



UNIVERSIDADE FEDERAL DO PARÁ
INSTITUTO DE CIÊNCIAS SOCIAIS APLICADAS
PROGRAMA DE PÓS GRADUAÇÃO EM ADMINISTRAÇÃO
MESTRADO ACADÊMICO EM ADMINISTRAÇÃO

NERIVANE DA SILVA MENDES

QUANDO AS MÁQUINAS PARECEM HUMANAS DEMAIS: Intrusividade, Estranheza
e Cinismo da Privacidade no uso de Chatbots de IA Generativa

BELÉM-PA
2026

NERIVANE DA SILVA MENDES

QUANDO AS MÁQUINAS PARECEM HUMANAS DEMAIS: Intrusividade, Estranheza
e Cinismo da Privacidade no uso de Chatbots de IA Generativa

Dissertação apresentada ao Programa de Pós-Graduação em Administração – PPGAD, do Instituto de Ciências Sociais Aplicadas, da Universidade Federal do Pará, como requisito parcial para obtenção do título de Mestre em Administração.

Área de concentração: Estratégia e Desempenho Organizacional.

Orientador: Prof. Dr. Thiago Poletto

BELÉM-PA
2026

Dados Internacionais de Catalogação na Publicação (CIP) de acordo com ISBD
Sistema de Bibliotecas da Universidade Federal do Pará
Gerada automaticamente pelo módulo Ficat, mediante os dados fornecidos pelo(a) autor(a)

M538q Mendes, Nerivane.
 Quando as máquinas parecem humanas demais: intrusividade,
 estranheza e cinismo da privacidade no uso de chatbots de IA
 generativa / Nerivane Mendes. — 2026.
 37 f. : il. color.

 Orientador(a): Prof. Dr. Thiago Poletto
 Dissertação (Mestrado) - Universidade Federal do Pará,
 Instituto de Ciências Sociais Aplicadas, Programa de Pós-
 Graduação em Administração, Belém, 2026.

 1. cinismo da privacidade. 2. chatbots. 3. vulnerabilidade
 percebida. 4. proteção da privacidade. 5. privacidade digital. I.
 Título.

CDD 006.3

NERIVANE DA SILVA MENDES

**QUANDO AS MÁQUINAS PARECEM HUMANAS DEMAIS: Intrusividade, Estranheza
e Cinismo da Privacidade no uso de Chatbots de IA Generativa**

Dissertação apresentada ao Programa de Pós-Graduação
em Administração – PPGAD, do Instituto de Ciências
Sociais Aplicadas, da Universidade Federal do Pará,
como requisito parcial para obtenção do título de Mestre
em Administração.

Data da aprovação: __/__/__
Conceito: _____

Banca examinadora:

Presidente: Prof Thiago Poletto (Orientador)
Titulação: Doutorado em Engenharia de Produção.
Universidade Federal de Pernambuco, UFPE, Brasil.

1º Examinador:
Professor Dr. Fernando de Assis Rodrigues
Titulação: Doutorado em Ciência da Informação
Universidade Estadual Paulista Júlio de Mesquita Filho, UNESP, Brasil.

2º Examinador Externo:
Professor Dr. Luis Hernan Contreras Pinochet
Titulação: Doutorado em Administração de Empresas.
Fundação Getulio Vargas - SP, FGV-SP

AGRADECIMENTOS

A Deus, primeiramente, por me conceder força, saúde, sabedoria e perseverança para seguir firme ao longo desta jornada, mesmo diante dos desafios e incertezas.

À minha família, por estar sempre presente, oferecendo apoio, incentivo, compreensão e paciência em todos os momentos, especialmente nos períodos mais intensos desta trajetória acadêmica.

Ao meu orientador, Professor Dr. Thiago Poletto, pela dedicação, competência e generosidade ao longo de todo o processo, desde a construção do projeto até a finalização desta dissertação. Seu apoio, confiança e orientação foram fundamentais para o meu crescimento acadêmico.

Aos demais professores do programa, em especial aos membros da banca examinadora, pelas valiosas contribuições, sugestões e reflexões, que enriqueceram significativamente este trabalho.

À minha turma de mestrado, pela parceria, companheirismo e apoio mútuo. Entre estudos, trocas e aprendizados, compartilhamos também muitas risadas, que tornaram essa trajetória ainda mais especial.

Aos meus colegas de trabalho da PROGEP/UFPA, pelo apoio constante, incentivo, e parceria, que tornaram essa caminhada mais leve e possível.

Às demais pessoas que estiveram ao meu lado desde a concepção da ideia de cursar o mestrado, que acompanharam o processo seletivo, contribuíram com debates, reflexões, e palavras de motivação. Sou imensamente grata a todos que, de alguma forma, acreditaram em mim e fizeram parte desta conquista.

Por fim, agradeço a mim mesma, pela coragem de sonhar e não desistir, mesmo diante das dificuldades. Esta conquista representa o resultado de muito esforço e dedicação.

RESUMO

Esta dissertação examina os antecedentes cognitivos do cinismo da privacidade em interações automatizadas, com foco específico na intrusividade percebida e na estranheza percebida no uso de chatbots. Além disso, investiga o papel mediador da vulnerabilidade percebida na relação entre o cinismo da privacidade e as intenções de proteção da privacidade. Adotou-se um delineamento de pesquisa quantitativo, realizado por meio de um questionário on-line aplicado a 257 usuários brasileiros que haviam interagido com *chatbots* baseados em inteligência artificial em suas atividades cotidianas nos seis meses anteriores à coleta de dados. Os dados foram analisados por meio da Modelagem de Equações Estruturais por Mínimos Quadrados Parciais (PLS-SEM). A intrusividade percebida exerceu um efeito mais substancial sobre o cinismo da privacidade do que a estranheza percebida, embora ambos os construtos tenham se mostrado preditores significativos. Além disso, a vulnerabilidade percebida mediou plenamente a relação entre o cinismo da privacidade e a intenção de proteção da privacidade, indicando que o cinismo, por si só, não conduz a comportamentos protetivos. Os resultados oferecem subsídios para o desenvolvimento de sistemas de atendimento ao cliente baseados em inteligência artificial, destacando a importância de equilibrar personalização, transparência e proteção de dados de forma compatível com as percepções e expectativas dos usuários. O estudo evidencia a relevância social do desenvolvimento de tecnologias que preservem a autonomia, a dignidade e a confiança dos usuários em ambientes digitais. Os resultados contribuem para os debates mais amplos sobre inteligência artificial ética e proteção de dados, ao apoiar o avanço de práticas de governança inclusivas e transparentes que auxiliem na mitigação da fadiga digital e da resignação em relação à privacidade.

Palavras-chave: cinismo da privacidade; chatbots; vulnerabilidade percebida; proteção da privacidade; privacidade digital.

ABSTRACT

This dissertation examines the cognitive antecedents of privacy cynicism in automated interactions, with a particular focus on perceived intrusiveness and perceived strangeness in chatbot use. Additionally, it explores the mediating role of perceived vulnerability in the relationship between privacy cynicism and intentions to protect privacy. This study adopted a quantitative research design. It was conducted via an online survey of 257 Brazilian users who had interacted with AI-based chatbots in their daily activities within 6 months of data collection. Data were examined using Partial Least Squares Structural Equation Modeling (PLS-SEM). Perceived intrusiveness exerted a more substantial effect on privacy cynicism than perceived strangeness, although both constructs were significant predictors. In addition, perceived vulnerability fully mediated the relationship between privacy cynicism and privacy protection intention, indicating that cynicism alone does not necessarily lead to protective behavior. The findings offer guidance for the design of AI-powered customer service systems, emphasizing the importance of balancing personalization, transparency, and data protection in ways that are consistent with users' perceptions and expectations. This study highlights the societal importance of developing technologies that safeguard user autonomy, dignity, and trust in digital environments. The findings contribute to broader debates on ethical AI and data protection, supporting the advancement of inclusive and transparent governance practices that help mitigate digital fatigue and privacy resignation.

Keywords: privacy cynicism; chatbots; perceived vulnerability; privacy protection; digital privacy.

LISTA DE FIGURAS

Figura 1 – Modelo teórico da pesquisa.....	17
Figura 2 – Teste do modelo estrutural	25

LISTA DE TABELAS

Tabela 1 – Base teórica dos construtos.....	13
Tabela 2 – Critérios para Avaliação dos Modelos de Mensuração e Estrutural.....	20
Tabela 3 – Dados demográficos da pesquisa.....	21
Tabela 4 – Cargas Fatoriais, Alfa de Cronbach, Confiabilidade Composta e AVE.....	22
Tabela 5 – Teste de validade discriminada (critério de Fornell-Larcker)	24
Tabela 6 – Teste de validade discriminada (razão heterotraço-monotraço — HTMT).....	24
Tabela 7 – Capacidade explicativa do modelo (R^2 , R^2 ajustado, f^2)	26
Tabela 8 – Resultados dos testes de hipóteses.....	27

SUMÁRIO

1 INTRODUÇÃO	10
2 FUNDAMENTAÇÃO TEÓRICA E HIPÓTESES	13
2.1 Intrusividade Percebida	13
2.2 Percepção de Estranheza	14
2.3 Cinismo da Privacidade, Vulnerabilidade Percebida e Intenção de Proteção da Privacidade.....	15
3 MATERIAIS E MÉTODOS	18
3.1 Coleta de dados e perfil da amostra.....	18
3.2 Instrumento de medida	19
4 RESULTADOS	21
4.1 Abordagem.....	21
4.2 Normalidade, Viés de Método Comum e Colinearidade.....	21
4.3 Teste do modelo de mensuração.....	22
4.4 Teste do modelo estrutural	25
5 DISCUSSÃO	28
5.1 Antecedentes do cinismo da privacidade.....	28
5.2 O papel da vulnerabilidade percebida.....	29
5.3 Mediação e implicações comportamentais	30
5.4 Implicações teóricas e práticas	31
6 CONCLUSÃO, LIMITAÇÕES E PESQUISAS FUTURAS	33
REFERÊNCIAS	35
APÊNDICE A – ESCALAS ADAPTADAS DOS CONSTRUTOS.	42
APÊNDICE B – TELAS E TEXTO VÍDEO QUESTIONÁRIO	43

1 INTRODUÇÃO

O avanço dos *chatbots* baseados em Inteligência Artificial (IA) tem remodelado as interações entre humanos e máquinas (Adamopoulou; Moussiades, 2020; Lim *et al.*, 2025; Pitardi; Marriott, 2021). A evolução de sistemas iniciais baseados em regras, como ELIZA e PARRY, para modelos generativos apoiados em Modelos de Linguagem de Grande Escala (*Large Language Models* – LLMs), possibilitou que sistemas conversacionais produzissem linguagem de forma mais natural e sensível ao contexto da interação (Radford *et al.*, 2019; Zhou *et al.*, 2020). Essa transformação tecnológica modificou a forma como os usuários percebem os riscos relacionados à privacidade, considerando a interação com *chatbots* (Dinev; Hart, 2006; Gnewuch *et al.*, 2017).

No contexto organizacional, a inteligência artificial generativa não configura uma simples contradição, mas uma tensão estrutural entre inovação tecnológica e responsabilidade informacional. *Chatbots* humanizados ampliam o engajamento, a personalização, a eficiência operacional e a capacidade de inovação nos serviços digitais (Chen *et al.*, 2022, 2025; Moriuchi; Murdy, 2024; Pinochet *et al.*, 2024). Contudo, os mesmos atributos que tornam essas ferramentas mais sofisticadas também intensificam preocupações relacionadas à transparência algorítmica, ao controle informacional e à proteção da privacidade. Estudos recentes indicam que essa combinação pode desencadear respostas psicológicas e éticas adversas (Hong, 2025; Ko *et al.*, 2025; Sohn *et al.*, 2025), manifestadas por meio da percepção de intrusividade e de sentimentos de estranheza nas interações digitais.

À medida que as distinções entre interação humana e artificial se tornam cada vez mais difusas, os usuários percebem a coleta e o uso de dados pessoais como invasivos (Edwards *et al.*, 2002; MacDorman, 2006; Mori *et al.*, 2012). A semelhança quase humana e a imprevisibilidade comportamental dos sistemas de IA causam uma percepção de desconforto emocional, intensificando o risco informacional (Liu; Wang, 2025). Nesse contexto, o cinismo da privacidade refere-se a um estado em que os usuários reconhecem esses riscos, mas percebem alternativas limitadas e continuam a utilizar o sistema apesar das preocupações (Choi *et al.*, 2018; Segijn *et al.*, 2024). Apesar da relevância, poucos estudos examinaram de forma integrada como respostas cognitivas e emocionais à IA humanizada moldam o cinismo da privacidade e as subsequentes intenções de autoproteção.

Este estudo tem como objetivo examinar como a intrusividade percebida e a percepção de estranheza em interações com *chatbots* baseados em inteligência artificial generativa influenciam o cinismo da privacidade e, subsequentemente, de que modo esse cinismo afeta as

intenções de proteção da privacidade dos usuários, por meio do papel mediador da vulnerabilidade percebida. Para isso, propõe-se e testa-se empiricamente um modelo teórico fundamentado na Teoria da Motivação para a Proteção (Protection Motivation Theory – PMT) (Liang; Xue, 2009; Maddux; Rogers, 1983), integrando antecedentes cognitivo-emocionais e mecanismos motivacionais que explicam as respostas dos usuários diante de ameaças percebidas à privacidade em interação humano-IA.

Com base no objetivo geral de analisar, à luz da Teoria da Motivação da Proteção, como respostas cognitivas e emocionais às interações com chatbots de IA generativa se articulam para explicar atitudes e intenções relacionadas à privacidade, este estudo é guiado pela seguinte questão de pesquisa:

QP: Como a intrusividade percebida e a percepção de estranheza afetam o cinismo da privacidade e por meio de quais mecanismos esse cinismo influencia as intenções dos usuários de proteger a privacidade ao interagir com *chatbots* baseados em IA generativa?

Embora pesquisas anteriores tenham examinado o cinismo da privacidade no âmbito de plataformas digitais e redes sociais (Lyu *et al.*, 2024), a principal contribuição teórica desta dissertação consiste na transposição e readequação conceitual desse fenômeno para o ecossistema das interações humanizadas mediadas por Inteligência Artificial. O estudo enriquece a literatura ao evidenciar que, no contexto de agentes conversacionais generativos, a gênese do cinismo transcende as preocupações informacionais convencionais, sendo deflagrada por uma conjunção multidimensional de fatores estético-afetivos (percepção de estranheza) e cognitivos (intrusividade percebida). Adicionalmente, a pesquisa estende a Teoria da Motivação para a Proteção (PMT). Em contraposição a uma aplicação estritamente empírica e linear, o modelo estrutural proposto incorpora o cinismo da privacidade como um estado afetivo-cognitivo antecedente à avaliação de ameaça. Desse modo, o trabalho comprova empiricamente que o cinismo, frequentemente caracterizado na literatura como um estado de inércia comportamental, atua, de fato, como um mecanismo preditor indireto, que induz intenções protetivas exclusivamente quando internalizado sob a forma de vulnerabilidade percebida. Por conseguinte, esta dissertação fornece subsídios teóricos para os debates contemporâneos acerca da ética algorítmica, da governança de dados e da preservação da autonomia do usuário em ambientes digitais de alta complexidade.

Ademais, o estudo alinha-se à Agenda 2030 das Nações Unidas, contribuindo diretamente para o alcance das metas dos Objetivos de Desenvolvimento Sustentável (ODS).

Destacam-se, em especial, as metas do Objetivo 9 (Indústria, Inovação e Infraestrutura) voltadas ao desenvolvimento de capacidades tecnológicas sustentáveis, e as metas do Objetivo 16 (Paz, Justiça e Instituições Eficazes). O trabalho dialoga com a Meta 16.6 (desenvolver instituições eficazes, responsáveis e transparentes) e a Meta 16.10 (assegurar o acesso à informação e proteger as liberdades fundamentais). Ao apoiar a inovação tecnológica responsável, a pesquisa fornece subsídios para reforçar mecanismos éticos e de governança no ecossistema digital.

O restante do estudo está dividido da seguinte forma: a seção 2 aborda a fundamentação teórica e hipóteses da pesquisa. A seção 3, materiais e métodos, descreve a coleta de dados e perfil da amostra, assim como o instrumento de pesquisa. Os resultados são detalhados na seção 4, enquanto a discussão e as implicações, teóricas e práticas, estão na seção 5. Por fim, na seção 6 é abordada a conclusão, limitações e pesquisas futuras.

2 FUNDAMENTAÇÃO TEÓRICA E HIPÓTESES

A Tabela 1 apresenta estudos que sustentam a relação entre intrusividade percebida, percepção de estranheza, vulnerabilidade percebida e cinismo da privacidade. Em conjunto, esses estudos ilustram como o comportamento realista dos sistemas de IA evocam desconforto, ambivalência e respostas defensivas nos usuários.

Tabela 1 – Base teórica dos construtos

Construto	Autores e aplicação do estudo
Intrusividade Percebida	Edwards, Hairong and Lee (2002) – estudo seminal na aplicação do termo “intrusividade percebida” na publicidade digital.
	Li, Edwards and Lee (2002) – discutiram a intrusividade percebida, explorando as respostas psicológicas dos usuários aos anúncios on-line.
	Okazaki, Li and Hirose (2009) – aplicaram o conceito de intrusão percebida no marketing móvel.
Percepção de Estranheza	Masahiro Mori (1970) – Uncanny Valley: estudo seminal sobre a percepção de estranheza frente a robôs e representações humanas quase perfeitas.
	David Hanson (2006) – pesquisas sobre robôs humanoides e a reação de estranheza provocada pelo realismo imperfeito.
	Karl MacDorman (2006) – expandiram o conceito do vale da estranheza na robótica e na interação humano-computador.
Vulnerabilidade Percebida	Dinev and Hart (2004, 2006) – aplicaram o conceito de vulnerabilidade percebida à privacidade e à adoção de tecnologias, destacando os riscos percebidos no ambiente on-line.
	Liang & Xue (2009) – aplicaram a vulnerabilidade percebida em modelos de ameaça, fundamentados na Teoria da Motivação para a Proteção, no contexto da segurança da informação.
	Bulgurcu, Cavusoglu and Benbasat (2010) – Estudo fundamental na área de segurança da informação, integrando a vulnerabilidade percebida como determinante da conformidade nas políticas de segurança.

Fonte: Elaborado pela autora, 2026.

2.1 Intrusividade Percebida

A intrusividade percebida refere-se à reação cognitiva e emocional imediata diante de estímulos invasivos que rompem as expectativas de controle em ambientes digitais (Li *et al.*, 2002; Okazaki *et al.*, 2009). Em contextos mediados por inteligência artificial, essa percepção emerge quando os usuários se tornam conscientes de que sistemas automatizados coletam, interpretam ou monitoram continuamente dados pessoais (Lau *et al.*, 2019). Essa sensação de invasão é ainda intensificada pela opacidade dos processos de vigilância algorítmica, que limita a compreensão dos usuários sobre como suas informações são processadas (Akdim; Casaló, 2023; Aw *et al.*, 2024).

De acordo com a Teoria da Motivação para a Proteção (Protection Motivation Theory – PMT), essa percepção de intrusividade atua como um fator de ameaça que desencadeia

processos de avaliação cognitiva, nos quais o usuário do *chatbot* de IA pondera os riscos de exposição e perda de controle em relação aos benefícios percebidos, como conveniência e personalização (Başer *et al.*, 2024; Fu *et al.*, 2023; Kawaf *et al.*, 2024; Meng; Liu, 2025; Schomakers *et al.*, 2022). Quando esse equilíbrio é percebido como desfavorável, os usuários tendem a adotar orientações defensivas, que se manifestam na forma de cinismo da privacidade (Kang; Shao, 2023; Lyu *et al.*, 2024; Nguyen; Le, 2025; Shao; Ho, 2025). Com base nesse raciocínio teórico, propõe-se a seguinte hipótese:

H1: A intrusividade percebida influencia positivamente o cinismo da privacidade em relação ao uso de *chatbots* baseados em IA generativa.

2.2 Percepção de Estranheza

A percepção de estranheza constitui uma resposta emocional de desconforto e incongruência que ocorre quando sistemas artificiais remetem familiaridade semelhante à humana, mas não conseguem atender plenamente às expectativas sociais relacionadas à privacidade e ao comportamento ético de uso de dados, como no caso de *chatbots* com linguagem quase natural (Langer; König, 2018; MacDorman, 2006; Mori *et al.*, 2012). Esse desconforto ultrapassa aspectos relacionados à aparência ou ao estilo de interação e intensifica as suspeitas dos usuários quanto às intenções subjacentes aos comportamentos das plataformas (Hyun Baek; Kim, 2023).

Em sistemas com linguagem humanizada, a personalização excessiva ou a formulação de perguntas inadequadas podem gerar um desalinhamento entre as normas esperadas de interação (clareza, impessoalidade ou previsibilidade) e as respostas recebidas (ambíguas, improvisadas ou afetivas). Esses desalinhamentos ampliam o desconforto e reduzem o controle percebido pelos usuários tanto sobre a interação quanto sobre o tratamento de seus dados pessoais (Ischen *et al.*, 2020; Maduku *et al.*, 2025). Pesquisas anteriores indicam que essa percepção favorece o desenvolvimento do cinismo da privacidade por meio de múltiplos mecanismos, incluindo a violação de expectativas sociais, a atribuição de intenções manipulativas e o aumento da consciência sobre riscos informacionais, frequentemente potencializados pela ansiedade tecnológica (Johnson; Verdicchio, 2017; Kshetri *et al.*, 2024; Suta *et al.*, 2020; Tai, 2020). Sob a perspectiva da gestão da informação, a estranheza atua como um gatilho estético e cognitivo, enfraquecendo a confiança na governança de dados e ampliando as percepções de risco. Nesse sentido, propõe-se a seguinte hipótese:

H2: A percepção de estranheza nas interações com *chatbots* de IA generativa influencia positivamente o cinismo da privacidade.

2.3 Cinismo da Privacidade, Vulnerabilidade Percebida e Intenção de Proteção da Privacidade

O cinismo da privacidade reflete uma orientação cognitiva e emocional caracterizada pela descrença e pelo sentimento de impotência diante das práticas organizacionais de coleta e uso de dados pessoais, fundamentada na percepção de que o controle individual é ineficaz e de que a exposição digital é inevitável (Lutz *et al.*, 2020). Evidências empíricas indicam que o cinismo da privacidade compromete a confiança e enfraquece o engajamento em comportamentos proativos de proteção de dados, devido a transparência limitada e exposição algorítmica (Segijn *et al.*, 2024; Segijn; Strycharz, 2025). A vulnerabilidade percebida corresponde à avaliação subjetiva da exposição a potenciais danos decorrentes do controle limitado sobre os dados pessoais (Dinev; Hart, 2004; Van Ooijen *et al.*, 2024). Quando os indivíduos questionam a eficácia das salvaguardas existentes de privacidade, tornam-se mais conscientes da vulnerabilidade e reconhecem o desequilíbrio informacional imposto pelas plataformas digitais (Agozie; Kaya, 2021; Min; Kim, 2015). Com base nesse raciocínio, propõe-se a seguinte hipótese:

H3: O cinismo da privacidade influencia positivamente a vulnerabilidade percebida em relação ao uso de *chatbots* baseados em IA generativa.

A vulnerabilidade percebida refere-se à avaliação subjetiva dos indivíduos quanto à sua exposição e ao risco decorrente de potenciais ameaças à privacidade dos dados pessoais, especialmente em ambientes digitais mediados por *chatbots* baseados em IA generativa. Quando os usuários acreditam que os dados pessoais são coletados, armazenados ou reutilizados de maneiras imprevisíveis, sua suscetibilidade percebida a danos aumenta, motivando a ativação de respostas autoprotetivas (Orszaghova e Blank, 2024). De acordo com a Teoria da Motivação para a Proteção (*Protection Motivation Theory* – PMT) (Rogers, 1975), a vulnerabilidade percebida constitui um componente central do processo de avaliação da ameaça, incentivando os indivíduos a adotarem comportamentos defensivos.

Nesse sentido, a vulnerabilidade funciona como o mecanismo psicológico por meio do qual a consciência abstrata dos riscos é traduzida em intenções concretas de proteção, como a limitação do compartilhamento de dados ou o aumento do controle sobre as configurações de

privacidade. Esse processo torna-se especialmente relevante nas interações com IA conversacional, nas quais a comunicação personalizada e empática pode, simultaneamente, ampliar o engajamento e intensificar as percepções de exposição (Gurung *et al.*, 2009; Ham, 2017; Mousavi *et al.*, 2020). Com base nesse raciocínio, propõe-se a seguinte hipótese:

H4: A vulnerabilidade percebida influencia positivamente as intenções dos usuários de proteger a privacidade ao interagir com *chatbots* baseados em IA generativa.

Diferentemente de uma postura de vigilância ativa, o cinismo da privacidade reflete um estado de distanciamento crítico, em que os usuários, embora conscientes das ameaças, não se percebem capazes de agir de forma eficaz diante delas. Embora esse ceticismo sinalize consciência dos riscos à privacidade, ele não motiva, de forma inerente, comportamentos protetivos (Acikgoz; Vega, 2022; Hoffmann *et al.*, 2016). O cinismo resulta em apatia ou desengajamento, levando os indivíduos a reconhecerem ameaças informacionais sem, contudo, adotarem ações para mitigá-las (Tang *et al.*, 2021; Tian *et al.*, 2022; Yu *et al.*, 2024).

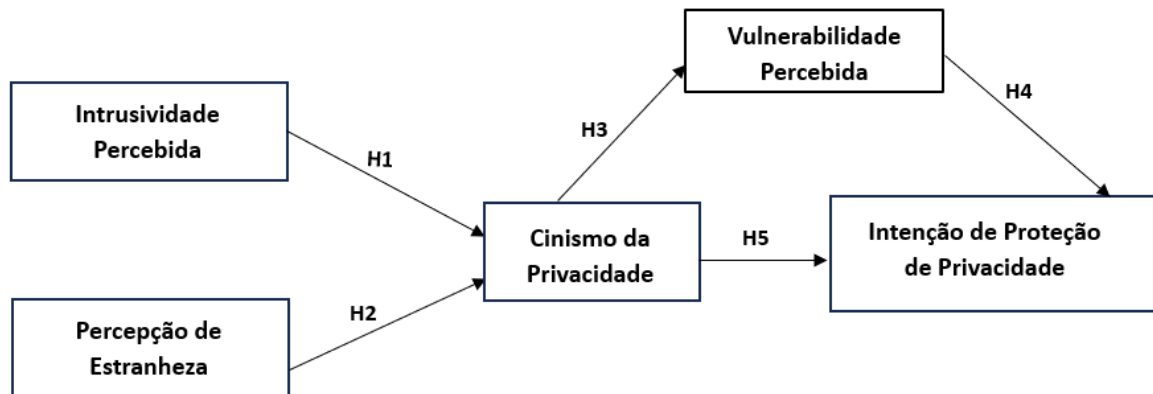
A vulnerabilidade percebida, portanto, representa o elo psicológico fundamental que determina se o cinismo permanecerá passivo ou se será convertido em ação. Sob a perspectiva da Teoria da Motivação para a Proteção, o comportamento autoprotetivo é desencadeado apenas quando os indivíduos percebem a ameaça como pessoalmente relevante e a si mesmos como suscetíveis a danos. Nesse sentido, o cinismo da privacidade leva a intenções protetivas somente quando acompanhado por um elevado senso de vulnerabilidade, levando os indivíduos a reavaliar decisões de consentimento ou a reduzir a interação com sistemas de IA generativa (Lee Kim *et al.*, 2025).

Na ausência dessa percepção de vulnerabilidade, o cinismo tende a se manifestar como desconfiança passiva, e não como mudança comportamental. Assim, a vulnerabilidade percebida atua como um mecanismo mediador que transforma o ceticismo cognitivo em intenção comportamental (Rashed Mohamed Al Humaid Alneyadi; Kassim Normalini, 2025). Com base nessa lógica teórica, propõe-se a seguinte hipótese:

H5: A vulnerabilidade percebida media integralmente a relação entre o cinismo da privacidade e a intenção de proteção da privacidade nas interações com *chatbots* baseados em IA generativa.

Com base no arcabouço conceitual apresentado, a Figura 1 ilustra o modelo teórico relacionado aos novos antecedentes do cinismo da privacidade.

Figura 1 – Modelo teórico da pesquisa



Fonte: Elaborado pela autora, 2026.

Os construtos abordados formam a base do modelo proposto neste estudo. A seguir, apresenta-se o delineamento metodológico que orientou a operacionalização dessas variáveis e a coleta dos dados empíricos.

3 MATERIAIS E MÉTODOS

3.1 Coleta de dados e perfil da amostra

Este estudo adota um delineamento de pesquisa quantitativo e transversal, uma vez que os dados foram coletados em um único corte temporal para examinar as relações entre os construtos teóricos propostos no modelo de pesquisa (Figura 1). Classifica-se como um estudo de natureza aplicada, pois busca gerar subsídios com relevância prática sobre as atitudes dos usuários em relação às práticas de privacidade nas interações com *chatbots* baseados em inteligência artificial (Prodanov; Freitas, 2013). Considerando a limitada evidência empírica sobre o cinismo da privacidade em contextos mediados por IA, o estudo caracteriza-se inicialmente como exploratório, com o objetivo de ampliar a compreensão dos fatores perceptuais que moldam as respostas relacionadas à privacidade. Simultaneamente, assume um caráter explicativo, na medida em que testa hipóteses e busca explicar os mecanismos causais e as relações de mediação entre as atitudes e as intenções comportamentais dos usuários.

A população-alvo foi composta por indivíduos brasileiros que haviam interagido com *chatbots* baseados em inteligência artificial em suas atividades cotidianas nos seis meses anteriores à aplicação do questionário, realizada em julho de 2025. Para assegurar a relevância das respostas, incluiu-se uma questão de triagem no início do instrumento, permitindo o prosseguimento apenas dos participantes que atendiam a esse critério. Essa delimitação temporal de seis meses foi intencionalmente estabelecida para garantir que as experiências estivessem recentes na memória dos participantes e refletissem o estágio atual das ferramentas de IA generativa. Ao focar exclusivamente em usuários ativos em suas rotinas diárias, o estudo delimita com precisão seu universo empírico, assegurando que as avaliações sobre intrusividade e estranheza sejam fundamentadas em vivências reais, práticas e contemporâneas, mitigando o risco de respostas baseadas em concepções abstratas ou defasadas sobre a tecnologia. Adotou-se uma estratégia de amostragem não probabilística por conveniência, comumente utilizada em estudos exploratórios sobre tecnologias emergentes e percepções dos usuários, em razão das limitações de acesso a um quadro amostral claramente definido.

Além disso, um vídeo¹ explicativo foi exibido no começo do questionário, simulando a interação típica com um *chatbot* e a respectiva solicitação de informações pessoais. A inserção desse material audiovisual teve o intuito de orientar os participantes sobre o processo

¹ Link do vídeo: <https://youtube.com/shorts/w0I-JMahTCQ%20>.

da pesquisa e ilustrar de forma concreta a situação de uso da tecnologia abordada. Essa estratégia metodológica foi rigorosamente adotada para mitigar potenciais vieses de interpretação e de memória. Ao garantir que todos os respondentes partissem de uma compreensão visual e prática padronizada do cenário interativo, minimizou-se o risco de respostas baseadas em premissas subjetivas ou assimétricas sobre o funcionamento da inteligência artificial, assegurando que as avaliações fossem direcionadas exatamente ao contexto do estudo.

O questionário continha 22 itens, adaptados de escalas de mensuração internacionalmente validadas, traduzidas e culturalmente ajustadas ao contexto brasileiro. Sua estrutura foi dividida em duas seções: (1) mensuração dos construtos e (2) caracterização sociodemográfica dos participantes. Antes da coleta completa dos dados, foi realizado um pré-teste com 31 respondentes, no período de 18 a 20 de fevereiro de 2025. Os resultados do pré-teste resultaram em ajustes pontuais na redação, eliminando interpretações divergentes e garantindo que as questões medissem adequadamente os construtos teóricos. Ao reduzir ruídos semânticos e minimizar vieses decorrentes de má interpretação das questões, o pré-teste contribuiu para o fortalecimento da validade interna do estudo, assegurando maior alinhamento entre os construtos teóricos e sua operacionalização empírica.

A adequação do tamanho da amostra foi avaliada por meio do software G*Power, considerando um modelo de regressão linear múltipla com cinco preditores, nível de significância de $\alpha = 0,05$, poder estatístico de 0,95 e tamanho de efeito médio ($f^2 = 0,15$), conforme Cohen (1988). A análise indicou um tamanho mínimo necessário de 138 observações. A amostra final foi composta por 257 respostas válidas, superando esse limiar e proporcionando poder estatístico suficiente para as análises subsequentes. Esse procedimento está em consonância com as práticas metodológicas recomendadas em pesquisas quantitativas no campo da gestão da informação (Doshi *et al.*, 2023).

3.2 Instrumento de medida

O instrumento de mensuração foi desenvolvido a partir da adaptação de itens provenientes de escalas previamente validadas (ver Apêndice A), bem como com base em estudos anteriores de Choi *et al.* (2018), Lau *et al.* (2019), Langer e König (2018), Dinev e Hart (2004) e Ko (2013). Todos os itens foram mensurados por meio de uma escala do tipo Likert de cinco pontos, variando de “discordo totalmente” (1) a “concordo totalmente” (5), amplamente utilizada em pesquisas sobre percepções de privacidade e atitudes relacionadas à tecnologia.

Os dados foram analisados por meio da Modelagem de Equações Estruturais por Mínimos Quadrados Parciais (*Partial Least Squares Structural Equation Modeling* – PLS-SEM), utilizando o software SmartPLS 4. A técnica PLS-SEM foi escolhida por sua adequação a pesquisas com enfoque preditivo e a modelos complexos envolvendo múltiplos construtos latentes, além do objetivo de testar a mediação completa do estudo. Todos os construtos foram especificados como reflexivos, em consonância com estudos anteriores sobre privacidade e adoção de tecnologias digitais, nos quais os indicadores observados são concebidos como manifestações das variáveis latentes subjacentes (Hair Jr *et al.*, 2014).

A Tabela 2 sintetiza os procedimentos estatísticos, os critérios de avaliação e os limiares de referência utilizados para a análise dos modelos de mensuração e estrutural.

Tabela 2 – Critérios para Avaliação dos Modelos de Mensuração e Estrutural

MODELO DE MENSURAÇÃO		
Procedimento	Parâmetro Utilizado	Referência
Consistência interna	Alfa de Cronbach	$\alpha \geq 0,70$
Consistência composta	rho_A e rho_C	$\geq 0,70$
Validade convergente	Comunalidade	AVE $\geq 0,7$ (ao nível dos indicadores)
		AVE $\geq 0,5$ (ao nível dos construtos)
Validade discriminante	Critério de Fornell-Larcker	Raiz quadrada da AVE deve ser maior do que as correlações entre os construtos
	HTMT (Heterotrait-Monotrait Ratio)	HTMT $< 0,90$
MODELO ESTRUTURAL		
Procedimento	Parâmetro Utilizado	Referência
Colinearidade	VIF (Variance Inflation Factor)	$1 < VIF < 5$
Poder explicativo	Coefficiente de determinação (R^2)	Pequeno: 0.02; Moderado: 0.13; Substancial: ≥ 0.26
Tamanho do efeito	f^2 (f-quadrado)	Pequeno ≥ 0.02 ; Médio ≥ 0.15 ; Grande ≥ 0.35
Significância estatística	Valores de p (via 5,000-sample bootstrapping)	$p \leq 0,05$

Fonte: Elaborado pela autora, 2026.

Com os procedimentos metodológicos definidos, parte-se para análise dos dados coletados, visando testar as hipóteses formuladas com base no modelo teórico proposto.

4 RESULTADOS

4.1 Abordagem

Conforme apresentado na Tabela 3, a amostra apresenta um perfil demográfico heterogêneo, abrangendo participantes de diferentes faixas etárias e níveis de escolaridade. Essa diversidade oferece uma base adequada para a análise das variações nas percepções e atitudes em relação ao uso de *chatbots* entre usuários com distintos perfis sociodemográficos.

Tabela 3 – Dados demográficos da pesquisa

	Item	Número	Frequência
Gênero	Masculino	142	55,2%
	Feminino	115	44,8%
Idade	18 – 24 anos	47	18,28%
	25 – 34 anos	84	32,68%
	35 – 44 anos	100	38,91%
	45 – 54 anos	26	10,1%
Nível educacional	Ensino médio completo	38	14,8%
	Graduação	58	22,6%
	Especialização	90	35 %
	Mestrado e/ou doutorado	71	27,6%

Fonte: Elaborado pela autora, 2026.

4.2 Normalidade, Viés de Método Comum e Colinearidade

Os resultados do coeficiente de Mardia indicaram um desvio em relação à normalidade multivariada, sustentando a adoção da abordagem PLS-SEM baseada em variância, uma vez que essa técnica não exige a normalidade na distribuição dos dados. Para avaliar a possível presença de viés de método comum, foi realizado o teste de fator único de Harman por meio de uma análise fatorial exploratória com o método de extração Principal Axis Factoring. Os resultados demonstraram que o primeiro fator explicou 28,96% da variância total, valor significativamente inferior ao limite de 50% sugerido por Podsakoff *et al.* (2003). Dessa forma, conclui-se que o viés de método comum não constitui uma preocupação relevante neste conjunto de dados.

Adicionalmente, os procedimentos diagnósticos procedimentais e estatísticos foram complementados pela avaliação da colinearidade entre os indicadores por meio do Fator de Inflação da Variância (Variance Inflation Factor – VIF), calculado para cada item e posteriormente promediado por construto. Os valores médios de VIF foram os seguintes:

Intrusividade Percebida (IP) = 1,908; Percepção de Estranheza (PE) = 1,908; Cinismo da Privacidade (CIN) = 1,000; Vulnerabilidade Percebida (VP) = 1,448; e Intenção de Proteção da Privacidade (IPP) = 1,448. Todos os valores ficaram bem abaixo do limite conservador de 5,0 recomendado por Hair *et al.* (2014), indicando que a multicolinearidade não representa uma preocupação no modelo proposto.

4.3 Teste do modelo de mensuração

O modelo de mensuração foi avaliado por meio da análise da confiabilidade, da validade convergente e da validade discriminante, em conformidade com as diretrizes estabelecidas para a aplicação do PLS-SEM. Conforme ilustrado na Tabela 4, as cargas fatoriais dos itens variaram de 0,712 a 0,929, indicando que todos os indicadores mantidos apresentaram forte associação com seus respectivos construtos latentes.

A confiabilidade da consistência interna foi confirmada por meio dos coeficientes alfa de Cronbach, que superaram o limite recomendado de 0,70 para todos os construtos: intrusividade percebida (IP = 0,788), percepção de estranheza (PE = 0,713), cinismo da privacidade (CIN = 0,851), vulnerabilidade percebida (VP = 0,753) e intenção de proteção da privacidade (IPP = 0,898).

A confiabilidade e a validade convergente dos construtos foram adicionalmente examinadas por meio das cargas fatoriais, do alfa de Cronbach, da confiabilidade composta (ρ_a e ρ_c) e da variância média extraída (*Average Variance Extracted* – AVE), conforme apresentado na Tabela 4, em consonância com as recomendações metodológicas estabelecidas para a modelagem de equações estruturais (Hair Jr *et al.*, 2014).

Tabela 4 – Cargas Fatoriais, Alfa de Cronbach, Confiabilidade Composta e AVE

	Cargas Externas	Alfa de Cronbach	Confiabilidade composta (ρ_a)	Confiabilidade composta (ρ_c)	Variância Média Extraída (AVE)
CIN	CIN1 0.845 CIN2 0.877 CIN4 0.908	0.851	0.869	0.909	0.769
IP	IP1 0.717 IP2 0.891 IP3 0.899	0.788	0.824	0.877	0.705

	Cargas Externas	Alfa de Cronbach	Confiabilidade composta (rho_a)	Confiabilidade composta (rho_c)	Variância Média Extraída (AVE)
IPP	IPP1 0.929 IPP2 0.921 IPP3 0.882	0.898	0.906	0.936	0.830
PE	PE2 0.748 PE4 0.764 PE5 0.747	0.713	0.718	0.797	0.567
VP	VP1 0.839 VP2 0.890 VP3 0.712	0.753	0.784	0.857	0.668

Fonte: Dados da pesquisa através do software SmartPLS 4.

Todos os valores de alfa de Cronbach e de confiabilidade composta superaram o limite recomendado de 0,70, confirmando a consistência interna e a estabilidade das escalas de mensuração. A validade convergente também foi estabelecida, uma vez que todos os valores de AVE foram superiores ao critério mínimo de 0,50 (Fornell; Larcker, 1981). De acordo com os critérios estabelecidos por Hair *et al.* (2014), coeficientes acima de 0,70 indicam elevada consistência interna, condição observada para todos os construtos analisados. Ademais, os índices de confiabilidade composta (rho_a e rho_c) excederam o limite de 0,70, reforçando a estabilidade e a confiabilidade das escalas utilizadas.

No que se refere à validade convergente, todas as AVEs apresentaram valores superiores a 0,50 (Fornell; Larcker, 1981), indicando adequação desse critério. Para atender a esse requisito, os itens PE1, PE3 e PE6 foram excluídos, uma vez que apresentaram cargas fatoriais inferiores a 0,40. Do ponto de vista teórico e contextual, a exclusão desses itens justifica-se pelo perfil da amostra. O item PE3 ('não sei exatamente como agir') não obteve aderência por se tratar de um público com alto nível educacional e com experiência prévia de uso de *chatbots* nos últimos seis meses. Da mesma forma, os itens PE1 e PE6, que mensuram um 'desconforto' mais visceral e direto, não refletiram a experiência atual dos usuários com a IA generativa. A retenção dos itens PE2, PE4 e PE5 demonstrou que, para este público, a percepção de estranheza não se manifesta como uma paralisia ou medo, mas sim como uma incerteza cognitiva e cautela em relação à imprevisibilidade da máquina. O construto intenção de proteção da privacidade (IPP) apresentou os indicadores mais robustos ($\alpha = 0,898$; rho_c = 0,936; AVE = 0,830), enquanto o construto percepção de estranheza (PE) apresentou o menor valor de AVE (0,567), permanecendo, contudo, acima do mínimo exigido.

O construto intenção de proteção da privacidade demonstrou os maiores índices de confiabilidade e validade, com alfa de Cronbach = 0,898, $\rho_c = 0,936$ e AVE = 0,830, indicando elevada consistência interna e forte validade convergente. Por outro lado, o construto estranheza percebida apresentou o menor valor de AVE (0,567). Ainda assim, esse resultado superou o limite mínimo recomendado, assegurando a validade convergente do construto.

A validade discriminante foi avaliada por meio dos critérios de Fornell-Larcker e HTMT. No primeiro, a raiz quadrada da variância média extraída (AVE) de cada construto superou suas correlações com os demais. A Tabela 5 sintetiza os resultados que sustentam a validade discriminante de todos os construtos analisados.

Tabela 5 – Teste de validade discriminada (critério de Fornell-Larcker)

	CIN	IP	IPP	PE	VP
CIN	0.877				
IP	0.827	0.840			
IPP	0.257	0.353	0.911		
PE	0.732	0.690	0.245	0.753	
VP	0.556	0.529	0.616	0.499	0.817

Fonte: Dados da pesquisa através do software SmartPLS 4.

O critério HTMT, considerado uma medida mais robusta para a avaliação da validade discriminante (Henseler *et al.*, 2015), também foi aplicado. Todos os valores de HTMT ficaram abaixo do limite conservador de 0,90, indicando adequada validade discriminante entre os construtos analisados (Tabela 6).

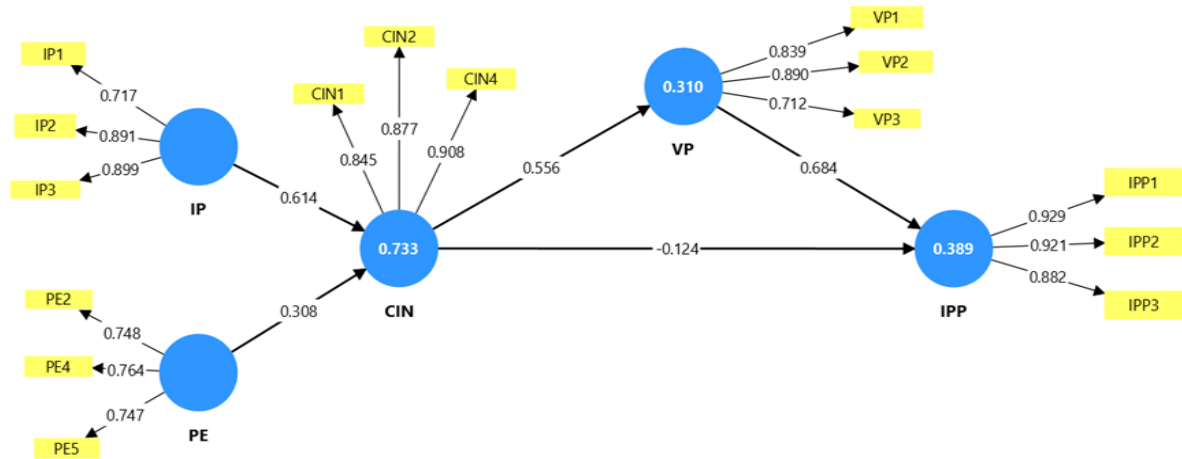
Tabela 6 – Teste de validade discriminada (razão heterotraço-monotraço — HTMT)

	CIN	IP	IPP	PE	VP
CIN					
IP	0.704				
IPP	0.355	0.226			
PE	0.624	0.428	0.451		
VP	0.545	0.613	0.550	0.388	

Fonte: Dados da pesquisa através do software SmartPLS 4.

4.4 Teste do modelo estrutural

Figura 2 – Teste do modelo estrutural



Fonte: Dados da pesquisa através do software SmartPLS 4.

Na etapa seguinte, foram realizadas a análise do modelo estrutural e o teste das hipóteses, conforme ilustrado na Figura 2. A qualidade do modelo foi avaliada com base no poder explicativo (R^2) e no tamanho do efeito (f^2), conforme apresentado na Tabela 7.

O coeficiente de determinação (R^2) indica a proporção da variância explicada em uma variável endógena. De acordo com Cohen (2009), valores de 0,02, 0,13 e 0,26 são considerados pequenos, moderados e substanciais, respectivamente. O construto cinismo da privacidade apresentou $R^2 = 0,733$, indicando que 73,3% de sua variância foi explicada pela intrusividade percebida e pela percepção de estranheza, o que sugere elevado poder explicativo. A vulnerabilidade percebida obteve $R^2 = 0,310$, indicando efeito moderado, enquanto a intenção de proteção da privacidade apresentou $R^2 = 0,389$, também correspondente a um efeito moderado.

Além da análise do R^2 , foi examinada a medida f^2 , que expressa o tamanho do efeito de uma variável exógena sobre uma variável endógena no modelo estrutural. Segundo os critérios de Cohen (2009), os valores de f^2 são classificados como pequenos (0,02), médios (0,15) ou grandes (0,35 ou superiores). A relação intrusividade percebida e o cinismo da privacidade apresentou $f^2 = 0,742$, indicando um efeito significativo e forte influência da intrusividade percebida sobre o cinismo. A relação percepção de estranheza e o cinismo da privacidade apresentou $f^2 = 0,187$, indicando impacto moderado e sugerindo que a percepção de estranheza também contribui para a formação desse construto, ainda que com menor intensidade.

A relação da variável cinismo da privacidade sob a vulnerabilidade percebida apresentou $f^2 = 0,448$, indicando efeito significativo e reforçando que o cinismo da privacidade constitui um fator relevante nas percepções de vulnerabilidade dos usuários. Ademais, a relação da variável vulnerabilidade percebida e da variável intenção de proteção da privacidade apresentou $f^2 = 0,530$, indicando contribuição substancial da vulnerabilidade percebida na explicação da intenção de proteção da privacidade. Por outro lado, a relação direta da variável cinismo da privacidade e da intenção de proteção da privacidade apresentou $f^2 = 0,017$, sugerindo um efeito desprezível e não significativo.

Tabela 7 – Capacidade explicativa do modelo (R^2 , R^2 ajustado, f^2)

	R-quadrado	R-quadrado ajustado
CIN	0.733	0.731
IPP	0.389	0.383
VP	0.310	0.306
Efeito (f^2)		
CIN -> IPP	0.017	
CIN -> VP	0.448	
IP -> CIN	0.742	
PE -> CIN	0.187	
VP -> IPP	0.530	

Fonte: Dados da pesquisa através do software SmartPLS 4.

A análise da significância estatística dos coeficientes de caminho foi realizada por meio do procedimento de *bootstrapping* com 5.000 amostragens, utilizando teste unilateral ao nível de significância de 0,05 (Hair Jr *et al.*, 2014). Esse procedimento possibilitou a avaliação da significância das relações estruturais propostas no modelo. Os coeficientes dos caminhos intrusividade percebida e cinismo da privacidade ($p = 0,000$), percepção de estranheza e cinismo da privacidade, cinismo da privacidade e vulnerabilidade percebida ($p = 0,000$) e vulnerabilidade percebida e intenção de proteção da privacidade ($p = 0,000$) mostraram-se estatisticamente significativos ($p < 0,01$), fornecendo evidências empíricas para as associações entre esses construtos. Por outro lado, o coeficiente da relação cinismo da privacidade e intenção de proteção da privacidade ($p = 0,059$) não apresentou significância estatística ($p > 0,05$), corroborando a interpretação de que a vulnerabilidade percebida atua como mediadora total entre o cinismo da privacidade e a intenção de proteção da privacidade.

Conforme apresentado na Tabela 8, a avaliação do modelo estrutural forneceu suporte empírico para todas as hipóteses formuladas (H1–H5). Os resultados revelaram efeitos estatisticamente significativos para as relações intrusividade percebida e cinismo da privacidade

($\beta = 0,614$; $p < 0,001$), percepção de estranheza e cinismo da privacidade ($\beta = 0,308$; $p < 0,001$), cinismo da privacidade e vulnerabilidade percebida ($\beta = 0,556$; $p < 0,001$) e vulnerabilidade percebida e intenção de proteção da privacidade ($\beta = 0,684$; $p < 0,001$). Esses achados demonstram que tanto a intrusividade percebida quanto a estranheza percebida contribuem para o aumento do cinismo da privacidade, o qual, por sua vez, está associado a níveis mais elevados de vulnerabilidade percebida e influencia a intenção de proteção da privacidade. Embora a relação direta entre o cinismo da privacidade e a intenção de proteção ($\beta = -0,124$; $p = 0,059$) não tenha sido significativa, esse resultado é consistente com a hipótese H5, que postula mediação total: o impacto do cinismo sobre a intenção de proteção ocorre exclusivamente por meio da vulnerabilidade percebida.

Tabela 8 – Resultados dos testes de hipóteses

	Amostra original (O)	Média da amostra (M)	Desvio padrão (STDEV)	Estatísticas T (O/STDEV)	Valores de p
CIN -> IPP	-0.124	-0.128	0.065	1.890	0.059
CIN -> VP	0.556	0.559	0.057	9.789	0.000
IP -> CIN	0.614	0.614	0.055	11.118	0.000
PE -> CIN	0.308	0.309	0.059	5.204	0.000
VP -> IPP	0.684	0.691	0.074	9.223	0.000

Fonte: Dados da pesquisa através do software SmartPLS 4.

Os resultados apresentados corroboram empiricamente as relações propostas entre os construtos, revelando padrões relevantes que serão discutidos no próximo tópico, à luz da literatura revisada.

5 DISCUSSÃO

5.1 Antecedentes do cinismo da privacidade

Os resultados confirmaram as hipóteses testadas, fornecendo suporte empírico ao modelo teórico e revelando interconexões significativas entre os construtos analisados. A intrusividade percebida emergiu como o antecedente mais influente do cinismo da privacidade, evidenciando como práticas excessivas ou insuficientemente justificadas de coleta de dados desencadeiam respostas cognitivas e emocionais. Essas percepções funcionam como indicadores de risco informacional, ampliando a desconfiança e enfraquecendo a confiança em ambientes digitais. Esse resultado é consistente com pesquisas anteriores que associam práticas intrusivas de coleta de dados à redução da confiança, à percepção de inadequação e às discussões mais amplas desses mecanismos em interações com IA generativa (Jan *et al.*, 2023; Kezer *et al.*, 2022; Raza *et al.*, 2024; Salah; Ayyash, 2025; Schomakers *et al.*, 2022; Shang; Yi, 2025; Van Den Broeck *et al.*, 2019).

Esses resultados indicam que o cinismo da privacidade abrange tanto avaliações racionais quanto respostas afetivas, estando enraizado na dissonância cognitiva dos usuários entre confiança e incerteza. Quando sistemas automatizados são percebidos simultaneamente como úteis e intrusivos, os usuários experimentam ambivalência emocional, o que contribui para o desenvolvimento do cinismo. Essa observação está alinhada aos modelos cognitivo-afetivos de uso da tecnologia, segundo os quais as decisões relacionadas à privacidade são influenciadas tanto por processos analíticos quanto intuitivos (Kezer *et al.*, 2022; Lyu *et al.*, 2024). Assim, o surgimento do cinismo funciona como um mecanismo de enfrentamento para aliviar o desconforto decorrente de atitudes conflitantes em relação aos sistemas de IA generativa.

Por outro lado, o sentimento de estranheza também contribuiu para o aumento do cinismo, ainda que em menor intensidade. Isso sugere que fatores como linguagem excessivamente humanizada, personalização exacerbada ou interações que geram desconforto e violam as expectativas dos usuários desempenham um papel relevante. A literatura indica que esse tipo de desconforto afeta negativamente a experiência do usuário, fragiliza a confiança na tecnologia e reduz a intenção de continuidade de uso (Dekkal *et al.*, 2024; Lee Jiang *et al.*, 2025; Raff *et al.*, 2024; Rajaobelina *et al.*, 2021). Contudo, como a estranheza opera de forma mais subjetiva e ambígua, sem envolver diretamente dados pessoais, o valor diagnóstico como indicador de risco mostra-se mais limitado.

Sob uma perspectiva sociotécnica, os resultados evidenciam que o cinismo da privacidade é um subproduto da tensão entre a eficiência algorítmica e a agência humana. Quando os sistemas de IA ultrapassam os limiares perceptivos dos usuários em termos de autonomia ou realismo, geram uma sensação implícita de vigilância e desequilíbrio de poder. Isso ressalta a importância do *design* responsável de sistemas de IA, no qual transparência, explicabilidade e autonomia do usuário sejam elementos centrais no desenvolvimento tecnológico (Carmody *et al.*, 2021; Pham *et al.*, 2024). Dessa forma, *designers* e organizações devem considerar como interfaces conversacionais equilibram a personalização e o controle percebido, reduzindo gatilhos de desconforto cognitivo e ceticismo em relação à privacidade.

5.2 O papel da vulnerabilidade percebida

No que se refere à terceira hipótese, os resultados indicam que o cinismo da privacidade está associado ao aumento da vulnerabilidade percebida. Usuários cínicos interpretam ações rotineiras dos sistemas como potencialmente ameaçadoras, o que intensifica sua sensação de vulnerabilidade. Estudos anteriores descrevem esse fenômeno como uma resposta à falta de clareza e controle, resultando em sentimento de impotência diante da coleta de dados (Lutz *et al.*, 2020; Van Ooijen *et al.*, 2024). A percepção de perda de controle atua como um estressor psicológico, reforçando a ideia de que violações são inevitáveis (Agozie; Kaya, 2021).

As percepções de vulnerabilidade são moldadas por fatores culturais e contextuais. Em economias emergentes, como o Brasil, os usuários frequentemente vivenciam o paradoxo de elevado engajamento digital aliado à baixa confiança institucional. Esse contexto amplia as percepções de assimetria informacional e impotência, aumentando a probabilidade de que o cinismo da privacidade se desenvolva como uma atitude defensiva generalizada em relação à tecnologia. Esses resultados evidenciam a necessidade de examinar construtos relacionados à privacidade em diferentes contextos socioculturais, em vez de presumir um padrão global uniforme de adoção tecnológica.

A quarta hipótese confirmou que a vulnerabilidade percebida é o preditor direto mais forte da intenção de adoção de comportamentos protetivos. Esse achado merece grande destaque interpretativo, pois sugere que os usuários só adotam medidas concretas e ativas de proteção quando se sentem diretamente ameaçados. Esse padrão está alinhado à Teoria da Motivação para a Proteção, segundo a qual a motivação para agir defensivamente depende intrinsecamente do reconhecimento de uma ameaça iminente e pessoal (Boerman *et al.*, 2021;

Maddux; Rogers, 1983). Estudos recentes reforçam a relevância da vulnerabilidade no desencadeamento de comportamentos defensivos (Nie *et al.*, 2024; Zhang; Zhang, 2024).

5.3 Mediação e implicações comportamentais

Por fim, a quinta hipótese revelou a mediação completa da vulnerabilidade percebida na relação entre cinismo e intenção de proteção. De forma mais didática, o modelo estrutural demonstrou que a relação direta entre o cinismo da privacidade e o comportamento de proteção foi insignificante ($p = 0,059$; $f^2 = 0,017$). Isso evidencia que um usuário tornar-se 'cínico' ou 'descrente' em relação à privacidade, por si só, não é suficiente para que ele passe a adotar medidas ativas, como fornecer dados incompletos ou alterar configurações. O cinismo isolado tende a paralisar o indivíduo na apatia e na resignação. Para que essa descrença teórica se transforme em uma ação protetiva real, é estritamente necessário o intermédio da vulnerabilidade percebida. Somente quando o cinismo faz com que o usuário perceba que aquela interação específica com o *chatbot* pode gerar danos reais à sua própria pessoa (elevando a suscetibilidade), é que a ponte psicológica para a ação é ativada. Assim, a mediação sustenta-se de forma integral: o cinismo não conduz diretamente à proteção; ele apenas engatilha a vulnerabilidade, que, então, motiva o comportamento de defesa.

Por fim, a quinta hipótese revelou a mediação completa da vulnerabilidade na relação entre cinismo e intenção de proteção. O cinismo, por si só, não induz à ação; ele intensifica principalmente a sensação de vulnerabilidade, que, quando percebida como ameaça pessoal, pode desencadear uma resposta protetiva. Esse resultado está em consonância com pesquisas que indicam que o cinismo frequentemente conduz à apatia e à resignação, em vez de comportamentos defensivos (Acikgoz; Vega, 2022; Hoffmann *et al.*, 2016). Para além das implicações teóricas e gerenciais, esses achados também contribuem para debates mais amplos sobre sustentabilidade digital e governança ética da tecnologia.

Ao evidenciar como percepções de intrusão e estranheza afetam a autonomia e a confiança dos usuários, o estudo contribui para a compreensão das respostas individuais às interações com IA generativa, fornecendo subsídios conceituais para debates normativos mais amplos sobre transparência, responsabilização e proteção de dados. Nesse sentido, os resultados dialogam com os princípios dos ODS 9 (Indústria, Inovação e Infraestrutura) e ODS 16 (Paz, Justiça e Instituições Eficazes), ao enfatizar a importância de uma inovação tecnológica inclusiva, transparente e alinhada aos direitos digitais. A contribuição do estudo para a agenda de governança ética e sustentabilidade digital reside, portanto, na elucidação de processos

psicológicos subjacentes ao comportamento do usuário, capazes de informar políticas públicas, diretrizes organizacionais e iniciativas institucionais, sem, contudo, estabelecer relações causais diretas entre os construtos analisados e esses objetivos macroestruturais.

5.4 Implicações teóricas e práticas

Este estudo amplia a compreensão teórica ao introduzir um arcabouço integrativo que transcende a aplicação linear da Teoria da Motivação para a Proteção (PMT), expandindo-a e contextualizando-a às características da inteligência artificial conversacional. Tradicionalmente, a PMT pressupõe que o comportamento de proteção é desencadeado diretamente pela avaliação cognitiva de uma ameaça. A contribuição deste trabalho consiste em integrar antecedentes perceptivos e emocionais específicos do uso de *chatbots* ao processo de proteção, estabelecendo-os como os fatores propulsores do cinismo da privacidade previamente à consolidação de uma ameaça tangível.

Ao evidenciar empiricamente que o efeito do cinismo sobre o comportamento protetivo é integralmente mediado pela vulnerabilidade percebida, o estudo articula a literatura de cinismo da privacidade à PMT de maneira inovadora. Esclarece-se, assim, um aspecto fundamental: o cinismo, de forma isolada, não suscita as respostas defensivas postuladas pela PMT. Ele atua como uma disposição latente e passiva que requer conversão em uma percepção concreta de risco pessoal (a avaliação de ameaça da PMT, manifestada como vulnerabilidade percebida) para, por conseguinte, impulsionar a intenção de proteção. Tal achado aprofunda o debate acerca do paradoxo da privacidade, revelando que a discrepância entre atitudes críticas e comportamentos protetivos na era da IA é elucidada pela incorporação desses novos mecanismos psicológicos intervenientes.

Dessa forma, o modelo proposto integra dimensões perceptivas, afetivas e cognitivas em uma estrutura explicativa unificada, oferecendo um avanço teórico ao conectar experiências de interação com IA generativa às dinâmicas psicológicas que sustentam decisões de proteção de dados.

As implicações práticas deste estudo orientam as organizações que utilizam *chatbots* baseados em IA generativa, enfatizando a necessidade de equilibrar o avanço tecnológico com uma governança ética eficaz. Para além da otimização da funcionalidade e da eficiência, os sistemas de comunicação mediados por IA devem ser concebidos a partir de uma filosofia centrada no ser humano, reconhecendo que cada interação algorítmica molda as respostas emocionais dos usuários, o senso de agência e a autonomia percebida.

Sob a perspectiva gerencial, o papel mediador da vulnerabilidade percebida revela duas direções centrais de atuação. Primeiramente, as organizações devem comunicar riscos e mecanismos de proteção com maior transparência, transformando noções abstratas de ética de dados em práticas concretas e centradas no usuário de responsabilidade digital. Em segundo lugar, devem capacitar os usuários por meio do design, desenvolvendo interfaces intuitivas que estimulem a autorregulação, como painéis de privacidade simplificados, mecanismos adaptativos de consentimento e níveis personalizáveis de exposição.

6 CONCLUSÃO, LIMITAÇÕES E PESQUISAS FUTURAS

A pesquisa aprofundou a compreensão do cinismo da privacidade ao propor e sustentar empiricamente um modelo integrativo que relaciona a intrusividade percebida, a estranheza percebida e a vulnerabilidade percebida nas interações com *chatbots* baseados em IA. Os resultados indicam que a intrusividade e a estranheza funcionam como antecedentes diretos do cinismo da privacidade. Ao mesmo tempo, a vulnerabilidade percebida emerge como o principal mecanismo mediador da relação entre o cinismo e a intenção de proteção da privacidade. Esse achado desafia a suposição de que o cinismo conduz diretamente a comportamentos protetivos, sugerindo, em vez disso, que a influência opera de forma indireta, por meio da internalização do risco percebido como vulnerabilidade pessoal.

Apesar das contribuições relevantes, o estudo apresenta limitações metodológicas que impõem restrições à generalização dos achados e abrem caminhos para pesquisas futuras. Primeiramente, a adoção de um delineamento de corte transversal limita a capacidade de inferir causalidade absoluta ao longo do tempo. Em segundo lugar, o uso de amostragem não probabilística por conveniência e o foco exclusivo em usuários com interação recente e perfil de alto letramento digital restringem a generalização para a população em geral. Consequentemente, as atitudes de não usuários, ou daqueles que abandonaram a tecnologia, não foram capturadas.

Como desdobramento para pesquisas futuras, recomenda-se explorar de forma mais profunda as variáveis sociodemográficas coletadas. A realização de uma Análise de Grupos Múltiplos (PLS-MGA) para avaliar como gênero, faixa etária e nível educacional atuam como moderadores no modelo estrutural pode enriquecer a discussão, revelando nuances comportamentais em diferentes subgrupos. Além disso, futuras investigações podem examinar moderadores contextuais, como a confiança nas instituições e variações culturais, que possam alterar a intensidade das relações observadas, destacando o valor de delineamentos de pesquisa longitudinais e transculturais.

Em síntese, a principal mensagem analítica deste estudo evidencia que o desenvolvimento acelerado de tecnologias de IA humanizadas gera reações complexas e paradoxais. A percepção de intrusividade algorítmica e a sensação de estranheza interacional alimentam intensamente o cinismo em relação à privacidade, empurrando o indivíduo para um estado sistêmico de descrença e fadiga. Contudo, a pesquisa prova de maneira conclusiva que a adoção de posturas de proteção ativas e defensivas apenas emerge quando esse cinismo latente se converte em uma percepção de vulnerabilidade pessoal imediata e tangível. Ao mapear essa

fronteira crítica entre o ceticismo passivo e a defesa ativa, evidencia-se que o futuro da adoção da inteligência artificial generativa não depende apenas de criar sistemas computacionalmente eficientes ou engajadores. Depende, fundamentalmente, do estabelecimento de uma arquitetura ética que respeite as fronteiras informacionais, mitigando a intrusividade percebida e devolvendo ao usuário o senso de agência e controle sobre a sua própria privacidade digital.

REFERÊNCIAS

- ACIKGOZ, Fulya; VEGA, Rodrigo Perez. The Role of Privacy Cynicism in Consumer Habits with Voice Assistants: A Technology Acceptance Model Perspective. **International Journal of Human–Computer Interaction**, v. 38, n. 12, p. 1138–1152, 21 jul. 2022.
- ADAMOPOULOU, Eleni; MOUSSIADES, Lefteris. Chatbots: History, technology, and applications. **Machine Learning with Applications**, v. 2, p. 100006, dez. 2020.
- AGOZIE, Divine Q.; KAYA, Tugberk. Discerning the effect of privacy information transparency on privacy fatigue in e-government. **Government Information Quarterly**, v. 38, n. 4, p. 101601, out. 2021.
- AKDIM, K.; CASALÓ, Luis V. Perceived value of AI-based recommendations service: the case of voice assistants. **Service Business**, v. 17, n. 1, p. 81–112, mar. 2023.
- AW, Eugene Cheng-Xi; LEONG, Lai-Ying; HEW, Jun-Jie; RANA, Nripendra P.; TAN, Teck Ming; JEE, Teck-Weng. Counteracting dark sides of robo-advisors: justice, privacy and intrusion considerations. **International Journal of Bank Marketing**, v. 42, n. 1, p. 133–151, 22 jan. 2024.
- BAŞER, Miraç Yücel; BÜYÜKBEŞE, Tuba; DURMAZ, Yakup. “Yes, It’s Cute, But How Can I Be Sure It’s Safe or Not?” Investigating the Intention to Use Service Robots in the Context of Privacy Calculus. **International Journal of Human–Computer Interaction**, v. 40, n. 20, p. 6151–6166, 17 out. 2024.
- BOERMAN, Sophie C.; KRUIKEMEIER, Sanne; ZUIDERVEEN BORGESIJUS, Frederik J. Exploring Motivations for Online Privacy Protection Behavior: Insights From Panel Data. **Communication Research**, v. 48, n. 7, p. 953–977, out. 2021.
- BULGURCU, Burcu; CAVUSOGLU, Hasan; BENBASAT, Izak. Information Security Policy Compliance: An Empirical Study of Rationality-Based Beliefs and Information Security Awareness. **MIS Quarterly**, v. 34, n. 3, p. 523–548, 2010.
- CARMODY, Jillian; SHRINGARPURE, Samir; VAN DE VENTER, Gerhard. AI and privacy concerns: a smart meter case study. **Journal of Information, Communication and Ethics in Society**, v. 19, n. 4, p. 492–505, 13 dez. 2021.
- CHEN, Qian; GONG, Yeming; LU, Yaobin; TANG, Jing. Classifying and measuring the service quality of AI chatbot in frontline service. **Journal of Business Research**, v. 145, p. 552–568, jun. 2022.
- CHEN, Qian; GONG, Yeming; LU, Yaobin; LUO, Xin Robert. The golden zone of AI’s emotional expression in frontline chatbot service failures. **Internet Research**, v. 35, n. 3, p. 1065–1103, 27 maio 2025.
- CHO, Hichang. Privacy helplessness on social media: its constituents, antecedents and consequences. **Internet Research**, v. 32, n. 1, p. 150–171, 18 jan. 2022.
- CHOI, Hanbyul; PARK, Jonghwa; JUNG, Yoonhyuk. The role of privacy fatigue in online privacy behavior. **Computers in Human Behavior**, v. 81, p. 42–51, 1 abr. 2018.

COHEN, Jacob. **Statistical power analysis for the behavioral sciences**. 2. ed., reprint ed. New York, NY: Psychology Press, 2009.

DEKKAL, Massilva; ARCAND, Manon; PROM TEP, Sandrine; RAJAABELINA, Lova; RICARD, Linha. Factors affecting user trust and intention in adopting chatbots: the moderating role of technology anxiety in insurtech. **Journal of Financial Services Marketing**, v. 29, n. 3, p. 699–728, 1 set. 2024.

DINEV, Tamara; AND HART, Paul. Internet privacy concerns and their antecedents - measurement validity and a regression model. **Behaviour & Information Technology**, v. 23, n. 6, p. 413–422, 1 nov. 2004.

DINEV, Tamara; HART, Paul. An Extended Privacy Calculus Model for E-Commerce Transactions. **Information Systems Research**, v. 17, n. 1, p. 61–80, mar. 2006.

DOSHI, Parinda; NIGAM, Priti; RISHI, Bikramjit. Customer values and patronage intention in social media networks: mediating role of perceived usefulness. **VINE Journal of Information and Knowledge Management Systems**, v. 55, n. 2, p. 404–420, 21 mar. 2023.

EDWARDS, Steven M.; LI, Hairong; LEE, Joo-Hyun. Forced Exposure and Psychological Reactance: Antecedents and Consequences of the Perceived Intrusiveness of Pop-Up Ads. **Journal of Advertising**, v. 31, n. 3, p. 83–95, out. 2002.

F. HAIR JR, Joe *et al.* Partial least squares structural equation modeling (PLS-SEM): An emerging tool in business research. **European Business Review**, v. 26, n. 2, p. 106–121, 4 mar. 2014.

FORNELL, Claes; LARCKER, David F. Evaluating Structural Equation Models with Unobservable Variables and Measurement Error. **Journal of Marketing Research**, v. 18, n. 1, p. 39, fev. 1981.

FU, Jialin; ZHANG, Jiaming; LI, Xihang. How do risks and benefits affect user' privacy decisions? An event-related potential study on privacy calculus process. **Frontiers in Psychology**, v. 14, p. 1052782, 16 fev. 2023.

GENWUCH, Ulrich; MORANA, Stefan; MAEDCHE, Alexander. **Towards Designing Cooperative and Social Conversational Agents for Customer Service**. 2017.

GURUNG, Anil; LUO, Xin; LIAO, Qinyu. Consumer motivations in taking action against spyware: an empirical investigation. **Information Management & Computer Security**, v. 17, n. 3, p. 276–289, 17 jul. 2009.

HAM, Chang-Dae. Exploring how consumers cope with online behavioral advertising. **International Journal of Advertising**, v. 36, n. 4, p. 632–658, 4 jul. 2017.

HANSON, David. Exploring the aesthetic range for humanoid robots. In: **Proceedings of the ICCS/CogSci-2006 long symposium: Toward social mechanisms of android science**. 2006. p. 39-42.

HENSELER, Jörg; RINGLE, Christian M.; SARSTEDT, Marko. A new criterion for assessing discriminant validity in variance-based structural equation modeling. **Journal of the**

Academy of Marketing Science, v. 43, n. 1, p. 115–135, jan. 2015.

HOFFMANN, Christian Pieter; LUTZ, Christoph; RANZINI, Giulia. Privacy cynicism: A new approach to the privacy paradox. **Cyberpsychology: Journal of Psychosocial Research on Cyberspace**, v. 10, n. 4, 1 dez. 2016.

HONG, Soo Jung. What drives AI-based risk information-seeking intent? Insufficiency of risk information versus (Un)certainly of AI chatbots. **Computers in Human Behavior**, v. 162, p. 108460, jan. 2025.

HYUN BAEK, Tae; KIM, Minseong. Is ChatGPT scary good? How user motivations affect creepiness and trust in generative artificial intelligence. **Telematics and Informatics**, v. 83, p. 102030, set. 2023.

ISCHEN, Carolin; ARAUJO, Theo; VOORVELD, Hilde; VAN NOORT, Guda; SMIT, Edith. Privacy Concerns in Chatbot Interactions. In: FØLSTAD, Asbjørn *et al.* (orgs.). Cham: **Springer International Publishing**, 2020.

JAN, Ihsan Ullah; JI, Seonggoo; KIM, Changju. What (de) motivates customers to use AI-powered conversational agents for shopping? The extended behavioral reasoning perspective. **Journal of Retailing and Consumer Services**, v. 75, p. 103440, nov. 2023.

JOHNSON, Deborah G.; VERDICCHIO, Mario. AI Anxiety. **Journal of the Association for Information Science and Technology**, v. 68, n. 9, p. 2267–2270, 1 set. 2017.

KANG, Weiyao; SHAO, Bingjia. The impact of voice assistants' intelligent attributes on consumer well-being: Findings from PLS-SEM and fsQCA. **Journal of Retailing and Consumer Services**, v. 70, p. 103130, jan. 2023.

KAWAF, Fatema; MONTGOMERY, Annaleis; THUEMMLER, Marius. Unpacking the privacy–personalisation paradox in GDPR-2018 regulated environments: consumer vulnerability and the curse of personalisation. **Information Technology & People**, v. 37, n. 4, p. 1674–1695, 6 maio 2024.

KEZER, Murat; DIENLIN, Tobias; BARUH, Lemi. Getting the privacy calculus right: Analyzing the relations between privacy concerns, expected benefits, and self-disclosure using response surface analysis. **Cyberpsychology: Journal of Psychosocial Research on Cyberspace**, v. 16, n. 4, 19 set. 2022.

KO, Hsiu-Chia. The determinants of continuous use of social networking sites: An empirical study on Taiwanese journal-type bloggers' continuous self-disclosure behavior. **Electronic Commerce Research and Applications**, v. 12, n. 2, p. 103–111, abr. 2013.

KO, Minjeong; LEE, Luri; KIM, Yunice YoungKyoung. The underlying mechanism of user response to AI assistants: from interactivity to loyalty. **Information Technology & People**, v. 38, n. 5, p. 2284–2304, 3 nov. 2025.

KSHETRI, Nir; DWIVEDI, Yogesh K.; DAVENPORT, Thomas H.; PANTELI, Niki. Generative artificial intelligence in marketing: Applications, opportunities, challenges, and research agenda. **International Journal of Information Management**, v. 75, p. 102716, abr. 2024.

LANGER, Markus; KÖNIG, Cornelius J. Introducing and Testing the Creepiness of Situation Scale (CRoSS). **Frontiers in Psychology**, v. 9, p. 2220, 16 nov. 2018.

LAU, Chloe K. H.; CHUI, Chi Fai Raymond; AU, Norman. Examination of the adoption of augmented reality: a VAM approach. **Asia Pacific Journal of Tourism Research**, v. 24, n. 10, p. 1005–1020, 3 out. 2019.

LEE, Chunsik; KIM, Junga; LIM, Joon Soo; SHIN, Donghee. Generative AI risks and resilience: How users adapt to hallucination and privacy challenges. **Telematics and Informatics Reports**, v. 19, p. 100221, set. 2025.

LEE, Jung-Chieh; JIANG, SiQi; TANG, Yuyin. Examining the impacts of artificial intelligence characteristics on user resistance to financial robo-advisors. **Information Technology & People**, 3 jun. 2025.

LI, Hairong; EDWARDS, Steven M.; LEE, Joo-Hyun. Measuring the Intrusiveness of Advertisements: Scale Development and Validation. **Journal of Advertising**, v. 31, n. 2, p. 37–47, jun. 2002.

LIANG, Huigang; XUE, Yajiong. Avoidance of Information Technology Threats: A Theoretical Perspective. **MIS Quarterly**, v. 33, n. 1, p. 71–90, 2009.

LIM, Joon Soo; SHIN, Donghee; LEE, Chunsik; KIM, Junga; ZHANG, Jun. The Role of User Empowerment, AI Hallucination, and Privacy Concerns in Continued Use and Premium Subscription Intentions: An Extended Technology Acceptance Model for Generative AI. **Journal of Broadcasting & Electronic Media**, v. 69, n. 3, p. 183–199, 27 maio 2025.

LIU, Fanjue; WANG, Rang. Fostering Parasocial Relationships with Virtual Influencers in the Uncanny Valley: Anthropomorphism, Autonomy, and a Multigroup Comparison. **Journal of Business Research**, v. 186, p. 115024, jan. 2025.

LUTZ, Christoph; HOFFMANN, Christian Pieter; RANZINI, Giulia. Data capitalism and the user: An exploration of privacy cynicism in Germany. **New Media & Society**, v. 22, n. 7, p. 1168–1187, jul. 2020.

LYU, Tu; GUO, Yulin; CHEN, Hao. Understanding people's intention to use facial recognition services: the roles of network externality and privacy cynicism. **Information Technology & People**, v. 37, n. 3, p. 1025–1051, 8 abr. 2024.

MACDORMAN, Karl F. **Subjective Ratings of Robot Video Clips for Human Likeness, Familiarity, and Eeriness: An Exploration of the Uncanny Valley**. 2006.

MADDUX, James E.; ROGERS, Ronald W. Protection motivation and self-efficacy: A revised theory of fear appeals and attitude change. **Journal of Experimental Social Psychology**, v. 19, n. 5, p. 469–479, 1 set. 1983.

MADUKU, Daniel K.; RANA, Nripendra P.; MPINGANJIRA, Misericordia; THUSI, Filile. Exploring the 'Dark Side' of AI -Powered Digital Assistants: A Moderated Mediation Model of Antecedents and Outcomes of Perceived Creepiness. **Journal of Consumer Behaviour**, v. 24, n. 3, p. 1194–1221, maio 2025.

MENG, Xiaoxiao; LIU, Jiaxin. "Talk to me, I'm secure": investigating information disclosure

to AI chatbots in the context of privacy calculus. **Online Information Review**, 26 fev. 2025.

MIN, Jinyoung; KIM, Byoungsoo. How are people enticed to disclose personal information despite privacy concerns in social network sites? The calculus between benefit and cost. **Journal of the Association for Information Science and Technology**, v. 66, n. 4, p. 839–857, 1 abr. 2015.

MORI, Masahiro; MACDORMAN, Karl; KAGEKI, Norri. The Uncanny Valley [From the Field]. **IEEE Robotics & Automation Magazine**, v. 19, n. 2, p. 98–100, jun. 2012.

MORIUCHI, Emi; MURDY, Samantha. The role of robots in the service industry: Factors affecting human-robot interactions. **International Journal of Hospitality Management**, v. 118, p. 103682, abr. 2024.

MOUSAVI, Reza; CHEN, Rui; KIM, Dan J.; CHEN, Kuanchin. Effectiveness of privacy assurance mechanisms in users' privacy protection on social networking sites from the perspective of protection motivation theory. **Decision Support Systems**, v. 135, p. 113323, ago. 2020.

NGUYEN, Tran Hung; LE, Xuan Cu. Artificial intelligence-based chatbots – a motivation underlying sustainable development in banking: standpoint of customer experience and behavioral outcomes. **Cogent Business & Management**, v. 12, n. 1, p. 2443570, 12 dez. 2025.

NIE, Fapeng; LI, Xiang; ZHOU, Chang. Impact of privacy regulation involving information collection on the ride-hailing market. **Omega**, v. 126, p. 103077, jul. 2024.

OKAZAKI, Shintaro; LI, Hairong; HIROSE, Morikazu. Consumer Privacy Concerns and Preference for Degree of Regulatory Control. **Journal of Advertising**, v. 38, n. 4, p. 63–77, dez. 2009.

ORSZAGHOVA, Eva; BLANK, Grant. Does the type of privacy-protective behaviour matter? An analysis of online privacy protective action and motivation. **Information, Communication & Society**, v. 27, n. 14, p. 2530–2547, 25 out. 2024.

PHAM, Thi Huyen; PHAN, Thuy-Anh; TRINH, Phuong-Anh; MAI, Xuan Bach; LE, Quynh-Chi. Information security risks and sharing behavior on OSN: the impact of data collection awareness. **Journal of Information, Communication and Ethics in Society**, v. 22, n. 1, p. 82–102, 4 mar. 2024.

PINOCHET, Luis Hernan Contreras; DE GOIS, Fernanda Silva; PARDIM, Vanessa Itacaramby; ONUSIC, Luciana Massaro. Experimental study on the effect of adopting humanized and non-humanized chatbots on the factors measure the intensity of the user's perceived trust in the Yellow September campaign. **Technological Forecasting and Social Change**, v. 204, p. 123414, jul. 2024.

PITARDI, Valentina; MARRIOTT, Hannah R. Alexa, *she's* not human but... Unveiling the drivers of consumers' trust in voice-based artificial intelligence. **Psychology & Marketing**, v. 38, n. 4, p. 626–642, abr. 2021.

PRODANOV, C. C.; FREITAS, E. C. DE. **Metodologia do trabalho científico: métodos e**

técnicas da pesquisa e do trabalho acadêmico. 2a ed. Novo Hamburgo: Universidade FEEVALE, 2013.

RADFORD, Alec *et al.* Language Models are Unsupervised Multitask Learners. 2019.
RAFF, Stefan; ROSE, Stefan; HUYNH, Tin. Perceived creepiness in response to smart home assistants: A multi-method study. **International Journal of Information Management**, v. 74, p. 102720, fev. 2024.

RAJAOBELINA, Lova; PROM TEP, Sandrine; ARCAND Manon; RICARD, Linha. Creepiness: Its antecedents and impact on loyalty when interacting with a chatbot. **Psychology & Marketing**, v. 38, n. 12, p. 2339–2356, 1 dez. 2021.

RASHED MOHAMED AL HUMAID ALNEYADI, Mohammed; KASSIM NORMALINI, Md. Intelligent Protection: A Study of the Key Drivers of Intention to Adopt Artificial Intelligence (AI) Cybersecurity Systems in the UAE. **Interdisciplinary Journal of Information, Knowledge, and Management**, v. 20, p. 003, 2025.

RAZA, Saqib *et al.* The Future of Learning: Building Trust and Transparency in AI Education. **Journal of Management Practices, Humanities and Social Sciences**, v. 8, n. 3, 2024.

SALAH, Omar Hasan; AYYASH, Mohannad Moufeed. Understanding user adoption of mobile wallet: extended TAM with knowledge sharing, perceived value, perceived privacy awareness and control, perceived security. **VINE Journal of Information and Knowledge Management Systems**, v. 55, n. 5, p. 1223–1250, 26 ago. 2025.

SCHOMAKERS, Eva-Maria; LIDYNIA, Chantal; ZIEFLE, Martina. The Role of Privacy in the Acceptance of Smart Technologies: Applying the Privacy Calculus to Technology Acceptance. **International Journal of Human–Computer Interaction**, v. 38, n. 13, p. 1276–1289, 9 ago. 2022.

SEGIJN, Claire M.; KIM, Eunah; VAN OOIJEN, Iris. The Role of Perceived Surveillance and Privacy Cynicism in Effects of Multiple Synced Advertising Exposures on Brand Attitude. **Journal of Current Issues & Research in Advertising**, v. 45, n. 4, p. 506–522, out. 2024.

SEGIJN, Claire M.; STRYCHARZ, Joanna. Transparency, Awareness, and Privacy Cynicism: Exploring Appropriate Dataveillance for Personalized Advertising Through Qualitative Interviews. **Journal of Interactive Advertising**, v. 25, n. 3, p. 272–284, 3 jul. 2025.

SHANG, Shanshan; YI, Tianyun. Search engine versus affective thought-guided chatbot for knowledge acquisition: the mediating roles of perceived social, expressive and diagnostic transparency and cognitive effort. **Internet Research**, 2025.

SHAO, Zhucheng; HO, Jessica Sze Yin. **Revealing the resistance of virtual streamers: intrusiveness, privacy disclosure and perceived justice.** 2025.

SOHN, Stefanie; LABRECQUE, Lauren; SIEMON, Dominik; MORANA, Stefan. Artificial intelligence versus human service agents: How their presence shapes consumer information privacy concerns. **Journal of Retailing**, p. S0022435925000247, abr. 2025.

SUTA, Prissadang *et al.* An Overview of Machine Learning in Chatbots. **International Journal of Mechanical Engineering and Robotics Research**, p. 502–510, 2020.

TAI, MichaelCheng-Tek. The impact of artificial intelligence on human society and bioethics. **Tzu Chi Medical Journal**, v. 32, n. 4, p. 339, 2020.

TANG, Jie; AKRAM, Umair; SHI, Wenjing. Why people need privacy? The role of privacy fatigue in app users' intention to disclose privacy: based on personality traits. **Journal of Enterprise Information Management**, v. 34, n. 4, p. 1097–1120, 15 jul. 2021.

TIAN, Xinluan; CHEN, Lina; ZHANG, Xiaojuan. The Role of Privacy Fatigue in Privacy Paradox: A PSM and Heterogeneity Analysis. **Applied Sciences**, v. 12, n. 19, p. 9702, 27 set. 2022.

VAN DEN BROECK, Evert; ZAROUALI, Brahim; POELS, Karolien. Chatbot advertising effectiveness: When does the message get through? **Computers in Human Behavior**, v. 98, p. 150–157, set. 2019.

VAN OOIJEN, Iris; SEGIJN, Claire M.; OPREE, Suzanna J. Privacy Cynicism and its Role in Privacy Decision-Making. **Communication Research**, v. 51, n. 2, p. 146–177, mar. 2024.

YU, Lu; ELE, Li; DU, Jian; WU Xiaobo. Protection or Cynicism? Dual Strategies for Coping with Privacy Threats. **Information Systems Frontiers**, 30 ago. 2024.

ZHANG, Xiaoxue; ZHANG, Zizhong. Leaking my face via payment: Unveiling the influence of technology anxiety, vulnerabilities, and privacy concerns on user resistance to facial recognition payment. **Telecommunications Policy**, v. 48, n. 3, p. 102703, abr. 2024.

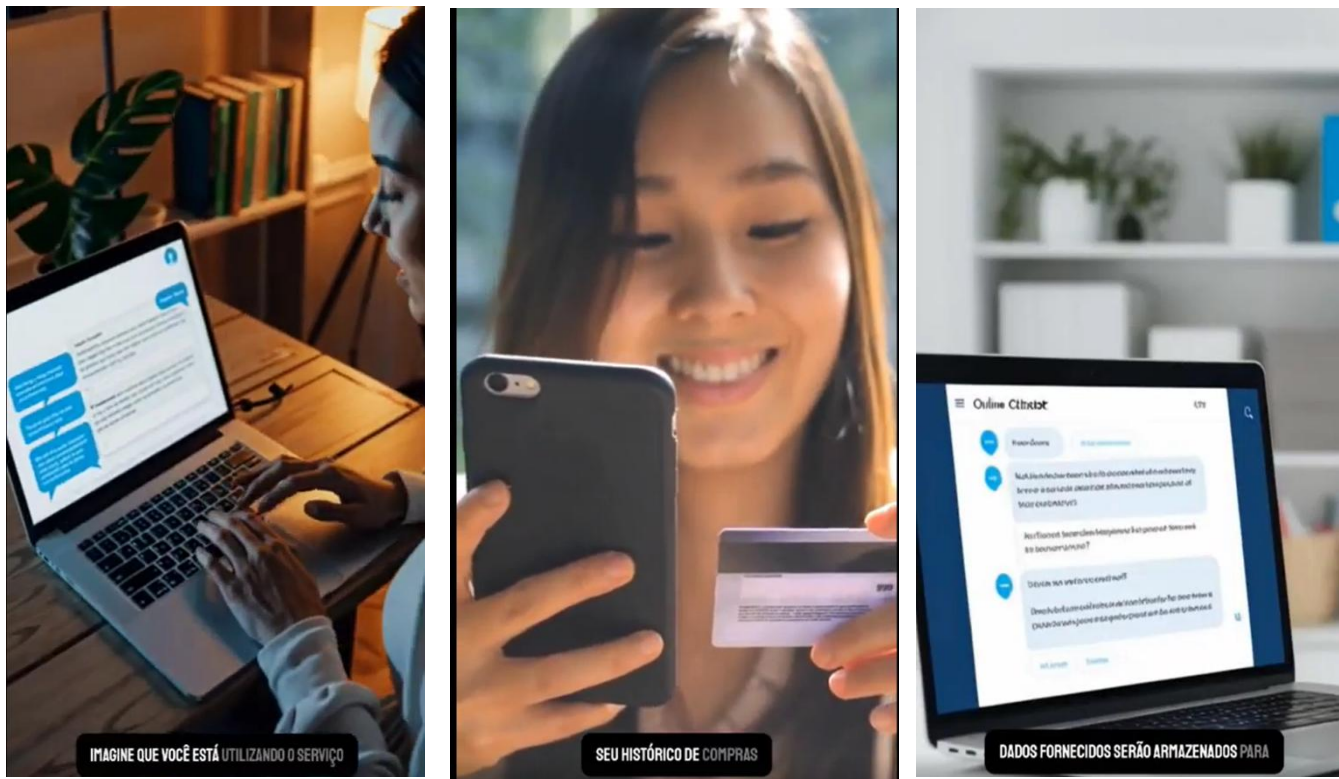
ZHOU, Li *et al.* The Design and Implementation of XiaoIce, an Empathetic Social Chatbot. **Computational Linguistics**, v. 46, n. 1, p. 53–93, mar. 2020.

APÊNDICE A – ESCALAS ADAPTADAS DOS CONSTRUTOS.

CINISMO DA PRIVACIDADE	
CIN1. Fiquei menos interessado em questões de privacidade on-line. CIN2. Fiquei mais descrente sobre se meus esforços para proteger minha privacidade ao interagir com <i>chatbots</i> são de alguma forma eficazes. CIN3. Muitas vezes duvido da importância das questões de privacidade on-line ao interagir com <i>chatbots</i> . CIN4. Estou menos motivado a proteger as informações pessoais fornecidas durante a interação com <i>chatbots</i> .	Adaptado de Choi et al. (2018)
INTRUSIVIDADE PERCEBIDA	
IP1. Ao interagir com <i>chatbots</i> , sinto que minhas interações estão sendo acompanhadas. IP2. Ao interagir com <i>chatbots</i> , percebo que minhas ações estão sendo monitoradas. IP3. Ao interagir com <i>chatbots</i> , tenho a impressão que eles registram mais informações do que eu esperava.	Adaptado de Lau et al., (2019)
PERCEPÇÃO DE ESTRANHEZA	
PE1. Ao interagir com o <i>chatbot</i> eu me senti desconfortável. PE2. Ao interagir com o <i>chatbot</i> , tive uma sensação de cautela. PE3. Ao interagir com o <i>chatbot</i> , não sei exatamente como agir. PE4. Ao interagir com o <i>chatbot</i> , não sabia exatamente o que esperar. PE5. Ao interagir com o <i>chatbot</i> , tenho a sensação de não compreender totalmente a experiência. PE6. Ao interagir com o <i>chatbot</i> , tive uma sensação estranha de desconforto.	Adaptado de Langer and König (2018)
VULNERABILIDADE PERCEBIDA	
VP1. Poderia estar sujeito a problemas de segurança maliciosos (por exemplo, perda de privacidade, roubo de identidade, vírus, hacking) usando <i>chatbots</i> . VP2. Eu sinto que minhas informações pessoais em <i>chatbots</i> podem ser disponibilizadas a indivíduos desconhecidos sem meu conhecimento. VP3. Eu sinto que minhas informações pessoais em <i>chatbots</i> podem ser disponibilizadas a empresas ou agências governamentais sem meu conhecimento.	Adaptado de Dinev and Hart (2004)
INTENÇÃO DE PROTEÇÃO DE PRIVACIDADE	
IPP1. Posso fornecer informações pessoais falsas durante o uso de <i>chatbots</i> . IPP2. Posso fornecer informações pessoais incompletas durante o uso de <i>chatbots</i> . IPP3. Estou disposto a fazer configurações de privacidade no uso de <i>chatbots</i> .	Adaptado de Ko (2013)

Fonte: Elaboração própria (2026).

APÊNDICE B – TELAS E TEXTO VÍDEO QUESTIONÁRIO



Texto: Imagine que você está utilizando o serviço de atendimento ao cliente de uma loja online e é atendido por um *chatbot*. Ele solicita informações pessoais, como nome, e-mail, número do pedido para localizar o seu histórico de compras. Durante a interação, o *chatbot* também pede seu número de telefone e informações sobre suas preferências de compra, justificando que esses dados serão usados para personalizar o atendimento e oferecer melhores recomendações no futuro. Além disso, o *chatbot* informa que as conversas e os dados fornecidos serão armazenados para melhorar o serviço. Contudo, a mensagem não esclarece como esses dados serão protegidos ou quem terá acesso a eles. Com base nesse cenário, eu peço que você responda as próximas perguntas, levando em consideração sua percepção sobre privacidade, segurança e transparência.