



**Maria Fernanda Araujo de Carvalho**

**Do Silêncio ao Clique: Como a Conveniência da  
Inteligência Artificial Está Transformando a Forma como  
Denunciamos nas Empresas**

**Dissertação de Mestrado**

Dissertação apresentada ao Programa de Pós-graduação  
em Administração de Empresas da PUC-Rio como  
requisito parcial para obtenção do título de Mestre em  
Administração de Empresas.

Orientador: Prof. Jorge Brantes Ferreira

Rio de Janeiro  
abril de 2025



**Maria Fernanda Araujo de Carvalho**

**Do Silêncio ao Clique: Como a Conveniência da  
Inteligência Artificial Está Transformando a Forma como  
Denunciamos nas Empresas**

Dissertação apresentada como requisito parcial para  
obtenção do grau de Mestre pelo Programa de Pós-  
Graduação em Administração de Empresas. Aprovada  
pela Comissão Examinadora abaixo

**Prof. Jorge Brantes Ferreira**

Orientador

Departamento de Administração- PUC-Rio

**Prof. Ricardo Limongi França Coelho**

Departamento de Administração - UFG

**Profa. Alessandra Baiocchi Antunes Corrêa**

Departamento de Administração - PUC-Rio

Rio de Janeiro, 28 abril de 2025

Todos os direitos reservados. É proibida a reprodução total ou parcial do trabalho sem autorização da universidade, do autor e do orientador.

### **Maria Fernanda Araujo de Carvalho**

Graduação em Administração pelo IBMEC em 2017. Formação em Compliance pelo Insper, em 2018. Head e Sócia de Novos Negócios e ESG em uma empresa de tecnologia e certificações ESG. Início das atividades no mestrado acadêmico em Administração de Empresas com ênfase em Marketing em março de 2023.

#### Ficha Catalográfica

Carvalho, Maria Fernanda Araujo de

Do silêncio ao clique : como a conveniência da inteligência artificial está transformando a forma como denunciemos nas empresas / Maria Fernanda Araujo de Carvalho ; orientador: Jorge Brantes Ferreira. – 2025.

100 f ; 30 cm

1. Dissertação (mestrado)–Pontifícia Universidade Católica do Rio de Janeiro, Departamento de Administração, 2025.

Inclui bibliografia

1. Administração – Teses. 2. Inteligência artificial. 3. Canal de denúncia. 4. Confiança. 5. Qualidade. 6. Conveniência. I. Ferreira, Jorge Brantes. II. Pontifícia Universidade Católica do Rio de Janeiro. Departamento de Administração. III. Título.

CDD: 658

## **Agradecimentos**

Ao meu pai, por todo o apoio e incentivo ao longo do mestrado e de minha vida.

À minha mãe por ser um exemplo de mulher e de profissional, e que me inspira todos os dias.

Ao meu irmão, por ir sempre além e me mostrar que nossos limites são, muitas vezes, crenças mentais.

Ao meu orientador por ter aceitado embarcar nessa jornada comigo, por toda a troca e todo o aprendizado ao longo do meu mestrado.

O presente trabalho foi realizado com apoio da Coordenação de Aperfeiçoamento de Pessoal de Nível Superior - Brasil (CAPES) - Código de Financiamento 001.

## Resumo

Carvalho, Maria Fernanda Araujo. **Do Silêncio ao Clique: Como a Conveniência da Inteligência Artificial Está Transformando a Forma como Denunciamos nas Empresas**. Rio de Janeiro, 2025. 100p. Dissertação de Mestrado – Departamento de Administração, Pontifícia Universidade Católica do Rio de Janeiro.

A maior difusão da inteligência artificial e o avanço contínuo das tecnologias digitais vêm transformando a forma como as pessoas interagem com sistemas automatizados em diferentes contextos, inclusive no ambiente corporativo. Nesse cenário de constante inovação, os canais de denúncia baseados em IA têm ganhado espaço como ferramentas estratégicas para a promoção da ética e do compliance nas organizações, oferecendo maior agilidade, segurança e acessibilidade no recebimento de relatos. Considerando as particularidades desses canais e o caráter sensível das informações ali transmitidas, torna-se fundamental compreender os fatores que influenciam a conveniência percebida do usuário e a sua disposição em utilizar essas ferramentas, muitas vezes em situações de fragilidades. O presente estudo tem como objetivo avaliar como a qualidade do serviço habilitado por IA, a confiança na ferramenta, o sacrifício percebido, a ansiedade e o risco percebido à privacidade se relacionam entre si, afetando a conveniência percebida ao utilizar canais de denúncia automatizados. Para isso, foi proposto um modelo conceitual que considera as relações entre esses construtos. A validação das hipóteses foi realizada por meio da modelagem de equações estruturais, com base em dados coletados junto a usuários que tiveram contato com IA. Os resultados indicam que a qualidade do serviço habilitado por IA tem efeito positivo sobre a confiança do usuário e contribui para a redução do sacrifício percebido. A ansiedade, por outro lado, mostrou-se um fator prejudicial, impactando negativamente a confiança, aumentando o sacrifício e elevando a percepção de risco à privacidade. Observou-se ainda que a confiança exerce influência positiva sobre a conveniência percebida, enquanto o sacrifício tem efeito negativo sobre essa mesma variável. O risco percebido à privacidade, no entanto, não demonstrou impacto significativo sobre a percepção de conveniência.

## Palavras-chave

Inteligência artificial, canal de denúncia, confiança, qualidade, conveniência.

## **Abstract**

Carvalho, Maria Fernanda Araujo. **From Silence to Click: How the Convenience of Artificial Intelligence Is Reshaping Corporate Whistleblowing Practices** Rio de Janeiro, 2025. 100p. MBA Dissertation – Departamento de Administração, Pontifícia Universidade Católica do Rio de Janeiro.

The widespread adoption of artificial intelligence and the continuous evolution of digital technologies have been reshaping how individuals interact with automated systems in various contexts, including corporate compliance environments. In this dynamic landscape of innovation, AI-powered whistleblowing platforms have gained prominence as strategic tools to foster organizational ethics and integrity, offering enhanced efficiency, data security, and accessibility in handling misconduct reports. Given the unique characteristics of these platforms and the sensitive nature of the information reported, it becomes crucial to understand the factors that shape the user experience and influence willingness to engage with such tools—often under emotionally charged and high-stakes conditions. This study aims to examine how perceived service quality, user trust in the AI-enabled platform, perceived burden (sacrifice), anxiety, perceived privacy risk, and perceived convenience interact to influence user behavior when utilizing automated whistleblowing channels. A conceptual model was developed to explore the interrelationships among these constructs. The proposed hypotheses were tested through structural equation modeling, based on data collected from individuals who had interacted with AI-driven whistleblowing tools in organizational contexts. Findings indicate that perceived service quality has a positive effect on user trust and contributes to lowering the perceived burden associated with using the platform. Anxiety, by contrast, emerged as a critical barrier, undermining trust, increasing perceived effort, and intensifying concerns about privacy risks. Moreover, trust was found to positively influence perceived convenience, while perceived burden had a negative impact on the same. Notably, perceived privacy risk did not show a statistically significant effect on convenience. These insights offer a deeper understanding of the psychological and experiential factors that enhance—or hinder—the effective adoption of AI-based whistleblowing systems, and provide actionable guidance for the development and refinement of such tools in corporate compliance settings.

## **Keywords**

Artificial intelligence, whistleblowing systems, trust, quality, convenience.

## Sumário

|                                                                     |    |
|---------------------------------------------------------------------|----|
| 1 Introdução                                                        | 12 |
| 1.1. Problema de Pesquisa                                           | 12 |
| 1.2. Objetivos                                                      | 14 |
| 1.2.1 Questões a investigar                                         | 14 |
| 1.3. Relevância do Estudo                                           | 14 |
| 1.4. Delimitações do Estudo                                         | 15 |
| 1.5. Organização do Estudo                                          | 16 |
| 2 Referencial Teórico                                               | 17 |
| 2.1. Conveniência do usuário habilitada pela IA                     | 17 |
| 2.2. Canais de denúncia corporativos                                | 18 |
| 2.3 Teoria do comprometimento-confiança                             | 19 |
| 2.3.1. Confiança                                                    | 20 |
| 2.3.2. Sacrifício percebido                                         | 22 |
| 2.3.3. Qualidade do serviço habilitada pela IA                      | 23 |
| 2.3.5. Ansiedade                                                    | 24 |
| 2.3.6. Privacidade                                                  | 25 |
| 2.4. Modelo Conceitual e hipóteses de pesquisa                      | 26 |
| 3 Metodologia da Pesquisa                                           | 35 |
| 3.1. Tipo de pesquisa                                               | 35 |
| 3.2. Operacionalização e definição das variáveis                    | 35 |
| 3.2.1. Definição operacional das variáveis                          | 36 |
| 3.2.2. Procedimentos de Tradução e Adaptação das Escalas Utilizadas | 36 |
| 3.2.3. Pré-teste do instrumento de pesquisa                         | 38 |
| 3.3. População e amostra                                            | 38 |
| 3.3.1. População                                                    | 38 |
| 3.3.2. Amostra                                                      | 38 |
| 3.4. Coleta de dados                                                | 38 |
| 3.4.1. O instrumento de coleta de dados                             | 38 |

|                                                          |    |
|----------------------------------------------------------|----|
| 3.4.2 A coleta de dados                                  | 39 |
| 3.5. Análise de dados                                    | 40 |
| 3.5.1. Validade e confiabilidade                         | 40 |
| 3.6. Análises estatísticas                               | 41 |
| 3.7. Limitações do método                                | 42 |
| 3.7.1. Limitações relacionadas ao critério de amostragem | 42 |
| 3.7.2. Limitações decorrentes da coleta de dados         | 42 |
| <br>                                                     |    |
| 4 Modelagem e Análise de dados                           | 43 |
| 4.1. Caracterização da amostra                           | 43 |
| 4.1.1. Análise da Variância de Método Comum              | 45 |
| 4.2. Análises e resultados                               | 45 |
| 4.2.1. Avaliação do modelo de mensuração                 | 45 |
| 4.2.2. Validade e confiabilidade dos construtos          | 46 |
| 4.2.3. Análise do modelo estrutural                      | 49 |
| 4.2.3.1. Ajuste no modelo proposto                       | 50 |
| 4.2.4. Teste das hipóteses de pesquisa                   | 51 |
| 4.3. Discussão dos resultados                            | 53 |
| 4.3.1. Ansiedade                                         | 53 |
| 4.3.2. Qualidade do serviço habilitado pela IA           | 54 |
| 4.3.3. Conveniência                                      | 55 |
| 4.3.4. Resultados gerais                                 | 56 |
| <br>                                                     |    |
| 5 Conclusões                                             | 58 |
| 5.1. Sumário do estudo                                   | 58 |
| 5.2. Conclusões e implicações                            | 60 |
| 5.2.1. Implicações teóricas                              | 60 |
| 5.2.2. Implicações práticas                              | 61 |
| 5.3. Limitações do estudo                                | 62 |
| 5.4. Sugestões para pesquisas futuras                    | 63 |
| <br>                                                     |    |
| 6 Referências Bibliográficas                             | 65 |
| <br>                                                     |    |
| Apêndice A – Questionário utilizado na pesquisa          | 84 |



## Lista de figuras

|                                                                              |    |
|------------------------------------------------------------------------------|----|
| Figura 2.1 - Modelo Conceitual Proposto com Hipóteses.....                   | 26 |
| Figura 4.1 - Modelo Proposto.....                                            | 50 |
| Figura 4.2 - Coeficientes Padronizados Estimados para o Modelo Proposto..... | 52 |

## Lista de tabelas

|                                                                                                             |    |
|-------------------------------------------------------------------------------------------------------------|----|
| Tabela 4.1 - Características da Amostra.....                                                                | 44 |
| Tabela 4.2 - Confiabilidade Composta e Variância Extraída Média (AVE) .....                                 | 48 |
| Tabela 4.3 - Índices de Ajuste do Modelo.....                                                               | 51 |
| Tabela 4.4 - Coeficientes Padronizados Estimados, Hipóteses e Significâncias<br>para o Modelo Proposto..... | 52 |

## Lista de quadros

|                                                         |    |
|---------------------------------------------------------|----|
| Quadro 3.1 - Definições operacionais das variáveis..... | 37 |
| Quadro 4.1 - Matriz de Correlação entre Construtos..... | 47 |
| Quadro 4.2 - Matriz de Validade Discriminante .....     | 49 |

# 1

## Introdução

### 1.1. Problema de Pesquisa

No contexto da disseminação da inteligência artificial e a crescente integração dela em ambientes corporativos, particularmente nos canais de denúncias, torna-se primordial examinar como a IA pode aprimorar a conveniência do usuário em fazer uma denúncia e assim aumentar as chances de uma situação ilícita ou antiética ser reportada.

Vale ressaltar a definição do Dicionário Michaelis (2025) para conveniência é simples e direta: “adequado; o que é de interesse para alguém; proveito, utilidade, vantagem.” Seiders et al. (2005), Berry, Seiders e Grewal (2002) e Kotler e Zaltman (1971) esclareceram que um serviço é conveniente quando dois fatores críticos são alcançados: otimização do tempo gasto e a minimização do esforço do usuário.

Os canais de denúncias servem como mecanismos vitais para promover o compliance e a transparência nas organizações. No entanto, Sari (2023) aponta que os usuários frequentemente encontram obstáculos, como a manutenção do anonimato, a compreensão do processo de denúncia e a diferenciação entre os tipos de relatos, o que pode dificultar a comunicação eficaz e o relato de irregularidades com precisão.

Sharma (2022) afirma que a IA pode aumentar a conveniência do usuário de diversas formas como, por exemplo, prover orientações personalizadas para os usuários, ajudando-os a navegar pelas complexidades do processo de denúncia e reclamação, garantindo que suas preocupações sejam devidamente categorizadas e endereçadas. E ainda, a IA pode operar por 24 horas, 7 dias por semana, manter uma comunicação com o usuário fiel aos padrões pré-estabelecidos, e, por fim, atender a uma grande demanda de solicitações, de diversos usuários simultâneos, aprendendo e alimentando o modelo através de cada uma delas.

Além disso, sistemas baseados em IA podem aumentar a precisão e a relevância das informações fornecidas aos usuários, promovendo um ambiente

mais acolhedor para os denunciantes (MOHAMED, 2024). No entanto, em uma perspectiva de mercado, o sucesso da implementação da IA nesses canais depende de conquistar a confiança dos usuários, que é influenciada por fatores como qualidade do serviço prestado pela IA e a percepção de confiabilidade do sistema que utiliza a IA (KUMAR, 2024).

Apesar dos benefícios potenciais, existem lacunas na compreensão de como a integração da IA pode aumentar efetivamente a conveniência do usuário nos canais de denúncias. A confiança dos usuários nos sistemas de IA é fundamental para sua aceitação e engajamento com essas tecnologias (KUMAR, 2024). Além disso, questões éticas relacionadas ao uso da IA devem ser abordadas para garantir que os usuários se sintam confortáveis ao utilizar esses canais segundo Rezwana e Maher (2022).

O trabalho de Lai (2017) estende o Modelo de Aceitação de Tecnologia ao incorporar a conveniência percebida como um fator-chave na formação de atitudes em relação ao aprendizado móvel. Isso corrobora a noção de que as experiências e intenções dos usuários são significativamente moldadas pela conveniência, demonstrando sua centralidade na aceitação das tecnologias. Além disso, a conveniência perpassa não apenas a predisposição intencional de uso, mas incorpora uma dimensão pragmática de valor percebido, conforme abordado pela Teoria do Valor do Consumidor, na qual o julgamento sobre uma oferta resulta de uma equação custo-benefício que enfatiza a redução do sacrifício por meio da economia de tempo e esforço (ZEITHAML, 1988).

Diferentemente dos modelos clássicos de aceitação de tecnologia, como o TAM de Davis (1989), que utilizam a intenção de uso como variável dependente final, baseada nos construtos de utilidade percebida e facilidade de uso para explicar o comportamento de adoção (DAVIS, 1989), as abordagens de service convenience defendem que a conveniência percebida impacta diretamente a efetivação do comportamento de consumo (BERRY, SEIDERS e GREWAL, 2002) e se mostra preditora robusta de satisfação e uso continuado de serviços digitais (SEIDERS, VOSS, GODFREY e GREWAL, 2007). Logo, ao considerar a conveniência como construto dependente final no modelo proposto neste trabalho, reforça-se a ênfase na experiência concreta do denunciante, cujo relato efetivo depende primordialmente de quão simples, célere e adaptado às suas necessidades se apresenta o sistema de denúncia assistido por IA.

Assim, chegamos à pergunta da pesquisa: Quais são os impactos da implementação de inteligência artificial nos canais de denúncia corporativos sobre os fatores de conveniência percebidos pelos usuários, incluindo qualidade do serviço, confiança na tecnologia e sacrifício percebido?

## **1.2. Objetivos**

O presente estudo tem como objetivo analisar antecedentes que influenciam a conveniência, ao utilizar ferramentas que integram a inteligência artificial ao processo de denúncias online em um canal corporativo.

### **1.2.1 Questões a investigar**

Diante do objetivo proposto, busca-se investigar, com base na teoria do compromisso de confiança (MORGAN; HUNT, 1994) e no modelo de qualidade do serviço (PARASURAMAN et al., 1994), de que forma, positiva ou negativa, a confiança, privacidade percebida e o sacrifício percebido influenciam a conveniência do serviço de denúncias habilitado por IA, avaliando também antecedentes como ansiedade e qualidade do serviço habilitado por IA. Bem como analisar os efeitos da privacidade na conveniência percebida pelo uso da IA (XU et al., 2011; LEE, RHA, 2016). Dessa forma, ao invés de posicionar a Intenção de Uso ou Adoção como desfecho final do modelo, este trabalho prioriza a conveniência percebida como construto dependente (BERRY, SEIDERS e GREWAL, 2002), pois ela representa de modo mais direto a experiência concreta do denunciante e, consequentemente, suas chances de efetivar o relato de irregularidades em ambientes corporativos assistidos por IA.

## **1.3. Relevância do Estudo**

A crescente popularização de plataformas de inteligência artificial de uso geral, como o ChatGPT, democratizou o acesso a ferramentas de IA, reduzindo custos e barreiras tecnológicas para organizações de todos os portes. Nesse cenário, a conveniência tornou-se um dos principais determinantes da adoção e do uso contínuo de sistemas digitais, sobretudo em processos críticos de *compliance*,

como canais de denúncias corporativas, em que a agilidade e a simplicidade de navegação podem reduzir o subrelato de práticas ilícitas e antiéticas.

Do ponto de vista teórico, este trabalho preenche lacunas importantes na literatura de aceitação de tecnologias ao deslocar o foco do desfecho tradicional – intenção de uso/adoção – para o construto de conveniência percebida como variável dependente central. Embora estudos clássicos, como o TAM (DAVIS, 1989) e o UTAUT (VENKATESH et al., 2003), reconheçam a importância da facilidade de uso e das condições facilitadoras, poucos se dedicam a mensurar em profundidade as múltiplas dimensões que compõem a conveniência em serviços digitais (BERRY, SEIDERS e GREWAL, 2002). No mesmo sentido, a incorporação dos constructos ansiedade (YANG et al., 2021) e privacidade (PATIL et al., 2024) adiciona originalidade à proposta, ao explorar como o receio de exposição de dados e o mal-estar emocional influenciam diretamente o julgamento de conveniência em canais de denúncia assistidos por IA.

Em termos de relevância gerencial, os resultados desta dissertação oferecem subsídios para provedores de soluções de *compliance* baseadas em IA e para corporações que avaliam sua adoção. A identificação dos antecedentes mais impactantes na conveniência percebida permitirá alinhar o design de interfaces, fluxos de atendimento e políticas de segurança de dados às expectativas e medos dos denunciantes, favorecendo maior adesão e eficácia na captura de informações críticas. Sendo assim, a contribuição prática deste trabalho reside na potencial melhoria dos mecanismos de transparência e *accountability* organizacional. Ao evidenciar os fatores que promovem uma experiência de denúncia mais conveniente e segura, a pesquisa contribui para ambientes corporativos eticamente mais responsáveis e resguardados de práticas nocivas. Assim, ao avaliar a integração da IA sob a ótica da conveniência, ansiedade e privacidade, o estudo não apenas amplia o arcabouço teórico existente, mas também traduz conhecimento acadêmico em diretrizes práticas capazes de fortalecer a cultura de *compliance* nas empresas.

#### **1.4. Delimitações do Estudo**

O estudo é delimitado a residentes do Brasil e adultos maiores de 18 anos.

O estudo ainda se concentrou na avaliação de uma única ferramenta que utiliza o ChatGPT como um *chatbot* para ajudar pessoas no momento de realizarem uma denúncia em um canal de denúncias corporativo específico.

### **1.5. Organização do Estudo**

O trabalho se organiza em capítulos, uma breve descrição deles se dá a seguir:

O primeiro capítulo apresenta os objetivos da pesquisa e destaca sua relevância para o estudo sobre como a integração da inteligência artificial nos canais de denúncias corporativos pode impactar a conveniência do usuário.

O segundo capítulo apresenta a revisão de literatura sobre o tema uso de inteligência artificial em canais de denúncia corporativos, e como alguns constructos influenciam a percepção do usuário seguido das hipóteses de pesquisa sugeridas. Esta etapa termina com a apresentação do modelo proposto por esse trabalho.

O capítulo 3 é dedicado à exposição da metodologia empregada neste estudo. São detalhados o delineamento da pesquisa, a caracterização da população de interesse, o processo de seleção da amostra, a construção do instrumento de pesquisa, o procedimento de coleta dos dados, as técnicas de análise de dados e as limitações da metodologia adotada.

O capítulo quatro apresenta os resultados do estudo. É realizado o teste dos modelos de mensuração e estrutural propostos, verificação das hipóteses de pesquisa e, então, discutidas as relações entre os construtos apresentados no referencial teórico.

O capítulo cinco finaliza o trabalho, apresentando suas principais conclusões e contribuições. Nessa etapa foram avaliadas as limitações do estudo e sugeridas pesquisas futuras.



## 2 Referencial Teórico

Este capítulo apresenta o referencial teórico utilizado no estudo. A partir de pesquisas em artigos científicos, teses e dissertações, é feita uma revisão de tópicos relevantes da literatura referentes a conceitos de conveniência, confiança, qualidade do serviço habilitado pela IA, ansiedade e privacidade e sacrifício percebido.

### 2.1. Conveniência do usuário habilitada pela IA

O conceito de conveniência no que diz respeito ao comportamento do consumidor remonta ao trabalho fundamental de Copeland (1923), que caracterizou os produtos de conveniência como aqueles que os consumidores compram com frequência e rapidez em locais de varejo de fácil acesso. A conveniência leva a um maior engajamento (ROY, BALAJI, SADEQUE, NGUYEN e MELEWAR, 2017; VAN DOORN et al., 2010).

Essa abordagem inicial, centrada no produto, evoluiu ao longo do tempo, com pesquisadores como Jian et al. (2013) identificando uma mudança de paradigma para uma perspectiva orientada ao serviço sobre a conveniência, conforme articulado por estudiosos como Seiders et al. (2005), Berry, Seiders e Grewal (2002) e Kotler e Zaltman (1971). Esses autores esclareceram que dois fatores críticos para a oferta de serviços convenientes aos usuários são a otimização do tempo gasto e a minimização do esforço do usuário.

A conveniência proporcionada pelos serviços que fazem uso de IA pode ser dividida em três dimensões principais. A primeira é a disponibilidade, já que o autoatendimento com suporte de IA está acessível 24 horas por dia, 7 dias por semana, permitindo que o serviço seja utilizado de qualquer lugar (WALCH, 2019). A segunda dimensão refere-se ao fornecimento de informações e suporte em tempo real ao longo da jornada do cliente (THIEL, 2019). E a terceira dimensão envolve agentes/robôs inteligentes (*bots*) que se utilizam da IA para iniciar interações de forma proativa, oferecer informações relevantes e apoiar os

usuários em cada ponto de contato. Isso possibilita que os usuários obtenham as respostas de que precisam no momento em que desejam, eliminando a necessidade de aguardar pela disponibilidade de um funcionário humano, o que contribui para a redução do tempo de resolução e para o aumento da satisfação dos usuários (WALCH, 2019). Chang et al. (2012) sugerem que o engajamento do usuário influencia a sua percepção de conveniência, fazendo com que a adoção da tecnologia se perpetue com o passar do tempo.

Com a pandemia do coronavírus (Covid-19), a conveniência associada à localização ganhou ainda mais relevância, considerando o contexto de isolamento social e distanciamento que o mundo enfrentou (MEYER, 2020). Os aspectos de conveniência relacionados à localização e à economia de tempo têm sido amplamente analisados, especialmente quanto ao impacto dos tempos de espera percebidos na experiência do usuário (ROY et al., 2017). De fato, a economia de tempo é um fator extensivamente investigado, com foco no efeito do tempo de espera percebido sobre a conveniência percebida do usuário (ROY et al., 2017).

Tecnologias baseadas em inteligência artificial (IA) melhoram a qualidade dos serviços, promovendo maior satisfação dos usuários (AMEEN et al., 2021). A intenção de reuso está ligada a conveniência percebida da tecnologia, fatores como facilidade de acesso aumentam o valor percebido pelos usuários (PHAM et al., 2018; TAN, 2022). Contudo, questões relacionadas à confiança, como preocupações com privacidade e segurança de dados em tecnologias de IA, reforçam a necessidade de equilibrar a conveniência com a segurança para atender às expectativas dos usuários e garantir a continuidade do uso dessas soluções (HASAN et al., 2021).

## **2.2. Canais de denúncia corporativos**

De acordo com Miceli e Near (1985), a denúncia pode ser definida como a divulgação de práticas ilegais, imorais ou ilegítimas para pessoas ou organizações que possam tomar medidas. Hedin e Mansson (2012) complementam que o objetivo da denúncia é corrigir ou interromper atos ilícitos, um processo que começa com a observação de irregularidades e posterior tomada de decisão da escolha entre agir, não agir ou se afastar da situação.

Caso a pessoa resolva agir, ela precisa ter meios para fazê-lo. A ferramenta de canal de denúncias online é um meio para receber essas informações. De acordo com o Relatório de 2024 da *Association of Certified Fraud Examiners*, 40% das denúncias foram recebidas via *web*. Assim, meios para facilitar a experiência do usuário durante o registro dessa denúncia se fazem essenciais. E ainda, de acordo com a *United Sentencing Commision*, é importante as empresas terem e divulgarem um sistema com mecanismos anônimos e confidenciais, para que colaboradores relatem ou busquem orientação sobre condutas criminosas sem medo de represálias.

Para Ruivo e Pires (2020) o regular funcionamento do canal de denúncias atua com um verificador da efetividade do programa de compliance em defesa da empresa.

No Brasil, diversas Leis como a Lei Federal Anticorrupção, CIPA e o Programa Emprega +Mulheres, LGPD e outras apontam o canal de denúncias corporativo como uma ferramenta de comunicação essencial para o desenvolvimento de um programa de compliance efetivo. Saad-Diniz e Marin (2021), destacam que as poucas pesquisas realizadas na área, ainda que produzidas no setor privado, dão conta da baixa maturidade do “*whistleblower*” (denunciante) no ambiente empresarial brasileiro.

Neste contexto, usar a inteligência artificial como um apoio ao denunciante no momento de realizar a sua denúncia se faz muito relevante de um ponto de vista didático e educacional ao usuário. Vale salientar a diferença entre o “*whistleblower*” e o delator. O primeiro não possui envolvimento com os fatos suspeitos e nem está engajado em uma associação criminosa, diferente do segundo (SAAD-DINIZ, MARIN, 2021).

Por fim, a ferramenta de canal de denúncias é uma ferramenta empresarial que é utilizada pelo usuário, ou seja, B2B2C (*business-to-business-to-consumer*). Apesar do contratante ser a empresa, quem vai utilizar a inteligência artificial para conseguir enviar a sua denúncia a empresa é o consumidor.

## 2.3 Teoria do comprometimento-confiança

Morgan e Hunt (1994) são os criadores da teoria do comprometimento-confiança. Eles enfatizam o papel central da confiança e do comprometimento no

desenvolvimento de relações entre compradores e vendedores. Essa teoria tem sido aplicada em diversos contextos, como o comportamento de envio online (REHMAN, BJATTI, MOHAMED e AYOUP, 2019).

A confiança, como elemento-chave na teoria do comprometimento-confiança de Morgan e Hunt (1994), é também essencial para o sucesso de serviços automatizados, por ser uma peça fundamental na relação entre humanos e automação (HENGGSTLER, ENKEL, e DUELLI, 2016). Além disso, Wang (2019) apontam a privacidade como um aspecto decisivo da confiança, já que consumidores buscam manter controle sobre como seus dados são utilizados pelas empresas. Siau e Wang (2018) e Ferrario, Loi e Vigano (2019) também sugerem que a confiança pode influenciar a relação entre fatores como qualidade de serviço e conveniência no uso de tecnologias baseadas em IA.

Segundo Morgan e Hunt (1994), os principais fatores contextuais que implicam no relacionamento duradouro são o comprometimento e a confiança. Isto ocorre, pois estes fatores conduzem a preservação de investimento nos relacionamentos, o que gera cooperação e auxilia a manutenção de vínculos de longo prazo. Comprometimento é considerado como um grande desejo de manter uma relação valiosa. Enquanto a confiança é a disposição de dependência para com o parceiro. A falta de comprometimento faz com que os relacionamentos não sejam duradouros. Dessa forma, comprometimento e confiança estão diretamente relacionados com os comportamentos de colaboração (FOCK, KOH, 2006), essenciais para a manutenção e expansão da IA.

### **2.3.1. Confiança**

A definição clássica de confiança é uma atitude de expectativa de que as vulnerabilidades de uma pessoa em uma situação de risco não serão exploradas (CORRITORE, KRACHER e WIEDENBECK, 2003). No contexto do comércio online, isso inclui confiar na marca, bem como na tecnologia (CORRITORE, 2003). No contexto da IA, estudos recentes mostram que a confiança é fundamental para garantir a aceitação, o progresso contínuo e o desenvolvimento dessa tecnologias (SIAU, WANG, 2018).

Um trabalho fundamental na área é o estudo de Lee e See (2004), que explora a confiança na automação a partir de múltiplas perspectivas, incluindo

pontos de vista organizacional, psicológico e sociológico. Eles argumentam que o contexto em que a automação é usada influencia significativamente a confiança e propõem princípios de design para fomentar a confiança apropriada em sistemas automatizados. Essa estrutura tem sido influente na formação de pesquisas subsequentes sobre confiança em IA enfatizando a necessidade de uma compreensão mais profunda de como os usuários percebem e interagem com sistemas automatizados. Quando algoritmos e lógica funcional são transparentes, a confiança tende a ser reforçada, afirmam os estudiosos.

Hengstler (2016) sugere que a introdução da tecnologia de IA no processo de serviço deve ser comunicada proativamente, começando nas primeiras etapas de difusão. Sua justificativa é que, quando os níveis de conhecimento são baixos, a comunicação pela marca tem uma maior chance de influenciar a aceitação social em relação a novas tecnologias.

A transparência e a comunicação eficaz sobre as capacidades da IA são fundamentais para estabelecer confiança (WINTER, CARUSI, 2022). A avaliação da confiabilidade dos sistemas de IA tem sido um foco central nas pesquisas recentes. Kaur (2022) fornece uma visão geral abrangente de diferentes abordagens para aumentar a confiança em IA incluindo estratégias de validação e verificação. Eles enfatizam que a confiabilidade não é apenas uma característica da tecnologia, mas um construto relacional que evolui através das interações. Isso está alinhado com os achados de Madhavan e Wiegmann (2007), que discutem como a confiança em sistemas automatizados se desenvolve de forma semelhante à confiança entre humanos, embora existam diferenças críticas nas reações dos usuários a conselhos automatizados versus humanos.

A natureza dinâmica da confiança em IA também é uma área significativa de investigação. Guo e Yang (2020) propõem que a confiança não é estática, mas evolui com base nas interações com sistemas de IA. Essa perspectiva é fundamental para entender como manter e aumentar a confiança ao longo do tempo, especialmente à medida que as tecnologias de IA continuam a se desenvolver e mudar. Além disso, Kumar (2024) enfatiza o papel da confiança na facilitação da colaboração eficaz entre humanos e IA, sugerindo que a dinâmica da confiança pode influenciar significativamente os processos de tomada de decisão.

Assim, a confiança se torna um elemento central na construção de relacionamentos eficazes entre humanos e sistemas de IA.

### **2.3.2. Sacrifício percebido**

O sacrifício percebido refere-se ao que os consumidores abrem mão para adquirir um produto ou serviço, incluindo custos monetários e não monetários, como tempo, esforço, envolvimento mental ou emoções negativas, como irritação e frustração (ZEITHAML, 1988). Estudos recentes ressaltam a importância de analisar os sacrifícios dos clientes ao utilizar serviços automatizados, especialmente quando há poucas opções de escolha (ANDRÉ et al., 2018).

A implementação de ferramentas de inteligência artificial em contextos de compliance tem recebido atenção significativa, especialmente no que diz respeito aos sacrifícios percebidos associados ao seu uso.

Embora certos sacrifícios sejam necessários para acessar um serviço, muitos custos não monetários, como perda de controle, privacidade, dinheiro, tempo ou esforço, além de emoções negativas, podem ser difíceis de mensurar (DE KERVILER, DEMOULIN e ZIDDA, 2016; MERISAVO et al., 2007; SHIN, LIN, 2016). Em serviços com IA, dois outros aspectos precisam ser considerados: a ausência de interação humana e o possível isolamento social, ambos podendo impactar negativamente a experiência do cliente (DAVENPORT et al., 2020).

A literatura enfatiza a relevância da interação humana e da prestação de serviços amigáveis (PHAM, AHAMMAD, 2017; PINTO et al., 2017). No entanto, serviços habilitados por IA introduzem interações modernas que demandam alta coordenação social, algo que novos usuários podem perceber como um sacrifício (CHRISTAKIS, 2019; DAVENPORT et al., 2020). Além disso, esses serviços podem implicar na perda de controle do usuário, pois frequentemente seguem uma estrutura rígida definida pela tecnologia, e não pelas necessidades de cada usuário, além de demandarem dados pessoais para operarem eficientemente (CHEATHAM, JAVANMARDIAN e SAMANDARI, 2019).

Uma das principais preocupações em relação às ferramentas de compliance baseadas em IA é a questão da transparência e responsabilização. Felmann et al. (2019) enfatizam a importância da transparência nos sistemas de IA observando que as normas legais devem ser integradas aos sistemas técnicos para garantir a

conformidade ética. A falta de transparência nos algoritmos de IA pode levar a vieses e consequências indesejadas, o que pode minar a confiança entre usuários e partes interessadas (FELZMANN et al., 2019). Blasch et al. (2021) apoiam ainda mais essa noção introduzindo a Tabela Multifonte de Scorecard de IA (MAST), que fornece uma lista de verificação para desenvolvedores e usuários promoverem transparência e confiança nos resultados de IA (BLASCH et al., 2021). Este scorecard é particularmente relevante em contextos de compliance, pois incentiva a adesão a padrões e políticas estabelecidas, garantindo que as ferramentas de IA operem dentro de limites éticos.

A percepção das ferramentas de IA também varia significativamente entre os usuários, influenciada pela familiaridade com a tecnologia, por exemplo. O estudo de Arowosegbe (2024) sobre IA generativa no ensino superior indica que a percepção dos usuários pode evoluir à medida que se tornam mais conscientes dos recursos e limitações das ferramentas. Essa percepção é crucial em contextos de compliance, onde a compreensão das capacidades operacionais da IA pode afetar a disposição dos usuários de adotar essas ferramentas.

A falta de assistência humana também pode criar desafios para clientes sem experiências prévias ou que precisam de mais tempo para se adaptar. Pesquisas mostram que muitos preferem um equilíbrio entre automação e interação humana (GAUVRIT, 2019). Além disso, a integração da IA em processos de compliance pode levar a preocupações com a segurança do emprego e o potencial de redução da supervisão humana, o que pode dificultar a aceitação do usuário (BHARGAVA et al., 2020).

### **2.3.3. Qualidade do serviço habilitada pela IA**

A qualidade do serviço é tradicionalmente definida como a diferença entre o serviço esperado e o percebido e avaliada pela forma como os clientes percebem as ofertas de serviço de uma marca (PARASURAMAN et al., 1994).

Essa conceitualização da qualidade do serviço tem suas raízes na teoria da desconfirmação de expectativas (COLLIER, BIENSTOCK, 2006), onde a avaliação da qualidade do serviço resulta de uma comparação entre a percepção do serviço recebido e as expectativas prévias sobre o que esse serviço deveria proporcionar (CHOI, LEE, LEE e SUBRAMANI, 2004). A qualidade do serviço

impulsionado por IA transformou significativamente a maneira como as indústrias interagem com os consumidores (YOUNES et al., 2021).

No contexto de ensinar e ajudar o consumidor a realizar a denúncia de forma correta, analisou-se IA no contexto de educação. Chaudhri et al. (2013) analisaram a aplicação da IA na educação e constataram que ela pode aumentar a eficácia dos professores e o engajamento dos alunos.

Da mesma forma, identificou-se que sistemas de IA têm o potencial de educar o público sobre políticas e oferecer mecanismos mais eficazes em ambientes com alta incerteza (MIKHAYLOV et al., 2018).

### **2.3.5. Ansiedade**

A ansiedade associada a inteligência artificial (IA) é, em grande parte, uma herança da ansiedade do uso de computadores, ou “tecnofobia” (LI e HUANG, 2020; WANG e WANG, 2022). Essa resistência à tecnologia, que se manifesta como um desconforto ou apreensão em relação a computadores e seus usos (PARASURAMAN, IGBARIA, 1990), se intensifica com a IA. Afinal, a IA representa uma nova fronteira tecnológica, com a capacidade de tomar decisões autônomas e complexas, o que gera incertezas e preocupações.

Epstein (1972) descreveu a ansiedade como um estado de alerta perturbador diante de ameaças percebidas ou um medo não resolvido. Segundo Spielberger (1966), a ansiedade pode ser dividida em dois tipos: a ansiedade “traço”, que é um traço de personalidade estável e duradouro, e a ansiedade “estado”, que é transitória e muda com o tempo.

A ansiedade tecnológica, como aquela relacionada ao uso de computadores, internet ou cursos online, geralmente se manifesta como um estado temporário, suscetível a mudanças conforme as condições variam (CAMBRE, COOK, 1985)

Uma das principais razões para a ansiedade atrelada ao uso da IA é a falta de transparência em seus processos de tomada de decisão (CLARKE, 2019). Sem entender como a IA chega a suas conclusões, é difícil confiar em suas decisões, especialmente quando elas podem ter um impacto significativo em nossas vidas. Outro ponto importante é a falta de conhecimento dos usuários quanto ao seu



melhor uso e aplicação, incluindo a construção de *prompts*, dado que é uma tecnologia relativamente recente.

Além disso, a IA pode perpetuar vieses presentes nos dados com os quais ela é treinada, levando a resultados injustos ou discriminatórios (ETZIONI e ETZIONI, 2017; RUSSELL et al., 2015; STAHL, WRIGHT, 2018).

### **2.3.6. Privacidade**

A privacidade é um conceito muito importante que está presente em diversas áreas como a filosofia, o direito, a sociologia, a política e a economia. Desde os tempos de Aristóteles, que diferenciava as esferas pública e privada da vida, se discute sobre privacidade (SWANSON, 1992). Ao longo do tempo, a literatura aprofundou o entendimento sobre os diversos tipos de danos à privacidade, o que levou à criação de leis para protegê-la (SOLOVE, 2006).

A privacidade da informação é um aspecto importante da privacidade como um todo e um tema debatido há muitos anos. Autores como Skinner et al. (2006) sugerem que a maioria das pessoas entende a privacidade como um direito humano, embora em diferentes contextos. Clarke (2002) chegou a definir quatro dimensões da privacidade: a da pessoa em si, a do seu comportamento, a de suas comunicações e a de seus dados. Hoje em dia, com a maioria das comunicações sendo digitais, a privacidade das comunicações e a dos dados se juntam no conceito de privacidade da informação.

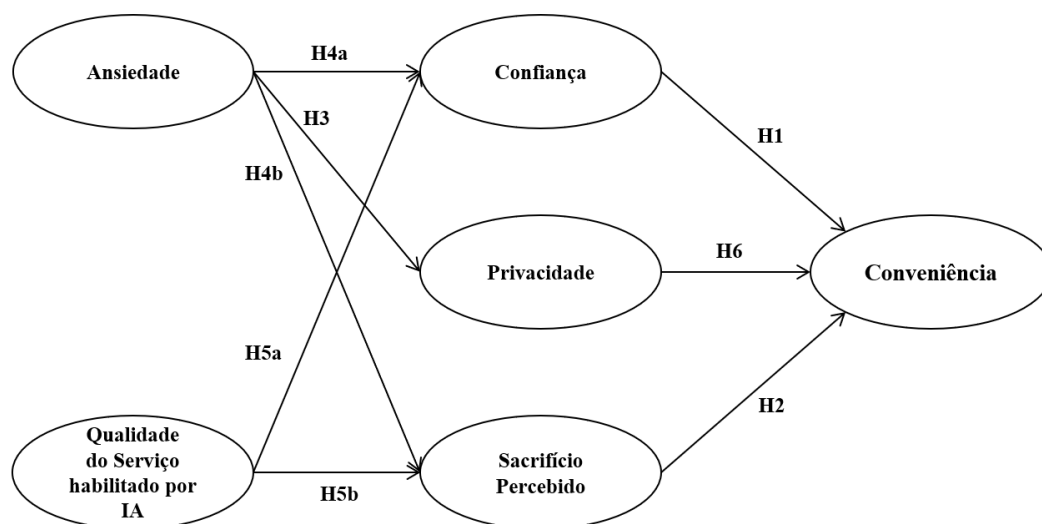
Com as tecnologias de informação e comunicação cada vez mais avançadas, a coleta, o armazenamento e a análise de dados nunca foram tão rápidos e volumosos (MALHOTRA et al. 2004). E ainda, há meios para que a coleta dos dados aconteça sem que o dono dos dados tome conhecimento (BÉLANGER, HILLER, 2006). Essa situação pode deixar o usuário ainda mais inseguro ao se deparar com uma ferramenta que use IA.

Um fator frequentemente estudado e que pode influenciar a forma como lidamos com a privacidade é o nosso nível de conhecimento sobre privacidade online. O paradoxo da privacidade pode ser explicado pela falta de experiência em primeira mão dos usuários com as consequências negativas de violações de privacidade, o que pode levá-los a subestimar os riscos à privacidade devido a sua inabilidade de gerenciá-lo (DIENLIN, TREPTE, 2015).

Estudos indicam que a falta de conhecimento sobre os riscos à privacidade e os direitos de privacidade de alguém, bem como a falta de conhecimento processual sobre como proteger a própria privacidade, pode reduzir as chances de que as preocupações dos indivíduos se traduzam em um comportamento que resulte na sua gestão de privacidade (DEBATIN et al., 2009; PARK, 2013).

## 2.4. Modelo Conceitual e hipóteses de pesquisa

Utilizando as bases das teorias de confiança e compromisso (MORGAN, HUNT, 1994) e o modelo de qualidade de serviço (PARASURAMAN et al., 1994), o modelo proposto (Figura 2.1) neste estudo apresenta uma abordagem para compreender como serviços habilitados por IA podem aprimorar a conveniência do usuário.



**Figura 2.1 - Modelo Conceitual Proposto com Hipóteses**

Confiança: A confiança, ao mitigar o risco percebido inerente a interações digitais, funciona como um fator catalisador que remove obstáculos cognitivos e operacionais, reduzindo o esforço mental exigido do usuário e acelerando o fluxo de decisão. Conforme definido por Corritore, Krackhardt e Wiedenbeck (2003), a confiança envolve a expectativa de que vulnerabilidades não serão exploradas, o que, em sistemas de IA, corresponde à crença de que o algoritmo atuará de forma previsível e segura. Essa crença gera um estado de calma e foco, em que o usuário

não se vê compelido a monitorar constantemente cada passo do processo, resultando em percepções mais elevadas de agilidade e simplicidade.

No mesmo sentido, Nadarzynski et al. (2019) indicam que interações com chatbots confiáveis — caracterizadas por respostas consistentes, coerentes e com “toque humano” — promovem um ambiente psicológico de conforto, no qual o usuário se sente liberto de hesitações e revisões constantes. Yang et al. (2022) corroboram essa relação, evidenciando relações fortes entre níveis de confiança e avaliações de conveniência, uma vez que indivíduos que confiam no sistema relatam experiências de uso mais fluídas e sem interrupções. Analogamente, Corrêa et al. (2023) e Glikson & Woolley (2020) mostram que a clareza das recomendações algorítmicas e a transparência dos processos internos aumentam a segurança percebida, desencadeando uma sensação de “caminho livre” para a execução de tarefas, sem a necessidade de buscas paralelas por validação externa. A adaptabilidade da IA às necessidades e preferências dos usuários também desempenha um papel fundamental na construção da confiança, conforme indicado em estudos sobre experiências em áreas como educação e finanças (PURNOMO et al., 2021; ODOGU, 2024).

Além disso, o estudo de Mou et al. (2024) aponta que, em um ciclo cumulativo, a confiança acumulada ao longo de múltiplas sessões de uso reforça a expectativa de conveniência futura, criando um efeito de “autoalimentação” onde cada interação bem-sucedida prepara o terreno para avaliações ainda mais positivas. Yin e Qiu (2021) confirmam que essa relação se fortalece em contextos de alta incerteza, como processos sensíveis de *compliance*, pois a confiança atua como um amortecedor contra temores de vazamento de informações e retaliações, possibilitando que o usuário avance com rapidez e segurança. Com base nesses argumentos, é proposta a seguinte hipótese:

**H1. A confiança tem um efeito direto positivo na conveniência percebida no uso de canais de denúncia habilitados por IA.**

Sacrifício Percebido: A percepção de sacrifício refere-se aos custos que o usuário atribui à interação com um sistema, incluindo esforço cognitivo, tempo despendido e riscos percebidos, como a exposição de dados pessoais ou a dependência tecnológica. Em serviços tradicionais, a presença humana media boa

parte da negociação de custos e benefícios, amortecendo barreiras iniciais (FØLSTAD et al., 2018). Contudo, em ambientes habilitados por IA, essa mediação é substituída por interfaces automatizadas que exigem maior adaptação e colaboração do usuário para interpretar respostas algorítmicas e gerenciar possíveis incertezas (HUANG, RUST, 2018). Quando os sacrifícios percebidos se elevam — por exemplo, frente à complexidade de fornecer informações sensíveis ou à necessidade de monitorar a precisão das recomendações —, a percepção de conveniência declina, pois o usuário passa a avaliar que o ganho de tempo ou de eficiência não compensa o desconforto ou o esforço adicional.

Em linha com essa lógica, Al-Okaily (2023) indicou que processos excessivamente demandantes reduzem o engajamento e minam o julgamento de utilidade de aplicações de IA. Sattu et al. (2024) corroboram essa relação, evidenciando que altos níveis de sacrifício percebido estão negativamente correlacionados com as avaliações de adoção e, por extensão, com indicadores de conveniência operacional. Adicionalmente, Rahman et al. (2021) e Arora et al. (2023) assinalam que, mesmo quando a qualidade do serviço é elevada, a sensação de vulnerabilidade — seja por receio de vazamento de dados ou de perda de controle — anula parte dos benefícios percebidos, gerando uma barreira psicológica ao uso fluido da plataforma.

Portanto, enuncia-se a seguinte hipótese:

**H2. O sacrifício percebido tem um efeito direto negativo na conveniência percebida no uso de canais de denúncia habilitados por IA.**

Ansiedade: Como já discorrido em seções anteriores, a resistência à tecnologia, que se manifesta como um desconforto ou apreensão em relação a computadores e seus usos (PARASURAMAN, IGBARIA, 1990), se intensifica com a IA.

Segundo Wang e Xiao (2025), a ansiedade induzida pela IA é multifacetada, sendo a violação da privacidade um fator predominante — principalmente entre indivíduos com diferentes níveis educacionais e gêneros; especificamente, níveis mais altos de ansiedade em relação à violação da privacidade foram observados em mulheres, em comparação com os homens.

Essa ansiedade pode ser agravada quando os usuários precisam compartilhar dados pessoais, pois o medo de uso indevido ou de acesso não autorizado a informações sensíveis intensifica a apreensão quanto à adoção de ferramentas de IA (KLEINE et al., 2023).

Pavlopoulos et al. (2024) apontam que esse temor atua como barreira central à adoção de soluções de IA, pois a percepção de risco à privacidade passa a sobrepor quaisquer benefícios funcionais ou de conveniência. Dessa forma, quanto maior o nível de ansiedade percebida pelo usuário, menor será sua avaliação da capacidade do sistema de IA em resguardar a confidencialidade e integridade de suas informações pessoais. Assim, é proposto:

**H3. A ansiedade tem um efeito direto negativo na privacidade no uso de canais de denúncia habilitados por IA.**

Conforme argumentado por Gado et al. (2021), a ansiedade elevada em relação à IA pode levar a comportamentos de evasão, resistência e desengajamento, especialmente em ambientes corporativos que adotam soluções automatizadas em processos sensíveis, como canais de denúncia. Nesses contextos, usuários ansiosos tendem a antecipar riscos, visualizar cenários de uso negativo e duvidar da eficácia ou imparcialidade dos sistemas, o que mina os fundamentos cognitivos e afetivos da confiança.

Wong et al. (2023) aprofundam essa discussão ao sugerirem que a percepção da IA como uma ameaça ao "eu psicológico" – seja em termos de autonomia, privacidade ou competência – intensifica a ansiedade e, por conseguinte, reduz significativamente a confiança nos sistemas. A percepção de ameaça ativa mecanismos defensivos que tornam os usuários mais cautelosos e menos propensos a delegar decisões ou confiar em interações automatizadas.

Adicionalmente, Karayianni (2025) confirma empiricamente que a ansiedade tecnológica não apenas afeta diretamente a confiança, mas também modera a relação entre confiança e intenção de uso, atuando como uma lente distorcida que contamina a avaliação racional e emocional das tecnologias. Dessa forma, mesmo em contextos nos quais a IA apresenta alto desempenho técnico ou forte reputação organizacional, a presença de ansiedade pode neutralizar ou reduzir o desenvolvimento de confiança.

Dessa forma, é proposta a seguinte hipótese:

**H4a. A ansiedade tem um efeito direto negativo na confiança no uso de canais de denúncia habilitados por IA.**

A ansiedade relacionada à tecnologia, especialmente em contextos envolvendo inteligência artificial (IA), exerce influência substancial sobre a forma como os indivíduos percebem os custos subjetivos associados ao uso dessas ferramentas. De acordo com Liu (2025), altos níveis de ansiedade em ambientes organizacionais tendem a intensificar a percepção de sacrifícios por parte dos usuários. Nesse cenário, indivíduos ansiosos passam a direcionar sua atenção de maneira desproporcional aos aspectos negativos da tecnologia — tais como complexidade operacional, perda de autonomia, risco de falhas ou consequências imprevistas — em detrimento da consideração dos benefícios proporcionados. Esse viés de atenção aos estressores reforça a sensação de esforço exigido para lidar com a tecnologia, configurando uma percepção de sacrifício mais acentuada.

Esse efeito é corroborado por Villanueva-Moya e Expósito (2023), que sustentam que os estados emocionais, particularmente a ansiedade, distorcem a análise subjetiva de custo-benefício realizada pelos indivíduos. Em contextos de ansiedade elevada, os custos percebidos tendem a ser superestimados, enquanto os benefícios são minimizados, o que compromete a racionalidade da avaliação e resulta em uma intensificação da percepção de sacrifício. Tal fenômeno pode ser particularmente problemático no caso da IA, dada a sua natureza opaca (black-box), que naturalmente gera sentimentos de incerteza e vulnerabilidade — catalisadores adicionais da ansiedade tecnológica.

Dessa forma, a literatura sugere que a ansiedade atua como um fator psicológico amplificador da percepção de sacrifícios, influenciando negativamente a experiência do usuário com sistemas de IA, sobretudo em contextos sensíveis como canais de denúncia, onde a percepção de risco já é elevada por natureza.

Com base nesses argumentos, formula-se a seguinte hipótese:

**H4b. A ansiedade tem um efeito direto positivo no sacrifício percebido no uso de canais de denúncia habilitados por IA.**

Qualidade do serviço habilitado pela IA: A confiança do usuário em sistemas de inteligência artificial está intimamente relacionada à percepção da qualidade dos serviços que esses sistemas oferecem. Tradicionalmente, qualidade de serviço é compreendida como o grau de discrepância entre as expectativas do consumidor e sua percepção real sobre a performance do serviço (PARASURAMAN et al., 1994). No contexto dos serviços mediados por IA, essa percepção torna-se ainda mais crítica, dado que a ausência de contato humano direto pode gerar insegurança ou ambiguidade quanto à confiabilidade da tecnologia.

Neste sentido, a qualidade percebida dos serviços de IA atua como um importante redutor de incertezas. Como afirmam Kerstan et al. (2023), a confiança na tecnologia depende diretamente do equilíbrio entre benefícios percebidos e riscos potenciais. Ou seja, quanto mais os usuários percebem que a IA oferece resultados consistentes, eficientes, personalizados e éticos, mais eles tendem a desenvolver confiança nesses sistemas. Essa confiança é fortalecida não apenas pela performance técnica, mas também pela capacidade da IA de operar de forma transparente, justa e imparcial (SARATCHANDRAN, 2019).

Adicionalmente, quando as organizações comunicam de maneira clara e eficaz os atributos positivos da IA — como sua precisão analítica, rapidez na resposta e capacidade de aprendizagem adaptativa — elas reduzem a ambivalência dos usuários e favorecem a construção de confiança (SATTU et al., 2024). Essa dinâmica é potencializada quando os usuários percebem que os sistemas estão alinhados com seus interesses e operam de modo confiável e previsível.

Assim, a qualidade do serviço habilitada por IA não apenas melhora a experiência do usuário, mas também fundamenta e fortalece a confiança, promovendo maior adesão às tecnologias e aceitação de sua aplicação em contextos sensíveis, como o atendimento automatizado em canais de denúncia. Desta forma:

**H5a. A qualidade do serviço habilitado por IA tem um efeito direto positivo na confiança no uso de canais de denúncia habilitados por IA.**

Por sua vez, a percepção de sacrifício — entendida como os custos ou esforços percebidos pelos usuários no uso de tecnologias, incluindo perda de controle, privacidade ou suporte humano — tende a ser atenuada quando a qualidade do serviço prestado por sistemas de IA é percebida como elevada. Isto é, quanto maior a percepção de qualidade do serviço habilitado por IA, menor a tendência de o indivíduo perceber o uso dessa tecnologia como um sacrifício.

Serviços de alta qualidade reduzem a frustração associada à ausência de interação humana, ao facilitar a navegação, reduzir o esforço cognitivo, aumentar a sensação de controle e garantir previsibilidade nos resultados (PARASURAMAN et al., 1994; HUANG, RUST, 2018). Nesse sentido, a qualidade percebida atua como um amortecedor das perdas subjetivamente atribuídas ao uso da IA, sobretudo em contextos sensíveis como o relato de denúncias, onde o suporte emocional humano poderia, em tese, ser insubstituível.

Além disso, ao proporcionar experiências mais eficazes, seguras e satisfatórias, os sistemas de IA tendem a mitigar preocupações ligadas à privacidade e ao uso indevido de dados — especialmente se houver clareza quanto aos mecanismos de segurança e à ética algorítmica envolvida (CHEATHAM et al., 2019). Quando os usuários percebem que os serviços são bem projetados, justos e funcionam como prometido, há uma redução significativa da sensação de exposição, vulnerabilidade e esforço excessivo.

Sattu et al. (2024) reforçam essa lógica ao indicarem que uma comunicação transparente sobre os atributos da IA — como confiabilidade e valor agregado — reduz a percepção de custo emocional e cognitivo associado ao seu uso. Desta forma, acredita-se que a qualidade do serviço habilitada por IA diminua a percepção de sacrifício, ao reduzir fricções, aumentar a previsibilidade e alinhar expectativas do usuário com os resultados concretos do serviço. Sendo assim:

**H5b. A qualidade do serviço habilitado por IA tem um efeito direto negativo no sacrifício percebido no uso de canais de denúncia habilitados por IA.**

Privacidade: No contexto da adoção de tecnologias baseadas em inteligência artificial (IA), a privacidade desponta como um fator crítico que influencia negativamente a percepção de conveniência. Essa relação ocorre, em



grande parte, devido ao *trade-off* entre proteção de dados pessoais e usabilidade tecnológica. Embora os sistemas de IA frequentemente promovam ganhos substanciais em termos de funcionalidade, automação e personalização, esses benefícios muitas vezes exigem, como contrapartida, a coleta intensiva de informações pessoais dos usuários — o que gera tensões em torno da privacidade (CLOAREC, 2020; MARTIN et al., 2017).

A teoria do cálculo da privacidade (*Privacy Calculus Theory*), conforme discutida por Princi e Krämer (2020), sustenta que os indivíduos realizam uma avaliação deliberada dos benefícios esperados (como conveniência, eficiência ou personalização) versus os custos percebidos (notadamente, riscos à privacidade e à segurança). Quando os riscos à privacidade são percebidos como elevados — por exemplo, pela coleta de dados sensíveis, pela ausência de transparência nos algoritmos ou pelo medo de uso indevido da informação — há uma tendência à retração na adoção da tecnologia, bem como à reavaliação da utilidade percebida dos serviços automatizados. Assim, a valorização da privacidade, ainda que legítima, atua como um redutor da percepção de conveniência, pois limita a fluidez, a agilidade e a personalização que os sistemas de IA poderiam proporcionar.

Nesse sentido, Lim e Shim (2022) discutem que, apesar da promessa de utilidade elevada oferecida por soluções baseadas em IA, a intensificação das preocupações com privacidade gera hesitação e incômodo nos usuários, especialmente quando a coleta de dados é percebida como invasiva, pouco transparente ou incontrolável. Tal hesitação compromete a fluidez da experiência de uso, minando a sensação de comodidade e eficiência — atributos centrais à noção de conveniência percebida.

Além disso, mesmo em contextos em que os usuários inicialmente priorizam a conveniência, como observam Afridi e Kumar (2024), experiências negativas com falhas de segurança ou vazamentos de dados tendem a alterar esse equilíbrio, levando a uma reavaliação retroativa da decisão de uso e, conseqüentemente, a uma queda na percepção de conveniência. Ou seja, a conveniência pode ser percebida como frágil ou ilusória quando o risco à privacidade é evidenciado ou experienciado diretamente.

Por fim, vale notar que o entendimento da privacidade como um direito fundamental (Skinner et al., 2006) reforça sua centralidade nos julgamentos dos

usuários. A ameaça a esse direito básico pode gerar sentimentos de vulnerabilidade e desconfiança, que afetam negativamente a experiência com a tecnologia, interferindo diretamente na percepção de que o uso da IA seja fácil, útil ou benéfico — elementos fundamentais da conveniência percebida.

Dessa forma, uma vez que preocupações com a privacidade introduzem barreiras cognitivas, emocionais e éticas que limitam o engajamento espontâneo, contínuo e despreocupado com tecnologias baseadas em IA, propõe-se que:

**H6. O risco de privacidade tem um efeito direto negativo na conveniência percebida no uso de IA.**

### **3**

## **Metodologia da Pesquisa**

Neste capítulo são apresentados os procedimentos metodológicos que foram utilizados para testar as hipóteses propostas na seção anterior. A seguir, serão descritos o tipo de pesquisa, escalas de mensuração de cada variável, o processo de amostragem, o procedimento de coleta de dados, a análise dos dados e as limitações do método.

### **3.1. Tipo de pesquisa**

Com o objetivo de testar as hipóteses propostas no estudo, foi realizada uma *cross-sectional survey* (PARASURAMAN et al., 2006), por meio de questionários auto-administrados via internet (AAKER et al., 2006), com uma amostra não probabilística da população de interesse.

### **3.2. Operacionalização e definição das variáveis**

O presente estudo faz uso de escalas já elaboradas e testadas pela literatura para a medição de todos os construtos envolvidos. O instrumento de pesquisa (Apêndice A) desenvolvido neste trabalho operacionaliza a medição dos seguintes construtos: confiança, sacrifício percebido, qualidade do serviço habilitado pela IA, ansiedade e privacidade e conveniência.

Foi utilizada uma escala do tipo Likert de cinco pontos, variando entre “discordo totalmente” e “concordo totalmente”. Essa escala demanda que os respondentes indiquem de 1 a 5 conforme a sua percepção do grau de concordância ou discordância com as afirmativas relacionadas a cada construto apresentado no questionário. Neste caso, 1 (um) significa a discordância total em relação a afirmativa, enquanto 5 (cinco) significa que o participante concorda totalmente com a frase apresentada e o 3 é considerado neutro, “não concorda nem discorda”.

### **3.2.1. Definição operacional das variáveis**

Essa pesquisa foi elaborada a partir dos modelos propostos por Wolfinbarger e Gilly (2003) e Chang e Wang, (2011) que tratam da qualidade dos serviços habilitados por IA; Wang et al. (2019) e Paparoidamis, Tran, e Leonidou, (2019) que tratam da confiança, Collier Sherrell (2010) que tratam da conveniência; Merisavo et al. (2007) Shin e Lin (2016) que tratam do sacrifício percebido; Xu et al. (2011) e Lee e Rha (2016) que abordam a privacidade e Compeau e Higgins (1999) que abordam a ansiedade.

Foi usada a escala Likert de 5 pontos variando entre concordo completamente e discordo completamente. As escalas já testadas pela literatura utilizadas no estudo, assim como as referências que as originaram, estão ilustradas no Quadro 3-1.

### **3.2.2. Procedimentos de Tradução e Adaptação das Escalas Utilizadas**

As escalas utilizadas no estudo foram originalmente desenvolvidas para a língua inglesa. Foi necessária a tradução e adaptação de cada escala para a língua portuguesa para que a pesquisa pudesse ser aplicada a respondentes brasileiros.

Como sugere o autor Sperber (2004), a adaptação de escalas para outras línguas deve passar por diversas etapas. Os itens do questionário proposto que estavam em inglês no texto original foram traduzidos pela autora e por profissional fluente em inglês independentemente, com objetivo de comparar as versões e obter a versão em português mais próxima da original. Os itens do questionário que já estavam no idioma português as questões sobre o constructo da ansiedade e da privacidade, foram apenas replicados no estudo.

| CONSTRUCTOS                             | ITENS DA ESCALA                                                                                                                                                                                                                                                                                                                                                                                                                                                                     | REFERÊNCIAS                                                |
|-----------------------------------------|-------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------------|------------------------------------------------------------|
| QUALIDADE DO SERVIÇO HABILITADA PELA IA | <p>A IA utilizada no canal de denúncia foi bem projetada.</p> <p>A IA utilizada no canal de denúncia fornece informações corretas.</p> <p>A IA utilizada no canal de denúncia é segura.</p> <p>A IA utilizada no canal de denúncia é atenciosa.</p>                                                                                                                                                                                                                                 | Wolfinbarger e Gilly (2003) e Chang e Wang, (2011)         |
| CONFIANÇA                               | <p>O desempenho da IA no canal de denúncia sempre atende às minhas expectativas.</p> <p>A IA do canal de denúncia possui bons recursos.</p> <p>A IA do canal de denúncia é confiável.</p> <p>A IA utilizada no canal de denúncia é honesta.</p> <p>Eu acredito na IA utilizada no canal de denúncia.</p> <p>Ao disponibilizar uma IA no canal de denúncia a empresa demonstra interesse em mim como cliente.</p>                                                                    | Wang et al. (2019) e Paparoidamis, Tran e Leonidou, (2019) |
| SACRIFÍCIO PERCEBIDO                    | <p>Fico preocupado(a) com a perda de controle ao usar a IA no canal de denúncia.</p> <p>Usar a IA no canal de denúncia consome muito tempo.</p> <p>Fico irritado(a) ao usar a IA no canal de denúncia.</p> <p>Eu fico preocupado(a) com a perda de privacidade ao usar a IA no canal de denúncia.</p> <p>Usar a IA no canal de denúncia um exige esforço adicional.</p> <p>Eu fico preocupado(a) com a falta de interação humana ao usar a IA no canal de denúncia.</p>             | Merisavo et al. (2007) Shin e Lin (2016)                   |
| ANSIEDADE                               | <p>Eu me sinto apreensivo em usar soluções de IA.</p> <p>Me assusta pensar que eu poderia perder uma grande quantidade de informação por ter utilizado a IA do canal de denúncias de forma incorreta, por exemplo se eu não for claro o suficiente.</p> <p>Eu hesito em usar uma solução de canal de denúncia que utilize IA por medo de erros que eu não conseguirei corrigir.</p> <p>Soluções de canal de denúncia que utilizam IA são de alguma forma intimidantes para mim.</p> | Compeau e Higgins (1999)                                   |
| PRIVACIDADE                             | <p>Preocupa-me que a IA do canal de denúncia partilhe as minhas informações pessoais de forma inadequada/não autorizada.</p> <p>Receio que IA do canal de denúncia resulte em uma invasão da minha privacidade.</p> <p>Preocupa-me a forma como as minhas ações enquanto usuário da IA do canal de denúncia possa ser controlada e monitorada.</p>                                                                                                                                  | Xu et al. (2011) e Lee e Rha (2016)                        |
| CONVENIÊNCIA                            | <p>A IA do canal de denúncia me permite usar o serviço sempre que eu escolher.</p> <p>A IA do canal de denúncia me permite usar o serviço em um momento conveniente.</p> <p>Eu valorizo a possibilidade de usar a IA do canal de denúncia no conforto de casa.</p>                                                                                                                                                                                                                  | Collier e Sherrell (2010)                                  |

**Quadro 3.1 - Definições operacionais das variáveis**

### **3.2.3. Pré-teste do instrumento de pesquisa**

Para a validação da tradução e do questionário, foi realizado um pré-teste com uma pequena amostra do público de interesse. Foi solicitado a 10 indivíduos que sinalizassem as dúvidas na compreensão do questionário ou que indicassem sugestões. As considerações apontadas foram ajustadas e um novo teste foi realizado com 10 indivíduos. A partir do resultado foi elaborado um questionário final que está apresentado no Apêndice A.

## **3.3. População e amostra**

### **3.3.1. População**

A população de interesse é formada por jovens e adultos brasileiros maiores de 18 anos que já utilizaram algum *chatbot* de inteligência artificial durante sua vida adulta.

### **3.3.2. Amostra**

A amostra estudada foi composta por brasileiros adultos de diversas regiões do país que já utilizaram alguma ferramenta de inteligência artificial ao longo de sua vida adulta.

O link para o questionário foi compartilhado em grupos de Whatsapp de *compliance* e LinkedIn. Ao todo, foram recebidas 235 respostas entre os dias 12 de março e 6 de abril de 2025. Destas algumas foram descartadas por se tratar de pessoas que nunca usaram uma ferramenta de IA ou ainda, que não compreenderam o que estava sendo proposto. Sendo assim, foram recebidos 214 formulários válidos.

## **3.4. Coleta de dados**

### **3.4.1. O instrumento de coleta de dados**

O instrumento de pesquisa foi um questionário estruturado (Apêndice A), composto por 26 afirmações referentes às variáveis: conveniência, confiança, sacrifício percebido, conveniência, qualidade do serviço habilitado pela IA,

ansiedade e privacidade. Além dessas dimensões, o questionário também apresentava itens para avaliar as variáveis demográficas dos respondentes e os hábitos de uso de ferramentas de inteligência artificial como ChatGPT, Gemini e outros semelhantes, proporcionando uma melhor compreensão da amostra envolvida na pesquisa.

### 3.4.2

#### A coleta de dados

O questionário da pesquisa foi composto por 33 itens: 3 itens abordavam a experiência com IA e o entendimento dos *prints* apresentados, 4 abordavam aspectos demográficos.

O entendimento dos *prints* foi essencial para o correto preenchimento do questionário. O objetivo era simular uma denúncia em uma plataforma de canal de denúncia que utiliza inteligência artificial (IA), proporcionando ao respondente uma experiência do serviço. Durante essa simulação, o participante pôde acompanhar cada interação do denunciante com a IA, bem como observar o comportamento da tecnologia e suas reações ao longo da observação das telas.

Após a análise das telas (cerca de 11 *prints*), os participantes responderam a uma pergunta de checagem, cujo intuito era verificar o grau de compreensão sobre o que haviam acabado de avaliar. Cinco respondentes indicaram não ter compreendido os *prints* e, por esse motivo, seus dados foram excluídos do presente estudo.

Os respondentes foram então direcionados para o questionário principal, o qual abordava os construtos investigados e a variável dependente. Foram feitas 4 perguntas para Qualidade, 6 para Confiança, 6 para o Sacrifício percebido, 4 para Ansiedade, 3 para Privacidade e 3 para Conveniência.

O questionário foi estruturado iniciando com perguntas que avaliavam os hábitos de uso de ferramentas de inteligência artificial. Em seguida foram apresentados os itens que medem os construtos propostos no modelo e por fim perguntas sobre características demográficas.

Os dados foram coletados por meio de uma *survey* online realizada em Março e Abril de 2025. O link para a pesquisa foi divulgado no LinkedIn e em grupos de discussão no WhatsApp relacionados ao tema de *compliance*

corporativo. O questionário era autoadministrado e os participantes foram voluntários interessados no tema.

### **3.5. Análise de dados**

Após a aplicação de todos os questionários, os dados foram reunidos para processamento estatístico em uma base de dados do SPSS (versão 23) e AMOS (versão 30), com o intuito de obter as análises univariadas e multivariadas dos dados obtidos.

A primeira etapa realizada foi uma limpeza dos dados para identificar e excluir informações ausentes e dados errôneos que não deveriam ser considerados na análise, o que resultou na eliminação de 21 registros. Dessa forma, a amostra válida totalizou 214 respostas.

Também foi feita uma análise das variáveis demográficas presentes no questionário para caracterizar a amostra pesquisada.

#### **3.5.1. Validade e confiabilidade**

Para o desenvolvimento do modelo de mensuração e a avaliação dos construtos empregados no instrumento de pesquisa, especialmente quanto à confiabilidade e validade, foi inicialmente realizada uma análise fatorial confirmatória (CFA) com os dados coletados. Conforme Hair et al. (2009) e Churchill (1979), a validade de um construto refere-se ao grau em que uma escala representa fielmente o construto latente que se propõe a medir, enquanto a confiabilidade diz respeito à consistência com que uma variável ou um conjunto de variáveis mede aquilo a que se destina.

A verificação da confiabilidade dos construtos foi feita por meio da confiabilidade composta (Composite Reliability – CR). De acordo com Hair et al. (2009) e Nunnally e Bernstein (1994), valores superiores a 0,8 são considerados satisfatórios, embora valores entre 0,7 e 0,8 também sejam aceitos como adequados.

No que diz respeito à validade dos construtos, foi utilizada a Variância Extraída Média (Average Variance Extracted – AVE) como medida da validade convergente. Segundo Hair et al. (2009), essa métrica expressa o quanto os diferentes indicadores de uma escala se relacionam com o mesmo construto.



Fornell e Larcker (1981), assim como Hair et al. (2009), recomendam uma AVE superior a 0,5 como indicativo de uma validade convergente satisfatória. A validade nomológica foi avaliada com o uso de uma matriz de correlação entre construtos, buscando-se verificar se os construtos se relacionam entre si conforme previsto pela teoria.

Posteriormente, procedeu-se à avaliação da validade discriminante, que demonstra o grau de distinção entre os construtos (HAIR et al., 2009). Para isso, comparou-se a AVE de cada construto com a variância compartilhada entre os pares de construtos, obtida pelo quadrado do coeficiente de correlação, conforme a proposta de Fornell e Larcker (1981). A validade discriminante é confirmada quando a AVE de cada construto é maior do que suas respectivas variâncias compartilhadas.

### 3.6. Análises estatísticas

As hipóteses formuladas neste estudo foram testadas por meio da modelagem de equações estruturais (SEM), utilizando-se o software AMOS. Essa técnica foi considerada adequada para responder às perguntas de pesquisa e testar as hipóteses propostas, pois permite estimar, de forma simultânea, as relações entre múltiplas variáveis independentes e dependentes (BAGOZZI; PHILLIP, 1982). A aplicação da SEM possibilita uma avaliação conjunta dos efeitos dos construtos no modelo proposto, evitando possíveis distorções que poderiam ocorrer caso as dependências entre as variáveis fossem analisadas isoladamente (HAIR et al., 2009).

Para verificar o ajuste dos modelos CFA e SEM, foram adotados índices comumente utilizados em pesquisas similares (CHANG & CHEN, 2009). Entre os índices considerados estão: o *Tucker-Lewis Index* (TLI), o *Comparative Fit Index* (CFI), o *Root Mean Squared Approximation Error* (RMSEA) e a estatística qui-quadrado do modelo (acompanhada dos respectivos graus de liberdade). Excetuando-se a estatística qui-quadrado, todos esses índices apresentam interpretação facilitada por estarem situados em uma escala contínua de 0 a 1 e por serem relativamente pouco influenciados pelo tamanho da amostra.

### **3.7. Limitações do método**

#### **3.7.1. Limitações relacionadas ao critério de amostragem**

A delimitação amostral consiste em brasileiros que afirmaram que compreenderam os *prints* apresentados e já utilizaram ferramentas com IA.

A amostra apresentou concentração de renda familiar estimada alta comparada com a média da população brasileira e escolaridade, a maioria dos respondentes possuíam Graduação, mestrado ou doutorado completos. Esses vieses devem ser considerados no resultado da pesquisa.

#### **3.7.2. Limitações decorrentes da coleta de dados**

As limitações, decorrentes da coleta de dados por meio de uma *survey* online, dizem respeito principalmente à atenção do respondente durante a participação na pesquisa, que levou cerca de 11 minutos para ser concluída.

Outra limitação observada refere-se à compreensão, por parte dos respondentes, das explicações previamente disponibilizadas por meio de textos e capturas de tela acerca da ferramenta analisada.

Por fim, apesar dos prints serem parte integrante do formulário de pesquisa e terem passado por testes de clareza durante a etapa de pré-teste, os respondentes da versão final do questionário podem não ter o completo entendimento acerca do uso da inteligência artificial no contexto da ferramenta de canal de denúncia apresentada.

## 4

### Modelagem e Análise de dados

Neste capítulo são avaliadas as propriedades estatísticas e métricas da amostra coletada, assim como os modelos de mensuração e estruturais propostos e também os resultados dos testes das hipóteses de pesquisa.

#### 4.1. Caracterização da amostra

O instrumento de pesquisa (Apêndice A) foi compartilhado em grupos de WhatsApp e no LinkedIn. O processo de coleta de dados foi realizado ao longo do mês de março e abril de 2025 e a participação dos respondentes foi voluntária e confidencial. Era esperado que os respondentes tivessem experiência com inteligência artificial (ChatGPT, Gemini e semelhantes) e entendessem os *prints* que estavam sendo apresentados em tela, como um exemplo do que a pesquisa iria tratar. Por essa razão, foram considerados apenas respondentes que afirmaram que entenderam os *prints* e já haviam usado alguma ferramenta de IA anteriormente, resultando em uma amostra final de 214 respostas válidas.

A Tabela 4.1 ilustra as características da amostra final segundo algumas estatísticas descritivas. Do total de 214 participantes, 127 pertencem ao gênero feminino (59,3%) e 87 ao gênero masculino (40,7%). A amostra apresentou uma grande variação em relação à diferença de idade, 18 a 85 anos com um desvio padrão máximo de 13,87, sendo a média de idade 44 anos. No que diz respeito à renda média familiar, 64,5% indicaram renda familiar mensal superior a dez salários-mínimos.

Considerando a experiência anterior com IA, do total de 214 respondentes, 42,5% (91) responderam que usam IA diariamente. 64 respondentes (29,9%) já fizeram denúncia em um canal de denúncias de corporativo e 139 (65%) nunca haviam feito, o que reforça a importância da compreensão dos *prints* apresentados durante a pesquisa.

**Tabela 4.1 - Características da Amostra**

| <b>GÊNERO</b>                                                             |       |     |
|---------------------------------------------------------------------------|-------|-----|
| Feminino                                                                  | 59,3% | 127 |
| Masculino                                                                 | 40,7% | 87  |
| <b>IDADE</b>                                                              |       |     |
| Média                                                                     | 44,7  |     |
| Desvio Padrão                                                             | 13,87 |     |
| Máximo                                                                    | 85    |     |
| Mínimo                                                                    | 18    |     |
| <b>ESCOLARIDADE</b>                                                       |       |     |
| Ensino Fundamental                                                        | 0,5%  | 1   |
| Ensino Médio                                                              | 2,8%  | 6   |
| Graduação                                                                 | 57,5% | 123 |
| Mestrado                                                                  | 35,5% | 76  |
| Doutorado                                                                 | 3,7%  | 8   |
| <b>RENDA FAMILIAR ESTIMADA</b>                                            |       |     |
| Acima de 10 salários-mínimos                                              | 64,5% | 138 |
| De 5 a 10 salários-mínimos                                                | 19,6% | 42  |
| De 2 a 5 salários-mínimos                                                 | 11,7% | 25  |
| De 1 a 2 salários-mínimos                                                 | 3,3%  | 7   |
| Até 1 salário-mínimo                                                      | 0,9%  | 2   |
| <b>EXPERIÊNCIA ANTERIOR COM IA</b>                                        |       |     |
| Usa Diariamente                                                           | 42,5% | 91  |
| Usa Semanalmente                                                          | 27,6% | 59  |
| Usa Mensalmente                                                           | 7,0%  | 15  |
| Usa Raramente                                                             | 13,6% | 29  |
| Nunca usou                                                                | 9,3%  | 20  |
| <b>EXPERIÊNCIA COM REALIZAÇÃO DE DENÚNCIAS</b>                            |       |     |
| Já fiz uma denúncia em um canal de<br>denúncia antes                      | 29,9% | 64  |
| Não me recordo se já fiz uma<br>denúncia em um canal de denúncia<br>antes | 5,1%  | 11  |
| Nunca fiz uma denúncia em um canal<br>de denúncia de uma empresa          | 65,0% | 139 |

#### **4.1.1. Análise da Variância de Método Comum**

De acordo com Podsakoff et al. (2003), variância comum do método (CMV) é a variância que está associada ao método de mensuração, e não aos construtos que as medidas representam. Para os autores, pesquisas nas quais os dados, tanto para a variável preditora quanto para a variável dependente, são obtidos da mesma pessoa, no mesmo contexto de mensuração, utilizando o mesmo contexto de itens e características similares nos itens, estariam suscetíveis ocorrência da CMV, o que poderia ocasionar erros sistemáticos de mensuração, inflacionando ou reduzindo as relações observadas entre os construtos.

Portanto, o teste de um-fator ou fator único de Harman (Podsakoff & Organ, 1986) foi utilizado, com o propósito de examinar até que ponto esse viés estava presente nos dados coletados no presente estudo. Os resultados da análise de componentes principais indicaram que nenhum dos fatores representava quase toda a variação (o fator que explicava mais conseguiu capturar apenas 36,31% da variação total).

Além disso, também foi utilizado o método do fator latente comum (*Common Latent Factor Method*), conforme descrito por Podsakoff et al. (2003) para avaliar a ocorrência de CMV. Os resultados mostraram que, ao incluir no modelo uma variável latente refletindo o método comum, a variância correspondente à variável latente do método comum foi apenas 15,19% da variância total.

Conclui-se então, conforme descrito por Podsakoff e Organ (1986) e Podsakoff et al. (2003), que não existem evidências de problemas relacionados à Variância Comum do Método para o presente estudo.

#### **4.2. Análises e resultados**

##### **4.2.1. Avaliação do modelo de mensuração**

No modelo inicial testado, com 26 indicadores, os índices de ajuste encontrados ficaram aquém dos valores recomendados pela literatura. Uma análise da matriz de covariância dos resíduos padronizados da CFA indicou diversos itens que poderiam estar contribuindo para ajuste do modelo abaixo do sugerido. A partir desse resultado, o modelo foi refinado e ajustado, levando à

eliminação de alguns itens que não estavam se adequando à estrutura dos construtos. Assim, foram eliminados 4 indicadores em três escalas diferentes, todas as outras permaneceram com os itens propostos originalmente. Os itens eliminados foram: o item 3 de Confiança, o item 1 de Ansiedade e os itens 1 e 4 de Sacrifício Percebido.

O modelo de mensuração final com 22 indicadores apresentou, então, bons índices de ajuste (SRMR = 0,054; RMSEA=0,056 com C.I. de 0,045 até 0,066; CFI= 0,944; IFI=0,945; TLI= 0,933;  $\chi^2=321,048$ , d.f.= 193,  $p<0,001$ ,  $\chi^2/d.f. = 1,663$ ), representando melhoria significativa em relação ao modelo inicial. Em conjunto, estes índices apontam que os dados se ajustam de maneira satisfatória ao modelo proposto (HU, BENTLER, 1999; SCHREIBER et al., 2006).

#### **4.2.2. Validade e confiabilidade dos construtos**

Segundo Hair et al., (2009), a validade de um construto é composta de quatro componentes: validade de face, validade nomológica, validade convergente e validade discriminante.

A validade da face pode ser definida como a consistência entre o conteúdo de cada item em relação ao construto por ele medido. No presente estudo, essa validade foi garantida durante o desenvolvimento do instrumento de pesquisa através da escolha de escalas utilizadas anteriormente na literatura, da tradução criteriosa dessas escalas do inglês para o português, da avaliação de cada item por pesquisadores experientes de comportamento do consumidor e dos pré-testes conduzidos com pequenas amostras da população de interesse.

A validade nomológica examina se as correlações entre os construtos da teoria de mensuração aplicada fazem sentido. A validade nomológica pode ser avaliada com o uso de uma matriz de correlação entre construtos em que é possível verificar se os construtos se relacionam entre si conforme previsto pela teoria. O Quadro 4.1 apresenta a matriz de correlação dos construtos estudados.

**Quadro 4.1 - Matriz de Correlação entre Construtos**

|                         | (1)   | (2)   | (3)   | (4)   | (5)   | (6)   |
|-------------------------|-------|-------|-------|-------|-------|-------|
| <b>Conveniência (1)</b> | 1     | -0,51 | 0,51  | 0,58  | -0,36 | -0,62 |
| <b>Ansiedade (2)</b>    | -0,51 | 1     | -0,44 | -0,46 | 0,45  | 0,61  |
| <b>Qualidade (3)</b>    | 0,51  | -0,44 | 1     | 0,63  | -0,41 | -0,48 |
| <b>Confiança (4)</b>    | 0,58  | -0,46 | 0,63  | 1     | -0,42 | -0,55 |
| <b>Privacidade (5)</b>  | -0,36 | 0,45  | -0,41 | -0,42 | 1     | 0,54  |
| <b>Sacrifício (6)</b>   | -0,62 | 0,61  | -0,48 | -0,55 | 0,54  | 1     |

Nota: Todas as correlações foram significativas ( $p < 0,01$ ).

Considerando que todas as correlações observadas estão alinhadas com a teoria utilizada (correlações positivas eram esperadas entre os construtos Qualidade/Confiança/Conveniência e negativas destes com os construtos Ansiedade/Privacidade/Sacrifício, com esses três últimos devendo apresentar correlação positiva entre si), pode-se concluir que os construtos empregados no estudo apresentam validade nomológica.

A fim de avaliar a consistência interna e a confiabilidade das escalas, foi adotada a medida de confiabilidade composta, a qual expressa o grau de consistência interna entre os indicadores que mensuram um mesmo fator (FORNELL & LARCKER, 1981). De acordo com os autores, valores superiores a 0,7 são considerados adequados para essa medida. Conforme evidenciado na Tabela 4.2, todas as escalas analisadas atendem aos critérios mínimos estabelecidos na literatura, apresentando valores de confiabilidade composta superiores a 0,7.

No que diz respeito à validade convergente, esta foi examinada por meio do cálculo da variância extraída média (*Average Variance Extracted* – AVE) para cada construto. Segundo Fornell e Larcker (1981), valores de AVE superiores a 0,50 são indicativos de validade convergente satisfatória. Os resultados obtidos nesta pesquisa estão disponíveis na Tabela 4.2, e todos os valores de AVE calculados situam-se entre 0,51 e 0,76, superando o patamar recomendado e comprovando, assim, a validade convergente das escalas utilizadas.

**Tabela 4.2 - Confiabilidade Composta e Variância Extraída Média (AVE)**

| <b>Escala</b>       | <b>Confiabilidade Composta</b> | <b>AVE</b> |
|---------------------|--------------------------------|------------|
| <b>Conveniência</b> | 0.81                           | 0,58       |
| <b>Ansiedade</b>    | 0.78                           | 0,54       |
| <b>Qualidade</b>    | 0.73                           | 0,51       |
| <b>Confiança</b>    | 0.83                           | 0,53       |
| <b>Privacidade</b>  | 0.89                           | 0,73       |
| <b>Sacrifício</b>   | 0.82                           | 0,54       |

A validade discriminante é confirmada quando todos os construtos apresentam variâncias extraídas superiores às respectivas variâncias compartilhadas (FORNELL & LARCKER, 1981), o que indica que os itens se relacionam mais fortemente com os construtos aos quais pertencem do que com os demais construtos do modelo. Para essa avaliação, Fornell e Larcker (1981) recomendam a comparação entre a variância extraída média (AVE) de cada construto e a variância compartilhada (obtida pelo quadrado do coeficiente de correlação) entre todos os pares de construtos.

A matriz utilizada para a análise da validade discriminante é apresentada no Quadro 4.2, em que a diagonal principal contém os valores de AVE de cada construto, enquanto as demais células mostram o quadrado dos coeficientes de correlação entre os pares de construtos. A partir da observação da matriz, verifica-se que, em todos os casos, as variâncias extraídas são superiores às variâncias compartilhadas, indicando a validade discriminante adequada dos construtos analisados.



**Quadro 4.2 - Matriz de Validade Discriminante**

|                         | (1)  | (2)  | (3)  | (4)  | (5)  | (6)  |
|-------------------------|------|------|------|------|------|------|
| <b>Conveniência (1)</b> | 0,58 | 0,26 | 0,26 | 0,34 | 0,13 | 0,38 |
| <b>Ansiedade (2)</b>    | 0,26 | 0,54 | 0,19 | 0,21 | 0,20 | 0,37 |
| <b>Qualidade (3)</b>    | 0,26 | 0,19 | 0,51 | 0,40 | 0,17 | 0,23 |
| <b>Confiança (4)</b>    | 0,34 | 0,21 | 0,40 | 0,53 | 0,18 | 0,3  |
| <b>Privacidade (5)</b>  | 0,13 | 0,2  | 0,17 | 0,18 | 0,73 | 0,29 |
| <b>Sacrifício (6)</b>   | 0,38 | 0,37 | 0,23 | 0,3  | 0,29 | 0,54 |

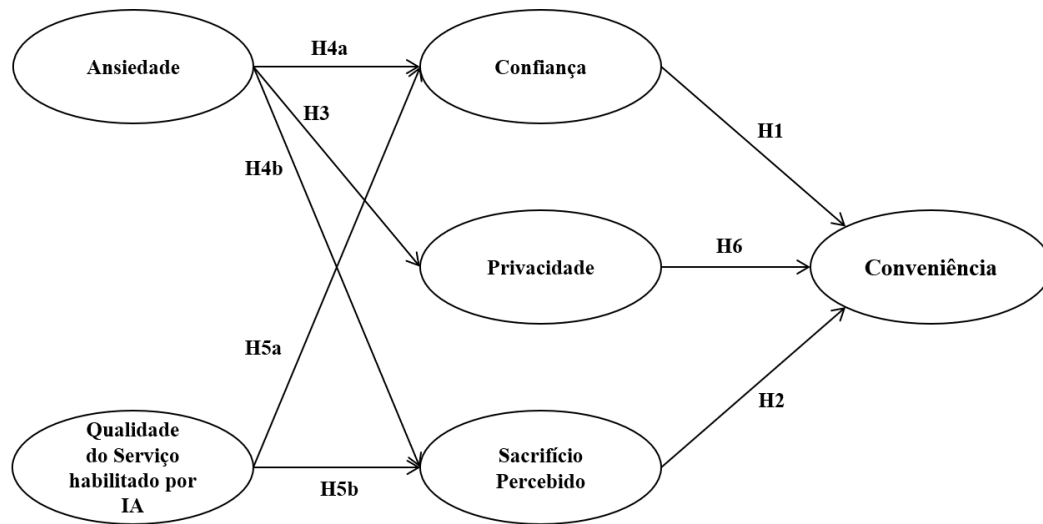
Dessa forma, a análise conjunta dos resultados obtidos na análise fatorial confirmatória indica que o modelo de mensuração proposto satisfaz os critérios esperados de confiabilidade, unidimensionalidade, validade de face, validade nomológica, validade convergente e validade discriminante. Com isso, torna-se viável a etapa seguinte do estudo: a investigação das relações entre os construtos latentes por meio da aplicação do modelo estrutural.

#### **4.2.3. Análise do modelo estrutural**

Para testar o modelo conceitual proposto e as hipóteses de pesquisa foi utilizada pela técnica de modelagem de equações estruturais (SEM), através do uso do software AMOS 30. Na modelagem de equações estruturais, a significância dos coeficientes estimados para as relações do modelo indica se cada hipótese formulada sobre as relações entre os construtos é confirmada ou não (BYRNE, 2010). Por fim, realiza-se a modelagem da estrutura proposta e sua adequação é avaliada com base nos índices de ajuste pertinentes.

#### 4.2.3.1. Ajuste no modelo proposto

O modelo proposto (Figura 4.1) foi avaliado por meio de diferentes índices de ajuste, conforme recomendado por Garver e Mentzer (1999) e Hair et al. (2009).



**Figura 4.1 - Modelo Proposto**

A estatística qui-quadrado obtida para o modelo foi estatisticamente significativa ( $\chi^2 = 272,124$ , d.f. = 161,  $p < 0,001$ ). Todavia, como sugerem alguns autores, é necessário que a interpretação seja realizada com cautela, pois é sensível ao tamanho da amostra e a violações da premissa de normalidade (BENTLER, 1990; HOOPER et al., 2008). Todos os outros índices indicaram um bom ajuste do modelo aos dados. A razão  $\chi^2/\text{d.f.}$  foi de 1,69, inferior ao valor de 3,0 sugerido por Byrne (2010). O SRMR (standardized root mean square residual) foi de 0,078, enquanto o RMSEA (root mean square error of approximation) foi de 0,057 (C. I. de 0,045 até 0,068), ambos abaixo de 0,08, conforme recomendado pela literatura (HU & BENTLER, 1999). Os índices de ajuste incrementais ficaram acima do patamar de 0,90 (BENTLER & BONNET, 1980), com um CFI (comparative fit index) de 0,944, um TLI (Tucker-Lewis index) de 0,934 e um IFI (incremental fit index) de 0,945. Estes resultados encontram-se resumidos na Tabela 4.3. Considerando-se os índices apresentados, pode-se concluir que o ajuste do modelo proposto é satisfatório.

**Tabela 4.3 - Índices de Ajuste do Modelo**

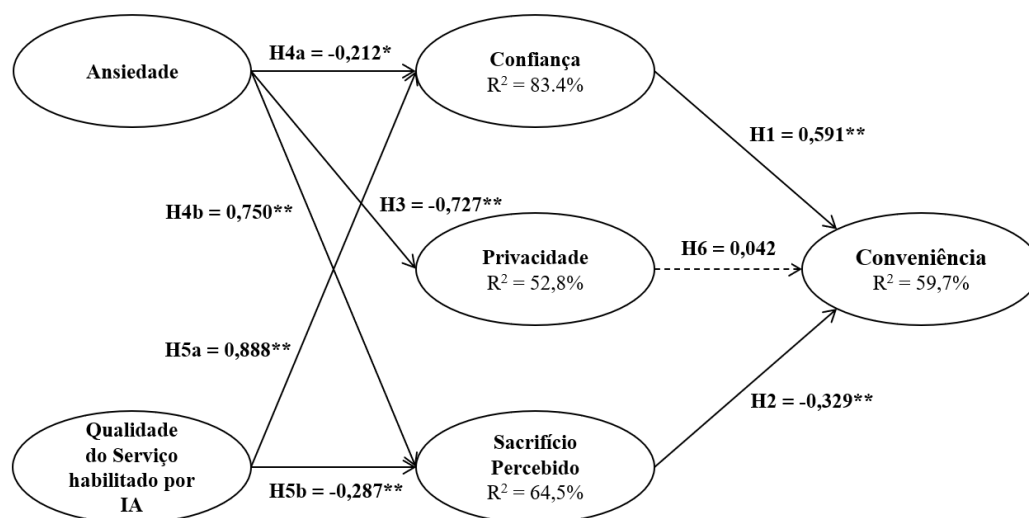
| <b>Índice de Ajuste</b> | <b>Modelo Proposto</b> | <b>Valor sugerido pela literatura</b> |
|-------------------------|------------------------|---------------------------------------|
| $\chi^2/\text{d.f.}$    | 1,69                   | $\leq 3$                              |
| <b>CFI</b>              | 0,944                  | $\geq 0,90$                           |
| <b>TLI</b>              | 0,934                  | $\geq 0,90$                           |
| <b>IFI</b>              | 0,945                  | $\geq 0,90$                           |
| <b>RMSEA</b>            | 0,067                  | $\leq 0,08$                           |
| <b>SRMR</b>             | 0,078                  | $\leq 0,08$                           |

#### **4.2.4. Teste das hipóteses de pesquisa**

A verificação das hipóteses de pesquisa foi conduzida por meio da análise da magnitude, direção e significância dos coeficientes padronizados estimados no modelo estrutural (BYRNE, 2010). Para que uma relação fosse considerada significativa, o p-valor associado ao teste t do coeficiente estimado deveria ser inferior ao nível de significância de 0,05 (BYRNE, 2010; HAIR et al., 2009). Os coeficientes estimados no modelo, juntamente com as hipóteses testadas e os respectivos níveis de significância, estão apresentados na Tabela 4-4 e ilustrados na Figura 4-2.

**Tabela 4.4 - Coeficientes Padronizados Estimados, Hipóteses e Significâncias para o Modelo Proposto**

| Relação Proposta               | Coeficiente Padronizado | <i>p-value</i> | Hipótese Verificada |
|--------------------------------|-------------------------|----------------|---------------------|
| H1: Confiança → Conveniência   | 0,591                   | < 0,001        | Sim                 |
| H2: Sacrifício → Conveniência  | -0,329                  | < 0,001        | Sim                 |
| H3: Ansiedade → Privacidade    | -0,727                  | < 0,001        | Sim                 |
| H4a: Ansiedade → Confiança     | -0,212                  | 0,002          | Sim                 |
| H4b: Ansiedade → Sacrifício    | 0,750                   | < 0,001        | Sim                 |
| H5a: Qualidade → Confiança     | 0,888                   | < 0,001        | Sim                 |
| H5b: Qualidade → Sacrifício    | -0,287                  | < 0,001        | Sim                 |
| H6: Privacidade → Conveniência | 0,042                   | 0,604          | Não                 |



**Figura 4.2 - Coeficientes Padronizados Estimados para o Modelo Proposto**

(\* indica  $p\text{-valor} < 0,01$ ; \*\* indica  $p\text{-valor} < 0,001$ )

Analizando os resultados apresentados na Tabela 4.4 e na Figura 4.2, é possível observar que foi obtido suporte confirmatório para 5 das 6 hipóteses de pesquisa apresentadas. Dessas cinco, quatro são significativas a um nível de 0,001, enquanto uma (o efeito da ansiedade sobre a conveniência) é significativa a um nível de 0,01.

### 4.3. Discussão dos resultados

#### 4.3.1. Ansiedade

Os achados relativos à ansiedade ajudam a compreender os mecanismos que regem a experiência do usuário em canais de denúncia assistidos por IA. No que tange à Hipótese H4a, o coeficiente padronizado de  $-0,212$  confirma um efeito direto e estatisticamente significativo da ansiedade sobre a confiança do usuário. Esse resultado corrobora estudos prévios (Gado et al., 2021; Wong et al., 2023; Karayianni, 2025) e reforça a perspectiva de que um estado de apreensão tecnológica diminui a crença na previsibilidade e integridade dos algoritmos de IA. Compreende-se que usuários ansiosos demandam garantias adicionais — como evidências de segurança e transparência — para restabelecer a confiança, o que sugere que iniciativas de design devem enfatizar indicadores de credibilidade (selos de auditoria, notificações de processamento ético, logs de conformidade) de modo a atenuar temores e fortalecer o vínculo de confiança.

Por sua vez, a Hipótese H4b revela um coeficiente padronizado de  $0,750$ , indicando que níveis elevados de ansiedade amplificam substancialmente o sacrifício percebido. Esse efeito, em consonância com Liu (2025) e Villanueva-Moya e Expósito (2023), aponta que a ansiedade atua como uma lente distorcedora na análise de custo-benefício, levando os usuários a atribuírem maior esforço cognitivo, gasto de tempo e percepção de risco aos processos de denúncia via IA. Na prática, torna-se imperativo reduzir a complexidade das interfaces e prover feedback contínuo — por exemplo, por meio de barras de progresso e confirmações passo a passo — para mitigar a sensação de esforço excessivo e preservar a eficiência da interação.

Finalmente, a Hipótese H3 apresentou um coeficiente de  $-0,727$ , evidenciando que a ansiedade exerce um impacto fortemente negativo sobre a privacidade percebida. Em linha com Wang & Xiao (2025), Kleine et al. (2023) e Pavlopoulos et al. (2024), usuários ansiosos percebem riscos à confidencialidade e ao controle de seus dados como mais elevados, prejudicando a avaliação de quão seguros são os canais de denúncia. Essa constatação corrobora a Teoria do Cálculo da Privacidade, segundo a qual consumidores ponderam utilidade versus riscos de exposição: à medida que a ansiedade inflama preocupações com violação de dados, mesmo benefícios utilitários passam a ser vistos como

insuficientes para compensar o desconforto. Desta forma, os três efeitos significativos sugerem que a gestão da ansiedade — por meio de treinamento, comunicação clara sobre políticas de privacidade e implementação de garantias técnicas — é condição essencial para aprimorar confiança, reduzir sacrifício e proteger a percepção de privacidade em sistemas de IA voltados para denúncias corporativas.

#### **4.3.2. Qualidade do serviço habilitado pela IA**

Os resultados referentes à Hipótese H5a indicam um impacto direto e substancial da qualidade do serviço habilitada por IA sobre a confiança dos usuários (0,888), confirmando e ampliando evidências de Cheatham et al. (2019) e Saratchandran (2019). Esse coeficiente indica que variações na percepção de qualidade — seja pela precisão das recomendações, pela velocidade de resposta, ou pela personalização do atendimento — se traduzem em ganhos de confiança. Em termos teóricos, esse achado reforça o conceito de disconfirmação de expectativas (PARASURAMAN et al., 1994), segundo o qual a superação das expectativas iniciais serve de alicerce para a construção de confiança cognitiva e afetiva. No âmbito dos serviços digitais, a IA, ao garantir consistência e transparência em múltiplos pontos de contato, reduz a incerteza do usuário, fomentando crenças de previsibilidade e integridade do sistema. Para gestores de canais de denúncia corporativos, isso sinaliza a necessidade de investir em algoritmos auditáveis, *feedback* em tempo real e interfaces que comuniquem claramente o valor agregado, de forma a consolidar relacionamentos de confiança duradouros.

Já a Hipótese H5b revela um efeito direto negativo e significativo da qualidade do serviço na percepção de sacrifício (−0,287), corroborando os achados de Kerstan et al. (2023) e Sattu et al. (2024). Embora o coeficiente seja de magnitude baixa, ele é suficiente para indicar que incrementos na qualidade do serviço mediado por IA — como fluxos de denúncia simplificados, navegação intuitiva e suporte proativo ao usuário — reduzem o esforço cognitivo, o tempo demandado e a sensação de perda de controle inerentes ao processo. Sob a ótica da *Service Convenience* (BERRY et al., 2002), essa diminuição de fricção operacional traduz-se em menor desconforto e menor percepção de “custo”

subjetivo, aumentando a eficiência percebida da ferramenta. Na prática, melhorias em usabilidade e clareza de instruções podem gerar reduções no sacrifício percebido pelo usuário, o que, por sua vez, favorece a adoção e o uso continuado.

Em conjunto, esses dois efeitos destacam a relevância da qualidade do serviço habilitada por IA como motor duplo de construção de confiança e de atenuação do sacrifício percebido. A forte relação positiva com a confiança sinaliza que esforços em manter padrões elevados de desempenho — suportados por transparência e consistência algorítmica — geram dividendos em credibilidade, enquanto o efeito negativo sobre o sacrifício ressalta a importância de minimizar atritos e barreiras cognitivas. Assim, organizações que visem maximizar a conveniência percebida em seus canais de denúncia devem priorizar estratégias de governança de dados, design centrado no usuário e comunicação clara de benefícios, garantindo que cada interação com o sistema mediado por IA seja simultaneamente segura, eficiente e de baixo esforço.

#### **4.3.3. Conveniência**

Os resultados evidenciam que a confiança possui o maior efeito (0,591) sobre a conveniência percebida no uso de sistemas de IA em canais de denúncia dentre os construtos testados no modelo proposto. Concordando com Lim e Seng (2024) e Yang et al. (2024), esta relação reforça a premissa de que, à medida que os usuários acreditam na previsibilidade, integridade e benevolência do algoritmo, sua percepção de conveniência sobre o uso da tecnologia habilitada por IA se potencializa. Do ponto de vista teórico, confirma também que sistemas adaptativos e transparentes geram um ciclo virtuoso de uso contínuo (CORRÊA et al., 2023; GLIKSON & WOOLLEY, 2020; PURNOMO et al., 2021; ODOGU, 2024; MOU et al., 2024; YIN & QIU, 2021). Tal achado reforça a necessidade de enfatizar, na interface e na comunicação institucional, evidências de auditoria independente, selos de privacidade e mensagens personalizadas que atestem a robustez do sistema de denúncias via IA, de modo a consolidar a confiança como alicerce da conveniência.

Em paralelo, a percepção de sacrifício mostrou-se significativamente inibidora da conveniência (−0,329), corroborando as proposições de Al-Okaily (2023), Sattu et al. (2024), Arora et al. (2023) e Rahman et al. (2021). Este

coeficiente negativo evidencia que qualquer esforço cognitivo adicional, incerteza sobre procedimentos ou tempo extra requerido para submeter uma denúncia reduz a avaliação de conveniência. Sob a ótica do modelo de *Service Convenience* (BERRY et al., 2002), sacrifícios elevados ampliam fricções operacionais, distorcem a relação custo-benefício e prejudicam a satisfação imediata. Assim, recomenda-se otimizar fluxos — por exemplo, por meio de assistentes virtuais guiados passo a passo, feedback visual de progresso e recuperação automática de dados — para diminuir o esforço e assegurar que cada interação seja percebida como de baixo custo mental e alto valor agregado.

Contrariamente às expectativas baseadas em Cloarec (2020), Martin et al. (2017), Lim & Shim (2022), Afridi & Kumar (2024) e Princi & Krämer (2020), a Hipótese H6 não foi confirmada ( $p = 0,604$ ). Embora Skinner et al. (2006) ressaltem que a privacidade é entendida por muitos como um direito fundamental, na amostra em questão ela não se mostrou um fator capaz de reduzir a conveniência percebida. Essa discrepância pode decorrer de um efeito de compensação, em que os benefícios funcionais e a urgência inerente ao processo de denúncia se sobrepõem às preocupações com privacidade — ao menos até que ocorram incidentes concretos de violação de dados. Além disso, é plausível que a confiança no sistema (potencialmente reforçada por altos índices de qualidade de serviço) possa mitigar as apreensões relativas à privacidade, tornando este construto menos relevante na avaliação de conveniência.

#### **4.3.4. Resultados gerais**

A análise integrada dos achados aponta que compreender a conveniência em profundidade implica reconhecer que ela não é gerada de forma isolada, mas resulta da sinergia entre a confiança no sistema e a qualidade percebida do serviço, ao mesmo tempo em que sofre as pressões de variáveis psicológicas como a ansiedade e o sacrifício subjetivo. A confiança atua como fundação, pois sem ela qualquer ganho funcional se dissipa em dúvidas e hesitações; a qualidade do serviço, por sua vez, molda a percepção de valor e de utilidade imediata, atenuando barreiras de uso e reduzindo o custo mental e temporal que a denúncia pode implicar.



Do ponto de vista prático, essas relações sugerem que maximizar a conveniência deve ir além de tornar a interface “amigável” ou de disparar alertas de segurança: é necessário estruturar toda a jornada do usuário de modo a construir, a cada clique, um ambiente de previsibilidade e suporte. Isso passa por articular mensagens claras de transparência (sobre tratamento de dados, fluxos de encaminhamento e confidencialidade), incorporar pontos de validação automatizada que reforcem a confiabilidade e projetar respostas adaptativas que demonstrem atenção às necessidades específicas de cada denunciante. O *feedback* imediato, os tutoriais contextuais e a disponibilização de indicadores visuais de progresso não somente diminuem o esforço requerido como comunicam cuidado institucional, o que fortalece a conveniência percebida.

Por fim, a ênfase na conveniência evidencia sua importância estratégica em ambientes corporativos, onde a eficácia dos canais de denúncia está diretamente ligada à saúde ética e ao *compliance* organizacional. Compreender profundamente essa variável significa, portanto, permitir que as empresas não apenas adotem tecnologias de IA como “mais um recurso”, mas repensem seus processos de atendimento, comunicação e governança de dados de forma integrada, garantindo que a experiência do denunciante seja percebida como segura, eficiente e com o menor dispêndio possível de tempo e esforço. Dessa forma, a conveniência deixa de ser um simples desdobramento teórico para se tornar um indicador de transformação cultural, sinalizando a capacidade de uma organização de ouvir, proteger e valorizar quem se dispõe a revelar condutas potencialmente danosas.

## **5**

### **Conclusões**

Este capítulo é uma síntese da pesquisa desenvolvida, seguida da análise das contribuições e dos impactos gerados pelo estudo, tanto sob a perspectiva teórica quanto prática. Por fim, são discutidas as limitações enfrentadas ao longo da investigação e indicadas possíveis direções para estudos futuros.

#### **5.1.**

##### **Sumário do estudo**

Esta pesquisa examinou a conveniência percebida pelos usuários no uso de inteligência artificial no contexto de canais de denúncia corporativos. Para isso, foram analisadas as relações entre os antecedentes confiança, privacidade, ansiedade, qualidade do serviço habilitado pela IA e sacrifício percebido e seu impacto na percepção de conveniência proporcionada pela IA. O objetivo foi aprofundar a compreensão dos fatores que influenciam a percepção de conveniência do usuário ao utilizar um canal de denúncia online no ambiente corporativo.

A revisão de literatura iniciou-se com uma visão geral sobre conveniência. Nesse contexto, revisou-se a teoria de Copeland (1923), Kotler e Zaltman (1971), Seiders et al. (2005), Van Doorn et al. (2010), Jian et al. (2013), e Roy et al. (2017), destacando a evolução do conceito de conveniência e as principais particularidades nas visões desses autores. Essa discussão foi aplicada ao contexto da inteligência artificial, apresentando o conceito de conveniência habilitada pela IA.

Em seguida, foi discutido sobre o contexto de aplicação da conveniência habilitada pela IA, um canal de denúncia corporativo. Foi definido o conceito de denúncia por Miceli e Near (1981), apresentando a sua importância (RUIVO, PIRES, 2020) e o contexto brasileiro legislativo que trata do tema.

Posteriormente, foi debatido sobre a Teoria do Comprometimento-Confiança, de Morgan e Hunt (1994) que aponta que esses são principais fatores contextuais que implicam no relacionamento duradouro. Wang (2019) adiciona que a privacidade é a um aspecto fundamental da confiança e Siau e Wong (2018) e Ferrario et al. (2019) adicionam que a confiança pode influenciar fatores como qualidade do serviço e conveniência no uso de tecnologias baseadas em IA.

Foram então aprofundadas algumas características e antecedentes da conveniência habilitada pela IA, detalhando construtos previamente identificados como relevantes na literatura (AMEEN et al., 2021; WALCH, 2019; THIEL, 2019; CHANG et al., 2012). Com o intuito de aprofundar o entendimento acerca da conveniência percebida pela IA, o presente estudo desenvolveu um modelo fundamentado na pesquisa de Ameen et al. (2021), com adaptações destinadas a testar o modelo de conveniência percebida, bem como novos constructos, como ansiedade e privacidade. Esses construtos foram incorporados com base nas contribuições de Compeau e Higgins (1999), Xu et al. (2011) e Lee e Rha (2016). Em seguida foi elaborado um instrumento de pesquisa composto por escalas pré-existentes, desenvolvidas e testadas anteriormente na literatura (WOLFINBARGER, GILLY, 2003; CHANG, WANG, 2011; WANG et al., 2019, PAPAROIDAMIS, TRAN, LEONIDOU, 2019; MERISAVO et al. (2007) SHIN, LIN, 2016; COMPEAU, HIGGINS, 1999; XU et al., 2011) e cuidadosamente traduzidas para o português. Outros construtos, já foram apresentados no idioma português e não foi necessária de tradução (LEE, RHA, 2016; COLLIER, SHERRELL, 2010) Realizou-se uma cross-sectional survey com uma amostra com 214 questionários válidos, respondidos por indivíduos da área de compliance.

O modelo de mensuração ajustado aos dados demonstrou a confiabilidade e a validade dos construtos avaliados por meio das escalas selecionadas, apresentando resultados consistentes com estudos prévios que utilizaram as mesmas escalas.

As maioria das hipóteses de pesquisa foram confirmadas por meio da modelagem de equações estruturais, com exceção de uma. O modelo proposto mostrou que a qualidade do serviço habilitado por IA tem efeito positivo na confiança do usuário e contribui para a redução do sacrifício percebido. A ansiedade, por outro lado, mostrou-se um fator prejudicial, ao impactar negativamente a confiança, aumentar a percepção de sacrifício e intensificar o

risco percebido à privacidade. A conveniência percebida foi influenciada positivamente pela confiança e negativamente pelo sacrifício, destacando o papel das emoções e da experiência subjetiva na avaliação da ferramenta. No entanto, o risco percebido à privacidade não apresentou efeito significativo sobre a conveniência.

## **5.2. Conclusões e implicações**

Os resultados e as relações identificadas na pesquisa configuram contribuições relevantes para o arcabouço teórico da teoria do comprometimento-confiança, ao aprofundarem a compreensão sobre a conveniência percebida no contexto da inteligência artificial. Além disso, o estudo apresenta implicações práticas para organizações interessadas na adoção de canais de denúncia baseados em IA, bem como para empresas fornecedoras dessas soluções que buscam aprimorar tecnologicamente seus produtos por meio da incorporação da IA.

### **5.2.1. Implicações teóricas**

As implicações teóricas deste estudo são contribuições relevantes para o campo da adoção de tecnologias baseadas em inteligência artificial (IA) em contextos organizacionais sensíveis, como os canais corporativos de denúncia. O presente trabalho contribui para a literatura ao propor e validar um modelo que articula fatores emocionais, funcionais e perceptuais com a conveniência percebida no uso de IA para fins organizacionais. Sendo assim, o estudo avança a compreensão teórica de como construtos como ansiedade, sacrifício percebido e risco de privacidade operam conjuntamente com a confiança e a qualidade do serviço para moldar a percepção de conveniência, um desfecho até então pouco explorado na literatura sob essa ótica.

Em primeiro lugar, ao evidenciar o papel central da confiança como o preditor mais relevante da conveniência percebida, este trabalho reforça, em linha com Glikson e Woolley (2020) e Yin e Qiu (2021), a necessidade de teorizações mais robustas sobre a confiança na inteligência artificial como variável mediadora e moderadora em contextos de risco ético. Além disso, a identificação de efeitos significativos da ansiedade e do sacrifício percebido sobre a confiança, privacidade e conveniência oferece uma contribuição relevante ao expandir a

literatura sobre barreiras emocionais à adoção tecnológica, sugerindo a necessidade de modelos que combinem aspectos da psicologia do uso de tecnologias com os princípios de design centrado no usuário.

A ausência de relação significativa entre risco de privacidade e conveniência também foi um resultado instigante. Embora esse achado esteja em aparente dissonância com a literatura tradicional (como MARTIN et al., 2017), ele sugere que, em certos contextos institucionais, o cálculo racional dos benefícios e da funcionalidade da IA pode temporariamente suplantar preocupações com privacidade, principalmente em cenários em que o usuário perceba a IA como uma solução institucional legitimada por uma promessa de segurança e confidencialidade.

Finalmente, este estudo contribui teoricamente ao reposicionar o construto conveniência percebida como um elemento central no processo de uso de IA em contextos de denúncia. Tradicionalmente tratada como consequência funcional da usabilidade, a conveniência aqui é elevada a um status teórico mais complexo, visto que emerge da interdependência de confiança, esforço percebido, afetos negativos e percepções de qualidade. Essa reformulação teórica permite que estudiosos do comportamento do consumidor passem a considerar a conveniência como uma variável crítica na modelagem da adesão a canais digitais em cenários de alto risco social e emocional.

Assim, os resultados encontrados não apenas respondem aos objetivos propostos e às lacunas identificadas na literatura, como também oferecem novas direções para a avaliação da adoção tecnológica, da confiança na automação por IA e da percepção de valor em serviços mediados por IA.

### **5.2.2. Implicações práticas**

Os resultados desta pesquisa oferecem direcionamentos importantes para empresas que operam canais de denúncia que baseados em IA. A percepção de qualidade do serviço desempenha um papel fundamental na construção da confiança do usuário e na redução do esforço percebido ao realizar uma denúncia. Assim, é essencial que essas ferramentas ofereçam uma experiência fluida, com interfaces intuitivas, respostas claras e bom desempenho técnico, garantindo que o

usuário se sinta amparado e compreendido ao longo do processo de realizar uma denúncia.

A ansiedade do usuário mostrou-se um fator de impacto negativo relevante, elevando a percepção de risco e comprometendo a confiança na ferramenta. Isso evidencia a importância de comunicar de forma transparente o funcionamento do canal, reforçando o anonimato, a segurança dos dados e a seriedade com que as denúncias são armazenadas. Estratégias como mensagens acolhedoras, explicações didáticas e a oferta de suporte humanizado – mesmo em plataformas automatizadas – ajudam a reduzir essa ansiedade e a criar um ambiente digital mais acolhedor.

Por fim, a confiança também se destacou como elemento-chave para a conveniência percebida. Usuários que confiam na ferramenta tendem a considerá-la mais prática e acessível, enquanto aqueles que percebem um alto sacrifício no uso enxergam menor conveniência. Portanto, investir em usabilidade, simplificação de etapas e clareza na comunicação é tão importante quanto a tecnologia em si. A privacidade, embora não tenha impacto direto na conveniência, segue sendo um ponto sensível na experiência do usuário e deve ser comunicada de forma clara e objetiva, reforçando a sensação de segurança e contribuindo para o engajamento efetivo com o canal de denúncia. Dessa forma, nenhuma mensagem importante deixará de ser enviada pelo usuário e a empresa poderá endereçar e tratar assuntos como assédio, fraude, violações a leis ambientais entre outros temas.

### **5.3. Limitações do estudo**

Uma limitação metodológica da pesquisa é sobre o processo de coleta de dados. Embora tenha havido diligência na apresentação de capturas de tela simulando a interface de um canal de denúncia com integração de Inteligência Artificial (IA) aos participantes, a heterogeneidade na interpretação das telas pode ter introduzido variabilidade na qualidade das respostas. Não obstante a exclusão de cinco questionários cujos respondentes indicaram incompreensão das imagens apresentadas, essa potencial disparidade interpretativa permanece como um fator a ser considerado. A fim de mitigar essa limitação, a aplicação do instrumento de coleta de dados imediatamente subsequente à experiência de uso de um canal de

denúncia efetivamente habilitado por IA seria metodologicamente interessante. Contudo, a escassez de exemplos concretos dessa tecnologia no mercado inviabilizou esse formato.

Em relação à validade externa dos resultados obtidos, e considerando que a amostra analisada compreende predominantemente indivíduos brasileiros, com uma parcela significativa (superior a 64%) apresentando renda familiar superior a 10 salários-mínimos, é importante ressaltar a possibilidade de que as relações identificadas nesta pesquisa possam não apresentar generalização irrestrita a todos os possíveis usuários.

Outra limitação da pesquisa é a avaliação de somente alguns construtos relacionados à conveniência percebida pelo uso da IA em decorrência da necessidade de manter o questionário dentro de um limite aceitável de tempo para resposta, 11 minutos no presente estudo. Apesar de terem sido testados construtos apontados como relevantes na literatura, a experiência do usuário não foi abordada na pesquisa e poderia ser tema de um trabalho futuro.

Essas limitações apontam que o tema é amplo e que ainda há espaço para outros trabalhos e discussões. Assim, apesar da possibilidade de outros testes, acredita-se que os achados dessa pesquisa sejam um importante passo para uma melhor compreensão da conveniência percebida do usuário ao realizar uma denúncia em um canal de denúncia corporativo habilitado pela IA.

#### **5.4. Sugestões para pesquisas futuras**

Considerando as limitações inerentes à presente investigação, a condução de estudos subsequentes com uma amostra de maior número e diversidade de perfil em relação àquela analisada – a exemplo da inclusão do público geral, de estratos socioeconômicos distintos e de uma menor faixa etária média, se apresentaria como uma via promissora para a validação e expansão do alcance dos resultados ora obtidos.

A aplicação de outras escalas e construtos teóricos ao modelo poderia contribuir para uma compreensão mais abrangente do processo de formação da conveniência percebida no contexto de canais de denúncia corporativos.

Além disso, realizar uma avaliação multigrupo ou mediação/moderação para uma visão mais aprofundada de algumas interseções como a escolaridade em relação ao uso de IA ou ainda a renda média familiar.

Por fim, em virtude das características e particularidades intrínsecas a cada aplicação da Inteligência Artificial (IA), a implementação do modelo proposto em outras ferramentas de gestão empresarial, tais como *Customer Relationship Management* (CRM) e *Enterprise Resource Planning* (ERP), pode resultar em conclusões e, potencialmente, implicações distintas.



## 6 Referências Bibliográficas

ANDRÉ, A. D. et al. The role of user experience in the adoption of automated services. **Journal of Service Research**, v. 21, n. 3, p. 267-284, 2018.

AFRIDI, M. F; KUMAR, K. S. Exploring data privacy and privacy paradox nature of consumers using block chain technology: application of SLR and bibliometric analysis tools. *Abhigyan*, 42(3), 181-205, 2024.

AKROUT, H.; NAGY, G. The role of trust in social media brand communities: A case study of Facebook brand pages. **Journal of Business Research**, Amsterdam, v. 92, p. 1-10, 2018.

AL-OKAILY, M. Does AIs usage matter in SMEs performance? An empirical investigation under the digital transformation revolution. *Information Discovery and Delivery*, 52(2), 125-137, 2023.

AMEEN, A.; et al. Artificial Intelligence and User Satisfaction: A Study of AI-Enabled Services. **Computers in Human Behavior**, v. 122, p. 106-115, 2021.

ARORA, A.; GUPTA, S.; DEVI, C.; WALIA, N. Customer experiences in the era of artificial intelligence (AI) in context to fintech: a fuzzy AHP approach. **International Journal**, 30(10), 4342-4369, 2023.

ASSOCIATION OF CERTIFIED FRAUD EXAMINERS. Report to the Nations 2024: global study on occupation fraud and abuse. 2024. Disponível em: <https://www.acfe.com/-/media/files/acfe/pdfs/rtnn/2024/2024-report-to-the-nations.pdf> Acesso em: 11 de fevereiro de 2025.

BAGOZZI, R. P.; PHILLIPS, L. Representing and Testing Organizational Theories: A Holistic Construal. **Administrative Science Quarterly**, v. 17, p. 459-489, 1982.

BENTLER, P. M. Comparative Fit Indices in Structural Models. **Psychological Bulletin**, v. 107, n. 2, p. 238-246, 1990.

BENTLER, P. M.; BONETT, D. G. Significance tests and goodness of fit in the analysis of covariance structures. **Psychological Bulletin**, 88(3), 588-606, 1980.

BERRY, L. L.; SEIDERS, K.; GREWAL, D. Understanding Service Convenience. **Journal of Marketing**, v. 66, July, pp. 1-17, 2002

BRASIL. Lei nº 12.846, de 1º de agosto de 2013. Dispõe sobre a responsabilização administrativa e civil de pessoas jurídicas pela prática de atos contra a administração pública, nacional ou estrangeira, e dá outras providências. **Diário Oficial da União Federativa do Brasil**, Brasília, DF, 2 de ago. 2013. Acesso em: 11 de fev. 2025.

BRASIL. Lei nº 13.709, de 14 de agosto de 2018. Dispõe sobre a proteção de dados pessoais e altera o Marco Civil da Internet. **Diário Oficial da União Federativa do Brasil**, Brasília, DF, 15 de ago. 2018. Acesso em: 11 fev. 2025.

BRASIL. Norma Regulamentadora nº 5 (NR-5). Estabelece a obrigatoriedade da Comissão Interna de Prevenção de Acidentes (CIPA) nas empresas, visando à prevenção de acidentes e doenças relacionadas ao trabalho. **Diário Oficial da União Federativa do Brasil**, Brasília, DF, 21 de set. 2022. Acesso em: 11 fev. 2025.

BYRNE, B. M. Structural Equation Modeling with AMOS: **Basic Concepts, Applications and Programming**. 2a ed. Routledge, NY, 2010.

BÉLANGER, F.; HILLER, J. S. A Framework for E-Government: Privacy Implications, **Business Process Management Journal** (12:1), pp. 48-60, 2006.

CAMBRE, M. A.; COOK, D. L. Computer anxiety: Definition, measurement, and correlates. **Journal of Educational Computing Research**, 1(1), 37–54, 1985.

CHANG, C.; YAN, C.; TSENG, J. Perceived convenience in an extended technology acceptance model: mobile technology and english learning for college students. **Australasian Journal of Educational Technology**, 28(5), 2012.

CHANG, H. H., CHEN, S. W. The impact of customer interface quality, satisfaction and switching costs on e-loyalty: Internet experience as a moderator. **Computers in Human Behavior**, 2927–2944, 2008.

CHAUDHRI, V. K.; LANE, H. C.; GUNNING, D.; ROSCHELLE, J. Applications of artificial intelligence to contemporary and emerging educational challenges. **Artificial Intelligence Magazine**, Intelligent Learning Technologies, 2(34), 4, 2013.

CHEATHAM, C. L.; JAVANMARDIAN, K.; SAMANDARI, H. The role of personal data in AI-enabled services: Implications for customer control. **Journal of Marketing Theory and Practice**, v. 27, n. 2, p. 123-138, 2019.

CHIESA, V.; FRATTINI, F. The role of trust in the innovation process: A study of the Italian manufacturing sector. **International Journal of Technology Management**, Bradford, v. 56, n. 1, p. 1-21, 2011.

CHOI, B.; LEE, C.; LEE, H.; SUBRAMANI, M. Effects of web retail service quality and product categories on consumer behavior: A research model and empirical exploration. **37th annual Hawaii international conference on system sciences**, 2004.

CHOI, J. A Study of Customer Acceptance of Artificial Intelligence Technology. **International Journal of E-Business Research**, Hershey, v. 19, n. 1, p. 1-20, 2023.

CHRISTAKIS, N. A. The social dynamics of AI: Implications for human interaction. **Artificial Intelligence**, v. 274, p. 1-11, 2019.

CHURCHILL, G. A. A Paradigm for Developing Better Measures of Marketing Constructs. **Journal of Marketing**, v. 16, p. 64-73, 1979.

CLARKE, I.; MILLER, S. G.; HONG, T. W.; SANDBERG, O. Protecting Free Expression Online with Freenet, **Internet Computing** (6:1), pp. 40-49, 2002.

CLARKE, R. Why the world wants control over artificial intelligence. **Computer Law & Security Review**, 35(4), 423– 433, 2019.

CLOAREC, J. The personalization–privacy paradox in the attention economy. *Technological Forecasting & Social Change*, 2020.

COLLIER, J. E.; BIENSTOCK, C. C. Measuring service quality in E-retailing. **Journal of Service Research**, 8(3), 260–275, 2006.

COMPEAU, D. R.; HIGGINS, C. A.; HUFF, S. Social Cognitive Theory and Individual Reactions to Computing Technology: A Longitudinal Study. **MIS Quarterly**, v. 23, n. 2, p. 145 – 158, 1999.

CONVENIÊNCIA, in Michaelis, Dicionário Brasileiro da Língua Portuguesa, 2025, <https://michaelis.uol.com.br/moderno-portugues/busca/portugues-brasileiro/conveni%C3%Aancia/#:~:text=1%20Qualidade%20do%20que%20%C3%A9,%20proveito%2C%20utilidade%2C%20vantagem>.

COPELAND, Melvin T. Relation of consumers' buying habits to marketing methods. **Harvard Business Review**, v. 1, n. 3, p. 282–289, 1923.

CORRITORE, C. L.; KRACHER, B.; WIEDENBECK, S. On-line trust: Concepts, evolving themes, and a model. **International Journal of Human-Computer Studies**, Amsterdam, v. 58, n. 6, p. 737-758, 2003.

CORRITORE, M.; KRACHER, B.; WIEDENBECK, S. The effects of cognitive style on trust in technology: A model of trust in e-commerce. **Proceedings of the 2003 American Conference on Information Systems**, 2003.

CORRÊA, N. K.; GALVÃO, C.; SANTOS, J. W.; PINO, C. D.; PINTO, E. P.; BARBOSA, C.; OLIVEIRA, N. Worldwide ai ethics: a review of 200 guidelines and recommendations for ai governance. **Patterns**, 4(10), 100857, 2023.

DAVENPORT, T. H.; GUHA, A.; GREWAL, D.; BRESSGOTT, T. How artificial intelligence will change the future of marketing. **Journal of the Academy of Marketing Science**, v. 48, n. 1, p. 24-42, 2020.

DAVIS, F. D. Perceived usefulness, perceived ease of use, and user acceptance of information technology. **MIS Quarterly**, v. 13, n. 3, p. 319–340, 1989.

DE KERVILER, G.; DEMOULIN, N. T.; ZIDDA, P. The impact of service automation on customer satisfaction: The role of perceived sacrifice. **Journal of Retailing and Consumer Services**, v. 31, p. 1-10, 2016.

DEBATIN, B.; LOVEJOY, J. P.; HORN, A.; HUGHES, B. N. Facebook and online privacy: Attitudes, behaviors, and unintended consequences. **Journal of Computer-Mediated Communication**, 15, 83–108, 2009.

DIENLIN, T.; TREPTE, S. Is the privacy paradox a relic of the past? An in-depth analysis of privacy attitudes and privacy behaviors. **European Journal of Social Psychology**, 45, 285–297, 2015.

ELBELTAGI, I.; AGAG, G. Trust and commitment in online retailing: A study of the impact of trust on customer commitment in the context of online shopping. **International Journal of Retail & Distribution Management**, Bradford, v. 44, n. 3, p. 294-309, 2016.

EPSTEIN, S. The nature of anxiety with emphasis upon its relationship to expectancy. In C. D. Spielberger (Ed.), **Anxiety: Current trends in theory and research** (Vol. 2, pp. 291–337). New York: Academic Press, 1972.

ETZIONI, A.; ETZIONI, O. Incorporating ethics into artificial intelligence. *The Journal of Ethics*, 21(4), 403–418, 2017.

FARAYOLA, O. Ethics and trustworthiness of AI for predicting the risk of recidivism: A systematic literature review. *Information*, Basel, v. 14, n. 8, p. 426, 2023.

FERRARIO, M.; LOI, M.; VIGANO`, C. Trust and the role of artificial intelligence in the customer experience. **Journal of Business Research**, Amsterdam, v. 100, p. 1-10, 2019.

FERREIRA, J. B. Aceitação e prontidão do consumidor para produtos de alta tecnologia: Elaboração e teste empírico do modelo CART para adoção de produtos de alta tecnologia. [s.l.] Tese (Doutorado em Administração) – **Universidade Federal do Rio de Janeiro – UFRJ, Instituto COPPEAD de Administração**, 2010.

FOCK, S. T.; KOH, H. C. Conceptualization of Trust and Commitment: Understanding the relationships between trust and commitment and the willingness to try Internet Banking Services. **International Journal of Business and Information**, v. 2, p. 194-208, 2006.

FORNELL, C.; LARCKER, D. F. Evaluating Structural Equation Models with Unobservable Variables and Measurement Error. *Journal of Marketing Research*, v. 18, fevereiro, p.39-50, 1981.

FØLSTAD, A.; BRANDTZAEG, P. B.; HAUGSTVEIT, I. M. What Makes Users Trust a Chatbot for Customer Service? An Exploratory Interview Study. **In Proceedings of the 2018 International Conference on Artificial Intelligence** (pp. 1-6), 2018.

GADO, S.; KEMPEN, R.; LINGELBACH, K.; BIPP, T. Artificial intelligence in psychology: how can we enable psychology students to accept and use artificial intelligence? **Psychology Learning and AMP Teaching**, 21(1), 37-56, 2021.

GARVER, M. S.; MENTZER, J. T. Logistics Research Methods: Employing Structural Equation Modeling to Test for Construct Validity. **Journal of Business Logistics**, v. 20, p. 33-57, 1999.

GAUVRIT, N. Balancing automation and human interaction in service delivery. **Journal of Service Research**, v. 22, n. 4, p. 459-474, 2019.

GAZWANI, A.; TAN, W. Artificial intelligence, financial anxiety and cashier-less checkouts: a Saudi Arabian perspective. **The International Journal of Bank Marketing**, v. 40, n. 2, p. 270-285, 2022.

GHAZIZADEH, M.; LEE, J. D.; BOYLE, L. N. Trust in technology: A review of the literature. *Human Factors*, Baltimore, v. 54, n. 4, p. 671-688, 2012.

GLIKSON, E.; WOOLLEY, A. W. Human trust in artificial intelligence: review of empirical research. **Academy of Management Annals**, 14(2), 627-660, 2020.

GUO, Y.; YANG, Y. Modeling and predicting trust dynamics in human–robot teaming: A Bayesian inference approach. **International Journal of Social Robotics**, Dordrecht, v. 12, n. 1, p. 1-15, 2020.

HAIR, J. F.; BLACK, W. C.; BABIN, B. J.; ANDERSON, R. E. *Multivariate Data Analysis*,. 7a ed., Upper Saddle River: Prentice-Hall, 2009.

HASAN, L.; et al. Consumer Trust and Perceived Risk in AI Technologies. **Journal of Business Research**, v. 124, p. 1-10, 2021.

HEDIN, U. C.; MÅNSSON, S. A. Whistleblowing processes in Swedish public organisations: Complaints and consequences. **European Journal of Social Work**, 15(2), 151–167, 2012.

HEIDENREICH, S.; SPIETH, P. Innovation-related communication and customer acceptance: The role of trust. **Journal of Business Research**, Amsterdam, v. 66, n. 8, p. 1-7, 2013.

HENGSTLER, M.; ENKEL, E.; DUELLI, S. Applied artificial intelligence and trust: The case of autonomous vehicles. **Technological Forecasting and Social Change**, New York, v. 105, p. 105-120, 2016.

HENGSTLER, M.; ENKEL, E.; DUELLI, S. Applied artificial intelligence and trust: The case of autonomous vehicles. **Technological Forecasting and Social Change**, New York, v. 105, p. 105-120, 2016.

HOLLEBEEK, L. D.; CHEN, T. Engaging consumers through artificially intelligent technologies: Systematic review, conceptual model, and further research. **Psychology and Marketing**, v. 41, n. 1, p. 62-78, 2024.

HOOPER, D.; COUGHLAN, J.; MULLEN, M. R. Structural Equation Modelling: Guidelines for Determining Model Fit. **The Electronic Journal of Business Research Methods**, v. 6, n. 1, p. 53 - 60, 2008.

HU, L. T.; BENTLER, P. M. Cutoff Criteria for Fit Indexes in Covariance Structure Analysis: Conventional Criteria Versus New Alternatives. **Structural Equation Modeling**, v. 6, n. 1, p. 1-55, 1999.

HUANG, M.H.; RUST, R. T. Artificial Intelligence in Service. **Journal of Service Research**, 21(2), 155-172, 2018.

JIAN, J. Y.; BISANTZ, A. M.; DRURY, C. G. Foundations for an empirically determined scale of trust in automated systems. **International Journal of Cognitive Ergonomics**, London, v. 4, n. 1, p. 53-71, 2000.



JIANG, L.; YANG, Z.; JUN, M. Measuring consumer perceptions of online shopping convenience, **Journal of Service Management**, Vol. 24 Iss 2 pp. 191 – 214, 2013.

KANG, J.; LOU, H. Exploring human-AI interaction in organizational contexts. **Journal of Applied Psychology**, v. 108, n. 2, p. 350-365, 2022.

KARAYIANNI, I.; KLIDAS, A.; KARAKITSOU, C. Between a bot and a hard place: student engagement, technology anxiety and ChatGPT trust in higher education, 2025.

KAUR, K.; SINGH, R.; GUPTA, S. Trustworthy artificial intelligence: A review. **ACM Computing Surveys**, New York, v. 54, n. 1, p. 1-35, 2022.

KEININGHAM, T. L.; AKSOY, L.; COOIL, B.; ANDREASSEN, T. W. A longitudinal examination of customer engagement and brand trust. **Journal of Service Research**, Thousand Oaks, v. 20, n. 1, p. 1-16, 2017.

KERSTAN, S.; BIENEFELD, N.; GROTE, G. Choosing human over AI doctors? How comparative trust associations and knowledge relate to risk and benefit perceptions of AI in healthcare. **Risk Analysis**, 44(4), 939-957, 2023.

KLEINE, A.; KOKJE, E.; LERMER, E.; GAUBE, S. Attitudes toward the adoption of 2 artificial intelligence enabled mental health tools among prospective psychotherapists: cross-sectional study. **JMIR Human Factors**, 10, e46859, 2023.

KOTLER, P.; ZALTMAN, G. Social marketing: an approach to planned social change. **Journal of Marketing**, v. 35, n. 3, p. 3-12, 1971

KUMAR, A. Trust's significance in human-AI communication and decision-making. **International Journal of Scientific Research in Engineering and Management**, New Delhi, v. 8, n. 1, p. 1-10, 2024.

KUMAR, R. Transparency and trust in AI-driven systems. *International Journal of Business Ethics*, v. 19, n. 4, p. 123-135, 2024.

LAI, P. C. Security as an extension to tam model: consumers' intention to use a single platform e-payment. ***Asia-Pacific Journal of Management Research and Innovation***, 13(3-4), 110-119, 2017.

LEE, J. D.; SEE, K. A. Trust in automation: Designing for appropriate reliance. ***Human Factors***, Baltimore, v. 46, n. 1, p. 50-80, 2004.

LI, J.; HUANG, J.S. Dimensions of artificial intelligence anxiety based on integrated fear acquisition theory. ***Technology in Society***, 63, 101410, 2020.

LIM, T. K; SENG, L. W. Determinants of artificial intelligence adoption in accounting among Malaysia small and medium-sized enterprises. ***Journal of Accounting, Business and Management***, 32(1), 90, 2024.

LIU, X; LI, Y. Examining the double-edged sword effect of ai usage on work engagement: the moderating role of core task characteristics substitution. ***Behavioral Sciences***, 15(2), 206, 2025.

MADHAVAN, P.; WIEGMANN, D. A. Effects of information source, pedigree, and reliability on operator interaction with decision support systems. ***Human Factors, Baltimore***, v. 49, n. 6, p. 1007-1018, 2007.

MALHOTRA, N. K.; KIM, S. S.; AGARWAL, J. Internet Users Information Privacy Concerns (IUIPC): The Construct, the Scale, and a Causal Model, ***Information Systems Research*** (15:4), pp. 336-355, 2004.

MARTIN, K. D.; ABHISHEK B.; PALMATIER, R. W. Data Privacy: Effects on Customer and Firm Performance. *Journal of Marketing*, 81 (1), 36–58, 2017.

MERISAVO, M. et al. The impact of digital marketing on customer satisfaction: A case study. ***Journal of Business Research***, v. 60, n. 3, p. 255-261, 2007.

MEYER, C. The Importance of Convenience in the COVID-19 Era. **Harvard Business Review**, 2020.

MICELI, M. P.; NEAR, J. P. Characteristics of organizational climate and perceived wrongdoing associated with whistle-blowing decisions. **Personnel Psychology**, 38(3), 525–544, 1985.

MIKHAYLOV, S. J.; ESTEVE, M.; CAMPION, A. Artificial intelligence for the public sector: Opportunities and challenges of cross-sector collaboration. *Philosophical Transactions - Royal Society Mathematical, Physical, and Engineering Sciences*, 376(2128), 20170357, 2018.

MOHAMED, S. The role of AI in enhancing whistleblowing systems. **Corporate Ethics Review**, v. 12, n. 1, p. 15-29, 2024.

MORGAN, R. M.; HUNT, S. D. The commitment-trust theory of relationship marketing. **Journal of Marketing**, Chicago, v. 58, n. 3, p. 20-38, jul. 1994.

MORGAN, R. M.; HUNT, S. D. The commitment-trust theory of relationship marketing. **Journal of Marketing**, v. 58, n. 3, p. 20-38, 1994.

MOROSAN, C.; CENGIZCI, A. Letting AI make decisions for me: an empirical examination of hotel guests' acceptance of technology agency. **International Journal of Contemporary Hospitality Management**, v. 35, n. 10, p. 3174-3195, 2023.

MOU, Y.; GONG, Y., e DING, Z. Complement or substitute? A study of the impact of artificial intelligence on consumers' resistance. **Marketing Intelligence and AMP; Planning**, 42(4), 647-665, 2024.

MURPHY, R. The loss of human control in AI-enabled services. **Journal of Business Research**, v. 78, p. 1-9, 2017.

NADARZYNSKI, T. et al. Acceptability of artificial intelligence (AI)-led chatbot services in healthcare: A mixed-methods study. **Digital Health**, 5, 1-12, 2019.

NADARZYŃSKI, Tomasz; WAGNER, Aleksandra; KWIATKOWSKA, Anna; ADAMSKA, Aneta; PŁUCIEŃ, Piotr. The role of the Facebook community in the treatment of chronic diseases. **Journal of Medical Internet Research**, v. 21, n. 12, 2019.

NIENABER, H.; SCHEWE, G. Trust in technology: The role of trust in the acceptance of technology. **Journal of Business Research**, Amsterdam, v. 67, n. 4, p. 1-7, 2014.

NJAMFA, A. The relationship between customer experience and brand trust: A study of the retail sector. **Journal of Retailing and Consumer Services**, Amsterdam, v. 45, p. 1-10, 2018.

NUNNALLY, J.; BERNSTEIN, I. Psychometric Theory. 3a ed., McGraw-Hill Humanities/Social Sciences/Languages, 1994.

ODOGU, T. K. Z. Artificial intelligence and cyber security: implications for e-trans and e-accounting in emerging economies. **African Journal of Accounting and Financial Research**, 7(4), 152-168, 2024.

PARASURAMAN, A.; ZEITHAML, V. A.; BERRY, L. L. Reassessment of expectations as a comparison standard in measuring service quality: Implications for further research. **Journal of Marketing**, v. 58, n. 1, p. 111-124, 1994.

PARASURAMAN, S.; IGBARIA, M. An examination of gender differences in the determinants of computer anxiety and attitudes toward microcomputers among managers. **International Journal of Man Machine Studies**, 32(3), 327– 340, 1990.

PARK, Y. J. Digital literacy and privacy behavior online. **Communication Research**, 40, 215–236, 2013.

PATIL, D.; RANE, N.; RANE, J. Acceptance of ChatGPT and generative artificial intelligence in several business sectors: key factors, challenges, and

implementation strategies, The Future Impact of ChatGPT on Several Business Sectors (pp. 201-241). **Deep Science Publishing**, 2024.

PAVLOPOULOS, A.; RACHIOTIS, T.; MAGLOGIANNIS, I. An overview of tools and technologies for anxiety and depression management using AI, 2024.

PHAM, M. T.; AHAMMAD, M. F. The role of human interaction in service quality: A systematic review. **Journal of Service Management**, v. 28, n. 2, p. 275-298, 2017.

PHAM, T. H.; TRAN, T. B.; MISRA, S.; MASKELIŪNAS, R.; DAMAŠEVIČIUS, R. Convenience, Perceived Value, and Repurchase Intention in Online Shopping. *Sustainability*, v. 10, n. 1, p. 156, 2018.

PINTO, J. C.; DELL'ERA, C.; VERGANTI, R.; BELLINI, E. The role of human interaction in the service experience: A qualitative study. **Journal of Service Management**, v. 28, n. 3, p. 401-420, 2017.

PODSAKOFF, P. M.; ORGAN, D. W. Self-Reports in Organizational Research: Problems and Prospects. **Journal of Management**, v. 12, n. 4, pp. 531-544, 1996.

PODSAKOFF, P. M.; MACKENZIE, S. B.; LEE, J-Y.; PODSAKOFF, N. P. Common Method Biases in Behavioral Research: A Critical Review of the Literature and Recommended Remedies. **Journal of Applied Psychology**, v. 88(5), p. 879-903, 2003.

PRINCI, E; KRÄMER, N. I spy with my little sensor eye: The effect of data-tracking and convenience on the intention to use smart technology. In **Proceedings of the Annual Hawaii International Conference on System Sciences**, 2020.

PURNOMO, M.; MAULINA, E.; WICAKSONO, A. R.; RIZAL, M. Implementasi technology acceptance model. **Jurnal Manajemen Dan Kewirausahaan**, 9(2), 173-181, 2021.

RAHMAN, M.; MING, T. H.; BAIGH, T. A.; SARKER, M. Adoption of artificial intelligence in banking services: an empirical analysis. **International Journal of Emerging Markets**, 18(10), 4270-4300, 2021.

REHMAN, M. U.; BHATTI, A.; MOHAMED, A. R.; AYOUP, A. The impact of trust on online purchase intentions: A study of online shopping behavior. **Journal of Retailing and Consumer Services**, Amsterdam, v. 50, p. 1-9, 2019.

REZWANA, F.; MAHER, K. Privacy and ethics in AI-enabled platforms. **Digital Ethics Journal**, v. 6, n. 3, p. 201-218, 2022.

ROY, A. et al. The Use of Restraint and Seclusion in Residential Treatment Care for Youth: A Systematic Review of Related Factors and Interventions. **Trauma Violence & Abuse**, v. 20, n. 1, p. 22-34, 2019.

ROY, S. K.; BALAJI, M. S.; SADEQUE, S.; NGUYEN, B.; MELEWAR, T. C. Customer Engagement in the Digital Age: The Role of Convenience. **Journal of Business Research**, v. 70, p. 1-10, 2017.

RUIVO, M.A.; PIRES, A.F. Limites do whistleblower – denúncia de crimes contra a administração pública, ilícitos administrativos e ações lesivas ao interesse público. **Revista Brasileira de Ciências Criminais**, São Paulo, v. 28, n. 174, p. 41- 69, 2020.

RUSSEL, S.; HAUERT, S; ALTMAN, R. B.; VELOSO, M. **Robotics: Ethics of artificial intelligence**. *Nature*, 521(7553), 415– 418, 2015.

SAAD-DINIZ, E.; MARIN, G. C. Criminalidade empresarial e programas de whistleblowing: defesa dos regimes democráticos ou mercancia de informações? **Revista eletrônica do CPJM**, v. 1, n. 1, p. 72-99, 2021.

SARATCHANDRAN, V. Artificial intelligence (AI): Ways AI is redefining the future of customer service, 2020.

SARI, B. Challenges in whistleblowing mechanisms and the role of AI. **Compliance Today**, v. 14, n. 2, p. 89-102, 2023.

SATTU, R.; DAS, S.; JENA, L. K. Should I adopt ai during talent acquisition? Evidence from HR professionals of Indian IT Organizations. **Journal of Organizational Effectiveness: People and Performance**, 11(4), 1005-1022, 2024.

SEIDERS, K.; VOSS, G. B.; GODFREY, A. L.; GALY, E. The role of convenience in service consumption. **Journal of Service Research**, 7(2), 129-142, 2005.

SEIDERS, K.; VOSS, G. B.; GODFREY, A. L.; GREWAL, D. SERVCON: Development and Validation of a Multidimensional Service Convenience Scale. **Journal of the Academy of Marketing Science**, v. 35, pp. 144-156, 2007.

SHARMA, V., Artificial Intelligence in fraud detection and personalization: transforming the landscape of security and user experience. **Journal of Artificial Intelligence & Cloud Computing**, v. 18, p. 45-60, 2022.

SHIN, D. H.; LIN, H. The impact of perceived sacrifice on customer satisfaction in mobile services. **Computers in Human Behavior**, v. 65, p. 1-10, 2016.

SIAU, K.; WANG, W. Building trust in artificial intelligence, machine learning, and robotics. **Journal of Database Management**, Hershey, v. 29, n. 1, p. 1-10, 2018.

SILVA, F. J.; FERREIRA, J. J.; SANTOS, J. Employees' perception of corporate social responsibility and performance: the mediating roles of job satisfaction,

organizational commitment and organizational trust. **Journal of Strategy and Management**, Bradford, v. 15, n. 3, p. 1-20, 2022.

SKINNER, G.; HAN, S.; CHANG, E. Defining and Protecting Meta Privacy: A New Conceptual Framework within Information Privacy, in **Proceedings of the 22nd International Conference on Data Engineering Workshops**, Atlanta, GA, April 3-8, 2006.

SMUHA, N. The EU approach to ethics guidelines for trustworthy artificial intelligence. **Computer Law Review International**, Oxford, v. 19, n. 4, p. 1-10, 2019.

SOLOVE, D.J. A taxonomy of privacy. *Univ. PA. Law Rev.* 154, 477–560, 2006.

SPIELBERGER, C. D. Theory and research on anxiety. In C. D. Spielberger (Ed.), **Anxiety and behaviour** (pp. 3–20). New York: Academic Press, 1966.

STAHL, B. C.; WRIGHT, D. Ethics and privacy in AI and big data: Implementing responsible research and innovation. **IEEE Security & Privacy**, 16(3), 26–33, 2018.

Sui, Y.; Zhang, H.; Wang, L. Impact of cultural tightness on consumers' preference for anthropomorphic AI services. **Psychology and Marketing**, 2024.

SWANSON, J.A. The public and private in Aristotle's political philosophy. Ithaca, NY: **Cornell University Press**, 1992.

TAN, S. Factors Influencing Consumer Purchase Intention: A Study of Convenience. **International Journal of Accounting & Finance in Asia Pacific**, v. 5, n. 2, p. 159-170, 2022.

THIEBES, S.; LINS, S.; HUBER, J. Trustworthy artificial intelligence: A systematic literature review. **Electronic Markets**, Berlin, Heidelberg, v. 30, n. 4, p. 1-18, 2020.



THIEL, J. Real-Time Information and Customer Support in AI Services. **Journal of Business Research**, v. 98, p. 1-10, 2019.

UNITED STATES SENTENCING COMMISSION. 2024 Guidelines Manual: Chapter 8 – Sentencing of Organizations. 2024. Disponível em: <https://www.ussc.gov/guidelines/2024-guidelines-manual/annotated-2024-chapter-8#8b21>. Acesso em: 11 de fevereiro de 2025.

VENKATESH, V.; MORRIS, M. G.; DAVIS, G. B.; DAVIS, F. D. User acceptance of information technology: toward a unified view. **MIS Quarterly**, v. 27, n. 3, p. 425–478, set. 2003.

VILLANUEVA-MOYA, L.; EXPÓSITO, F. How do women and men perceive the sacrifice of leaving work for their families? A cost–benefit Analysis. **Personality and Social Psychology Bulletin**, 51(4), 495-513, 2023.

WALCH, K. The Role of AI in Enhancing Service Convenience. **Journal of Service Management**, v. 30, n. 4, p. 487-507, 2019.

WANG, C.; XIAO, A. Anxiety induced by artificial intelligence (AI) painting: an investigation based on the fear acquisition theory. **Psychological Trauma: Theory, Research, Practice, and Policy**, 2025.

WANG, Y,Y.; Wang, Y.S. Development and validation of an artificial intelligence anxiety scale: An initial application in predicting motivated learning behavior. **Interactive Learning Environments**, 30(4), 619–634, 2022.

WANG, Y.; TAJVIDI, M.; LIN, X.; HAJLI, N. The role of trust in social media in building customer commitment to online retailers. **Journal of Retailing and Consumer Services**, Amsterdam, v. 50, p. 1-10, 2019.

WANG, Y.; WANG, Y.; LIU, Y. The impact of trust on customer commitment in group buying: The mediating role of customer satisfaction. **Journal of Business Research**, Amsterdam, v. 69, n. 8, p. 2990-2995, 2016.

WINTER, S.; CARUSI, K. 'If You're Going to Trust the Machine, Then That Trust Has Got to Be Based on Something'. **Science & Technology Studies**, London, v. 35, n. 1, p. 1-20, 2022.

WONG, L.; TAN, G. W.; OOI, K.; DWIVEDI, Y. K. The role of institutional and self in the formation of trust in artificial intelligence technologies. *Internet Research*, 34(2), 343-370, 2023.

XU, H.; LUO, X.; CARROLL, J. M.; ROSSON, M. B. The personalization privacy paradox: An exploratory study of decision-making process for location-aware marketing. **Decision Support Systems**, 51(1), 42-52, 2011.

YANG, X.; FANG, Y. CHEN, H.; ZHANG, T.; YIN, X.; MAN, J.; LÜ, M. Global, regional and national burden of anxiety disorders from 1990 to 2019: results from the global burden of disease study 2019. **Epidemiology and Psychiatric Sciences**, 30, 2021.

YANG, Y.; LUO, J.; LAN, T. An empirical assessment of a modified artificially intelligent device use acceptance model—from the task-oriented perspective. **Frontiers in Psychology**, 13, 2023.

YANG, Y.; ZHANG, Y.; WANG, Y. Attitudes of Chinese Cancer Patients Toward the Clinical Use of Artificial Intelligence. **Patient Preference and Adherence**, London, v. 13, p. 1-10, 2019.

YAU, M.; WONG, K.; LEE, T. Artificial Intelligence Marketing (AIM) for Enhancing Customer Relationships. **Applied Sciences**, Basel, v. 11, n. 18, p. 1-20, 2021.

YIN, J.; QIU, X. AI technology and online purchase intention: structural equation model based on perceived value. **Sustainability**, 13(10), 5671, 2021.

YOUNES, S. S.; ALHARBI, A. H.; ABOELDAHAB, M. M. Measuring the Impact of an AI-Enabled Mobile Application for University Students. **Mobile Information Systems**, 4141576, 2021.

ZEITHAML, V. A. Consumer perceptions of price, quality, and value: A means-end model and synthesis of evidence. **Journal of Marketing**, v. 52, n. 3, p. 2-22, 1988.

ZHANG, Y.; BILGIHAN, A.; KANDAMPULLY, J.; LU, Y. The role of trust in building brand loyalty in online communities. **Journal of Retailing and Consumer Services**, Amsterdam, v. 45, p. 1-10, 2018.

## Apêndice A – Questionário utilizado na pesquisa

### Pesquisa sobre o uso da inteligência artificial em canal de denúncia

Prezado(a),

Me chamo Maria Fernanda Araujo de Carvalho, mestranda em Administração no Programa de Pós-Graduação da Pontifícia Universidade Católica do Rio de Janeiro (PUC-Rio), e estou realizando a coleta de dados para a minha dissertação de mestrado, sob orientação do Prof. Dr. Jorge Brantes Ferreira.

Na dissertação, irei **analisar como a inteligência artificial (ex. ChatGPT) está transformando a experiência do usuário ao realizar uma denúncia em um canal de denúncia de uma empresa.**

A sua participação consiste em responder de forma fidedigna o questionário **até o final**. Os dados serão analisados de forma consolidada, **garantindo seu anonimato**.

O tempo médio de preenchimento do questionário é de **10 minutos**.

Você poderá entrar em contato pelo e-mail [mariacarvalho1507@gmail.com](mailto:mariacarvalho1507@gmail.com) para esclarecer eventuais dúvidas.

**Certamente a sua participação fará total diferença na consolidação da minha pesquisa, muito obrigada pelo seu tempo!**

Maria Fernanda  
Mestranda no Programa de Pós-Graduação em Administração – PUC-Rio

---

\* Indica uma pergunta obrigatória

#### Experiência Anterior com Ferramentas de inteligência Artificial \*

Você já utilizou alguma ferramenta de inteligência artificial? Por exemplo ChatGPT, Copilot, Gemini ou outros?

- ☐ Diariamente
- ☐ Semanalmente
- ☐ Mensalmente
- ☐ Raramente
- ☐ Nunca usei

### Exemplo de Ferramenta de IA para recebimento de denúncias

Agora, você analisará algumas capturas de tela. O objetivo é explorar como a tecnologia de inteligência artificial (IA), como o ChatGPT, auxilia o denunciante no processo de realização de uma denúncia no canal de denúncia da Empresa X.

A ferramenta de inteligência artificial atua como um guia, ajudando o denunciante a entender, por exemplo, qual a diferença entre fazer uma denúncia e fazer uma reclamação contra a Empresa X, como elas poderiam ser feitas de forma anônima ou identificada, entre outros.

Durante a análise das telas, você será estimulado(a) a avaliar se a inteligência artificial (ex. ChatGPT) está transformando a experiência do usuário durante a realização de uma denúncia no canal da denúncia da Empresa X. Posteriormente, será convidado a responder algumas perguntas.

#### Descrição da Ferramenta e como ela utiliza o ChatGPT:

A ferramenta foi criada com base no ChatGPT com objetivo de agir como um canal de denúncia. Ela foi treinada com dados específicos para atuar da melhor forma possível no auxílio e suporte a realização de uma denúncia ou ainda uma reclamação.

**Consumidores e Funcionários podem realizar uma denúncia ou reclamação contra a Empresa X através do canal de denúncias.**

O fácil acesso e o uso correto do canal de denúncia é fundamental para que as informações cheguem até as empresas de forma ágil e com o maior nível de detalhes possível facilitando a condução da investigação dos relatos.

**Por favor, selecione uma das opções abaixo: \***

- ☐ Nunca fiz uma denúncia em um canal de denúncia de uma empresa
- ☐ Não me recordo se já fiz uma denúncia em um canal de denúncia antes
- ☐ Já fiz uma denúncia em um canal de denúncia antes

## Leitura dos diálogos

Neste momento você deverá analisar algumas capturas de tela de uma interação de um cliente com uma Ferramenta de Inteligência Artificial usada no Canal de Denúncia de uma empresa.

**Leia com atenção toda a conversa de exemplo, para entender bem como funciona essa Inteligência Artificial.**

### Passo 1

**EMPRESA X** *Powered by* [Marca do canal de denúncias]



#### Olá! Vamos iniciar o seu relato?

Eu sou Ana, a IA do Canal de Denúncias e vou te ajudar a enviar a sua denúncia ou reclamação hoje.

Não se preocupe, seus dados estão seguros comigo!

Vamos lá?

Iniciar

### Passo 2

**EMPRESA X** *Powered by* [Marca do canal de denúncias]



Bem vindo/a!

Este é um meio seguro para você reportar a sua mensagem.

Primeiro me fale, você gostaria de fazer uma denúncia ou uma reclamação?

- > O que é uma denúncia
- > O que é uma reclamação

## Passo 3

**EMPRESA X** *Powered by* [Marca do canal de denúncias]**Você diz:**

O que é uma denúncia?

**Ana está digitando...**

## Passo 4

**EMPRESA X** *Powered by* [Marca do canal de denúncias]**Você diz:**

O que é uma denúncia?



Uma denúncia é quando você informa algo errado ou grave que está acontecendo na Empresa X. Isso pode ser algo ilegal ou que vai contra as regras e valores da empresa. Por exemplo, **Assédio moral ou sexual** (quando alguém sofre humilhações ou investidas indesejadas) e **Fraudes e corrupção** (como pagar ou receber dinheiro de forma errada para conseguir vantagens).

---

## Passo 5

**EMPRESA X** *Powered by* [Marca do canal de denúncias]

**Você diz:**

Quero fazer uma denúncia



Claro, você quer se identificar e deixar seu nome, email e telefone para a Empresa X ou gostaria de fazer de forma anônima, sem deixar nenhum dado seu?

**Você diz:**

Quem terá acesso aos meus dados?

## Passo 6

**EMPRESA X** *Powered by* [Marca do canal de denúncias]



A Empresa X terá acesso aos seus dados.

Caso você resolva se identificar, ela terá acesso ao seu nome, email e telefone.

>Quero me identificar

> Quero fazer minha denúncia de forma anônima sem compartilhar meus dados

**Você diz:**

Quero me identificar



## Passo 7

**EMPRESA X** *Powered by* [Marca do canal de denúncias]

Entendi, vamos fazer essa denúncia de forma identificada. Por favor, me conte seu nome, email e telefone.

**Você diz:**

Maria Silva, [maria@email.com](mailto:maria@email.com), (01)000-0000



Obrigada! Agora me conte o que aconteceu e quando isso ocorreu. Não esqueça de deixar claro qual foi o local do ocorrido e se outras pessoas presenciaram a situação com você. Se você tiver evidências, não deixe de me informar também!

## Passo 8

**EMPRESA X** *Powered by* [Marca do canal de denúncias]**Você diz:**

Eu sou cliente da loja de vocês no shopping há mais de 10 anos. Esse sábado eu fui abordada por um segurança novo que pediu para ver a minha sacola de compras com as peças que eu havia acabado de comprar. Eu questionei o porquê e ele foi muito grosseiro e agressivo em suas palavras, me disse que eu não deveria questionar e apenas mostrar logo a minha sacola de compras. Ele tirou tudo de dentro dela, deixou as roupas no chão e me entregou a sacola vazia. Falou que eu estava liberada. Nunca fui tão maltratada.

## Passo 9

**EMPRESA X** *Powered by* [Marca do canal de denúncias]

Maria, para que possamos tratar o quanto antes a sua denúncia, poderia me informar a data, horário e qual Bairro foi ocorrido? Por fim, você teria o nome do segurança que te abordou? E poderia compartilhar conosco a vendedora que te atendeu?

**Você diz:**

Sim, a vendedora é a Julia, foi sábado dia 15/02/25, estive na loja do Bairro Y entre 13h e 14h. O nome do segurança é José.

## Passo 10

**EMPRESA X** *Powered by* [Marca do canal de denúncias]

Ok, entendi. Por favor, veja se está correto:

Maria Silva, telefone (01) 000-000 e email [maria@email.com](mailto:maria@email.com), é cliente da loja do Bairro Y há mais de 10 anos. Dia 15/02/25, sábado, entre 13h e 14h ela esteve na loja. Ao se direcionar para a porta para ir embora, Maria foi abordada pelo segurança José, ainda dentro da loja, que pediu para ver a sua sacola de compras. Maria questionou o porquê e ele foi muito grosseiro e agressivo em suas palavras, disse que ela não deveria questionar e apenas mostrar logo a sua sacola de compras. Maria afirma que José tirou tudo de dentro da sacola, deixou as roupas no chão e a entregou a sacola vazia. Então, falou que Maria estava liberada. Maria afirma que nunca foi tão maltratada.

## Passo 11

**EMPRESA X** *Powered by* [Marca do canal de denúncias]**Você diz:**

Sim, o que você descreveu está correto.



Certo, obrigada pela confirmação. Sua denúncia foi enviada a Empresa X, o seu protocolo é o **0987ZY9OLYY**.

Por favor, anote em um local seguro, através dele você vai poder acompanhar o andamento da sua denúncia.

Ótimo dia e até logo!

**Finalizar**

Confirma se você já realizou a leitura dos diálogos? \*

- ☐ Sim, confirmo que li com atenção os prints.
- ☐ Não, não entendi o que foi apresentado nos prints.

Após a interação com a IA do canal de denúncia, solicitamos que você responda ao questionário abaixo.

Por favor, leia cada afirmação com atenção e avalie-as de acordo com a sua experiência. As respostas devem ser dadas com base na seguinte escala de 1 a 5:

- 1: Discordo totalmente
- 2: Discordo parcialmente
- 3: Não concordo nem discordo
- 4: Concordo parcialmente
- 5: Concordo totalmente

**Ao responder, tenha em mente a experiência vivida durante a interação com a IA do canal de denúncia.**

A IA utilizada no canal de denúncia foi bem projetada \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

A IA utilizada no canal de denúncia fornece informações corretas \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

A IA utilizada no canal de denúncia é segura. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

A IA utilizada no canal de denúncia é atenciosa

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

O desempenho da IA no canal de denúncia sempre atende às minhas expectativas.

\*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

A IA do canal de denúncia possui bons recursos. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

A IA do canal de denúncia é confiável. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

A IA utilizada no canal de denúncia é honesta. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Eu acredito na IA utilizada no canal de denúncia \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Ao disponibilizar uma IA no canal de denúncia a empresa demonstra interesse em mim como cliente. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

A IA do canal de denúncia me permite usar o serviço sempre que eu escolher. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

A IA do canal de denúncia me permite usar o serviço em um momento conveniente.

\*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Eu valorizo a possibilidade de usar a IA do canal de denúncia no conforto de casa.

\*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Eu fico preocupado(a) com a perda de privacidade ao usar a IA no canal de denúncia. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Usar a IA no canal de denúncia um exige esforço adicional. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Fico preocupado(a) com a perda de controle ao usar a IA no canal de denúncia. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Usar a IA no canal de denúncia consome muito tempo. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Fico irritado(a) ao usar a IA no canal de denúncia. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente



Eu fico preocupado(a) com a falta de interação humana ao usar a IA no canal de denúncia. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Eu me sinto apreensivo em usar soluções de IA. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Me assusta pensar que eu poderia perder uma grande quantidade de informação por ter utilizado a IA do canal de denúncias de forma incorreta, por exemplo se eu não for claro o suficiente. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Eu hesito em usar uma solução de canal de denúncia que utilize IA por medo de erros que eu não conseguirei corrigir. \*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Soluções de canal de denúncia que utilizam IA são de alguma forma intimidantes para mim.

\*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Preocupa-me que a IA do canal de denúncia partilhe as minhas informações pessoais de forma inadequada/não autorizada.

\*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Receio que IA do canal de denúncia resulte em uma invasão da minha privacidade.

\*

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Preocupa-me a forma como as minhas ações enquanto usuário da IA do canal de \*  
denúncia possa ser controladas e monitoradas.

- ☐ Discordo totalmente
- ☐ Discordo parcialmente
- ☐ Não concordo nem discordo
- ☐ Concordo parcialmente
- ☐ Concordo totalmente

Seção 5 de 6

Dados Demográficos



O questionário está chegando ao fim. Agradecemos por seguir conosco, seguem as últimas perguntas.

Qual o seu gênero? \*

- ☐ Mulher
- ☐ Homem
- ☐ Outro

Qual a sua idade? \*

Texto de resposta curta

Qual a sua faixa de renda mensal? \*

(Somando a renda de todos os familiares que vivem na sua residência)

Para fins de análise estatística, gostaríamos de conhecer sua faixa de renda familiar mensal. Por favor, considere a renda de todos os membros da sua família que residem na mesma casa. As informações fornecidas serão tratadas de forma confidencial. (campo obrigatório)

- ☐ Até 1 salário mínimo
- ☐ De 1 a 2 salários mínimos
- ☐ De 2 a 5 salários mínimos
- ☐ De 5 a 10 salários mínimos
- ☐ Acima de 10 salários mínimos

Qual o seu nível de escolaridade mais alto? \*

- ☐ Ensino fundamental
- ☐ Ensino médio
- ☐ Ensino superior
- ☐ Mestrado
- ☐ Doutorado

Seção 6 de 6

Obrigada!



**Muito obrigada por dedicar seu tempo para responder ao questionário.** Suas contribuições são extremamente valiosas para minha pesquisa!

Até logo!