem de limiares ajuda a balancear eficiência (menos falsos positivos) com proteção contra distorções relevantes (menos falsos negativos); o *trade-off* será considerado pela organização em vista das consequências benéficas que a abordagem trará.

Em ambos os casos, pode-se começar com bases menores e curadas, expandindo gradualmente escopo e variáveis à medida que processos e controles de governança se consolidam; a adoção deve contemplar infraestrutura proporcional, políticas de controle interno, gestão de vieses presentes nos dados e resultados e treinamento da equipe para interpretação crítica dos *outputs*. Dessa forma, a contabilidade passa a combinar modelos preditivos com critérios profissionais e regulatórios, integrando explicabilidade, custo e qualidade da informação como eixos de decisão e ampliando sua contribuição estratégica para a gestão, a confiabilidade dos relatórios e a supervisão dos riscos.

5.1 Limitações

Apesar dos resultados positivos obtidos, esta pesquisa apresenta algumas limitações importantes. Primeiramente, os experimentos foram conduzidos com conjuntos de dados públicos, o que pode não refletir com exatidão a complexidade e a heterogeneidade dos dados utilizados pelas organizações num ambiente real. Além disso, a representatividade temporal e setorial desses dados não foi confrontada com a realidade contábil brasileira, o que limita a generalização externa dos achados. Outro fator que limita a generalização é a proveniência incerta/ambígua dos dados do conjunto usado para estimativa de inadimplência.

Outra limitação refere-se à escolha restrita de algoritmos. Embora *random forest* e *XGBoost* sejam métodos amplamente reconhecidos por sua robustez, a ausência de outras abordagens, restringe a possibilidade de comparação de desempenho e pode limitar a abrangência dos achados.

Ademais, o estudo não contempla questões práticas relacionadas à implementação dos

modelos em contextos organizacionais reais, como infraestrutura de dados, segurança da informação, viabilidade econômica e resistência cultural, fatores relevantes para a adoção da tecnologia no ambiente contábil. Também não foram aprofundados questões de integração no fluxo de trabalho, ou foram considerados os custos de implantação nem as demandas de operação contínua, como monitoramento, atualização e auditoria dos modelos.

Outra limitação reside na não exploração de modelos não supervisionados, semissupervisionados ou de aprendizado por reforço, que poderiam ser úteis para detecção de padrões em contextos com dados não rotulados, como análises textuais ou categorização automatizada de documentos. Por fim, A avaliação dos modelos foi centrada em métricas estatísticas tradicionais, como acurácia, F1-score, precisão, recall e curva ROC, sem considerar diretamente o impacto qualitativo do uso de ML na rotina e na tomada de decisão de profissionais contábeis. Faltou, ainda, avaliar calibração de probabilidades e definição de *thresholds* orientada a custo, bem como relatar incerteza.

5.2 Sugestões Para Pesquisas Futuras

Recomenda-se a realização de estudos com dados proprietários de empresas, o que permitiria testar a robustez e a adaptabilidade dos modelos a contextos reais, mais desafiadores e sensíveis. Essa abordagem possibilitaria a avaliação mais fiel dos desafios técnicos e institucionais enfrentados na adoção do aprendizado de máquina no cotidiano das organizações. Além disso, seria enriquecedor comparar os modelos utilizados com algoritmos de outras naturezas, como redes neurais profundas, métodos não supervisionados (ex.: *clustering*), modelos híbridos e técnicas de aprendizado por reforço, ampliando o escopo metodológico e explorando outras dimensões do *machine learning* na contabilidade.

Também se sugerem estudos que abordem a integração entre modelos preditivos e sistemas contábeis tradicionais (ERPs), bem como os impactos sobre a governança algorítmica, auditoria contínua, conformidade regulatória e segurança dos dados. Investigações futuras poderiam contemplar demandas de operação contínua, controles internos de dados ou requisitos de conformidade para uso dos modelos.

Outra frente de pesquisa relevante consiste em investigar modelos voltados à explicabilidade e causalidade, além da previsão. Estudos que busquem compreender relações entre variáveis, medir efeitos e sustentar decisões explicativas. Também, novas pesquisas poderiam explorar o uso da tecnologia, ou da IA em geral, na contabilidade gerencial, no gerenciamento de inventário, planejamento financeiro e escrituração contábil.

Finalmente, considerando o papel crescente do contador como analista de dados, recomenda-se o aprofundamento pesquisas sobre formação, capacitação e aceitação profissional da IA no ambiente contábil brasileiro Tais investigações podem subsidiar políticas de reestruturação curricular nos cursos de ciências contábeis e fomentar estratégias organizacionais de desenvolvimento humano, visando uma transição tecnológica mais eficaz e sustentável.

REFERÊNCIAS

MARTINS, Eliseu. Contabilidade de Custos. 9. ed. Atlas, 2003.

ABDI, Meyad Diriba; DOBAMO, Habtamu Abafoge; BAYU, Kefiyalew Belachew. Exploring current opportunity and threats of artificial intelligence on small and medium enterprises accounting function; evidence from south west part of ethiopia, oromiya, jimma and snnpr, bonga. **Academy of Accounting and Financial Studies Journal**, v. 25, n. 2, p. 1-11, 2021. Disponível em: <https://www.researchgate.net/profile/Kefiyalew-Belachew/publication/350431218>. Acesso em: 8 mai. 2025.

AGOSTINI, Carla; CARVALHO, Joziane T. de. A Evolução da Contabilidade: seus avanços no Brasil e a Harmonização com as Normas Internacionais. **Instituto de Ensino Superior Tancredo de Almeida Neves. Armário de Produção**, v. 1, 2012. Disponível em: <https://www.academia.edu/8230010/>. Acesso em: 20 dez. 2024.

AHMED, Israr; AZIZ, Abdul. Dynamic Approach for Data Scrubbing Process. **International Journal on Computer Science and Engineering**, Paquistão, v. 2, n. 2, p. 416-423, 2010. Disponível em: <https://www.academia.edu/download/77620663/IJCSE10-02-02-53.pdf>. Acesso em: 17 jan. 2025.

AHMED, Salman et al. FLF-LSTM: A novel prediction system using Forex Loss Function. **Applied Soft Computing**, v. 97, p. 106780, 2020. Disponível em: <https://doi.org/10.1016/j.asoc.2020.106780>. Acesso em: 15 fev. 2025.

ALIFERIS, Constantin; SIMON, Gyorgy. Overfitting, underfitting and general model overconfidence and under-performance pitfalls and best practices in machine learning and AI. Artificial intelligence and machine learning in health care and medical sciences: **Best practices and pitfalls**, EUA, p. 477-524, mar./2024. Disponível em: https://link.springer.com/content/pdf/10.1007/978-3-031-39355-6_10.pdf. Acesso em: 25 jan. 2025.

ALTMAN, Edward I. Financial ratios, discriminant analysis and the prediction of corporate bankruptcy. **The journal of finance**, v. 23, n. 4, p. 589-609, 1968. Disponível em: <https://doi.org/10.2307/2978933>. Acesso em: 19 fev. 2025.

ANAND, Vic et al. Predicting profitability using machine learning. **Available at SSRN 3466478**, 2019. Disponível em: <https://dx.doi.org/10.2139/ssrn.3466478>. Acesso em: 29 mar. 2025.

ARRUDA, D.C.S.; GOMES, E.Z.; SANTOS, Cleston Alexandre. Uma Na Análise Da Percepção Dos Profissionais Da Área De Contabilidade Do Município De Corumbá-Ms Sobre O Sped. **Revista Científica Semana Acadêmica**. Fortaleza, 2013, n. 42, nov./2013. Disponível em: https://semanaacademica.org.br/system/files/artigos/artigo_sped.pdf. Acesso em: 08 jan. 2025.

ASHTIANI, Matin N.; RAAHEMI, Bijan. Intelligent fraud detection in financial statements using machine learning and data mining: a systematic literature review. **Ieee Access**, v. 10, p.

72504-72525, 2021. Disponível em:
<https://ieeexplore.ieee.org/iel7/6287639/6514899/09481913.pdf>. Acesso em: 4 abr. 2025

ASSAF NETO, Alexandre. **Estrutura e análise de balanços: um enfoque econômico-financeiro**. 2020. p. 105-106.

ATHEY, Susan; IMBENS, Guido W. Machine learning methods that economists should know about. **Annual Review of Economics**, v. 11, n. 1, p. 685-725, 2019. Disponível em:
<https://doi.org/10.1146/annurev-economics-080217-053433>. Acesso em: 26 jan. 2025.

AUTOR, David H. Why are there still so many jobs? The history and future of workplace automation. **Journal of economic perspectives**, v. 29, n. 3, p. 3-30, 2015. Disponível em:
<http://doi.org/10.1257/jep.29.3.3>. Acesso em: 6 mai. 2025.

AYAD, Meryem; MEZOUARI, Said El; KHARMOUM, Nassim. Impact of Machine Learning on the Improvement of Accounting Information Quality. **International Conference On Advanced Intelligent Systems For Sustainable Development**, v. 637, n. 1, p. 501-514, 14 jun./2023. Disponível em: http://doi.org/10.1007/978-3-031-26384-2_43. Acesso em: 14 jan. 2025.

AYALA, Jordan et al. Technical analysis strategy optimization using a machine learning approach in stock market indices. **Knowledge-Based Systems**, v. 225, 107119, 2021. Disponível em: <https://doi.org/10.1016/j.knosys.2021.107119>. Acesso em: 11 fev. 2025.

BAO, Wang; LIANJU, Ning; YUE, Kong. Integration of unsupervised and supervised machine learning algorithms for credit risk assessment. **Expert Systems with Applications**, v. 128, p. 301-315, 2019. Disponível em: <https://doi.org/10.1016/j.eswa.2019.02.033>. Acesso em: 27 fev. 2025.

BAO, Yang et al. Detecting accounting fraud in publicly traded US firms using a machine learning approach. **Journal of Accounting Research**, v. 58, n. 1, p. 199-235, 2020. Disponível em: <https://papers.ssrn.com/sol3/Delivery.cfm?abstractid=2670703>. Acesso em: 3 abr. 2025.

BARBOZA, Flavio; KIMURA, Herbert; ALTMAN, Edward. Machine learning models and bankruptcy prediction. **Expert systems with applications**, v. 83, p. 405-417, 2017. Disponível em: <https://doi.org/10.1016/j.eswa.2017.04.006>. Acesso em: 21 fev. 2025.

BASAK, Suryoday et al. Predicting the direction of stock market prices using tree-based classifiers. **The North American Journal of Economics and Finance**, v. 47, p. 552-567, 2019. Disponível em: <https://doi.org/10.1016/j.najef.2018.06.013>. Acesso em: 12 fev. 2025.

BAY, Stephen et al. Large scale detection of irregularities in accounting data. **Sixth International Conference on Data Mining (ICDM'06). IEEE**, 2006. p.75-86. Disponível em: <https://doi.org/10.1109/ICDM.2006.93>. Acesso em: 5 abr. 2025.

BEAVER, William H. Financial ratios as predictors of failure. **Journal of accounting research**, p. 71-111, 1966. Disponível em: <https://doi.org/10.2307/2490171> Acesso em: 19 fev. 2025.

BENBYA, Hind; DAVENPORT, Thomas H.; PACHIDI, Stella. Artificial intelligence in organizations: Current state and future opportunities. **MIS Quarterly Executive**, v. 19, n. 4, 2020. Disponível em: <https://papers.ssrn.com/sol3/Delivery.cfm?abstractid=3741983>. Acesso em: 27 dez. 2024.

BENGIO, Yoshua; LECUN, Yann; HINTON, Geoffrey. Deep learning for AI. **Communications of the ACM**, v. 64, n. 7, p. 58-65, 2021. Disponível em: <https://dl.acm.org/doi/abs/10.1145/3448250>. Acesso em: 26 jan. 2025.

BERTOMEU, J. *et al.* Using machine learning to detect misstatements. **Review of Accounting Studies**, EUA, v. 26, n. 26, p. 468-519, 2020a. Disponível em: <https://link.springer.com/article/10.1007/s11142-020-09563-8>. Acesso em: 20 fev. 2025.

BERTOMEU, Jeremy. Machine learning improves accounting: discussion, implementation and research opportunities. **Review of Accounting Studies**, v. 25, n. 3, p. 1135-1155, 2020. Disponível em: <https://doi.org/10.1007/s11142-020-09554-9>. Acesso em: 17 mar. 2025.

BOOKS NGRAM VIEWER. **Machine learning mentions**. Disponível em: https://books.google.com/ngrams/graph?content=big+data&year_start=1950&year_end=2022. Acesso em: 19 dez. 2024.

BOSE, Sudipta; DEY, Sajal Kumar; BHATTACHARJEE, Swadip. Big Data, Data Analytics and Artificial Intelligence in Accounting: An Overview. **Handbook of Big Data Methods**, p. 32-51, mar/2022. Disponível em: <https://doi.org/10.4337/9781800888555.00007>. Acesso em: 19 dez. 2024.

BROWN, Richard. **A History of Accounting and Accountants**. 2. ed. Cosimo Classics, 2006.

BROWNLOW, Josh *et al.* Data and analytics-data-driven business models: A Blueprint for Innovation. **Cambridge Service Alliance**, v. 7, n. Fevereiro, p. 1-17, 2015. Disponível em: <https://doi.org/10.13140/RG.2.1.2233.2320>. Acesso em: 9 jan. 2025.

BÜHLMANN, Peter. Bagging, boosting and ensemble methods. Handbook of computational statistics: **Concepts and methods**, p. 985-1022, 2012. Disponível em: https://www.econstor.eu/bitstream/10419/22204/1/31_pb.pdf. Acesso em: 2 fev. 2025.

BURKOV, Andriy. **The Hundred-Page Machine Learning Book**. 1. ed. EUA: Andriy Burkov, 2019. p. 1-161.

CAO, Longbing. Data science: a comprehensive overview. **ACM Computing Surveys (CSUR)**, v. 50, n. 3, p. 1-42, 2017. Disponível em: <https://dl.acm.org/doi/pdf/10.1145/3076253>. Acesso em: 12 fev. 2025.

CAWLEY, Gavin C.; TALBOT, Nicola L.c. On over-fitting in model selection and subsequent selection bias in performance evaluation. **The Journal of Machine Learning Research**, [S.L], v. 11, n. 1, p. 2079-2107, ago./2010. Disponível em: <https://www.jmlr.org/papers/volume11/cawley10a/cawley10a.pdf>. Acesso em: 25 jan. 2025.

CHAGAS, Mário Francisco et al. Tecnologia na contabilidade: Uma análise dos sistemas fiscais, trabalhistas e contábeis. **Diálogos em Contabilidade: Teoria e Prática**, v. 1, n. 1, 2013. Disponível em: <http://periodicos.unifacef.com.br/dialogoscont/article/download/1224/934>. Acesso em: 19 dez. 2024.

Chapman, C.; Chua, W. F. Technology-driven integration, automation, and standardization of business processes: implications for accounting. **Management Accounting in the Digital Economy**. Nova York: Oxford University Press. v. 1, p. 74-94, 2003.

CHATZIS, Sotirios P. et al. Forecasting stock market crisis events using deep and statistical machine learning techniques. **Expert systems with applications**, v. 112, p. 353-371, 2018. Disponível em: <https://doi.org/10.1016/j.eswa.2018.06.032>. Acesso em: 16 fev. 2025.

CHEN, Tianqi; GUESTRIN, Carlos. Xgboost: A scalable tree boosting system. **Proceedings of the 22nd acm sigkdd international conference on knowledge discovery and data mining**. 2016. p. 785-794. Disponível em: <https://dl.acm.org/doi/pdf/10.1145/2939672.2939785> . Acesso em: 3 fev. 2025.

CHEN, Wei et al. Mean–variance portfolio optimization using machine learning-based stock price prediction. **Applied soft computing**, v. 100, p. 106943, 2021. Disponível em: <https://doi.org/10.1016/j.asoc.2020.106943>. Acesso em: 13 fev. 2025.

CHU, X. *et al.* Data Cleaning: Overview and Emerging Challenges. **Proceedings of the 2016 international conference on management of data**, USA, v. 16, n. 1, p. 2201-2206, jun./2016. Disponível em: <https://doi.org/10.1145/2882903.2912574>. Acesso em: 17 jan. 2025.

CLOUD GOOGLE. **O que é aprendizado não supervisionado?** 2024. Disponível em: <https://cloud.google.com/discover/what-is-unsupervised-learning?hl=pt-BR>. Acesso em: 16 jan. 2025.

COELHO, Felipe Fernandes; DE LIMA AMORIM, Daniel Penido; DE CAMARGOS, Marcos Antônio. Analisando métodos de machine learning e avaliação do risco de crédito. **Revista Gestão & Tecnologia**, v. 21, n. 1, p. 89-116, 2021. Disponível em: <https://revistagt.fpl.emnuvens.com.br/get/article/view/2089/1198>. Acesso em: 2 mar. 2025.

Comitê de Pronunciamentos Contábeis (CPC). **Pronunciamento Técnico CPC 23: Políticas Contábeis, Mudança de Estimativa e Retificação de Erro**. Brasília, 2022. Disponível em: https://s3.sa-east-1.amazonaws.com/static.cpc.aatb.com.br/Documentos/296_CPC_23_rev%2020.pdf. Acesso em: 12 mar. 2025.

Conselho Federal de Contabilidade (CFC). **NBC TA 540 (R2): Auditoria de Estimativas Contábeis, Inclusive do Valor Justo, e Divulgações Relacionadas**. Brasília, 2019. Disponível em: [https://www1.cfc.org.br/sisweb/SRE/docs/NBCTA540\(R2\).pdf](https://www1.cfc.org.br/sisweb/SRE/docs/NBCTA540(R2).pdf). Acesso em: 12 mar. 2025.

Conselho Federal de Contabilidade (CFC). **NBC TG Estrutura Conceitual: Estrutura conceitual para elaboração e divulgação de relatório contábil-financeiro**. Brasília, 2019. Disponível em: <https://www1.cfc.org.br/sisweb/SRE/docs/NBCTGEC.pdf>. Acesso em: 10 mar. 2025.

COVERT, Ian; LUNDBERG, Scott M.; LEE, Su-In. Understanding global feature contributions with additive importance measures. **Advances in Neural Information Processing Systems**, v. 33, p. 17212-17223, 2020. Disponível em: <http://dx.doi.org/10.48550/arXiv.2004.00668>. Acesso em: 26 jan. 2025.

CREPALDI, S. A.; CREPALDI, G. S. **Auditoria contábil**. Grupo Gen-Atlas, 2016.

DASH, Sushree Sasmita; NAYAK, Subrat Kumar; MISHRA, Debahuti. A review on machine learning algorithms. **Intelligent and Cloud Computing: Proceedings of ICICC 2019**, Volume 2, p. 495-507, 2020. Disponível em: http://doi.org/10.1007/978-981-15-6202-0_51. Acesso em: 2 fev. 2025.

DE ARAUJO, Marcelo Henrique; CORNACCHIONE, Edgard. Reflexões sobre o uso de inteligência artificial na contabilidade gerencial: oportunidades, desafios e riscos. **Revista de Contabilidade e Organizações**, v. 18, p. e231688-e231688, 2024. Disponível em: <https://www.revistas.usp.br/rco/article/download/231688/210408>. Acesso em: 7 mai. 2025.

DE IUDÍCIBUS, Sérgio. **Manual de contabilidade societária: aplicável a todas as sociedades, de acordo com as normas internacionais e do CPC**. Ed. Atlas, 2018.

DECHOW, Niels; GRANLUND, Markus; MOURITSEN, Jan. Management control of the complex organization: relationships between management accounting and information technology. **Handbooks of management accounting research**, v. 2, p. 625-640, 2006. Disponível em: [https://doi.org/10.1016/S1751-3243\(06\)02007-4](https://doi.org/10.1016/S1751-3243(06)02007-4). Acesso em: 08 jan. 2025.

DECHOW, Patricia M.; DICHEV, Ilia D. The quality of accruals and earnings: The role of accrual estimation errors. **The accounting review**, v. 77, n. s-1, p. 35-59, 2002. Disponível em: <https://doi.org/10.2308/accr.2002.77.s-1.35>. Acesso em: 29 mar. 2025.

DELOITTE. **Becoming an AI-fueled organization: Deloitte's State of AI in the Enterprise**. Disponível em: <https://www.deloitte.com/au/en/services/consulting/perspectives/becoming-ai-fueled-organisation.html>. Acesso em: 10 jan. 2025.

DING, Kexing et al. Machine learning improves accounting estimates: Evidence from insurance payments. **Review of accounting studies**, v. 25, n. 3, p. 1098-1134, 2020. Disponível em: <https://papers.ssrn.com/sol3/Delivery.cfm?abstractid=3253220>. Acesso em: 17 mar. 2025.

EMMANUEL, T. *et al.* A survey on missing data in machine learning. **Journal of Big Data**, [S.L], v. 8, n. 1, p. 1-37, out./2021. Disponível em: <https://doi.org/10.1186/s40537-021-00516-9>. Acesso em: 20 jan. 2025.

EY. **Audit innovation**. Disponível em: https://www.ey.com/en_gl/services/audit/innovation. Acesso em: 11 jan. 2025.

FAWCETT, Tom; PROVOST, Foster. Data science and its relationship to big data and data-driven decision making. **Big data**, v. 1, n. 1, p. 51-59, 2013. Disponível em: <https://www.liebertpub.com/doi/pdf/10.1089/big.2013.1508>. Acesso em: 20 dez. 2024.

FISCHER, Alice E.; GRODZINSKY, Frances S. **The Anatomy Of Programming Languages**. 1. ed. Nova Jersey: Prentice Hall, 1992. p. 3-5. Acesso 18 jan. 2025.

FISCHER, T., KRAUSS, C. Deep learning with long short-term memory networks for financial market predictions. **European journal of operational research**, v. 270(2), p. 654-669, 2018. Disponível em: <https://doi.org/10.1016/j.ejor.2017.11.054>. Acesso em: 12 fev. 2025.

FREUND, Yoav; SCHAPIRE, Robert E. Experiments with a new boosting algorithm. **icml**. 1996. p. 148-156. Disponível em: <https://cseweb.ucsd.edu/~yfreund/papers/boostingexperiments.pdf>. Acesso em: 2 fev. 2025.

FREY, Carl Benedikt; OSBORNE, Michael A. The future of employment: How susceptible are jobs to computerisation?. **Technological forecasting and social change**, v. 114, p. 254-280, 2017. Disponível em: <https://doi.org/10.1016/j.techfore.2016.08.019>. 15 mai. 2025.

FRIEDMAN, Jerome H. Stochastic gradient boosting. **Computational statistics & data analysis**, v. 38, n. 4, p. 367-378, 2002. Disponível em: [https://doi.org/10.1016/S0167-9473\(01\)00065-2](https://doi.org/10.1016/S0167-9473(01)00065-2). Acesso em: 3 fev. 2025.

GAO, Hanyao et al. Machine learning in business and finance: a literature review and research opportunities. **Financial Innovation**, v. 10, n. 1, p. 86, 2024. Disponível em: <http://doi.org/10.1186/s40854-024-00629-z>. Acesso em: 10 fev. 2025.

GAVER, Jennifer J.; PATERSON, Jeffrey S. Do insurers manipulate loss reserves to mask solvency problems? **Journal of Accounting and Economics**, v. 37, n. 3, p. 393-416, 2004. Disponível em: <https://doi.org/10.1016/j.jacceco.2003.10.010>. Acesso em: 13 mar. 2025.

GHODDUSI, Hamed; CREAMER, Germán G.; RAFIZADEH, Nima. Machine learning in energy economics and finance: A review. **Energy Economics**, v. 81, p. 709-727, 2019. Disponível em: <https://doi.org/10.1016/j.eneco.2019.05.006>. Acesso em: 11 fev. 2025.

GOGAS, Periklis; PAPADIMITRIOU, Theophilos. Machine learning in economics and finance. **Computational Economics**, v. 57, p. 1-4, 2021. Disponível em: <http://doi.org/10.1007/S10614-021-10094-W>. Acesso em: 10 fev. 2025.

GUO, Li; SHI, Feng; TU, Jun. Textual analysis and machine leaning: Crack unstructured data in finance and accounting. **The Journal of Finance and Data Science**, v. 2, n. 3, p. 153-170, 2016. Disponível em: <https://doi.org/10.1016/j.jfds.2017.02.001>. Acesso em: 15 abr. 2025.

HAJEK, Petr; HENRIQUES, Roberto. Mining corporate annual reports for intelligent detection of financial statement fraud—A comparative study of machine learning

methods. **Knowledge-Based Systems**, v. 128, p. 139-152, 2017. Disponível em: <https://doi.org/10.1016/j.knosys.2017.05.001>. Acesso em: 9 abr. 2025.

Hashem, I. A. T., Yaqoob, I., Anuar, N. B., Mokhtar, S., Gani, A., e Khan, S. U. The rise of “Big Data” on cloud computing: Review and open research issues. **Information systems**, 47, p. 98-115, 2015. Disponível em: <http://dx.doi.org/10.1016/j.is.2014.07.006>. Acesso em: 20 dez. 2024.

HASTIE, Trevor; TIBSHIRANI, Robert; FRIEDMAN, Jerome. **The Elements of Statistical Learning: Data Mining, Inference, and Prediction**. 2. ed. [S.L]: Springer, 2009. p.30-31, 260, 607. Acesso em: 22 jan. 2025.

HERRERA, Salvador Robles; CEBERIO, Martine; KREINOVICH, Vladik. When is deep learning better and when is shallow learning better: qualitative analysis. *International Journal of Parallel, Emergent and Distributed Systems*, v. 37, n. 5, p. 589-595, 2022. Disponível em: <https://doi.org/10.1080/17445760.2022.2070748>. Acesso em: 26 jan. 2025.

HOANG, Daniel; WIEGRATZ, Kevin. Machine learning methods in finance: Recent applications and prospects. **European Financial Management**, v. 29, n. 5, p. 1657-1701, 2023. Disponível em: <https://onlinelibrary.wiley.com/doi/full/10.1111/eufm.12408>. Acesso em: 25 jan. 2025.

HOODA, Nishtha; BAWA, Seema; RANA, Prashant Singh. Fraudulent firm classification: a case study of an external audit. **Applied Artificial Intelligence**, v. 32, n. 1, p. 48-64, 2018. Disponível em: <https://doi.org/10.1080/08839514.2018.1451032>. Acesso em: 8 abr. 2025.

Hooda, Nishtha. Audit Data Dataset. UCI Machine Learning Repository. 2018. <https://doi.org/10.24432/C5930Q>.

HOSAKA, Tadaaki. Bankruptcy prediction using imaged financial ratios and convolutional neural networks. **Expert systems with applications**, v. 117, p. 287-299, 2019. Disponível em: <http://www.isc.meiji.ac.jp/~hosaka/img/file.pdf>. Acesso em: 21 fev. 2025.

HOU, Xuechen. Design and application of intelligent financial accounting model based on knowledge graph. **Mobile Information Systems**, v. 2022, n. 1, p. 8353937, 2022. Disponível em: <https://onlinelibrary.wiley.com/doi/pdf/10.1155/2022/8353937>. Acesso em: 12 mai. 2025.

HU, Hanxin; SUN, Ting. The applications of machine learning in accounting and auditing research. **Encyclopedia of finance**, 2021. p. 1-21. Disponível em: https://link.springer.com/rwe/10.1007/978-3-030-73443-5_91-1. Acesso em: 9 mar. 2025.

IBM. **O que é aprendizado de máquina (ML)?** 2021. Disponível em: <https://www.ibm.com/br-pt/topics/machine-learning>. Acesso em: 13 jan. 2025.

IBRAHIM, A. E. A; ELAMER, Ahmed A.; EZAT, Amr Nazieh. The convergence of big data and accounting: innovative research opportunities. **Technological Forecasting and Social Change**, v. 137, n. 1, dez. 2021. Disponível em: <https://pure.port.ac.uk/ws/portalfiles/portal/42853653/Ibrahim.pdf>. Acesso em: 19 dez. 2024.

IFAC. **Why Accountants Must Embrace Machine Learning**. Disponível em: <https://www.ifac.org/knowledge-gateway/discussion/why-accountants-must-embrace-machine-learning>. Acesso em: 8 jan. 2025.

ISLAM, Md Saiful; HOSSAIN, Emam. Foreign exchange currency rate prediction using a GRU-LSTM hybrid network. **Soft Computing Letters**, v. 3, p. 100009, 2021. Disponível em: <https://doi.org/10.1016/j.socl.2020.100009>. Acesso em: 15 fev. 2025.

JABBAR, Haider Khalaf; KHAN, Rafiqul Zaman. Methods to avoid over-fitting and under-fitting in supervised machine learning (comparative study). **Computer Science, Communication and Instrumentation Devices**, India, v. 70, n. 10.3850, p. 978-981, dez./2014. Disponível em: https://doi.org/10.3850/978-981-09-5247-1_017. Acesso em: 20 jan. 2025.

JONES, Stewart. Corporate bankruptcy prediction: a high dimensional analysis. **Review of Accounting Studies**, v. 22, p. 1366-1422, 2017. Disponível em: <http://doi.org/10.1007/s11142-017-9407-1>. Acesso em: 17 jan. 2025.

JOSEPH, V. Roshan. Optimal ratio for data splitting. **Statistical Analysis and Data Mining**, N.Y., v. 15, n. 4, p. 531-538, abr./2022. Disponível em: <https://doi.org/10.1002/sam.11583>. Acesso em: 24 jan. 2025.

JOVIC, Alan; BRKIC, Karla; BOGUNOVIC, Nikola. A review of feature selection methods with applications. **International Convention on Information and Communication Technology, Electronics and Microelectronics (MIPRO)**, Croácia, v. 38, n. 1, p. 1200-1205, jul./2015. Disponível em: <https://ieeexplore.ieee.org/abstract/document/7160458/>. Acesso em: 20 jan. 2025.

KIRKOS, Efstathios; SPATHIS, Charalambos; MANOLOPOULOS, Yannis. Data mining techniques for the detection of fraudulent financial statements. **Expert systems with applications**, v. 32, n. 4, p. 995-1003, 2007. Disponível em: <http://delab.csd.auth.gr/papers/ESWA07ksm.pdf>. Acesso em: 3 abr. 2025

KORATAMADDI, Prahlad et al. Market sentiment-aware deep reinforcement learning approach for stock portfolio allocation. **Engineering Science and Technology, an International Journal**, v. 24, n. 4, p. 848-859, 2021. Disponível em: <https://doi.org/10.1016/j.jestch.2021.01.007>. Acesso em: 14 fev. 2025.

KOTSIANTIS, Sotiris et al. Predicting fraudulent financial statements with machine learning techniques. **Hellenic Conference on Artificial Intelligence**. Berlin, Heidelberg: Springer Berlin Heidelberg, 2006. p. 538-542. Disponível em: https://link.springer.com/chapter/10.1007/11752912_63. Acesso em: 6 abr. 2025.

LAMBERT, Richard; LEUZ, Christian; VERRECCHIA, Robert E. Accounting information, disclosure, and the cost of capital. **Journal of accounting research**, v. 45, n. 2, p. 385-420, 2007. Disponível em: <https://doi.org/10.1111/j.1475-679X.2007.00238.x>. Acesso em: 10 mar. 2025.

Lashine, S.H.; Mohamed, E.K.A. (2003), "Accounting knowledge and skills and the challenges of a global business environment", **Managerial Finance**, Vol. 29 No. 7, pp. 3-16. <https://doi.org/10.1108/03074350310768319>. Acesso em: 20 dez. 2024.

LECUN, Yann; BENGIO, Yoshua; HINTON, Geoffrey. Deep learning. **Nature**, v. 521, n. 7553, p. 436-444, mai./2015. Disponível em: <https://hal.science/hal-04206682/document>. Acesso em: 12 jan. 2025.

LEWELLEN, Jonathan; RESUTEK, Robert J. Why do accruals predict earnings? **Journal of Accounting and Economics**, v. 67, n. 2-3, p. 336-356, 2019. Acesso em: 29 mar. 2025.

LI, Feng et al. Textual analysis of corporate disclosures: A survey of the literature. **Journal of accounting literature**, v. 29, n. 1, p. 143-165, 2010. Disponível em: <https://www.cuhk.edu.hk/acy2/workshop/20110215FengLI/Paper1.pdf>. Acesso em: 23 abr. 2025.

LI, Yelin et al. The role of text-extracted investor sentiment in Chinese stock price prediction with the enhancement of deep learning. **International Journal of Forecasting**, v. 36(4), 2020. Disponível em: <https://doi.org/10.1016/j.ijforecast.2020.05.001>. Acesso em: 13 fev. 2025.

LI, Yu. Credit risk prediction based on machine learning methods. **2019 14th international conference on computer science & education (ICCSE)**. IEEE, 2019. p. 1011-1013. Disponível em: <https://doi.org/10.1109/ICCSE.2019.8845444>. Acesso em: 28 fev. 2025.

LIMA, Lemonier Barbosa de. O uso de técnicas de Machine Learning para melhorar a prevenção à fraude. 2022. Disponível em: <https://repositorio.unb.br/handle/10482/44797>. Acesso em: 8 mar. 2025.

LIU, Yi et al. Applying machine learning algorithms to predict default probability in the online credit market: Evidence from China. **International Review of Financial Analysis**, v. 79, p. 101971, 2022. Disponível em: <https://doi.org/10.1016/j.irfa.2021.101971>. Acesso em: 27 fev. 2025.

LUNDBERG, Scott M. et al. From local explanations to global understanding with explainable AI for trees. **Nature machine intelligence**, v. 2, n. 1, p. 56-67, 2020. Disponível em: <https://www.nature.com/articles/s42256-019-0138-9>. Acesso em: 26 jan. 2025.

LUNDBERG, Scott M.; LEE, Su-In. A unified approach to interpreting model predictions. **Advances in neural information processing systems**, v. 30, 2017. Disponível em: https://papers.nips.cc/paper_files/paper/2017/file/8a20a8621978632d76c43dfd28b67767-Paper.pdf. Acesso em: 18 jul. 2025.

MA, Yilin; HAN, Ruizhu; WANG, Weizhong. Portfolio optimization with return prediction using deep learning and machine learning. **Expert Systems with Applications**, v. 165, p. 113973, 2021. Disponível em: <https://doi.org/10.1016/j.eswa.2020.113973>. Acesso em: 13 fev. 2025.

MACIEIRA, F. C. Além da Fiscalização Tradicional: Utilizando Inteligência Artificial para

Seleção de Amostra de Auditoria. 2025. Disponível em: <https://hdl.handle.net/10438/37047>. Acesso em: 30 jul. 2025.

MALANDRI, Lorenzo et al. Public mood–driven asset allocation: The importance of financial sentiment in portfolio management. **Cognitive Computation**, v. 10, n. 6, p. 1167-1176, 2018. Disponível em: <https://doi.org/10.1007/s12559-018-9609-2>. Acesso em: 15 fev. 2025.

MARTINS, Pablo Luiz *et al.* Tecnologia e sistemas de informação e suas influências na gestão e contabilidade. IX SEGeT, 2012. Disponível em: <https://www.aedb.br/seget/arquivos/artigos12/28816533.pdf>. Acessado em: 08 jan. 2025.

MATHUR, Gauri. Data science vs data analytics: Unpacking the differences. **IBM**. 2023. Disponível em: <https://www.ibm.com/think/topics/data-science-vs-data-analytics>. Acesso em: 19 dez. 2024.

MATIN, Rastin et al. Predicting distresses using deep learning of text segments in annual reports. **Expert systems with applications**, v. 132, p. 199-208, 2019. Disponível em: <https://arxiv.org/pdf/1811.05270>. Acesso em: 25 fev. 2025.

MAYR, Andreas et al. The evolution of boosting algorithms. **Methods of information in medicine**, v. 53, n. 06, p. 419-427, 2014. Disponível em: <https://arxiv.org/pdf/1403.1452>. Acesso em: 2 fev. 2025.

MEIR, Yuval et al. Efficient shallow learning as an alternative to deep learning. **Scientific Reports**, v. 13, n. 1, p. 5423, 2023. Disponível em: <https://www.nature.com/articles/s41598-023-32559-8.pdf>. Acesso em: 26 jan. 2025.

MIT SLOAN. **Machine learning, explained**. Disponível em: <https://mitsloan.mit.edu/ideas-made-to-matter/machine-learning-explained>. Acesso em: 19 dez. 2024.

MOGHADDAM, Arya Hadizadeh; MOMTAZI, Saeedeh. Image processing meets time series analysis: Predicting Forex profitable technical pattern positions. **Applied Soft Computing**, v. 108, p. 107460, 2021. Disponível em: <https://doi.org/10.1016/j.asoc.2021.107460>. Acesso em: 15 fev. 2025.

MOLNAR, Christoph. **Interpretable machine learning**. Lulu. com, 2020. p. 10; 241.

MONARD, Maria Carolina; BARANAUSKAS, José Augusto. Conceitos sobre aprendizado de máquina. **Sistemas inteligentes-Fundamentos e aplicações**, v. 1, n. 1, p. 32, 2003.

MUCCI, Tim. Overfitting vs. underfitting. **IBM**, 2024. Disponível em: <https://www.ibm.com/think/topics/overfitting-vs-underfitting>. Acesso em: 20 jan. 2025.

MÜLLER, Andreas C.; GUIDO, Sarah. **Introduction to machine learning with Python: a guide for data scientists**. " O'Reilly Media, Inc.", 2016. p.35, 73-88.

NASTESKI, Vladimir. An overview of the supervised machine learning methods. **HORIZONS B**, v. 4, n. 1, p. 51-62, dez./2017. Disponível em:

<http://dx.doi.org/10.20544/HORIZONS.B.04.1.17.P05>. Acesso em: 14 jan. 2025.

NAZARETH, Noella; REDDY, Y. V. R. Financial applications of machine learning: A literature review. **Expert Systems with Applications**, v. 213, n. 1, p. 1-33, jun./2023. Disponível em: <https://doi.org/10.1016/j.eswa.2023.119640>. Acesso em: 16 jan. 2025.

NETFLIX RESEARCH. **Machine Learning Platform. 2023**. Disponível em: <https://research.netflix.com/research-area/machine-learning-platform>. Acesso em: 20 fev. 2025.

NIELSEN, Didrik. **Tree boosting with xgboost-why does xgboost win" every" machine learning competition?** 2016. Dissertação de Mestrado. NTNU. Disponível em: <https://pzs.dstu.dp.ua/DataMining/boosting/bibl/Didrik.pdf>. Acesso em: 3 fev. 2025.

OUYANG, Zi-sheng; YANG, Xi-te; LAI, Yongzeng. Systemic financial risk early warning of financial market in China using Attention-LSTM model. **The North American Journal of Economics and Finance**, v. 56, p. 101383, 2021. Disponível em: <https://doi.org/10.1016/j.najef.2021.101383>. Acesso em: 17 fev. 2025.

OZBAYOGLU, A. M., GUDELEK, M. U., & SEZER, O. B. (2020). Deep learning for financial applications: A survey. **Applied Soft Computing Journal**. Vol. 1, No. 45-56, p. 228. Disponível em: <https://arxiv.org/pdf/2002.05786>. 12 fev. 2025.

PAIVA, Felipe Dias et al. Decision-making for financial trading: A fusion approach of machine learning and portfolio selection. **Expert systems with applications**, v. 115, p. 635-655, 2019. Disponível em: <https://doi.org/10.1016/j.eswa.2018.08.003>. Acesso em: 13 fev. 2025.

PASUPA, Kitsuchart; SUNHEM, Wisuwat. A comparison between shallow and deep architecture classifiers on small dataset. **2016 8th International Conference on Information Technology and Electrical Engineering (ICITEE)**. IEEE, 2016. p. 1-6. Disponível em: <https://doi.org/10.1109/ICITEED.2016.7863293>. Acesso em: 26 jan. 2025.

PAULA, Daniel Abreu Vasconcellos de et al. Estimating credit and profit scoring of a Brazilian credit union with logistic regression and machine-learning techniques. **RAUSP Management Journal**, v. 54, p. 321-336, 2019. Disponível em: <https://www.scielo.br/j/rmj/a/MQSND9HPXDqpRdm7QTswnxm/>. Acesso em: 7 mar. 2025.

PEROLS, Johan. Financial statement fraud detection: An analysis of statistical and machine learning algorithms. Auditing: **A Journal of Practice & Theory**, v. 30, n. 2, p. 19-50, 2011. Disponível em: <https://doi.org/10.2308/ajpt-50009>. Acesso em: 8 mar. 2025.

PETROPOULOS, Anastasios et al. Predicting bank insolvencies using machine learning techniques. **International Journal of Forecasting**, v. 36, n. 3, p. 1092-1113, 2020. Disponível em: <https://doi.org/10.1016/j.ijforecast.2019.11.005>. Acesso em: 25 fev. 2025.

PICASSO, Andrea et al. Technical analysis and sentiment embeddings for market trend prediction. **Expert Systems with Applications**, v. 135, p. 60-70, 2019. Disponível em: <https://sentic.net/market-trend-prediction.pdf>. Acesso em: 15 fev. 2025.

PRINCE, Simon J.d.. **Understanding Deep Learning**. 1. ed. EUA: The MIT Press, 2023. p. 1-22.

PROBST, Philipp; WRIGHT, Marvin N.; BOULESTEIX, Anne-Laure. Hyperparameters and tuning strategies for random forest. **Wiley Interdisciplinary Reviews: data mining and knowledge discovery**, v. 9, n. 3, p. e1301, 2019. Disponível em: <https://arxiv.org/pdf/1804.03515>. Acesso em: 1 fev. 2025.

RANTA, Mikko; YLINEN, Mika; JÄRVENPÄÄ, Marko. Machine Learning in Management Accounting Research: Literature Review and Pathways for the Future. **European Accounting Review**, EUR, v. 32, n. 3, p. 607-636, nov. 2022. Disponível em: <https://www.tandfonline.com/doi/pdf/10.1080/09638180.2022.2137221>. Acesso em: 14 jan. 2025.

REIS, Carolina et al. Assessing the drivers of machine learning business value. **Journal of Business Research**, v. 117, p. 232-243, 2020. Disponível em: <https://doi.org/10.1016/j.jbusres.2020.05.053>. Acesso em: 23 abr. 2025.

REZAEI, Zabihollah et al. Continuous auditing: Building automated auditing capability. **Auditing: A Journal of Practice & Theory**, v. 21, n. 1, p. 147-163, 2002. Disponível em: <https://doi.org/10.2308/aud.2002.21.1.147>. Acesso em: 3 abr. 2025.

RIDZUAN, Fakhithah; ZAINON, W. M. N. W. A Review on Data Cleansing Methods for Big Data. **Procedia Computer Science**, Malásia, v. 161, p. 731-738, nov./2019. Disponível em: <https://doi.org/10.1016/j.procs.2019.11.177>. Acesso em: 17 jan. 2025.

SAMITAS, Aristeidis; KAMPOURIS, Elias; KENOURGIOS, Dimitris. Machine learning as an early warning system to predict financial crisis. **International Review of Financial Analysis**, v. 71, p. 101507, 2020. Disponível em: <https://doi.org/10.1016/j.irfa.2020.101507>. Acesso em: 16 fev. 2025.

SAMUEL, Arthur Lee. Some Studies in Machine Learning Using the Game of Checkers. **IBM Journal of Research and Development**, EUA, v. 3, n. 3, p. 210-229, jul./1959. Disponível em: <http://www.cs.virginia.edu/~evans/greatworks/samuel1959.pdf>. Acesso em: 10 jan. 2025.

SARKER, Iqbal H. Machine learning: Algorithms, real-world applications and research directions. **SN computer science**, v. 2, n. 3, p. 1-21, mar./2021. Disponível em: <https://link.springer.com/content/pdf/10.1007/s42979-021-00592-x.pdf>. Acesso em: 20 fev. 2024.

SEKAR, Maris. **Machine learning for auditors**. Apress, 2022. p. 91.

SEMIROMI, Hamed Naderi; LESSMANN, Stefan; PETERS, Wiebke. News will tell: Forecasting foreign exchange rates based on news story events in the economy calendar. **The North American Journal of Economics and Finance**, v. 52, p. 101181, 2020. Disponível em: <https://doi.org/10.1016/j.najef.2020.101181>. Acesso em: 16 fev. 2025.

SHEHAB, M. et al. Machine learning in medical applications: A review of state-of-the-art methods. **Computers in Biology and Medicine**, Amst, v. 145, n. 2022, p. 1-38, jun./2022. Disponível em: <https://doi.org/10.1016/j.combiomed.2022.105458>. Acesso em: 20 fev. 2024.

SHI, Yanling. The impact of artificial intelligence on the accounting industry. **Cyber Security Intelligence and Analytics**. Springer International Publishing, 2020. p.971-978. Disponível em: http://doi.org/10.1007/978-3-030-15235-2_129. Acesso em: 5 mai. 2025.

SHMUELI, Galit. To explain or to predict? **Statistical Science**, v. 25, n. 3, p. 289–310, 2010. Disponível em: <https://doi.org/10.1214/10-STS330>. Acesso em: 24 abr. 2025.

SHRESTHA, Yash Raj; BEN-MENACHEM, Shiko M.; VON KROGH, Georg. Organizational decision-making structures in the age of artificial intelligence. **California management review**, v. 61, n. 4, p. 66-83, 2019. Disponível em: <https://journals.sagepub.com/doi/abs/10.1177/0008125619862257>. Acesso em: 15 jun. 2025.

ŚLUSARCZYK, Beata. Industry 4.0: Are we ready? **Polish Journal of Management Studies**, v. 17, 2018. Disponível em: <https://bibliotekanauki.pl/articles/406054.pdf>. Acesso em: 25 abr. 2025.

SONG, Xin-Ping et al. Application of machine learning methods to risk assessment of financial statement fraud: Evidence from China. **Journal of Forecasting**, v. 33, n. 8, p. 611-626, 2014. Disponível em: <https://doi.org/10.1002/for.2294>. Acesso em: 8 mar. 2025.

TAN, Zheng; YAN, Ziqin; ZHU, Guangwei. Stock selection with random forest: An exploitation of excess return in the Chinese stock market. **Heliyon**, v. 5, n. 8, 2019. Disponível: [https://www.cell.com/heliyon/pdf/S2405-8440\(19\)35970-5.pdf](https://www.cell.com/heliyon/pdf/S2405-8440(19)35970-5.pdf). Acesso em: 13 fev. 2025.

THEOBALD, Oliver. **Machine Learning For Absolute Beginners: A Plain English Introduction**. 2. ed. Scatterplot Press, 2017. p.1-168.

THEODORAKOPOULOS, Leonidas; THANASAS, Georgios; HALKIOPOULOS, Constantinos. Implications of Big Data in Accounting: Challenges and Opportunities. **Emerging Science Journal**, Grécia, v. 8, n. 3, p. 1-14, jun./2024. Disponível em: <https://www.ijournalse.org/index.php/ESJ/article/download/2228/pdf>. Acesso em: 19 dez. 2024.

THOMSON REUTERS. **2024 generative AI in professional services**. 2024b. Disponível em: https://www.thomsonreuters.com/content/dam/ewp-m/documents/thomsonreuters/en/pdf/reports/tr4322226_rgb.pdf. Acesso em: 10 jan. 2025.

THOMSON REUTERS. **How do different accounting firms use AI?** 2024a. Disponível em: <https://tax.thomsonreuters.com/blog/how-do-different-accounting-firms-use-ai/>. Acesso em: 9 jan. 25.

TSAI, Chih-Fong; HSU, Yu-Feng; YEN, David C. A comparative study of classifier ensembles for bankruptcy prediction. **Applied Soft Computing**, v. 24, p. 977-984, 2014.

Disponível em: <https://doi.org/10.1016/j.asoc.2014.08.047>. Acesso em: 7 mar. 2025.

UCOGLU, Derya. Current machine learning applications in accounting and auditing. **PressAcademia Procedia**, EUA, v. 12, n. 1, p. 1-7, dez./2020. Disponível em: <https://doi.org/10.17261/Pressacademia.2020.1337>. Acesso em: 11 jan. 2025.

VEGANZONES, David; SÉVERIN, Eric. An investigation of bankruptcy prediction in imbalanced datasets. **Decision Support Systems**, v. 112, p. 111-124, 2018. Disponível em: <https://doi.org/10.1016/j.dss.2018.06.011>. Acesso em: 24 fev. 2025.

VEZHNEVETS, Alexander; BARINOVA, Olga. Avoiding boosting overfitting by removing confusing samples. **European Conference on Machine Learning**. Berlin, Heidelberg: Springer Berlin Heidelberg, 2007. p.430-441. Disponível em: http://doi.org/10.1007/978-3-540-74958-5_40. Acesso em: 4 fev. 2025.

WANG, Sun-Chong. Artificial neural network. **Interdisciplinary computing in java programming**. Boston, MA: Springer US, 2003. p. 81-100. Disponível em: https://doi.org/10.1007/978-1-4615-0377-4_5. Acesso em: 6 fev. 2025.

WECKS, Janik Ole; VOSHAAR, Johannes; ZIMMERMANN, Jochen. Using Machine Learning to Address Individual Learning Needs in Accounting Education. **Available at SSRN 4648223**, 2023. Disponível em: <http://dx.doi.org/10.2139/ssrn.4648223>. Acesso em: 18 jun. 2025.

WENG, B., AHMED, M. A., MEGAHED, F. M. Stock market one-day ahead movement prediction using disparate data sources. **Expert Systems with Applications**, v. 79, p. 153-163, 2017. Disponível em: <https://doi.org/10.1016/j.eswa.2017.02.041>. Acesso em: 12 fev. 2025.

WHITE, Halbert. Economic prediction using neural networks: The case of IBM daily stock returns. **ICNN**. 1988. v. 2, p. 451-458, 1988. Disponível em: <https://doi.org/10.1109/ICNN.1988.23959>. Acesso em: 10 fev. 2025.

WRIGHT, Scott A.; SCHULTZ, Ainslie E. The rising tide of artificial intelligence and business automation: Developing an ethical framework. **Business Horizons**, v. 61, n. 6, p. 823-832, 2018. Disponível em: <https://doi.org/10.1016/j.bushor.2018.07.001>. Acesso em: 25 abr. 2025

WU, Desheng; MA, Xiyuan; OLSON, David L. Financial distress prediction using integrated Z-score and multilayer perceptron neural networks. **Decision Support Systems**, v. 159, p. 113814, 2022. Disponível em: <https://doi.org/10.1016/j.dss.2022.113814>. Acesso em: 26 fev. 2025.

XAVIER, L. M.; CARRARO, W. B. W. H.; RODRIGUES, A. T. L. Indústria 4.0 E Avanços Tecnológicos da Área Contábil: Perfil, Percepções e Expectativas Dos Profissionais. **ConTexto - Contabilidade em Texto, Porto Alegre**, v. 20, n. 45, 2020. Disponível em: <https://seer.ufrgs.br/index.php/ConTexto/article/view/97774>. Acesso em: 19 dez. 2024.

YAN, Hongju; OUYANG, Hongbing. Financial time series prediction based on deep learning. **Wireless Personal Communications**, v. 102, p. 683-700, 2018. Disponível em: <https://link.springer.com/article/10.1007/s11277-017-5086-2>. Acesso em: 13 fev. 2025.

YAO, Jianrong; ZHANG, Jie; WANG, Lu. A financial statement fraud detection model based on hybrid data mining methods. **2018 international conference on artificial intelligence and big data (ICAIBD)**. IEEE, 2018. p. 57-61. Disponível em: <https://doi.org/10.1109/ICAIBD.2018.8396167>. Acesso em: 8 mar. 2025.

YING, Xue. An Overview of Overfitting and its Solutions. **Journal of Physics**, [S.L], v. 1168, n. 2, p. 22022, fev./2019. Disponível em: <https://doi.org/10.1088/1742-6596/1168/2/022022>. Acesso em: 20 jan. 2025.

YU, Xiaojiao. Machine learning application in online lending risk prediction. **arXiv preprint arXiv:1707.04831**, 2017. Disponível em: <https://arxiv.org/pdf/1707.04831>. Acesso em: 28 fev. 2025.

ZHANG, Yingying et al. The impact of artificial intelligence and blockchain on the accounting profession. **Ieee Access**, v. 8, p. 110461-110477, 2020. Disponível em: <https://ieeexplore.ieee.org/iel7/6287639/8948470/09110603.pdf>. Acesso em: 10 mai. 2025.

ZORIČÁK, Martin et al. Bankruptcy prediction for small-and medium-sized companies using severely imbalanced datasets. **Economic Modelling**, v. 84, p. 165-176, 2020. Disponível em: <https://doi.org/10.1016/j.econmod.2019.04.003>. Acesso em: 25 fev. 2025.

APÊNDICE A – INADIMPLÊNCIA - ALGORITMO BASE PRINCIPAL (RF)

```

import pandas as pd
from sklearn.model_selection import (train_test_split, GridSearchCV)
from sklearn.ensemble import RandomForestClassifier
from sklearn.metrics import (
    accuracy_score, precision_score, recall_score, f1_score,
    roc_auc_score, confusion_matrix, ConfusionMatrixDisplay)
import matplotlib.pyplot as plt
import seaborn as sns
import shap

# CARREGAMENTO DOS DADOS
df = pd.read_csv("DATASETS/IBM Late Payment Histories/WA_Fn-UseC_-Accounts-
Receivable.csv")
# Features originais do conjunto de dados
print("=== Formação original do conjunto === \n", df.columns)

# PREPROCESSAMENTO + FEATURE ENGINEERING
# Criar variável (Y) binária.
df["Late"] = (df["DaysLate"] > 0).astype(int) # feature criada a partir de dias em atraso
# Transformar features temporais em data object
df["InvoiceDate"] = pd.to_datetime(df["InvoiceDate"])
df["DueDate"] = pd.to_datetime(df["DueDate"])
df["SettledDate"] = pd.to_datetime(df["SettledDate"])
# Criar features específicas derivadas e converter features categóricas
df["InvoiceDay"] = df["InvoiceDate"].dt.day
df["InvoiceMonth"] = df["InvoiceDate"].dt.month
df["InvoiceYear"] = df["InvoiceDate"].dt.year
df["DueDateDay"] = df["DueDate"].dt.day
df["DueDateMonth"] = df["DueDate"].dt.month
df["DueDateYear"] = df["DueDate"].dt.year
df["PaperlessBill"] = pd.Categorical(df["PaperlessBill"], categories=["Paper", "Electronic"],
ordered=True)

```

```

# Binnary encoding
df = pd.get_dummies(df, columns=["Disputed", "PaperlessBill"], drop_first=True)
df["countryRaw"] = df["countryCode"]

# One hot encoding
df = pd.get_dummies(df, columns=["countryCode"])
df.rename(columns={
    "countryCode_391": "Argentina", "countryCode_406": "Brasil", "countryCode_770":
    "Uruguai",
    "countryCode_818": "Paraguai", "countryCode_897": "Mexico"
}, inplace=True)

print("\n=== Novo conjunto ===\n", "\n".join(df.columns))

# Matriz de correlação inicial das features
plt.figure(figsize=(16, 12))
sns.heatmap(df.corr(numeric_only=True), annot=True, cmap="coolwarm", linewidths=0.5)
plt.title("Matriz de Correlação das Features")
plt.tight_layout()
plt.show()

# Novas features criadas
# Feat. referente ao total da fatura disputada
df["DisputedAmount"] = df["InvoiceAmount"] * df["Disputed_Yes"]
# Feat. referente a quantia média devida por cliente
AvgAmount_Customer = df.groupby("customerID")["InvoiceAmount"].mean()
df["AvgAmount_byCustomer"] = df["customerID"].map(AvgAmount_Customer)
# Feat. referente a proporção de faturas disputadas por cliente
dispute_rate = df.groupby("customerID")["Disputed_Yes"].mean()
df["DisputeRate_byCustomer"] = df["customerID"].map(dispute_rate)
# Total de faturas por devedor
invoice_totals = df.groupby("customerID").size()
df["TotalInvoices_byCustomer"] = df["customerID"].map(invoice_totals)
print("\n=== dataset + features criadas ante o treino === \n", df.columns)

```

```

# Definir features
features= df.drop(columns=[
    "invoiceNumber", "DueDate", "Late", "PaperlessDate"
]).columns.tolist()

# — Seleção e encode de features
X = df[features].copy()
y = df["Late"]

# Amostra de treino e teste
X_train, X_test, y_train, y_test = train_test_split(
    X, y, test_size=0.2, random_state=42, stratify=y)

# Criar dataframe a partir do subconjunto de treino
train_df = df.loc[X_train.index]

# Calcular estatísticas no treino e Z-score da média de cada país
df["countryRaw"].head()
country_train_means = train_df.groupby("countryRaw")["DaysToSettle"].mean()
country_zscores = (country_train_means - train_df["DaysToSettle"].mean()) /
train_df["DaysToSettle"].std()

# Adicionar feature CountryZScore ao dataframe original
df["CountryZScore"] = df["countryRaw"].map(country_zscores)

# Atribuir aos conjuntos de treino (e teste) a nova feature
X_train["CountryZScore"] = df.loc[X_train.index, "CountryZScore"]
X_test["CountryZScore"] = df.loc[X_test.index, "CountryZScore"]

# Calcular médias por cliente a partir dos dados de treino
avg_days_to_pay = train_df.groupby("customerID")["DaysToSettle"].mean()
avg_days_late = train_df.groupby("customerID")["DaysLate"].mean()

# Mapear para todos os dados (com base nos valores aprendidos no treino)
df["AvgDaysToPay_byCustomer"] = df["customerID"].map(avg_days_to_pay)
df["AvgDaysLate_byCustomer"] = df["customerID"].map(avg_days_late)

# Adicionar as novas features ao treino/teste

```

```

X_train["AvgDaysToPay_byCustomer"] = df.loc[X_train.index,
"AvgDaysToPay_byCustomer"]
X_test["AvgDaysToPay_byCustomer"] = df.loc[X_test.index,
"AvgDaysToPay_byCustomer"]
X_train["AvgDaysLate_byCustomer"] = df.loc[X_train.index, "AvgDaysLate_byCustomer"]
X_test["AvgDaysLate_byCustomer"] = df.loc[X_test.index, "AvgDaysLate_byCustomer"]
# Remover colunas que causariam data leakage ou features usadas no cálculo
X_train.drop(columns=["countryRaw", "DaysToSettle",
"customerID", "DaysLate", "SettledDate",
"InvoiceDate"], inplace=True)
X_test.drop(columns=["countryRaw", "DaysToSettle",
"customerID", "DaysLate", "SettledDate",
"InvoiceDate"], inplace=True)

print("\n=== conjunto de treino ===\n", "\n".join(X_train.columns))

# TREINAMENTO DO MODELO
hyperparameters = {
    'n_estimators' : [200, 300, 400, 500],
    'max_depth' : [None, 5, 10, 15],
    'min_samples_split': [10, 12, 15, 17, 20],
    'min_samples_leaf': [6, 7, 9, 10, 12, 15],
    'class_weight': ["balanced"]
}
grid = GridSearchCV(
    RandomForestClassifier(),
    param_grid=hyperparameters,
    scoring="roc_auc",
    cv=10,
    n_jobs=-1,
    verbose=1
)
grid.fit(X_train, y_train)
print("Melhores parâmetros:", grid.best_params_)

```

```

print("Melhor ROC-AUC (CV):", grid.best_score_)
best_rf = grid.best_estimator_

# PREDIÇÃO E MÉTRICAS DE DESEMPENHO
y_pred = best_rf.predict(X_test)
y_proba = best_rf.predict_proba(X_test)[:,1]
#y_pred= (y_proba >= 0.48).astype(int)

metrics = {
    "Accuracy" : accuracy_score(y_test, y_pred),
    "Precision" : precision_score(y_test, y_pred),
    "Recall" : recall_score(y_test, y_pred),
    "F1 Score" : f1_score(y_test, y_pred),
    "ROC AUC" : roc_auc_score(y_test, y_proba)
}
print("==== Model Performance ====")
for k,v in metrics.items():
    print(f'{k:>9}: {v:.4f} ')

# MATRIZ DE CONFUSÃO
cm = confusion_matrix(y_test, y_pred)
disp = ConfusionMatrixDisplay(cm, display_labels=["No prazo", "Atrasado"])
disp.plot(cmap="Blues")
plt.title("Matriz de confusão")
plt.xlabel("Valor previsto")
plt.ylabel("Valor real")
plt.show()

# SHAP
explainer = shap.Explainer(best_rf)
shap_vals = explainer(X_test)
sv = shap_vals.values[:,1]
# importancia global
shap.summary_plot(sv, X_test, plot_type="bar")

```

```
# Resumo da distribuição (positivo vs negativo)  
shap.summary_plot(sv, X_test)
```

APÊNDICE B – INADIMPLÊNCIA – RECORTE DE TREINO (XGB)

```

import pandas as pd
from collections import Counter
from sklearn.model_selection import (train_test_split, GridSearchCV)
from xgboost import XGBClassifier
from sklearn.metrics import (
    accuracy_score, precision_score, recall_score,
    f1_score, roc_auc_score, confusion_matrix, ConfusionMatrixDisplay)
import matplotlib.pyplot as plt
import seaborn as sns
import shap

[...]

# TREINAMENTO DO MODELO
# Cálculo do scale_pos_weight
counts = Counter(y_train)
scale = counts[0] / counts[1]

hyperparameters = {
    "n_estimators": [300, 400, 500, 600],
    "max_depth": [3, 5, 7, 10],
    "learning_rate": [0.01],
    "subsample": [0.8, 1.0],
    "colsample_bytree": [0.8, 1.0],
    "min_child_weight": [3, 4, 5, 6, 7],
    "scale_pos_weight": [scale, scale * 1.25]
}

grid = GridSearchCV(
    estimator=XGBClassifier(objective="binary:logistic", eval_metric="logloss"),
    param_grid=hyperparameters,
    scoring="precision", # roc_auc também foi utilizado
    cv=10,

```

```
n_jobs=-1,  
verbose=1  
)  
grid.fit(X_train, y_train)  
print("Melhores parâmetros:", grid.best_params_)  
print("Melhor ROC-AUC (CV):", grid.best_score_)  
best_m = grid.best_estimator_  
  
[...]
```

APÊNDICE C – AUDITORIA – ALGORITMO PRINCIPAL (RF)

```

import pandas as pd
from sklearn.model_selection import train_test_split, GridSearchCV
from sklearn.ensemble import RandomForestClassifier
from sklearn.metrics import (accuracy_score, roc_auc_score, confusion_matrix,
                             precision_score, recall_score, f1_score, ConfusionMatrixDisplay)

import shap
import matplotlib.pyplot as plt
import seaborn as sns

# # Carregamento e pré-processamento
df = pd.read_csv("DATASETS/Audit Data/audit_data.csv")
df.rename(columns={"PROB": "Loss_score", "Prob": "History_score"}, inplace=True)

features = ["PARA_A", "PARA_B", "Money_Value", "numbers", "Sector_score",
            "District", "Loss_score", "History_score"]

# Multiplicando scores por 10
df[["Loss_score", "History_score"]] = df[["Loss_score", "History_score"]] * 10
print(df[features].head(5))

# Removendo linhas com valores faltantes
df.dropna(axis=0, how="any", inplace=True)

# Matriz de correlação
plt.figure(figsize=(10, 8))
sns.heatmap(df[features].corr(numeric_only=True), annot=True, cmap='coolwarm',
            linewidths=0.5)
plt.title("Matriz de Correlação das Features")
plt.tight_layout()
plt.show()

X = df[features]
y = df["Risk"]

```

```

# Split treino/teste
X_train, X_test, y_train, y_test = train_test_split(
    X, y, test_size=0.3, random_state=42, stratify=y
)

# Treinar RandomForest com GridSearch
hyperparameters = {
    "n_estimators": [100, 150, 200, 300],
    "max_depth": [None, 5, 10, 15, 20],
    "min_samples_split": [2, 3, 4, 5, 6],
    "min_samples_leaf": [2, 3, 4, 5, 6],
    "class_weight": ["balanced"]
}
grid = GridSearchCV(
    RandomForestClassifier(),
    param_grid=hyperparameters,
    scoring='roc_auc',
    cv=10,
    n_jobs=-1,
    verbose= 1
)

grid.fit(X_train, y_train)
print("Melhores hiperparâmetros:", grid.best_params_)
print("Melhor ROC-AUC (CV):", grid.best_score_)
best_rf = grid.best_estimator_

# Avaliação
#y_pred = best_rf.predict(X_test)
y_proba = best_rf.predict_proba(X_test)[:, 1]
y_pred= (y_proba >= 0.49).astype(int)

# Métricas de desempenho

```

```

metrics = {
    "Accuracy:": accuracy_score(y_test, y_pred),
    "Precision:": precision_score(y_test, y_pred),
    "Recall:": recall_score(y_test, y_pred),
    "F1 Score:": f1_score(y_test, y_pred),
    "ROC AUC:": roc_auc_score(y_test, y_proba)
}

print("==== Model Performance TEST====")
for k,v in metrics.items():
    print(f'{k:>9}: {v:.4f} ')

# Matriz de confusão
cm = confusion_matrix(y_test, y_pred)
disp = ConfusionMatrixDisplay(cm, display_labels=["Não-Fraude", "Fraude"])
disp.plot(cmap="Blues")
plt.title("Matriz de confusão")
plt.xlabel("Valor previsto")
plt.ylabel("Valor real")
plt.show()

# SHAP
explainer = shap.Explainer(best_rf)
shap_vals = explainer(X_test)
sv = shap_vals.values[:, :, 1]
# importancia global
shap.summary_plot(sv, X_test, plot_type="bar")
# Resumo da distribuição (positivo vs negativo)
shap.summary_plot(sv, X_test)

```

APÊNDICE D – AUDITORIA – RECORTE DE TREINO (XGB)

```

import pandas as pd
from collections import Counter
from sklearn.model_selection import train_test_split, GridSearchCV
from xgboost import XGBClassifier
from sklearn.metrics import (accuracy_score, roc_auc_score, confusion_matrix,
                             precision_score, recall_score, f1_score, ConfusionMatrixDisplay)

import shap
import matplotlib.pyplot as plt
import seaborn as sns

[...]

# Treinar XGBoost com GridSearch
# Cálculo do scale_pos_weight
counts = Counter(y_train)
scale = counts[0] / counts[1]

hyperparameters = {
    "n_estimators" : [200, 300, 400, 500],
    "max_depth" : [3, 5, 7, 10],
    "learning_rate" : [0.01],
    "subsample": [0.8, 1.0],
    "colsample_bytree": [0.8, 1.0],
    "min_child_weight": [3, 4, 5, 6, 7],
}

grid = GridSearchCV(
    XGBClassifier(objective="binary:logistic"),
    param_grid=hyperparameters,
    scoring="roc_auc",
    cv=10,
    n_jobs=-1,
    verbose= 1

```

)

```
grid.fit(X_train, y_train)
```

```
print("Melhores hiperparâmetros:", grid.best_params_)
```

```
print("Melhor ROC-AUC (CV):", grid.best_score_)
```

```
best_rf = grid.best_estimator_
```

```
[...]
```