

**MESTRADO EM
MARKETING**

**TRABALHO FINAL DE MESTRADO
DISSERTAÇÃO**

**INTELIGÊNCIA ARTIFICIAL GENERATIVA EM PUBLICIDADES:
INFLUÊNCIA NA CONFIANÇA NA MARCA**

GIULLIA SCARPA SCHMEISKE DA SILVA

JUNHO – 202

MESTRADO EM
MARKETING

TRABALHO FINAL DE MESTRADO
DISSERTAÇÃO

INTELIGÊNCIA ARTIFICIAL GENERATIVA EM PUBLICIDADES:
INFLUÊNCIA NA CONFIANÇA NA MARCA

GIULLIA SCARPA SCHMEISKE DA SILVA

ORIENTAÇÃO: PROFESSORA DOUTORA SUSANA CATARINA DE
JESUS FERNANDES DOS SANTOS

JUNHO - 2026

MESTRADO EM

MARKETING

TRABALHO FINAL DE MESTRADO

DISSERTAÇÃO

**INTELIGÊNCIA ARTIFICIAL GENERATIVA EM PUBLICIDADES:
INFLUÊNCIA NA CONFIANÇA NA MARCA**

GIULLIA SCARPA SCHMEISKE DA SILVA

JURY:

PRESIDENTE: PROFESSORA DOUTORA HELENA DO CARMO
MILAGRE MARTINS GONÇALVES

ORIENTADORA: PROFESSORA DOUTORA SUSANA CATARINA DE
JESUS FERNANDES DOS SANTOS

VOGAIS: PROFESSOR DOUTOR RICARDO MARINO FRANCISCO
RODRIGUES

JUNHO - 2026

AGRADECIMENTOS

Mudar de país aos 19 anos foi a coisa mais difícil que já fiz. Encarei o desconhecido, desbravei um país novo, sem amigos, sem família, apenas com o desejo de conquistar meu futuro. Em meio à saudade avassaladora de casa, eu me encontrei, cresci, fiz amigos, me apaixonei, terminei uma licenciatura e mudei de cidade. Agora, cinco anos depois, Portugal também é sinônimo de casa. Ao olhar para trás, me pergunto de onde tirei tanta coragem, mas é justamente disso que mais me orgulho. Eu poderia ter escolhido o caminho mais fácil, menos dolorido, com mais comodidades, mas nunca saberia o tamanho da minha força, da minha resiliência e da minha capacidade de enfrentar qualquer coisa. O caminho não foi fácil, mas eu garanto que valeu a pena. Por isso, o primeiro agradecimento vai para a Giullia de 19 anos, que deu o primeiro passo para que hoje, a Giullia de 23 pudesse estar aqui, concluindo mais um sonho.

Gostaria de dedicar esta dissertação aos meus pais, tias e avós, porque, sem eles, nada disto seria possível. Esta dissertação não é só minha, é o reflexo de cada sacrifício e de cada lágrima de saudade para que eu pudesse viver meu sonho. Dedico aos meus pais, que sob muito sol, me fizeram chegar até aqui pela sombra. Viver o que eu vivo é um privilégio, e eu sei que não é fácil. Por isso, agradeço do fundo do coração por me permitirem trilhar meu caminho, este é o ato mais puro de amor, pelo qual serei eternamente grata. Aos meus avós, que sempre acreditaram e investiram no meu potencial, eu tenho maior orgulho de ser neta de vocês. Às minhas tias, que são exemplos cotidianos de mulheres fortes e batalhadoras. Quero agradecer também ao meu namorado, Oliver, você é meu porto seguro no mar aberto que é a vida. Nossa cumplicidade, dedicação e horas na biblioteca juntos vão refletir em um futuro brilhante. Ao meu irmão, Pedro, que, mesmo sendo mais novo, é quem busco em momentos de crise, à minha melhor amiga, Gabi, que nunca soltou a minha mão e me mostrou o que é amizade verdadeira, e aos meus amigos, obrigada por serem a minha família onde quer que eu vá.

Por fim, mas longe de ser menos importante, gostaria de agradecer à minha orientadora, Susana Santos. Professora, sem você esta tese não teria sido possível. O seu compromisso e cuidado com nossos trabalhos mudou completamente esta experiência. Particularmente, vou lembrar sempre deste momento com carinho, pois em nenhum momento me senti à deriva, você esteve sempre lá, nos guiando, presente, e por isso sou extremamente grata. Muito obrigada por toda a empatia e consideração que teve comigo nestes últimos meses, e por ter tornado tudo mais leve.

DISCLOSURE DO USO DE IA

No desenvolvimento da presente dissertação foram utilizadas ferramentas de inteligência artificial como instrumentos auxiliares de apoio ao processo de revisão textual e à gestão das referências bibliográficas. As principais ferramentas utilizadas foram o *ChatGPT*, o *Elicit* e o *Anara*. O *ChatGPT* foi utilizado para apoiar a melhoria da coesão e clareza textual do estudo, bem como a concordância gramatical em português de Portugal. O *Elicit* e o *Anara*, por sua vez, prestaram apoio à organização dos artigos científicos selecionados nas bases de dados consultadas, bem como à coordenação das citações e referências. Em nenhum momento estas ferramentas substituíram a análise crítica, a interpretação de dados ou a construção argumentativa presente nesta investigação. O seu uso ocorreu de forma consciente e responsável, com objetivo de aumentar a eficiência do processo de revisão e formatação do texto, sem comprometer a responsabilidade autoral, originalidade e integridade acadêmica da presente dissertação, a autora assume total responsabilidade pelo conteúdo do trabalho.

RESUMO

A crescente utilização da inteligência artificial generativa na criação publicitária tem reconfigurado os processos de comunicação das marcas. Contudo, as respostas dos consumidores à sua sinalização é heterogênea, tornando cada vez mais relevante compreender os seus efeitos no comportamento dos consumidores. O presente estudo, teve como objetivo analisar de que forma o *disclosure* do uso de IA generativa em publicidades reconfigura a confiança dos consumidores na marca. Com base na *Assemblage Theory*, a confiança foi conceptualizada como um sistema emergente composto pela interação entre marca, consumidor e tecnologia, sendo operacionalizada através de três propriedades: autenticidade percebida, transparência percebida e competência tecnológica. A investigação adotou uma abordagem quantitativa, com recurso a uma experiência fatorial 3x3 entre sujeitos, cominando três marcas, ASICS, Ferrero Rocher e Rexona, com três condições de *disclosure*: sem *disclosure*, com *disclosure* IA Criação e *disclosure* de IA Sugestão. A amostra final foi composta por 363 respostas válidas, analisadas através de ANOVA, MANOVA, regressão linear múltipla e modelos PROCESS. Os resultados demonstram que o *disclosure* de IA não afeta diretamente a confiança na marca. Contudo, a condição IA Criação produziu efeitos ambivalentes: aumentou significativamente a transparência percebida, sobretudo as dimensões de Disclosure e Clareza, mas reduziu a originalidade percebida, fragilizando parcialmente a autenticidade. A mediação paralela, revelou que a IA Criação exerceu um efeito indireto negativo sobre a confiança por meio da autenticidade percebida. A competência tecnológica manteve-se estável entre as condições, embora tenha sido um preditor positivo da confiança. Por fim, a familiaridade com a IA moderou apenas o efeito do *disclosure* sobre a transparência. Conclui-se que o *disclosure* de IA generativa atua como força reconfiguradora e paradoxal do sistema da confiança, tornando a marca mais transparente, mas não necessariamente mais confiável quando a autenticidade é fragilizada.

Palavras-chave: Comportamento do Consumidor; Confiança na Marca; Rótulo do Uso de IA Generativa; Autenticidade da Marca; Transparência

ABSTRACT

The growing use of generative artificial intelligence in advertising creation has been reshaping brand's communication process. However, consumer response to AI disclosure remains heterogeneous, therefore making it increasingly important to understand its effects on consumer behavior. This study aims to analyze how disclosing generative AI involvement in advertisements reconfigures consumers' trust in brands. Drawing on Assemblage Theory, brand trust is conceptualized as an emergent system formed through the interaction between brand, consumer, and technology, and operationalized through three properties: perceived authenticity, perceived transparency, and technological competence. This study adopted a quantitative approach, using a 3x3 between-subjects factorial experiment that combined three brands, ASICS, Ferrero Rocher, and Rexona, with three disclosure conditions: no disclosure, AI Creation disclosure, and AI Suggestion disclosure. The final sample consisted of 363 valid responses, analyzed through ANOVA, MANOVA, multiple linear regression, and PROCESS models. The results show that AI disclosure does not directly affect brand trust. However, the AI Creation condition produced ambivalent effects: it significantly increased perceived transparency, particularly the Disclosure and Clarity dimensions, while reducing perceived originality, thereby partially weakening authenticity. Parallel mediation analysis revealed that AI Creation had a negative indirect effect on brand trust through perceived authenticity. Technological competence remained stable across conditions, although it was a positive predictor of trust. Finally, familiarity with AI moderated only the effect of disclosure on transparency. The study concludes that generative AI disclosure acts as a paradoxical reconfiguring force within the trust system, making the brand more transparent but not necessarily more trustworthy when authenticity is weakened.

Keywords: Brand Trust; Consumer Behavior; Generative AI Disclosure; Brand Authenticity; Transparency

LISTA DE SIGLAS E ABREVIATURAS

ANOVA	<i>Análise de Variância</i>
APC	<i>Análise de Componentes Principais</i>
CTR	<i>Click-Through-Rate</i>
IA	<i>Inteligência Artificial</i>
KMO	<i>Kaiser-Meyer-Olkin</i>
LLM	<i>Large-Language-Model</i>
MANOVA	<i>Análise Multivariada de Variância</i>
e-WOM	<i>Electronic Word-Of-Mouth</i>
WOM	<i>Word-Of-Mouth</i>

ÍNDICE

AGRADECIMENTOS	IV
DISCLOSURE DO USO DE IA.....	V
RESUMO	VI
ABSTRACT.....	VII
LISTA DE SIGLAS E ABREVIATURAS.....	VIII
1. INTRODUÇÃO	1
1.1 CONTEXTUALIZAÇÃO TEÓRICA	1
1.2 RELEVÂNCIA DO ESTUDO	2
1.3 OBJETIVO DO ESTUDO	3
1.4 ESTRUTURA DO ESTUDO.....	3
2. REVISÃO DE LITERATURA	4
2.1 IA GENERATIVA E A PUBLICIDADE	4
2.1.1 OS DIFERENTES TIPOS DE USO DA IA GENERATIVA NA PUBLICIDADE	5
2.1.2 O <i>DISCLOSURE</i> DO USO IA GENERATIVA NAS PUBLICIDADES E OS SEUS EFEITOS	5
2.2 CONFIANÇA NA MARCA.....	7
2.3 INTEGRAÇÃO TEÓRICA: CONFIANÇA COMO <i>ASSEMBLAGE</i> EMERGENTE	8
2.3.1 AUTENTICIDADE PERCEBIDA COMO FORÇA EXPRESSIVA DA CONFIANÇA	10
2.3.2 TRANSPARÊNCIA COMUNICACIONAL COMO FORÇA NORMATIVA DA CONFIANÇA.....	12
2.3.3 COMPETÊNCIA TECNOLÓGICA COMO FORÇA MATERIAL DA CONFIANÇA	13
2.4 FAMILIARIDADE COM A IA COMO MODERADORA DOS EFEITOS DO <i>DISCLOSURE</i>	15
2.4.1 FAMILIARIDADE E A AUTENTICIDADE PERCEBIDA	16
2.4.2 FAMILIARIDADE E A TRANSPARÊNCIA PERCEBIDA	16
2.4.3 FAMILIARIDADE E A COMPETÊNCIA TECNOLÓGICA	17
3. MODELO CONCEPTUAL.....	17
4. METODOLOGIA	18
4.1 TIPO DE ESTUDO	18
4.2 POPULAÇÃO E AMOSTRAGEM	19
4.3 RECOLHA DE DADOS	19
4.4 QUESTIONÁRIO.....	20
4.5 ESCALAS DE MEDIDA.....	21
4.6 TRATAMENTO E ANÁLISE PRELIMINAR DOS DADOS	21
5. ANÁLISE DE DADOS.....	22
5.1 CARACTERIZAÇÃO DA AMOSTRA	22
5.2 FIABILIDADE, CONSISTÊNCIA INTERNA E CRIAÇÃO DE ÍNDICES SINTÉTICOS.....	23
5.3 ANÁLISE DESCRITIVA DOS ÍNDICES	23
5.4 TESTE DE HIPÓTESES	24
5.4.1 EFEITO DO <i>DISCLOSURE</i> DE IA NA CONFIANÇA DA MARCA	24
5.4.2 EFEITO DO <i>DISCLOSURE</i> NA AUTENTICIDADE PERCEBIDA	25
5.4.3 O IMPACTO DA AUTENTICIDADE PERCEBIDA NA CONFIANÇA DA MARCA	25
5.4.4 EFEITO DO <i>DISCLOSURE</i> NA TRANSPARÊNCIA PERCEBIDA	26
5.4.5 O IMPACTO DA TRANSPARÊNCIA PERCEBIDA NA CONFIANÇA DA MARCA.....	27
5.4.6 O EFEITO DO <i>DISCLOSURE</i> NA COMPETÊNCIA TECNOLÓGICA	28
5.4.7 O IMPACTO DA COMPETÊNCIA TECNOLÓGICA NA CONFIANÇA DA MARCA	29
5.4.8 EFEITO MODERADOR DA FAMILIARIDADE COM A IA NOS EFEITOS DO <i>DISCLOSURE</i>	29
5.4.8.1 EFEITO MODERADOR DA FAMILIARIDADE COM A IA E A AUTENTICIDADE.....	29
5.4.8.2 EFEITO MODERADOR DA FAMILIARIDADE COM A IA E A TRANSPARÊNCIA	30
5.4.8.2 EFEITO MODERADOR DA FAMILIARIDADE COM A IA E A COMPETÊNCIA TECNOLÓGICA	31
5.4.9 INTEGRAÇÃO DO MODELO	31
5.4.9.1 MEDIAÇÃO PARALELA	31

5.4.9.2 MODERAÇÃO MEDIADA.....	32
6. DISCUSSÃO E CONCLUSÃO	33
6.1 DISCUSSÃO.....	33
6.2 CONCLUSÃO	34
6.3 CONTRIBUTOS ACADÊMICOS	35
6.4 CONTRIBUTOS PRÁTICOS.....	36
6.5 LIMITAÇÕES DO ESTUDO E SUGESTÕES DE INVESTIGAÇÃO FUTURA	36
7. REFERÊNCIAS	37
8. ANEXOS	43
ANEXO A - ANÚNCIOS PRESENTES NOS QUESTIONÁRIOS	43
ANEXO B – QUESTIONÁRIO ONLINE	44
ANEXO C - ESCALAS DE MEDIDA	49
ANEXO D - AMOSTRA POR SUBGRUPO	50
ANEXO E - PERFIL SOCIODEMOGRÁFICO DA AMOSTRA	50
ANEXO F - PCA E ANÁLISE DE FIABILIDADE	52
ANEXO G – ANÁLISE DESCRITIVA DOS ITENS	53
ANEXO H – TESTE DA HIPÓTESE 1: PRESSUPOSTOS	54
ANEXO I - TESTE DA HIPÓTESE 1: TWO-WAY ANOVA	54
ANEXO J – TESTE HIPÓTESE 2: PRESSUPOSTOS.....	55
ANEXO K - TESTE HIPÓTESE 2: TWO-WAY ANOVA	55
ANEXO L- REGRESSÃO LINEAR MÚLTIPLA – H3, H5 E H7	57
ANEXO M – TESTE HIPÓTESE 4: PRESSUPOSTOS.....	59
ANEXO N - TESTE HIPÓTESE 4: TWO-WAY ANOVA	59
ANEXO O – TESTE HIPÓTESE 6: PRESSUPOSTOS	61
ANEXO P – TESTE HIPÓTESE 6: TWO-WAY ANOVA	61
ANEXO Q – TESTE HIPÓTESE H8A: PROCESS MODELO 1	62
ANEXO R – TESTE HIPÓTESE H8B: PROCESS MODEL 1	63
ANEXO S - TESTE HIPÓTESE H8C: PROCESS MODEL 1	64
ANEXO T – MEDIAÇÃO PARALELA: PROCESS MODELO 4	64
ANEXO U- MODERAÇÃO MEDIADA: PROCESS MODELO 7	66

ÍNDICE DE FIGURAS

FIGURA 1 - MODELO CONCEPTUAL	18
------------------------------------	----

ÍNDICE DE TABELAS

TABELA 1 - MODELO DE REGRESSÃO LINEAR MÚLTIPLA MACRO PARA CONFIANÇA	26
---	----

1. INTRODUÇÃO

1.1 Contextualização teórica

A rápida ascensão da inteligência artificial (IA) generativa tem transformado profundamente os processos de criação e comunicação das marcas, bem como a própria indústria do marketing (Hartmann et al., 2025). Ferramentas como DALL-E 3 e ChatGPT permitem gerar imagens, textos e campanhas publicitárias a partir de comandos textuais, redefinindo os limites entre criação humana e tecnológica (Feuerriegel et al., 2024). Baseada em modelos de *deep learning* treinados em grandes volumes de dados, a IA generativa é capaz de identificar padrões e produzir conteúdos diversos, desde publicações em redes sociais até imagens e elementos criativos (Feuerriegel et al., 2024; Grewal et al., 2025). O impacto desta tecnologia já é visível no mercado. Campanhas como “A.I Ketchup”, da Heinz, alcançaram mais de 850 milhões de impressões globais e reconhecimento internacional (The One Club For Creativity, 2023). Num contexto de forte expansão económica, o uso de IA no marketing tem projeções de retorno global superiores a 107 mil milhões de dólares até 2028 e podendo atingir 1.3 trilhões até 2030 (Bloomberg, 2023; Statista, 2025). Paralelamente, a IA generativa reduz custos, acelera a produção de conteúdo e aumenta a personalização das campanhas, reforçando o seu potencial estratégico no marketing (Ammanath et al., 2024; Grewal et al., 2025; Hartmann et al., 2025).

Apesar destes benefícios, a adoção da IA generativa é acompanhada por incertezas relevantes. Cerca de 30% dos líderes de negócios e tecnologia reconhecem ressalvas quanto à sua aplicação prática (Ammanath et al., 2024), enquanto preocupações associadas à ética, transparência, privacidade e opacidade operacional podem comprometer a confiança dos consumidores (Feuerriegel et al., 2024; Glikson & Woolley, 2020; Grewal et al., 2025). Neste contexto, com o aumento da utilização de IA na criação de publicidades torna-se fundamental compreender as atitudes dos consumidores face a estes conteúdos (Campbell et al., 2022; Chen et al., 2024). Além disto, à medida que a distinção entre conteúdos artificiais e criações humanas se torna cada vez mais ténue (To et al., 2025), reforça-se a relevância ética do *disclosure*, para que o consumidor possa tomar decisões informadas acerca da mensagem publicitária (Chen et al., 2024). Ainda assim, as respostas dos consumidores variam conforme o contexto de uso e as suas crenças individuais, refletindo a complexidade deste fenómeno (Zhu et al., 2022).

1.2 Relevância do estudo

A credibilidade de uma publicidade está associada à fonte da informação e às características da mensagem, sendo a avaliação do público-alvo determinante para sua eficácia (Chen et al., 2024). No contexto da IA generativa, a aceitação dos conteúdos depende da confiança cognitiva, isto é, da percepção de capacidade, competência e confiabilidade atribuída à tecnológica, em que tende a ser mais elevada em tarefas técnicas e objetivas e menos favorável em contextos hedônicos e criativos (Chen et al., 2024; Glikson & Woolley, 2020). Neste cenário estudos recentes mostram que o *disclosure* do uso de IA generativa nas publicidades pode produzir efeitos paradoxais: embora possa reforçar a transparência e a percepção ética da marca, tende a reduzir a autenticidade percebida e, em alguns casos, a competência percebida (Baek et al., 2024; Chen et al., 2024; To et al., 2025; Wu et al., 2025). Adicionalmente, o uso de IA pode ser interpretado não como inovação, mas como uma falta de esforço criativo, comprometendo a confiança quando percebido como inautêntico (Campbell et al., 2022; To et al., 2025).

Apesar destes avanços, permanece uma lacuna relevante na compreensão destes impactos em sistemas mais amplos do comportamento do consumidor, como a confiança na marca (Chen et al., 2024; Cillo & Rubera, 2025; Rajavi et al., 2019). Considerando que a confiança constitui um antecedente fundamental da lealdade dos consumidores, torna-se fundamental investigar como ela pode ser reconfigurada face ao uso de IA generativa na publicidade (Chaudhuri & Holbrook, 2001; Khamitov et al., 2024; Rajavi et al., 2019). Além disto, a heterogeneidade das respostas ao *disclosure* sugere que a familiaridade com a IA desempenha um papel crítico, uma vez que a exposição crescente a estas tecnologias reduz a incerteza e reconfigura as heurísticas usadas na avaliação de autenticidade, transparência e competência tecnológica, podendo atenuar ou intensificar os seus efeitos sobre a confiança da marca (Glikson & Woolley, 2020; Rehman & Fareed, 2025; Topsakal, 2025).

A relevância deste tema estende-se para além da teoria. Chaudhuri e Holbrook (2001) destrincham que a confiança na marca constitui a base cognitiva do *brand equity*, ao desencadear efeitos comportamentais e psicológicos, como a lealdade comportamental e a lealdade atitudinal, que traduzem-se em maiores quotas de mercado e maior disposição para pagar preços *premium* (Chaudhuri & Holbrook, 2001; Khamitov et al., 2024; Rajavi et al., 2019). Dado o papel fundamental da confiança no sucesso da marca (Rajavi et al.,

2019), compreender os efeitos do *disclosure* da IA generativa em publicidades torna-se crucial, sobretudo porque anúncios são um dos principais instrumentos de comunicação e construção da marca (Chaudhuri & Holbrook, 2001; Khamitov et al., 2024). Do ponto de vista da gestão, esta compreensão igualmente essencial, já que o *disclosure* pode influenciar a eficácia persuasiva das campanhas e deve ser estrategicamente alinhado ao posicionamento da marca e ao tipo de tarefa comunicacional, podendo gerar percepções negativas de autenticidade, criatividade ou congruência, ao mesmo tempo que impacta a construção e proteção de ativos de confiança no longo prazo (Ali et al., 2025; Cillo & Rubera, 2025; Wu et al., 2025).

1.3 Objetivo do estudo

Com o objetivo de aprofundar a literatura sobre confiança na marca, esta investigação procura compreender os potenciais impactos do uso de IA generativa em publicidades na relação entre marca e consumidor, respondendo à seguinte questão: “Como é que o uso de inteligência artificial (IA) generativa na publicidade reconfigura a confiança dos consumidores na marca?”. Para isto, adota-se a lente ontológica de Manuel DeLanda (2006) aplicada ao contexto de mercado por Canniford e Bajde, (2015), para compreender a confiança na marca como sistema emergente e multifacetado (i.e., *assemblage*) resultante das interações entre marca, consumidor e tecnologia (Duffek et al., 2025). Assim, o objetivo central do estudo é testar empiricamente como as *propriedades emergentes* (autenticidade percebida, transparência percebida e competência tecnológica) estabilizam ou desestabilizam a confiança na marca no contexto do *disclosure* do uso de IA generativa em publicidades. Especificamente pretende-se: analisar os efeitos do *disclosure* da IA sobre as propriedades emergentes; examinar como essas percepções influenciam a confiança na marca; verificar se a familiaridade com a IA modera a relação entre o *disclosure* e as propriedades emergentes; e por fim, interpretar os resultados à luz da *Assemblage Theory* (DeLanda, 2006).

1.4 Estrutura do estudo

O presente estudo organiza-se em sete capítulos; (1) Introdução; (2) Revisão de Literatura; (3) Modelo Conceptual; (4) Metodologia; (5) Análise dos Resultados dos Questionários; (6) Discussão; e por fim (7) Conclusão. O primeiro capítulo, é composto pela Introdução do estudo, desenvolvendo a sua contextualização teórica, relevância, o objetivo central e objetivos específicos da dissertação e a estruturação do trabalho. O segundo capítulo engloba a Revisão de Literatura, fornecendo uma visão detalhada do

corpo teórico e fundamentando a investigação. O capítulo do Modelo Conceptual, demonstra a estruturação dos objetivos e hipóteses que serão estudadas. Posteriormente, a metodologia aplicada neste estudo é descrita, seguida pelas análises realizadas e discussões sobre os resultados provenientes do questionário online aplicado, finalizando com as conclusões da dissertação e suas limitações.

2. REVISÃO DE LITERATURA

2.1 IA generativa e a publicidade

Inteligências artificiais (IA) generativas são modelos de caráter auto supervisionados, cujo treino consiste em retirar partes dos dados e prever os elementos em falta, permitindo, ao final do processo, gerar múltiplos resultados plausíveis para preencher estas lacunas (Cillo & Rubera, 2025). Estruturadas por modelos complexos, como *large language models* (LLM), são treinadas em volumes extensos e heterogêneos de dados (*inputs*), o que lhes permite produzir novos *outputs* baseadas em bases de dados pré existentes (Cillo & Rubera, 2025; Feuerriegel et al., 2024). Assim, as inteligências artificiais podem ser diferenciadas tanto pela natureza de seus *inputs* quando dos seus *outputs* (Grewal et al., 2025).

As IA Generativas operam por diferentes modalidades (e.g., texto, imagens, dados numéricos), e por meio do reconhecimento de padrões (*deep neural network*), são capazes de gerar novos conteúdos inéditos e diversificados (Feuerriegel et al., 2024). Esta capacidade ultrapassa os limites das IA analíticas, tradicionalmente focadas em tarefas objetivas e preditivas (Grewal et al., 2025). Ao viabilizar a criação automatizada de conteúdos sintéticos, incluindo imagens e vídeos ultrarrealistas, expande os limites da produção criativa e altera estes processos antes exclusivamente humanos (Campbell et al., 2022). No contexto publicitário, estas tecnologias introduzem a possibilidade de produzir conteúdos originais e escaláveis em larga escala, revolucionando a indústria do marketing e os seus modelos tradicionais de criação de valor (Chen et al., 2024; Feuerriegel et al., 2024; Hartmann et al., 2025).

Contudo, apesar dos ganhos operacionais e de eficiência, os decisores ainda se mostram receosos quanto à adoção e ao controlo desta tecnologia (Ammanath et al., 2024). Tais reservas decorrem dos riscos e limitações que comprometem a sua legitimidade e confiança, como alucinação de respostas, infração de direitos autorais, a reprodução de vieses problemático e a opacidade no tratamento e armazenamento de

dados (Cillo & Rubera, 2025; Feuerriegel et al., 2024; Grewal et al., 2025). Ademais, a opacidade operacional e diretrizes vagas de privacidade, reduzem a confiança nesta tecnologia impedindo a sua implementação e a sua aceitação (Glikson & Woolley, 2020).

2.1.1 Os diferentes tipos de uso da IA generativa na publicidade

Mesmo face a estes desafios, a IA generativa tem sido integrada em várias fases do processo publicitário, sendo igualmente crucial para a recolha de *insights* dos consumidores (Wu et al., 2025). Wu et al. (2025) destacam que, para além da criação do conteúdo e medição de performance, a IA é capaz de automatizar as decisões de *targeting* e posicionamento do anúncio, aumentando a sua eficácia ao manipular grandes volumes de dados comportamentais que dificilmente poderiam ser tratados apenas por humanos. Do mesmo modo, modelos recentes conseguem gerar automaticamente diferentes elementos publicitários, desde *copy* (e.g., sistemas de LLM), até imagens e porta-vozes sintéticos (e.g., sistemas de *text-to-image*) (Wu et al., 2025). O estudo de Hartmann et al. (2025) ilustra este carácter disruptivo, ao demonstrar que *banners* publicitários criados por IA generativas obtiveram uma performance 50% superior em *clicks-to-rate* (CTR) em comparação aos *banners* desenvolvidos com imagens de *stock* de alta qualidade, reduzindo significativamente o custo de criação de conteúdo e simultaneamente aumento o volume de produção (Cillo & Rubera, 2025; Hartmann et al., 2025).

Ademais, Grewal et al. (2025) apontam que, à medida que a utilização da IA generativa no marketing se intensifica, os seus resultados criativos e decisórios tendem a ser refinados por meio de testes de larga escala ao longo do processo publicitário. Ainda assim a sua utilização suscita preocupações relativas à autonomia do consumidor, à ocultação de processos internos e às implicações éticas de persuasão (Shi & Jiang, 2026). Por isto, o debate acerca do direito do consumidor em saber quando a IA generativa está presente nas publicidades tem ganhado relevância entre académicos e agências reguladoras (Chen et al., 2024; Glikson & Woolley, 2020), sobretudo porque a sinalização da utilização da IA pode afetar as perceções de inovação, autenticidade e credibilidade, critérios centrais para a confiança do consumidor na marca (Baek et al., 2024; Chen et al., 2024; Shi & Jiang, 2026).

2.1.2 O *disclosure* do uso IA generativa nas publicidades e os seus efeitos

Os rápidos avanços da IA generativa já permitem criar *deepfakes* e publicidades sintéticas ultrarrealistas que ofuscam as fronteiras entre o conteúdo artificial e humano (Campbell et al., 2022). Neste sentido, plataformas digitais e organizações regionais

caminham para a implementação de rótulos de *disclosure* sobre conteúdos gerados por IA (e.g., imagens, vídeos e músicas), como estratégia para educar os usuários sobre a presença de IA generativa (Baek et al., 2024; Campbell et al., 2022; Hartmann et al., 2025). Wu et al. (2025) defendem a crucialidade do *disclosure*, não apenas para reforçar a confiança na marca, mas também por contribuir para a literacia digital dos consumidores e para o seu direito à informação, permitindo-lhes reconhecer as tecnologias subjacentes às publicidade e proteger-se de táticas de manipulação (Baek et al., 2024; Cillo & Rubera, 2025; Wu et al., 2025).

Contudo os efeitos do *disclosure* não são unívocos, variando conforme o contexto, a categoria do produto e das percepções psicológicas individuais acerca da IA generativa (Shi & Jiang, 2026). Quando utilizada para fins predominantemente funcionais (e.g., descrição e especificações de produtos) e devidamente sinalizada, a autenticidade e voz da marca não são prejudicadas, sugerindo que a transparência ética pode ser estrategicamente utilizada em contextos utilitários e de baixo teor emocional (Kirkby et al., 2023). De forma similar, quando a IA generativa é percecionada como competente na resolução de tarefas complexas, o *disclosure* pode aumentar a intenção de compartilhamento e WOM (Shi & Jiang, 2026).

Ainda assim, a literatura também aponta para efeitos adversos. O uso indiscriminado dos rótulos de IA, pode reduzir a credibilidade de anúncios de cunho social (Baek et al., 2024), enfraquecer a imagem percebida e autenticidade da marca (Ali et al., 2025; To et al., 2025) e diminuir a intenção de uso em processos de recrutamento, quando as indicações da IA são percebidas como sintéticas e pouco humanas (Shi & Jiang, 2026). Ademais, quando a IA generativa é identificada, mas não explicitamente declarada, a confiança na marca pode ser ainda mais minada, pois os consumidores tendem a interpretar esta omissão como tentativa de engano (Shi & Jiang, 2026).

Esta ambiguidade evidencia que o debate e torno da sinalização da IA generativa na publicidade não deve constranger-se à decisão de implementar ou não o *disclosure*, mas antes centrar-se em como calibrar a sua adequação face às normas éticas e ao uso consciente da tecnologia (Baek et al., 2024; Cillo & Rubera, 2025; Grewal et al., 2025; Shi & Jiang, 2026; Wu et al., 2025). Importa notar, contudo, que este efeito paradoxal permanece insuficientemente explorado, sobretudo no que concerne aos mecanismos psicológicos que explicam a sua influência sobre a confiança ativados pela a sinalização da IA generativa nas publicidades (Shi & Jiang, 2026). Uma explicação plausível advém do fenómeno conhecido como “aversão à IA” ou “aversão ao algoritmo”, segundo o qual

os indivíduos, ao tomarem consciência do uso da tecnologia, tendem a avaliar tanto o conteúdo quanto a marca como menos humanos e mais artificiais, desencadeando respostas de desconfiança e rejeição (Dang & Liu, 2024; Shi & Jiang, 2026).

2.2 Confiança na marca

Mayer et al. (1995) definem confiança como a disposição do indivíduo em aceitar vulnerabilidade diante das ações de outro ator, baseando-se apenas na expectativa de que este cumprirá com suas obrigações, mesmo sem o controle direto sobre o resultado. No marketing, a confiança na marca é tradicionalmente compreendida como a crença do consumidor na capacidade e na disposição da marca em cumprir com os compromissos que comunica (Chaudhuri & Holbrook, 2001; Khamitov et al., 2024; Rajavi et al., 2019). Assim, tais convicções emergem do processo relacional entre consumidor e marca ao longo de interações e experiências acumuladas (Portal et al., 2019), sustentadas por avaliações cognitivas acerca dos seus atributos expressivos e materiais (Khamitov et al., 2024; Rajavi et al., 2019). Neste contexto, a lente ontológica de DeLanda (2006) aplicado ao contexto de mercado por Canniford e Bajde (2015), permitem compreender a confiança como um fenómeno emergente de interações entre *componentes heterogêneos*, cujas relações podem se estabilizar ou se fragilizar.

A literatura converge ao indicar que esse julgamento relacional se organiza em torno de duas dimensões centrais: *integridade* e *competência* (Khamitov et al., 2024; Portal et al., 2019). Nesse sentido, a integridade corresponde às percepções do consumidor sobre a essência moral da marca, isto é, se ela é percebida como honesta, ética e alinhada com os seus valores, reforçando a credibilidade das suas promessas e ações quando estas são interpretadas como autênticas e verdadeiras (Portal et al., 2019). Já a competência refere-se a percepção da capacidade da marca para entregar desempenho e valor de forma consistente, sendo reforçada por sinais como qualidade percebida, valor agregado, investimentos de marketing, publicidade e inovação (Khamitov et al., 2024; Portal et al., 2019; Rajavi et al., 2019). Assim, enquanto a integridade reduz a incerteza por meio de inferências de autenticidade e transparência ética, a competência sinaliza a capacidade de cumprimento e fiabilidade (Khamitov et al., 2024; Rajavi et al., 2019).

A confiança na marca constitui, por isto, um elo fundamental na construção de relações duradouras entre consumidor e marca, estando associada a efeitos atitudinais como lealdade e satisfação, e a efeitos comportamentais, como a intenção de compra, a disposição para pagar preços superiores e *word-of-mouth* (WOM) (Chaudhuri &

Holbrook, 2001; Khamitov et al., 2024). Deste modo, o presente estudo compreende a confiança na marca como um sistema emergente constituído por atores independentes: marca, consumidores e tecnologia, cuja configuração resulta de interações contínuas (Canniford & Bajde, 2015; DeLanda, 2006). Neste quadro, das interações destes atores, *propriedades emergentes* caracterizam os antecedentes de confiança que derivam das dimensões de *integridade* e *competência*, operacionalizados como autenticidade percebida, transparência percebida e competência tecnológica, refletindo respectivamente as forças expressiva, normativa e material (Duffek et al., 2025; Khamitov et al., 2024; Portal et al., 2019)

2.3 Integração teórica: confiança como *assemblage* emergente

DeLanda (2006) define os fenómenos sociais como *assemblages* resultantes das interações de componentes heterogêneos que se interligam sem perder a sua autonomia, podendo estabilizar o sistema (*territorialização*) ou fragilizá-lo (*desterritorialização*). Aplicada à confiança na marca, esta perspectiva permite compreendê-la não como uma qualidade inerente da marca, mas como um fenómeno relacional e emergente, desenvolvido ao longo do alinhamento entre as expectativas dos consumidores quanto à capacidade e integridade da marca em cumprir as suas promessas (Chaudhuri & Holbrook, 2001; Khamitov et al., 2024; Rajavi et al., 2019). Neste sentido, a confiança resulta de interações contínuas entre consumidores, marcas e tecnologia, especialmente em contextos digitais e nas redes sociais, onde estes atores se articulam de forma cada vez mais intensa (Canniford & Bajde, 2015; Duffek et al., 2025). À luz desta lente ontológica, este estudo traduz os componentes heterogêneos da confiança na marca, consumidores e tecnologia, e sustenta que, quando interagem de forma consistente e alinhada, as propriedades emergentes da confiança, autenticidade, transparência percebida e competência tecnológica, se estabelecem e *territorializam* o sistema de confiança (Duffek et al., 2025; Khamitov et al., 2024; Rajavi et al., 2019).

Em termos ontológicos, DeLanda (2006) distingue dimensões expressivas e materiais dos *assemblages*, isto é, a combinação de elementos simbólicos e capacidades práticas. Aplicando ao contexto de mercado, Canniford e Bajde (2015) argumentam que as marcas funcionam como arranjos sociotécnicos, nos quais significados, normas e componentes tecnológicas se articulam continuamente e orientam o julgamento do consumidor. Neste estudo, a autenticidade percebida corresponde a um dos componentes da dimensão expressiva da confiança, por refletir percepções de genuinidade e

credibilidade construídas ao longo das comunicações e experiências acumuladas do consumidor com a marca (Duffek et al., 2025; Kirkby et al., 2023; Portal et al., 2019). Ademais, a transparência percebida opera como a dimensão normativa, na medida em que a clareza ética do processo comunicacional, nomeadamente por via do *disclosure*, reduz incertezas e reforça o “direito de saber” dos consumidores em contextos de conteúdos sistéticos (Kirkby et al., 2023; Mittelstadt, 2019; Wu et al., 2025). Por fim, a competência tecnológica traduz a dimensão material do *assemblage*, por envolver inferências sobre a capacidade da marca de mobilizar novas tecnologias como a IA generativa de forma eficaz e adequada para sustentar o seu desempenho e cumprir suas promessas (Glikson & Woolley, 2020; Khamitov et al., 2024; Rajavi et al., 2019). Embora possam ser analisadas separadamente, estas devem ser compreendidas como partes interligadas do sistema da confiança (Mayer et al., 1995).

Nesta perspectiva, o *assemblage* da confiança territorializa-se quando a marca é percebida como autêntica, comunica de forma transparente e demonstra competência na utilização da tecnologia, reduzindo a incerteza e reforçando a disposição do consumidor em confiar (Khamitov et al., 2024; Portal et al., 2019; Rajavi et al., 2019). Contudo, os *assemblages* são inerentemente instáveis e podem ser desterritorializados quando mudanças contextuais salientam tensões entre as suas dimensões expressivas, normativas e materiais (Canniford & Bajde, 2015; DeLanda, 2006). É precisamente neste ponto que a lente de DeLanda (2006) revela-se particularmente adequada, ao permitir captar reconfigurações simultâneas e até contraditórias que modelos lineares tendem a tratar de forma unilateral (Baek et al., 2024; To et al., 2025; Wu et al., 2025).

Deste modo, o *disclosure* do uso de IA generativa em publicidades atua como uma força reconfiguradora do sistema da confiança devido à sua natureza paradoxal. Por um lado, ao revelar os processos de criação do conteúdo, pode *territorializar* a dimensão normativa da confiança, reforçando a transparência e a legitimidade ética (Glikson & Woolley, 2020; Kirkby et al., 2023). Por outro, pode *desterritorializar* as dimensões expressiva e material ao ativar percepções negativas de empenho humano, originalidade e adequação da IA a tarefas criativas, reduzindo assim, a autenticidade e competência tecnológicas percebidas (Baek et al., 2024; Glikson & Woolley, 2020; To et al., 2025). Como a confiança na IA generativa não é homogênea, o *disclosure* desencadeia um sistema de ganhos e perdas, no qual os consumidores avaliam a congruência entre o uso da tecnologia, a essência simbólica da marca e o apelo do anúncio (Grewal et al., 2025). Quando este alinhamento falha, ou quando se ativa a “aversão à IA”, os indivíduos

tendem a rejeitar o conteúdo e a desacreditar a marca pelo uso da tecnologia (Dang & Liu, 2024; Shi & Jiang, 2026; To et al., 2025). Desta forma, propõe-se:

H1: O *disclosure* de IA generativa em publicidades impacta negativamente a confiança na marca.

2.3.1 Autenticidade percebida como força expressiva da confiança

Na literatura de *branding*, a autenticidade é frequentemente definida como a percepção de que uma marca é genuína e credível, construída por meio de uma comunicação consistente e das experiências dos consumidores, sendo elemento essencial para a formação da confiança (Ali et al., 2025; Rajavi et al., 2019). Contudo, trata-se de uma propriedade inerentemente instável, que exige esforço contínuo de marketing para ser sustentada, pois incongruências percebidas podem erodir a confiança construída ao longo do tempo (Duffek et al., 2025). Neste sentido, o *disclosure* do uso de IA generativa em publicidades tende a alterar os julgamentos sobre a mensagem e sobre o emissor, pois os consumidores avaliam o conteúdo com base em inferências acerca de autoria, intenção e empenho criativo (Kirkby et al., 2023). Este efeito pode ser explicado tanto pela percepção de artificialidade do conteúdo gerado (Zhu et al., 2022), quanto pela ativação da “aversão à IA”, na qual a publicidade passa a ser avaliada negativamente por carecer de essência humana (Dang & Liu, 2024). Além disto, como a autenticidade é uma condição central para identificação e o vínculo relacional com a marca, ela tende a ser fragilizada quando a IA é percebida como incapaz de expressar nuances culturais e emocionais próprias de contextos humanos (Portal et al., 2019; Tokita, 2024).

Ainda assim, a ocultação da presença da IA generativa nas publicidades não deve ser entendida como antídoto para os potenciais efeitos adversos do *disclosure* (Kirkby et al., 2023). Quando descoberta, esta omissão pode ser interpretada como violação da integridade, tornando-se ainda mais corrosiva para a confiança (Shi & Jiang, 2026). Assim, num contexto em que distinção entre a produção humana e artificial torna-se progressivamente opacas (To et al., 2025), o *disclosure* assume relevância normativa ao informar os consumidores e contribuir para o letramento tecnológico dos indivíduos (Wu et al., 2025), podendo também reforçar a integridade da marca (Khamitov et al., 2024; Rajavi et al., 2019), ao sinalizar um compromisso explícito com a transparência ética nas comunicações (Kirkby et al., 2023; Shi & Jiang, 2026).

Todavia, as crescentes exigências de transparências podem instaurar um paradoxo no *assemblage* da confiança (Chen et al., 2024; Kirkby et al., 2023; Mittelstadt, 2019; Wu

et al., 2025). Se, por um lado, o *disclosure territorializa* a força normativa ao alinhar a marca às normas regulatórias e ao direito à informação (Baek et al., 2024; Cillo & Rubera, 2025), por outro lado pode *desterritorializar* as forças expressiva e material ao ativar inferências negativas sobre a autenticidade e, em certos contextos, sobre a competência tecnológica, sobretudo quando o uso de IA é interpretado como sinal de menor empenho criativo ou inadequação a tarefa (Chen et al., 2024; Glikson & Woolley, 2020; To et al., 2025). A literatura já documenta este efeito em diferentes contextos. Em anúncios de cunho social, o *disclosure* reduziu a credibilidade e intenções *call-to-action* (Baek et al., 2024). Já no luxo, diminui a autenticidade e qualidade percebida por contrariar expectativas de artesanato e empenho (To et al., 2025), e no setor da restauração, reduziu a genuinidade e originalidade percebidas, especialmente entre consumidores que valorizam marcas mais humanas e calorosas (Ali et al., 2025; Portal et al., 2019). Assim, estabelece-se:

H2: O *disclosure* do uso de IA nas publicidades impacta negativamente a autenticidade percebida da marca.

Propõe-se também que a autenticidade percebida, por sua vez, irá influenciar a confiança na marca (Ali et al., 2025). Portal et al. (2019) argumentam que uma marca autêntica é aquela que age em consonância com os seus valores, traduzindo-os nas suas operações e comunicações, e demonstrando continuamente comprometimento com as suas promessas. Tal coerência configura a marca como uma entidade íntegra, competente e digna de confiança (Portal et al., 2019). Concomitantemente, Ali et al. (2025) reforçam esta ligação ao sustentar que a autenticidade constitui um dos principais motores de lealdade à marca, especialmente dentre os indivíduos que valorizam a sinceridade e um toque humano nas comunicações.

Considerando que os consumidores buscam marcas com quem possam se identificar e construir conexões, o desenvolvimento destes relacionamentos tende a ocorrer de forma análoga aos relacionamentos sociais (Ali et al., 2025). Assim, empresas percebidas como autênticas tendem a conquistar consumidores ao demonstrarem consistentemente “calor humano”, originalidade, boas intenções e o alinhamento dos seus valores nas suas comunicações (Ali et al., 2025; Chaudhuri & Holbrook, 2001; Portal et al., 2019). Dessa forma, a autenticidade pode ser entendida como a dimensão expressiva da confiança, pois uma comunicação coerente e consistente sinaliza integridade e credibilidade, reduzindo a incerteza por meio do compromisso contínuo com seus valores e atitudes declaradas (Kirkby et al., 2023; Portal et al., 2019). Khamitov et al. (2024)

corroboram essa perspectiva ao demonstrar que os antecedentes associados à integridade exercem maior influência no desenvolvimento da confiança do que aqueles relacionados exclusivamente à competência da marca (Khamitov et al., 2024). Assim propõe-se:

H3: A autenticidade percebida impacta positivamente na confiança na marca.

2.3.2 Transparência comunicacional como força normativa da confiança

A transparência comunicacional refere-se ao grau em que a marca comunica de maneira clara e intencional, a origem, os processos e as intenções de suas mensagens (Kirkby et al., 2023; Schnackenberg et al., 2021). Na produção do conteúdo criativo, esta clareza tende a reforçar percepções éticas e a confiança ao sinalizar honestidade institucional (Kirkby et al., 2023). Neste cenário, o *disclosure* de IA generativa em publicidades torna-se crucial por assegurar o “direito de saber” do consumidor e mitigar assimetrias informacionais em um ambiente marcado por desinformação e opacidade entre conteúdos humanos e sintéticos (Kirkby et al., 2023; Wu et al., 2025).

Contudo, este processo é paradoxal. Ao passo que reforça a transparência da marca, o rótulo sinaliza o uso de táticas de persuasivas mediadas por tecnologia (Shi & Jiang, 2026). Baek et al. (2024) demonstram que o *disclosure* ativa o “conhecimento persuasivo” dos indivíduos, tornando-os mais céticos quanto à credibilidade do anúncio e elevando suas barreiras cognitivas de defesa. Este efeito é agravado pela opacidade estrutural dos sistemas de IA, cujos processos de funcionamento são complexos e pouco acessíveis à população geral, contribuindo para erosão da confiança em sistemas algorítmicos (Glikson & Woolley, 2020). Riscos associados a enviesamento, alucinações, fabricação de informação e indefinições em diretrizes de privacidade e segurança (Feuerriegel et al., 2024; Grewal et al., 2025; Mittelstadt, 2019) alimentam resistência e percepções negativas em relação ao algoritmo (Baek et al., 2024), sobretudo quando os consumidores percebem a opacidade operacional e ética (Glikson & Woolley, 2020; Sundar & Kim, 2019).

Adicionalmente, os efeitos do *disclosure* variam conforme o tipo de tarefa atribuído à IA. Quando associada a tarefas objetivas, como o posicionamento de anúncios, a tecnologia pode ativar heurísticas positivas de objetividade e capacidade, elevando o *word-of-mouth* e as suas percepções de competência (Glikson & Woolley, 2020; Wu et al., 2025). Porém, quando vinculada à criação de conteúdo criativo, tende a ativar percepções de inadequação ou limitação da IA, reduzindo a intenção de compartilhamento, tendo em vista que os consumidores não a consideram plenamente

capaz de desempenhar tarefas criativas e emocionalmente sensíveis (Chen et al., 2024; Cillo & Rubera, 2025). Assim, ainda que o *disclosure* se apresente como mecanismo normativo de transparência, a sua sinalização pode não ser interpretada como aumento efetivo da transparência percebida, mas como evidência de mediação tecnológica opaca e potencialmente manipulativa (Baek et al., 2024; Shi & Jiang, 2026; Wu et al., 2025). Logo propõe-se:

H4: O *disclosure* de IA nas publicidades impacta negativamente a transparência percebida da marca.

Ainda assim, apesar dos efeitos adversos associados ao *disclosure*, a transparência comunicacional da marca permanece fundamental para o desenvolvimento da confiança cognitiva dos consumidores (Glikson & Woolley, 2020), pois omissões ou ambiguidades nesta comunicação podem prejudicar as atitudes do público (Schnackenberg et al., 2021). Segundo Glikson e Woolley (2020), explicitar claramente as lógicas operacionais da marca, do desenvolvimento à criação publicitária, reduz a incerteza e alinha expectativas às capacidades percebidas. Portanto, diante do avanço das regulações internacionais, a transparência deve ser comunicada estrategicamente, instruindo os consumidores sem comprometer o valor de marca (Cillo & Rubera, 2025) ou reduzir as intenções de compartilhamento (Wu et al., 2025).

Essa recepção ao *disclosure* tende a ser mais favorável quando a IA generativa é utilizada em anúncios de apelo racional, visto que os consumidores aceitam melhor recomendações de cunho utilitário feitas pela tecnologia (Chen et al., 2024), ou em tarefas de baixo teor emocional, que preservam a imagem e capacidade da marca (Kirkby et al., 2023). Ademais, a transparência na utilização de IA generativa pode atuar como mecanismo normativo de sustentação da confiança (Glikson & Woolley, 2020; Portal et al., 2019; Wu et al., 2025). Quando adaptado ao contexto estratégico, o *disclosure* legitima processos internos, reduz incertezas e pode elevar as percepções de inovação e curiosidade (Shi & Jiang, 2026). Logo, propõe-se:

H5. A transparência percebida impacta positivamente na confiança na marca.

2.3.3 Competência Tecnológica como força material da confiança

Na literatura de confiança na marca, a competência refere-se à capacidade operacional da organização em cumprir com suas promessas, sendo o desempenho pilar fundamental da confiança (Portal et al., 2019; Rajavi et al., 2019). Neste estudo, a competência tecnológica percebida reflete a crença do consumidor de que a marca é capaz

de mobilizar tecnologias, como a IA generativa, de forma eficaz e apropriada para executar tarefas para sustentar as suas promessas, ativando a percepção de inovação como motor deste vínculo (Glikson & Woolley, 2020). Ao empregar uma tecnologia avançada, este estímulo pode, despertar curiosidade ao sinalizar esta novidade (Shi & Jiang, 2026).

Contudo, a confiança em IA ancora-se sobretudo em componentes cognitivos, isto é, a aceitação dos seus resultados depende das inferências de competência que os indivíduos formulam sobre ela (Glikson & Woolley, 2020). Sundar e Kim (2019) argumentam que heurísticas associadas à máquina moldam atitudes e julgamentos durante as interações com os sistemas automatizados. Assim, quando a presença da tecnologia se torna saliente por via do *disclosure*, podem ser ativados estereótipos de objetividade e neutralidade em contextos de tarefas práticas. Empiricamente, as evidências confirmam que a confiança na IA generativa é calibrada em função da natureza da tarefa realizada, na medida em que as percepções de confiabilidade e competência tendem a ser superiores em tarefas técnicas e analíticas (e.g., previsão de vendas e cálculos) e inferiores ao realizar tarefas sociais, morais ou criativas, frequentemente associadas a cognições afetivas consideradas inerentemente humanas (Chen et al., 2024; Glikson & Woolley, 2020).

Deste modo, a IA generativa expressa a dimensão material da confiança na marca ao oferecer altos níveis de inovação tecnológica, mas também introduz riscos éticos e simbólicos (Chen et al., 2024; Cillo & Rubera, 2025; Feuerriegel et al., 2024). Se, por uma lado, transforma a indústria publicitária por meio da automatização da produção de conteúdo (Campbell et al., 2022; Grewal et al., 2025), por outro, o *disclosure* da sua presença pode levar os consumidores a avaliar criticamente a capacidade da IA para desempenhar atividades criativas e emocionalmente carregadas (Chen et al., 2024; Glikson & Woolley, 2020; Wu et al., 2025). Assim, quando há o alinhamento entre as heurísticas de capacidade e características da tarefa, o *assemblage* da confiança tende a estabilizar-se, porém quando a IA generativa é percebida como inadequada para executar tarefas emocionais ou criativas, a confiança pode ser *desterritorializa* (Glikson & Woolley, 2020). Com base nisto, fundamenta-se a seguinte hipótese:

H6. O *disclosure* do uso de IA generativa impacta negativamente a competência tecnológica da marca.

As definições tradicionais de confiança definem a competência da marca como a sua capacidade de cumprir com as tarefas prometidas (Mayer et al., 1995). Neste sentido, Khamitov et al. (2024) argumentam que organizações que sinalizam capacidades técnicas robustas por meio de alta qualidade percebida, inovação e desempenho consistente

constroem uma base cognitiva sólida para o desenvolvimento da confiança. Assim, a introdução de novas ferramentas, como a IA generativa nos processos de marketing pode sinalizar o compromisso com a qualidade, eficiência e diferenciação, reforçando a dimensão de competência da confiança (Khamitov et al., 2024).

Evidências empíricas sugerem que o uso contextualizado da IA generativa nas publicidades pode aumentar a competência tecnológica percebida (Shi & Jiang, 2026). Shi et al. (2026) demonstram que, quando porta-vozes artificiais são devidamente sinalizados, as intenções de compra aumentam, uma vez que a transparência tecnológica impacta positivamente a capacidade percebida e a imagem de inovação da marca. De maneira semelhante, publicidades geradas e sinalizadas com IA generativa podem ser avaliadas mais favoravelmente, elevando as intenção de compartilhamento (Chen et al., 2024). Desta forma, demonstrar maestria na utilização de tecnologias avançadas, por meio de comunicações claras, *performance* consistente e investimentos visíveis, estabelece caminhos para converter sistemas inovadores em confiança durável dos consumidores (Khamitov et al., 2024). Assim propõe-se:

H7: A competência tecnológica impacta positivamente na confiança da marca.

2.4 Familiaridade com a IA como moderadora dos efeitos do *disclosure*

A confiança dos indivíduos na IA generativa é fundamental para seu emprego e para a aceitação dos resultados por ela gerados (Glikson & Woolley, 2020). Neste processo, a familiaridade com a tecnologia desempenha papel central, uma vez que as percepções dos consumidores sobre as capacidades da IA são fortemente influenciadas pelo seu grau de conhecimento e pelas experiências prévias de uso (Topsakal, 2025). Segundo Topsakal (2025), níveis mais elevados de familiaridade tendem a reduzir a incerteza, aumentar a facilidade de utilização e reforçar a confiança e a adoção da tecnologia em diferentes contextos. Adicionalmente, a aceitação da IA depende de percepções de transparência, capacidade e autenticidade, que também são moldadas pela experiência acumulada com a tecnologia (Rehman & Fareed, 2025). Sob esta ótica, o uso repetido de rótulos de IA nas publicidades pode gerar mudanças sistemática nas respostas dos consumidores à sinalização (Baek et al., 2024; Campbell et al., 2022; Rehman & Fareed, 2025). Sundar e Kim (2019) argumentam que indivíduos com “heurísticas de máquina”, acreditando que os algoritmos são confiáveis e objetivos, tendem a perceber positivamente a sinalização tecnológica. Consequentemente, os efeitos negativos do *disclosure* sobre a credibilidade e as atitudes em relação à marca tendem a ser atenuados

entre consumidores que acreditam na capacidade da IA de realizar tarefas complexas, mas reforçados na ausência dessa percepção (Wu et al., 2025).

2.4.1 Familiaridade e a autenticidade percebida

No que concerne à autenticidade percebida, a familiaridade com IA tende a moderar os efeitos adversos do *disclosure* sobre esta propriedade emergente (Rehman & Fareed, 2025). Consumidores com baixa familiaridade tendem interpretar a sua sinalização sobretudo como evidência de que o conteúdo não foi criado por humanos, ativando o fenómeno de aversão ao algoritmo e percepções de artificialidade, menor empenho criativo e reduzida genuinidade (Ali et al., 2025). Entretanto, à medida que a familiaridade aumenta por meio da exposição contínua à tecnologia, a incerteza e a percepção de novidade diminuem, permitindo que os consumidores compreendam melhor o que a IA é capaz ou incapaz de realizar (Rehman & Fareed, 2025; Topsakal, 2025). Neste contexto, a IA pode ser legitimada como coautora criativa, através da transparência dos processos e posicionada como sofisticação tecnológica (Baek et al., 2024; Rehman & Fareed, 2025; Wu et al., 2025). Logo, propõe-se:

H8a: Quanto maior a familiaridade com a IA, menor são os efeitos negativos do *disclosure* na autenticidade percebida.

2.4.2 Familiaridade e a transparência percebida

O *disclosure* de IA generativa pode funcionar como sinalizador de transparência, mas o seu impacto positivo depende da familiaridade do consumidor com a tecnologia (Baek et al., 2024; Topsakal, 2025). Entre os consumidores com baixa familiaridade, a rotulação tende ativar o “conhecimento persuasivo” e a “aversão a IA”, tornando-os mais céticos quanto a tecnologia, além de salientar a opacidade estrutural dos sistemas mais do que a transparência operacional (Baek et al., 2024). Embora reconheçam que um sistema algorítmico está em uso, a baixa literacia tecnológica dificulta a compreensão do seu funcionamento, elevando a incerteza e a desconfiança, podendo comprometer a efetividade e credibilidade do anúncio (Baek et al., 2024; Wu et al., 2025). Por outro lado, entre consumidores com maior familiaridade, construída pela exposição contínua à estes sistemas, os efeitos negativos do *disclosure* tendem a ser atenuados, pois estes compreendem melhor a tecnologia e interpretam a sinalização como comunicação honesta sobre os processos e regras de utilização (Rehman & Fareed, 2025; Topsakal, 2025). Nestes casos, o *disclosure* pode ser percecionado como alinhamento regulatório e

transparência comunicacional, reforçando positivamente e calibrando a confiança na marca (Rehman & Fareed, 2025; Topsakal, 2025). Estabelece-se então a hipótese:

H8b: Quanto maior a familiaridade com a IA, menor são os efeitos negativos do *disclosure* na transparência percebida.

2.4.3 Familiaridade e a competência tecnológica

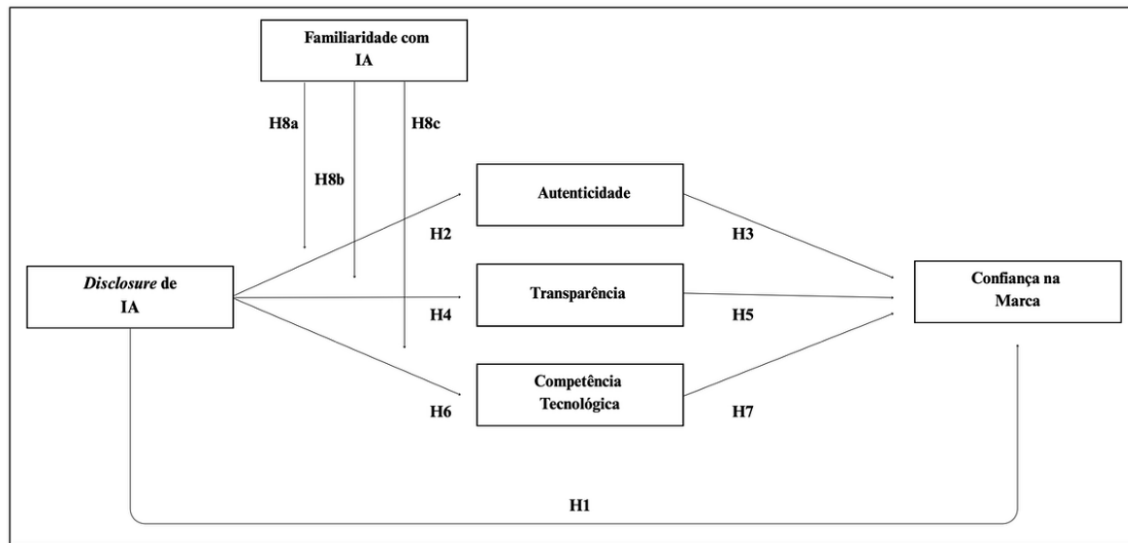
A familiaridade com a IA também influencia a forma como o *disclosure* é interpretado em termos de competência tecnológica, isto é, se será percebido como sinal de sofisticação ou como indícios de risco e incompetência da marca (Glikson & Woolley, 2020). Glikson e Woolley (2020) argumentam que indivíduos com pouco contacto e experiência com a IA tendem a não possuir modelos mentais claros sobre o seu funcionamento, percebendo-a como opaca ou pouco confiável. Neste caso, o *disclosure* pode aumentar a incerteza e reduzir as percepções de competência tecnológica, posicionando a marca como experimental ou imprudente, em vez de inovadora. Em contraste, o uso repetido de ferramentas de IA aumenta as percepções utilitárias e normaliza a IA como decisor eficiente (Topsakal, 2025). Assim, entre consumidores mais familiares ou com “heurísticas de máquina”, a IA é percebida como objetiva, precisa e segura. Neste segmento, o *disclosure* tende a reforçar a competência tecnológica da marca ao sinalizar investimentos e capacidade de adoção de tecnologias inovadoras, sobretudo quando os consumidores acreditam que a IA generativa pode desempenhar tarefas complexas (Sundar & Kim, 2019; Wu et al., 2025). Com base nisto, fundamenta-se a seguinte hipótese:

H8c: Quanto maior a familiaridade com a IA, menor são os efeitos negativos do *disclosure* na competência tecnológica.

3. MODELO CONCEPTUAL

A concepção do modelo conceptual do estudo (Figura 1) deu-se a partir da adaptação da *Assemblage Theory* de DeLanda (2006), aplicada ao contexto de mercado por Canniford e Bajde (2015). Ademais, as dimensões do sistema da confiança na marca foram baseadas nas investigações realizadas por Khamitov et al. (2024) e Rajavi et al. (2019). A estrutura do modelo conceptual, em conjunto com as *propriedades emergentes* e dimensões, foram adaptadas e relacionadas conforme as ligações estabelecidas destrinchadas anteriormente.

Figura 1 - Modelo conceptual



Fonte: Adaptação própria com base em DeLanda (2006), Khamitov et al. (2024), Rajavi et al. (2019) e Canniford e Bajde (2015).

4. METODOLOGIA

4.1 Tipo de estudo

A presente dissertação assume um propósito explanatório, visando testar empiricamente as hipóteses derivadas a partir do enquadramento teórico previamente estabelecido. Neste sentido, adota-se, uma abordagem dedutiva, adequada à análise das relações causais entre o *disclosure* de IA generativa e a confiança da marca (Saunders et al., 2023). Para a recolha de dados, foi utilizado um mono-método quantitativo, operacionalizado através de um questionário online. O inquérito foi escolhido por permitir a mensuração padronizada das percepções dos indivíduos, bem como a comparação sistemática entre os grupos (Saunders et al., 2023). A partir deste, foi conduzido uma experiência do tipo *between-subjects* no qual os participantes foram separados aleatoriamente atribuídos a condições experimentais diferentes (Saunders et al., 2023). Especificamente, os indivíduos foram expostos a uma publicidade das marcas ASICS, Ferrero Rocher e Rexona, que poderiam ou não incluir o *disclosure* do uso da IA. De forma a aumentar o realismo experimental e a validade externa, foi desenvolvido um cenário hipotético que simula um contexto de exposição publicitária real, permitindo captar as percepções dos indivíduos de forma mais naturalista (Viglia et al., 2021). A experiência tem como objetivo analisar de que forma a presença do *disclosure* de IA influencia as percepções de autenticidade, transparência e competência tecnológica nas

diferentes marcas, e como estas impactam a confiança dos consumidores. Considerando as limitações temporais da investigação, a recolha de dados assume um carácter *cross-sectional*, sendo realizada num único momento (Saunders et al., 2023). Por fim, ao adotar uma investigação objetiva da realidade, e independente do investigador, este estudo enquadra-se na filosofia positivista (Saunders et al., 2023).

4.2 População e amostragem

A população-alvo deste estudo é composta por indivíduos com idade igual ou superior a 18 anos, de ambos os géneros, com diferentes perfis sociodemográficos, abrangendo participantes residentes em diferentes contextos geográficos e com experiências distintas relativamente às marcas utilizadas na experiência. Esta definição teve como objetivo captar uma maior diversidade de percepções sobre a confiança na marca em contextos de *disclosure* do uso IA generativa. A amostra utilizada é do tipo não-probabilística por conveniência. Esta opção justifica-se pela viabilidade temporal da investigação, permitindo uma recolha eficiente num curto período de tempo, nomeadamente através da distribuição do questionário em redes sociais (Saunders et al., 2023). Importa ainda pontuar que, por não refletir a população geral, esta amostra esta sujeita a possíveis vieses.

4.3 Recolha de dados

A definição do problema de investigação e a elaboração das hipóteses deste estudo foram efetuadas a partir da recolha de dados secundários, sobretudo por artigos científicos. Posteriormente, com objetivo de empiricamente responder às questões e objetivos desta investigação, a recolha de dados primários estabeleceu-se por meio de um inquérito, traduzido em um questionário autoadministrado e mediado pela internet (Saunders et al., 2023), através da plataforma *Qualtrics XM*. O questionário foi divulgado por meio das redes sociais digitais (*Instagram e Facebook*) da investigadora, bem como de familiares e amigos, entre os dias 1 e 26 de maio de 2026. Complementarmente, realizou-se uma recolha de dados presencial através de visitas a diversos polos universitários na região de Lisboa. Ao término deste período, obteve-se um total de 427 respostas, das quais 363 foram consideradas válidas. Ademais, em conformidade com todas às questões éticas, o objetivo do estudo foi claramente comunicado aos inqueridos, e o anonimato e confidencialidade dos dados foram garantidos em todo o processo. Por fim, utilizou-se a ferramenta *G*Power 3.1* para determinar o tamanho da mínimo da amostra necessária para esta experiência (Viglia et al., 2021). Verificou-se que o

necessário seriam 196 inquiridos para atingir poder estatístico suficiente (tamanho de efeito médio, $f = 0,25$; $\alpha = 0,05$; $poder = 0,80$) para uma experiência de 3x3, que totalizam 9 grupos distintos.

4.4 Questionário

A experiência foi conduzida por meio de um questionário estruturado que, através da randomização automática da plataforma, expôs cada participante a uma das nove condições experimentais de igual probabilidade (Anexo A). Este desenho 3x3 resultou do cruzamento de três marcas distintas com três níveis de *disclosure* para o mesmo anúncio: ausência de rótulo, IA como criadora e IA como sugestão. A utilização de três marcas diferentes foi adotada como um refinamento metodológico, visando garantir que as percepções sobre a IA não fossem restritas a um único cenário, mas sim verificadas em diferentes categorias de produtos e o tipo de tarefa realizada pelo algoritmo (hedônica x objetiva). Os rótulos aplicados são respectivamente: “Anúncio criado por IA” e “Anúncio sugerido por IA” sinalizando respectivamente o caráter de cada atividade. Essa distinção ocorre porque os consumidores tendem a confiar mais na IA em tarefas objetivas, como o *targeting* do anúncio, e menos em tarefas criativas e hedônicas, como a criação de anúncios (Chen et al., 2024; Glikson & Woolley, 2020).

As peças publicitárias foram desenvolvidas no *ChatGPT* com base em anúncios reais, utilizando o *input* “Crie um anúncio desta marca, com base neste anúncio pré-existente. Tem que parecer realista e mimicar o estilo do anúncio da marca e o seu logo”. A plataforma foi escolhida pela sua eficiência e fidelidade visual em relação aos originais do *Instagram*. Posteriormente, as imagens foram editadas no *Canva* para inserção dos rótulos de *disclosure* e dos perfis das marcas ASICS, Ferrero Rocher e Rexona, simulando a interface real da rede social para maximizar o realismo experimental. A seleção estratégica destas marcas pautou-se na sua elevada notoriedade (Kirkby et al., 2023), associada a uma neutralidade atitudinal, evitando líderes de mercado. Esta escolha visa mitigar vieses decorrentes de lealdades extremas ou opiniões polarizadas que poderiam influenciar a formação de atitude, assegurando que as respostas reflitam o impacto da IA e não a identificação prévia com a marca (Chaudhuri & Holbrook, 2001).

O questionário foi organizado em 12 secções para mensurar as variáveis do modelo conceptual. O fluxo inicia-se com o consentimento informado e um filtro de exclusão sobre o uso de redes sociais. A terceira secção contextualiza o cenário do anúncio, o qual é adaptado pelo sistema conforme a condição experimental. Após a

visualização, seguem-se as secções que medem a confiança na marca (4^a) e as suas propriedades emergentes: autenticidade, transparência e competência tecnológica (5^a a 7^a). As etapas finais verificam a visualização da manipulação (8^a), a atenção do participante (9^a), a familiaridade com a IA generativa (10^a) e as atitudes pré-existentes face à marca (11^a). O questionário encerra-se com a recolha de dados sociodemográficos na 12^a secção, totalizando 22 questões. O pré-teste foi administrado para 15 participantes que atendiam os requisitos de participação desta experiência. Após o feedback foram realizadas as devidas alterações e melhorias no questionário (Anexo B).

4.5 Escalas de medida

Para mensurar as variáveis deste estudo, utilizaram-se escalas previamente validadas na literatura, as quais foram traduzidas e adaptadas ao contexto desta investigação. Assim, a confiança na marca foi operacionalizada através de quatro itens adaptada do estudo de Chaudhuri e Holbrook (2001). Para a variável da autenticidade percebida, aplicaram-se 11 itens baseados na escala de Portal et al. (2019). Por sua vez, a transparência comunicacional, foi medida por meio de uma escala composta por 12 itens, desenvolvida por Schnackenberg et al. (2021) e adaptada à comunicação da marca. No que toca à competência tecnológica, esta variável foi aferida através de cinco itens adaptados do estudo de Portal et al. (2019). Por fim, a familiaridade com a IA generativa, que atua como a variável moderadora do modelo conceptual, foi mensurada a partir da escala de Topsakal (2025), composta por quatro itens. Todas as métricas foram avaliadas mediante uma escala do tipo Likert de sete pontos, variando entre 1 (Discordo Totalmente) e 7 (Concordo Totalmente). O conjunto completo das escalas encontra-se disponível para consulta no Anexo C.

4.6 Tratamento e análise preliminar dos dados

Após a recolha, os dados foram tratados no *IBM SPSS Statistics 30*. Das 427 respostas obtidas, 363 foram consideradas válidas após quatro etapas de limpeza, que resultaram na exclusão de 64 inquiridos por não cumprirem os critérios globais, deixarem o questionário incompleto ou falharem nos *attention/manipulation checks*. Atendendo ao desenho experimental fatorial 3×3 (3 marcas × 3 tipos de *disclosure*), a amostra distribuiu-se de forma equilibrada, com uma média de aproximadamente 40 indivíduos por condição (Anexo D). Este volume supera o limiar mínimo de 20 observações por célula, assegurando a viabilidade estatística do modelo (Sarstedt & Mooi, 2019). Para viabilizar as análises, foram operacionalizadas duas variáveis categóricas de agrupamento: a

variável "Tipo_Disclosure", que identifica o estímulo visual exposto (1 = Sem disclosure; 2 = Com disclosure de Criação; 3 = Com disclosure de Sugestão), e a variável "Tipo_Marca", que codifica os cenários experimentais (1 = ASICS; 2 = Ferrero Rocher; 3 = Rexona). Relativamente à distribuição geográfica, as respostas isoladas de residentes na Bélgica, Alemanha, Países Baixos e Tailândia (N=4) foram agregadas na categoria "Outros" devido à sua baixa representatividade. Por fim, a validação dos valores máximos e mínimos confirmou a consistência interna e a correta inserção de toda a matriz de dados.

5. ANÁLISE DE DADOS

5.1 Caracterização da amostra

A amostra final é constituída por 363 participantes, com predominância do género feminino (55,1%), seguido pelo masculino (44,4%) e por indivíduos não binários (N=2). Relativamente ao perfil etário, observa-se uma forte concentração de jovens, com destaque para a faixa dos 18 aos 24 anos (68,9%), seguida pelos grupos dos 25 aos 34 anos (16,5%; N=60), apenas com uma pequena parcela acima dos 35 aos (13,6%). A nível de habilitações literárias, a maioria possui Licenciatura (56,2%), sendo o Mestrado o segundo nível mais frequente (24%), seguido por Pós-graduação (11,3%) e 12º ano (8,5%). Em linha com o perfil etário, a situação profissional é liderada por estudantes (57,3%). Entre a população ativa, registam-se trabalhadores por conta de outrem (28,7%), trabalhadores-estudantes (9,9%) e trabalhadores por conta própria (3%), tendo 3 participantes omitido a resposta. Geograficamente, os respondentes residem maioritariamente em Portugal (59,2%) e no Brasil (36,6%). Devido às disparidades cambiais, o rendimento foi segmentado regionalmente: na subamostra europeia, 39,7% não auferem rendimentos e 11,30% recebem entre 1001€ e 1500€. Já na subamostra latina, 12,4% não possuem rendimentos, destacando-se a faixa entre R\$10,886 e R\$14,191 (8,5%).

No âmbito do desenho experimental (Anexo E), os indivíduos foram distribuídos de forma equilibrada pelas marcas testadas: ASICS (N=114), Ferrero Rocher (N=119) e Rexona (N=130). Todas as marcas registaram 100% de conhecimento prévio. Os índices de familiaridade foram igualmente elevados, fixando-se em 99,2% para a Ferrero Rocher, 98,2% para a ASICS e 96,9% para a Rexona. Quanto à experiência de consumo anterior, verificou-se compra prévia na Rexona (93,8%) e na Ferrero Rocher (89,1%), enquanto a ASICS apresentou uma divisão exata na sua subamostra (50%). Por fim, a opinião prévia

(escala de 1 a 7) revelou-se altamente positiva e homogênea, liderada pela Ferrero Rocher ($M=6,13$), seguida de perto pela Rexona ($M=6,12$) e pela ASICS ($M=5,86$), validando o uso de estímulos consolidados e com forte reputação.

5.2 Fiabilidade, consistência interna e criação de índices sintéticos

Realizou-se uma Análise de Componentes Principais (APC) de cada variável do estudo, seguida por uma análise do alpha de Cronbach de todos os itens as compunham (Anexo F). Posteriormente a esta confirmação, foram criados índices sintéticos globais, a partir do cálculo da média de seus respectivos itens, bem como dimensões sintéticas para as variáveis aplicáveis.

Tendo em vista que o coeficiente alpha Cronbach é extremamente sensível escalas com mais de 10 itens (Pallant, 2020), a Análise de Componentes Principais foi utilizada para identificar a correlação interna das variáveis, bem como agrupar em fatores menores permitindo uma análise minuciosa (Sarstedt & Mooi, 2019). Deste modo, após a verificação do critério *Kaiser-Meyer-Olkin* (KMO), todas as variáveis apresentaram um valor superior 0,5, constatando a adequação dos dados (Sarstedt & Mooi, 2019). Adiante, acerca do Teste da Esfericidade de Barlett, todas as escalas obtiveram um $p<0,05$, indicando a correlação significativa dos itens (Pallant, 2020). No que toca as escalas “Autenticidade Percebida” e “Transparência Percebida”, ambas indicaram um Eigenvalue <1 , identificando 4 e 3 fatores, respectivamente. Para além da análise da fiabilidade interna dos índices globais criados (escalas multi-itens), foi realizado também o cálculo do alpha de Cronbach para estas dimensões identificadas. Este varia entre 0 e 1, em que apenas valores acima de 0,7 são significativos e devem ser utilizados (Pallant, 2020). Assim, verificou-se que os alphas dos índices globais se encontravam entre 0,708 e 0,890, apresentando uma elevada consistência interna, e para as dimensões sintéticas variaram entre 0,777 e 0,885. Embora o Fator 4 da Autenticidade tenha apresentado um valor de alpha de Cronbach inferior ao limiar recomendado ($\alpha=0,558$), optou-se por não construir uma dimensão sintética autónoma para este fator. Contudo, os respectivos itens não foram excluídos da análise, dado que contribuíam para o índice global de Autenticidade, cuja consistência interna se revelou adequada, com alpha superior a 0,7.

5.3 Análise descritiva dos índices

Através da análise das estatísticas descritivas associadas aos índices globais, observa-se que a Confiança na Marca destaca-se com o valor média mais elevado da

amostra ($M=5,58$), seguindo-se o construto da Competência Tecnológica ($M=5,35$) e, com valores muito próximos, a variável Autenticidade ($M=5,21$). Contudo, o índice que reuniu a média menos expressiva foi a Familiaridade com a IA ($M=4,36$), antecedido pela Transparência ($M=5,17$). Importa referir, que todos os índices se posicionam acima do ponto médio da escala de Likert ($M= 4$), evidenciando uma tendência geral de percepção positiva dos inquiridos face aos estímulos avaliados. Adicionalmente verifica-se que a Familiaridade com a IA gerou maior heterogeneidade nas respostas ($DP=1,179$), contrastando com o forte consenso da Autenticidade ($DP=0,611$) (Anexo G).

5.4 Teste de hipóteses

5.4.1 Efeito do *disclosure* de IA na confiança da marca

A Hipótese 1 postulava que o *disclosure* do uso de IA generativa em publicidades teria um impacto negativo na confiança. Para este teste foi realizada uma ANOVA bifatorial, tendo a confiança na marca como variável dependente e o tipo de *disclosure* e o tipo de marca como fatores independentes. Embora os testes de normalidade tenham indicado desvios significativos à normalidade nas três condições ($p<0,001$), o teste de Levene não foi significativo ($p= 0,186$) (Anexo H), confirmando a homogeneidade da variância e robustez do teste paramétrico devido a amostra ser >30 (Pallant, 2020).

As médias indicaram níveis de confiança ligeiramente superiores na condição “sem *disclosure* de IA” ($M = 5,66$) em comparação com as condições “com *disclosure* de IA criação” ($M = 5,53$) e “com *disclosure* de IA sugestão” ($M = 5,57$). Contudo, esta diferença não foi estatisticamente significativa ($p = 0,322$). A interação entre tipo de *disclosure* e marca também não foi significativa ($p = 0,268$), indicando que o efeito do *disclosure* não variou significativamente entre as marcas. Verificou-se, porém, um efeito principal significativo da marca sobre a confiança ($F(2, 354) = 7,773$, $p < 0,001$, $\eta^2 = 0,042$). As comparações *Post-Hoc* com correção de Bonferroni mostraram que a ASICS ($M=5,38$) apresentou níveis de confiança significativamente inferiores aos da Ferrero Rocher ($M= 5,65$; $p = 0,007$) e da Rexona ($M=5,7$; $p < 0,001$), não havendo diferenças significativas entre Ferrero Rocher e Rexona ($p= 1,000$). Deste modo, a H1 não é suportada, visto que o tipo de *disclosure* não teve impacto estatístico significativo direto na confiança na marca. Contudo, de maneira independente, verificou-se que a confiança varia em função da marca avaliada. Os resultados podem ser consultados no Anexo I.

5.4.2 Efeito do *disclosure* na autenticidade percebida

A Hipótese 2 (H2) postulava que o *disclosure* de IA exerceria um impacto negativo na autenticidade percebida da marca. Ao nível macro, realizou-se uma ANOVA bifatorial sobre a Autenticidade Global. O pressuposto de normalidade foi violado ($p < 0,001$), mas a robustez do teste foi assegurada pelo tamanho da amostra (Pallant, 2020), tendo o Teste de Levene validado a homogeneidade das variâncias ($p = 0,772$) (Anexo J). Os resultados indicaram que nem o efeito principal do *disclosure* ($p = 0,102$), nem o efeito da marca ($p = 0,339$) ou a sua interação ($p = 0,723$) foram significativos, não suportando a H2 a nível macro (Anexo K).

A fim de aprofundar a análise das subdimensões Benevolência, Continuidade e Originalidade, a MANOVA foi realizada. Dada a violação da homogeneidade das matrizes de covariância (Teste de Box, $p < 0,001$), utilizou-se a estatística do Traço de Pillai devido à sua robustez (Sarstedt & Mooi, 2019), enquanto o Teste de Levene validou a homogeneidade individual de todas as subdimensões ($p < 0,477$). Ao nível micro, a MANOVA revelou um efeito principal do *disclosure* significativo (Traço de Pillai = 0,039; $F(6,706) = 2,342$; $p = 0,030$; $\eta^2 = 0,02$). O Teste de Efeitos Entre-Sujeitos demonstrou neutralidade para a Benevolência ($p = 0,446$) e Continuidade ($p = 0,254$), mas um impacto significativo na Originalidade ($F(2,354) = 5,860$; $p = 0,003$; $\eta^2 = 0,032$). O teste *Post-Hoc* de Bonferroni para a Originalidade indicou que a condição sem *disclosure* superou significativamente ($M = 5,31$) a condição IA de Criação ($M = 4,88$; $p = 0,002$), não existindo diferenças significativas face à IA de Sugestão ($p = 0,086$) ou entre os dois tipos de aviso ($p = 0,635$). Consequentemente, a H2 obteve suporte parcial, confirmando-se que o aviso de IA de Criação prejudica especificamente o pilar da Originalidade, reduzindo a perceção de originalidade da marca. Ademais, verificou-se um efeito significativo do tipo de marca ($p < 0,001$), impactando significativamente apenas a dimensão da Originalidade ($p = 0,004$). O teste de Bonferroni para Originalidade demonstrou que a marca Ferrero Rocher foi considerada significativamente mais original ($M = 5,3$) em relação a ASICS ($M = 4,89$; $p = 0,004$). A Rexona não apresentou diferenças significativas em relação ASICS ($p = 0,675$) ou a Ferrero Rocher ($p = 0,099$) (Anexo K).

5.4.3 O impacto da autenticidade percebida na confiança da marca

A Hipótese 3 postulava que a Autenticidade Percebida teria um impacto positivo na confiança da marca. Para este teste, realizou-se primeiramente à um nível macro, uma Regressão Linear Múltipla, incluindo os preditores globais da “Autenticidade”,

“Transparência” e “Competência tecnológica” e “Confiança da Marca” como variável dependente. A matriz de correlações (Pearson) indicou associações positivas significantes entre a confiança e todas as variáveis independentes ($p < 0,001$). O modelo de regressão foi estatisticamente significativo ($p < 0,001$) explicando 23,8% (R^2) a variância da confiança na marca (Tabela 1), com níveis de colinearidade aceitáveis com o VIF entre 1,224 e 1,940 (Anexo L). Relativamente à H3, os resultados indicaram um efeito positivo e estatisticamente significativo da Autenticidade ($B = 0,374$; $\beta = 0,326$; $t = 5,157$; $p < 0,001$). Assim controlando as demais variáveis, a autenticidade revelou-se o preditor mais forte da confiança na marca. Deste modo, a H3 foi suportada à nível macro.

Na segunda etapa, foi realizada novamente uma regressão linear múltipla, incluindo as subdimensões da Autenticidade (e.g., Benevolência, Continuidade e Originalidade) e as da Transparência (e.g., Disclosure, Clareza e Exatidão), mantendo as demais variáveis. A correlação de Pearson demonstrou correlação de todas as variáveis com a confiança da marca ($p < 0,001$), com destaque para a dimensão da Continuidade ($r=0,454$), Competência Tecnológica ($r=0,417$) e Originalidade ($r=0,319$) como preditores de maior magnitude. O modelo micro foi significativo ($p < 0,001$) explicando 26,1% da variância da confiança e não indicando problemas de multicolinearidade (VIF). Entre as dimensões da autenticidade, apenas a Continuidade apresentou um efeito positivo sobre a confiança na marca ($B = 0,294$; $\beta = 0,288$; $t = 4,761$; $p < 0,001$), já a Benevolência ($p=0,561$) e Originalidade ($p=0,177$), não apresentaram efeitos significativos. Assim a H3 é suportada parcialmente a nível micro, sugerindo que a confiança é explicada sobretudo pela peceção de continuidade da marca (Anexo L).

Tabela 1 - Modelo de Regressão Linear Múltipla Macro para Confiança

VD	N	R	R^2	R^2 Ajustado	F	df	p	Durbin-Watson
Confiança	363	0,488	0,238	9,231	37,345	3,359	<0,001	1,749

Fonte: Elaboração Própria

5.4.4 Efeito do disclosure na transparência percebida

A Hipótese 4 estipulava que o *disclosure* de IA exerceria um impacto negativo na transparência percebida. Ao nível macro, realizou-se uma ANOVA a dois fatores sobre a Transparência Global. Embora a normalidade tenha sido violada ($p < 0,001$), o tamanho da amostra assegurou a robustez da análise ($N=363$), tendo o Teste de Levene baseado na mediana validado a homogeneidade das variâncias ($p=0,063$) (Anexo M). As estatísticas descritivas indicaram que a condição com *disclosure* IA Criação apresentou a média mais

elevada de transparência percebida ($M= 5,52$), em comparação com a condição sem *disclosure* de IA ($M= 5,02$) e com a condição com *disclosure* IA Sugestão ($M= 4,95$). Os resultados da ANOVA revelaram um efeito principal significativo do tipo de *disclosure* ($F(2,354) = 22,788$; $p < 0,001$; $\eta^2 = 0,114$). No entanto, nem o efeito da marca ($p = 0,521$) nem a interação entre marca e *disclosure* ($p = 0,643$), foram significativos. O teste *Post-Hoc* de Bonferroni rejeitou a premissa de um impacto negativo: a condição IA Criação aumentou significativamente a transparência ($M=5,52$), face ao cenário sem *disclosure* ($M=5,02$; $p<0,001$), enquanto a condição IA Sugestão permaneceu estatisticamente neutra ($M=4,95$; $p=1,000$). Assim, embora o efeito do *disclosure* tenha sido significativo, ele ocorreu na direção contrária à prevista pela H4, a IA de Criação aumentou, e não reduziu, a transparência, não suportando a hipótese (Anexo N).

Para desconstruir este impacto ao nível micro, procedeu-se a uma MANOVA com as subdimensões Disclosure, Clareza e Exatidão (Anexo N). Dado o Teste de Box significativo ($M=91,86$; $p<0,001$), interpretou-se através do Traço de Pillai. Este indicou um efeito significativo do *disclosure* (Traço de Pillai = 0,304; $F(6,706) = 21,068$; $p < 0,001$; $\eta^2 = 0,152$), mas não do tipo de marca ($p = 0,693$), nem da interação entre eles ($p = 0,978$). As análises univariadas indicaram que o tipo de *disclosure* teve efeito significativo apenas sobre as dimensões “Disclosure” ($F(2,354) = 55,582$; $p < 0,001$; $\eta^2 = 0,239$), e “Clareza” ($F(2,354) = 9,049$; $p < 0,001$; $\eta^2 = 0,049$). Os contrastes *Post-Hoc* de Bonferroni mostraram que na dimensão “Disclosure” a condição IA Criação apresentou média significativa superior ($M=5,36$) à condição sem *disclosure* ($M= 4,17$; $p < 0,001$) e à condição com *disclosure* IA Sugestão ($M= 3,95$; $p < 0,001$). Na dimensão “Clareza” observou-se um padrão semelhante. A condição com *disclosure* IA Criação apresentou média significativamente superior ($M= 5,88$) à condição sem *disclosure* de IA ($M = 5,53$; $p < 0,001$) e à condição com *disclosure* IA Sugestão ($M= 5,58$; $p = 0,002$). A comparação entre a condição sem *disclosure* e a condição IA Sugestão não foi significativa ($p = 1,000$). Logo a análise micro indica que o *disclosure* de IA de Criação reforçou a perceção da comunicação da marca como explícita e clara, mas não alterou a perceção de exatidão da informação, por esta não apresentar resultados significantes ($p = 0,904$).

5.4.5 O impacto da transparência percebida na confiança da marca

A Hipótese 5 (H5), apontava que a Transparência percebida impactaria positivamente a confiança na marca. Como visto na H3, a análise macro desta hipótese foi realizada através de uma Regressão Linear Multipla, em que as variáveis globais

Autenticidade, Transparência e Competência Tecnológica são as variáveis independentes e a Confiança na Marca a variável dependente. Tal como apontado na Tabela 1, o modelo de regressão é estatisticamente significativo ($p < 0,001$) e explica cerca de 23,8% (R^2) a variância da confiança na marca. No que toca a H5, a Transparência Percebida global não apresentou um efeito estatisticamente significativo ($p = 0,263$). Embora a correlação bivariada entre a transparência e a confiança tenha sido positiva e significativa ($r = 0,249$), este efeito deixou de ser significativo quando controladas a autenticidade e a competência tecnológica. Assim, a H5 não foi suportada a nível macro. A nível micro, este padrão também é percebido. Para aprofundar a análise foi realizado uma segunda regressão múltipla para as dimensões de Autenticidade e Transparência. O modelo micro (Anexo L) foi estatisticamente significativo ($F(7,355) = 17,910$; $p < 0,001$) explicando 26,1% da variância da Confiança na Marca Global. Contudo, nenhuma das subdimensões da Transparência apresentou efeitos significativos sobre a confiança da marca: Disclosure ($p = 0,380$), Clareza ($p = 0,739$) e Exatidão ($p = 0,838$). Isto sugere que a Transparência percebida não é um motor da confiança, mas sim um dever da marca, mesmo demonstrando correlação entre as dimensões e a confiança.

5.4.6 O efeito do *disclosure* na competência tecnológica

A Hipótese 6 estipulava que o *disclosure* de IA exerceria um impacto negativo na competência tecnológica percebida. Para o seu teste foi realizada uma ANOVA bifatorial, considerando a Competência Tecnológica como variável dependente e o tipo de *disclosure* e o tipo de marca como fatores independentes (Anexo O e P). O pressuposto de normalidade foi violado ($p < 0,001$), mas a robustez do teste foi assegurada pelo tamanho da amostra (Pallant, 2020). As estatísticas descritivas indicaram médias muito próximas entre as três condições experimentais: sem disclosure de IA ($M = 5,35$; $DP = 0,621$), com disclosure IA Criação ($M = 5,29$; $DP = 0,692$) e com disclosure IA Sugestão ($M = 5,40$; $DP = 0,691$). O Teste de Levene não foi significativo ($p = 0,945$), confirmando a homogeneidade das variâncias. Os resultados da ANOVA indicaram que o efeito principal do tipo de *disclosure* não foi estatisticamente significativo ($p = 0,384$). Do mesmo modo, não se verificou efeito significativo da marca ($p = 0,255$), nem da interação entre marca e *disclosure* ($p = 0,612$). O modelo corrigido também não foi significativo ($p = 0,481$). Deste modo, A H6 não foi suportada. Os resultados indicam que o *disclosure* de IA não reduziu significativamente a competência tecnológica da marca.

5.4.5 O impacto da competência tecnológica na confiança da marca

A Hipótese 7 propunha que a Competência Tecnológica da marca exerceria uma influência direta e positiva sobre os seus níveis de Confiança. O Teste desta hipótese foi realizado através da Regressão Linear Múltipla, isolando o seu peso preditivo diante das variáveis Autenticidade e Transparência (Anexo L). A análise da correlação de Pearson indicou uma associação positiva, moderada e estatisticamente muito significativa entre a Competência Tecnológica e a Confiança ($r=0,417$; $p<0,001$). Os pressupostos de colinearidade ($VIF=1,940$) e a estatística de Durbin-Watson ($DW=1,749$) foram validados. O modelo de regressão foi estatisticamente significativo ($p>0,001$) (Tabela 1). A Competência Tecnológica apresentou um impacto positivo e estatisticamente válido sobre a variável confiança ($B = 0,181$; $\beta = 0,174$; $t = 2,709$; $p = 0,007$), indicando que mesmo controlando as demais variáveis, a competência tecnológica contribui significativamente para explicar a confiança na marca. Portanto, a H7 foi suportada.

5.4.8 Efeito moderador da familiaridade com a IA nos efeitos do *disclosure*

Para avaliar a Familiaridade com a IA como uma variável moderadora, utilizou-se o PROCESS macro Modelo 1 de Hayes e Little (2018). Para mitigar a multicolinearidade e garantir robustez estatística, o *software* aplicou automaticamente a centralização das médias (*mean-centering*) às variáveis contínuas, além do método de *bootstrapping* com 5000 reamostragens (IC 95%).

5.4.8.1 Efeito moderador da familiaridade com a IA e a autenticidade

A Hipótese 8a (H8a) previa que a familiaridade com IA atenuaria os efeitos negativos do *disclosure* na Autenticidade Percebida. Para testá-la, aplicou-se o Modelo 1 do PROCESS, definindo o tipo de *disclosure* como variável independente, e a condição “Sem *disclosure*” como referência, a Familiaridade com IA como moderadora e o tipo de marca como covariável. Globalmente, o modelo não foi significativo para a Autenticidade ($p= 0,193$), assim como o termo de interação ($p= 0,341$), não revelando moderação para as condições IA Criação ($p= 0,937$) ou IA Sugestão ($p= 0,208$). Também, não emergiram diferenças significativas na percepção de autenticidade entre as marcas ($p>0,194$).

No nível micro, o padrão manteve-se, a interação entre *disclosure* e familiaridade não foi significativa para as dimensões de Benevolência ($p= 0,102$) e Continuidade ($p= 0,442$). Já para a Originalidade, embora o modelo tenha sido significativo ($R^2= 0,065$; $F(7,355)= 3,528$; $p= 0,001$), com efeitos diretos negativos da IA Criação ($B= -0,427$; $p=$

0,001) e da IA Sugestão ($B = -0,256$; $p = 0,046$), a interação com a familiaridade permaneceu não significativa ($p = 0,402$). Deste modo, a H8a não foi suportada, evidenciando que a familiaridade com IA não atenua os efeitos do *disclosure* sobre a autenticidade, quer ao nível global ou em suas subdimensões. Contudo, os resultados micro reforçam que o rótulo de IA reduz especificamente a percepção de originalidade, independentemente do nível de familiaridade dos consumidores. Por fim, apenas na dimensão da Originalidade a marca Ferrero Rocher ($p = 0,03$), apresentou efeito significativo positivo em relação a Rexona (i.e., referência), enquanto a ASICS não diferiu desta última ($p = 0,244$; Anexo Q).

5.4.8.2 Efeito moderador da familiaridade com a IA e a transparência

A Hipótese 8b (H8b) previa que a Familiaridade com IA moderaria o efeito do *disclosure* na Transparência Percebida. Utilizando o mesmo arranjo metodológico da H8a no Modelo 1 do PROCESS, o modelo macro para a Transparência Global foi significativo ($R^2 = 0,176$; $F(7,355) = 10,855$; $p < 0,001$). Verificou-se um efeito direto positivo da IA Criação ($B = 0,537$; $p < 0,001$), e uma efeito interação significativa entre *disclosure* e familiaridade ($\Delta R^2 = 0,031$; $F(2,355) = 6,631$; $p = 0,002$). A análise dos efeitos condicionais indicou que o impacto positivo da IA Criação aumentou com a familiaridade com IA, variando de $B = 0,374$ ($p = 0,003$), em baixa familiaridade, para $B = 0,697$ ($p < 0,001$), em alta familiaridade. Em contrapartida, a condição IA Sugestão não apresentou efeitos significativos nos diferentes níveis do moderador.

No nível micro, a moderação concentrou-se na dimensão Disclosure, cujo modelo foi significativo ($R^2 = 0,322$; $F(7,355) = 24,085$; $p < 0,001$), bem como a sua interação ($\Delta R^2 = 0,055$; $F(2,355) = 14,415$; $p < 0,001$). Os efeitos condicionais mostraram que o efeito positivo da IA Criação aumentou progressivamente com a familiaridade, de $B = 0,800$ (i.e., baixa) para $B = 1,303$ (i.e., média) e $B = 1,705$ (i.e., alta), todos com $p < 0,001$. Para IA Sugestão, apesar da interação específica negativa ($B = -0,253$; $p = 0,034$), os efeitos condicionais não foram significativos. Nas dimensões Clareza ($p = 0,201$), e Exatidão ($p = 0,463$), não se verificou moderação significativa. Deste modo, a H8b obteve suporte parcial, visto que a familiaridade moderou o efeito do *disclosure* sobre a Transparência Global e, especificamente, sobre a dimensão Disclosure. Contudo, este efeito ocorreu pela amplificação do impacto positivo da condição IA Criação, e não pela atenuação de um efeito negativo. Por fim, os tipos de marca não geraram diferenças significativas nas percepções de transparência global ($p > 0,375$) ou dimensional ($p > 0,1$; Anexo R).

5.4.8.2 Efeito moderador da familiaridade com a IA e a competência tecnológica

A Hipótese 8c (H8c) previa que a Familiaridade com a IA moderaria a relação entre o tipo de *disclosure* e a Competência Tecnológica, atenuando os eventuais efeitos negativos do *disclosure*. Para testar esta hipótese, aplicou-se o Modelo 1 do PROCESS, tendo o Tipo de Disclosure como variável independente, a Familiaridade com IA como moderadora e o Tipo de Marca como covariável. Os resultados indicaram que o modelo global não foi estatisticamente significativo ($p = 0,588$), sugerindo baixa capacidade explicativa sobre a Competência Tecnológica Percebida. O teste da interação entre o tipo de *disclosure* e Familiaridade com IA também não foi significativo ($p = 0,924$), indicando que a familiaridade não moderou a relação entre *disclosure* e competência tecnológica. Adicionalmente, não foram observados efeitos diretos significativos das condições IA Criação ($p = 0,552$), ou IA Sugestão, ($p = 0,453$), em comparação com a condição sem *disclosure*. Assim, a H8c não foi suportada. Os resultados sugerem que a competência tecnológica percebida permaneceu estável entre as condições experimentais, não sendo significativamente influenciada pelo tipo de *disclosure* nem pelo nível de familiaridade dos consumidores com IA, ou pelo tipo de marca ($p > 0,161$), (Anexo S).

5.4.9 Integração do modelo

5.4.9.1 Mediação paralela

Para integrar o modelo conceptual, testou-se uma mediação paralela via Modelo 4 do PROCESS (Hayes & Little, 2018), definindo o Tipo de Disclosure como variável independente, a Confiança na Marca como variável dependente e a Autenticidade Percebida, a Transparência Percebida e a Competência Tecnológica como mediadores paralelos. A condição “Sem *disclosure*” foi utilizada como categoria de referência, sendo o tipo de marca incluído como covariável. A condição IA Criação apresentou um efeito negativo significativo sobre a Autenticidade Percebida ($B = -0,171$; $p = 0,033$), e um efeito positivo significativo sobre a Transparência Percebida ($B = 0,503$; $p < 0,001$), mas não afetou a Competência Tecnológica ($p = 0,504$). No modelo final, a Confiança na Marca foi explicada significativamente conjunto de preditores ($R^2 = 0,273$; $F(7,355) = 19,045$; $p < 0,001$). Entre os mediadores, apenas a Autenticidade Percebida ($B = 0,321$; $p < 0,001$), e a Competência Tecnológica ($B = 0,218$; $p = 0,001$), exerceram efeitos positivos significativos sobre a confiança, a Transparência Percebida não apresentou impacto ($p =$

0,235). Quanto as covariáveis, a ASICS registrou menor percepção de confiança em relação à Rexona (i.e., referência; $B = -0,281$; $p = 0,03$), ao passo que a Ferrero Rocher não diferiu da última ($p = 0,879$). Os efeitos diretos do *disclosure* sobre a confiança não foram significativos, tanto para IA Criação ($p = 0,271$), quanto para IA Sugestão ($p = 0,275$).

A análise de *Bootstrapping* (5000 subamostras) revelou que, dentre os caminhos indiretos, apenas o mediado pela Autenticidade foi significativo (Efeito = $-0,055$; IC 95% [$-0,120$; $-0,004$]), indicando que a condição IA Criação reduziu a confiança na marca indiretamente ao mitigar a autenticidade. Os efeitos indiretos via Transparência e Competência Tecnológica, bem como todos os caminhos associados à IA de Sugestão, não foram significativos. Assim, a mediação paralela oferece suporte parcial ao modelo conceptual, confirmando a Autenticidade Percebida como principal mecanismo explicativo da relação entre o *disclosure* de IA e a confiança na marca (Anexo T).

5.4.9.2 Moderação mediada

Para testar a integração final do modelo, aplicou-se o Modelo 7 do PROCESS (Hayes & Little, 2018), com o tipo de *disclosure* como variável independente e a condição “Sem *disclosure*” como referência, a Confiança na Marca como variável dependente, a Autenticidade Percebida, a Transparência Percebida e a Competência Tecnológica como mediadores paralelos. A Familiaridade com IA como moderadora dos caminhos entre o *disclosure* e os mediadores, definindo o tipo de marca como covariável. Os resultados mostraram que a Familiaridade com IA não moderou significativamente o efeito do *disclosure* sobre a Autenticidade ($p = 0,341$), nem sobre a Competência Tecnológica ($p = 0,924$). Em contrapartida, emergiu uma moderação significativa no caminho entre *disclosure* e Transparência ($\Delta R^2 = 0,031$; $F(2,355) = 6,631$; $p = 0,002$). O efeito positivo da condição IA Criação expandiu-se com a elevação da familiaridade, variando de $B = 0,374$ ($p = 0,003$), em baixa familiaridade, para $B = 0,697$ ($p < 0,001$), em alta familiaridade. A IA Sugestão não apresentou efeitos condicionais significativos.

No modelo final de predição da confiança na marca, a Autenticidade ($B = 0,321$; $p < 0,001$), e a Competência Tecnológica ($B = 0,218$; $p = 0,001$), apresentaram efeitos positivos significativos, enquanto a Transparência permaneceu não significativa ($p = 0,235$). Os efeitos diretos do *disclosure* sobre a confiança não foram significativos tanto para IA Criação, ($p = 0,271$) quanto para IA Sugestão ($p = 0,275$). A análise dos efeitos indiretos condicionais indicou que nenhum dos índices de mediação moderada foi significativo, com todos os intervalos de confiança incluindo zero. Assim, embora a

Familiaridade com IA tenha moderado o efeito da IA Criação sobre a Transparência, esse efeito não se converteu em impacto indireto sobre a Confiança na Marca, não se verificando evidências de moderação mediada no modelo integrado (Anexo U)

6. DISCUSSÃO E CONCLUSÃO

6.1 Discussão

O presente estudo teve como objetivo compreender de que modo o *disclosure* do uso de IA generativa em publicidades reconfigura a confiança dos consumidores na marca. Partindo da *Assemblage Theory* (DeLanda, 2006), a confiança na marca foi entendida como um sistema emergente composto pela articulação entre marca, consumidor e tecnologia (Canniford & Bajde, 2015), no qual, a autenticidade, a transparência e a competência tecnológica, atuam como propriedades emergentes que estabilizam ou fragilizam o sistema perante mudanças contextuais (i.e., presença do *disclosure* de IA) (Khamitov et al., 2024; Mayer et al., 1995; Shi & Jiang, 2026). Os resultados demonstram que o *disclosure* não possui um efeito direto isolado sobre a confiança (H1). Contudo, esta ausência de linearidade não anula seu impacto. Os resultados sugerem que o *disclosure* atua como uma força paradoxal que reconfigura o *assemblage* ao afetar as suas propriedades emergentes, alinhando-se à perspectiva de que os consumidores manifestam avaliações diversas face a sinalização (Shi & Jiang, 2026).

O principal achado reside na confirmação parcial desta lógica paradoxal. A condição do *disclosure* de IA de Criação (tarefa hedônica), aumentou significativamente a transparência percebida, sobretudo nas subdimensões associadas ao *disclosure* e à clareza (H4), corroborando que a sinalização pode elevar as percepções de responsabilidade ética da marca (Kirkby et al., 2023). Em contrapartida, esta mesma condição reduziu a subdimensão “Originalidade” da autenticidade (H2), convergindo com Ali et al. (2025), ao demonstrar que o rótulo de IA pode comprometer genuinidade percebida. Assim, ao tornar o uso da tecnologia visível, o *disclosure* reforça a dimensão normativa da confiança, mas pode fragilizar a sua dimensão expressiva, ligada à autoria e à criatividade (Ali et al., 2025; Glikson & Woolley, 2020).

Ademais, embora reforçada pelo *disclosure*, a transparência não se traduziu em maior confiança na marca (H5), contrariando a expectativas de Rehman e Fareed (2025). Em contraste, a autenticidade revelou-se o preditor mais forte da confiança (H3), seguida pela competência tecnológica (H7). Este cenário aproxima-se de Khamitov et al. (2024), ao sugerir que atributos de integridade e autenticidade podem exercer maior influência na

confiança do que as dimensões puramente funcionais. Assim, a transparência associada ao *disclosure* parece ser interpretada pelos consumidores como um dever ético da marca (To et al., 2025), mas revela-se insuficiente para sustentar a confiança se a autenticidade for mitigada. Notou-se ainda que as avaliações de originalidade e confiança variam conforme a marca testada, indicando que a experiência prévia do consumidor molda estas percepções (Chaudhuri & Holbrook, 2001).

A análise de mediação paralela reforça esta interpretação, a condição “IA de Criação” não reduziu diretamente a confiança, mas exerceu um efeito indireto negativo mediado pela perda de autenticidade. Isto sugere que o consumidor não penaliza automaticamente o uso de IA generativa, mas sim quando este degrada o caráter genuíno da marca. Já a competência tecnológica não variou entre as condições (H6), contrariando a premissa de que o rótulo expandiria o senso de inovação (Khamitov et al., 2024; Shi & Jiang, 2026). Mesmo estável, essa propriedade manteve-se como preditor positivo da confiança, evidenciando que o público valoriza marcas competentes, ainda que a sinalização não potencialize esta percepção.

Por fim, a familiaridade com a IA apresentou moderação limitada. Ela não atenuou os efeitos negativos do *disclosure* na autenticidade (H8a), nem na competência tecnológica (H8c), mas amplificou o efeito positivo da “IA Criação” sobre a transparência (H8b). Assim, consumidores familiarizados decodificam melhor o valor informativo do rótulo, avaliando-o como comunicação honesta, em consonância com Topsakal (2025) e Rehman e Fareed (2025). Todavia, como a transparência não se converteu em maior confiança (H5), esta maior clareza não foi suficiente para elevar a confiança da marca. Assim, à luz da *Assemblage Theory* (DeLanda, 2006), os resultados expõe a instabilidade inerente do *assemblage* da confiança. O *disclosure* de IA generativa, *territorializa* a dimensão normativa, ao reforçar a transparência, mas *desterritorializa* parcialmente a dimensão expressiva, ao reduzir a originalidade. Como a competência tecnológica atua como força material estável, conclui-se que o sistema da confiança é reconfigurado não por impactos diretos do *disclosure*, mas pelas tensões internas que ele introduz entre as propriedades emergentes da marca.

6.2 Conclusão

A presente investigação buscou compreender de que forma o *disclosure* do uso de IA generativa em publicidades reconfigura a confiança dos consumidores na marca. Os resultados mostram que o rótulo de IA não afeta diretamente a confiança, mas atua de

forma indireta e paradoxal sobre as propriedades que a caracterizam. A lente ontológica da *Assemblage Theory* (DeLanda, 2006), permite compreender estas reconfigurações simultâneas e paradoxais do sistema da confiança. Neste estudo, o *disclosure* da IA como criadora do anúncio *territorializou* a dimensão normativa ao reforçar a transparência, mas *desterritorializou* parcialmente a dimensão expressiva, ao fragilizar a autenticidade, sobretudo por meio da originalidade. Assim, na presença do *disclosure* da IA generativa da publicidade, a confiança não é afetada diretamente pelo rótulo, mas indiretamente pelo seu impacto sobre a autenticidade. Conclui-se assim, que o *assemblage* da confiança parece-se estabilizar sobretudo quando a marca é percebida como autêntica e tecnologicamente competente, em que embora a transparência seja relevante enquanto dimensão normativa entre marca e consumidor, não se revelou como motor direto da confiança, sendo antes interpretada como dever ético.

6.3 Contributos académicos

O presente estudo oferece uma nova abordagem à literatura da confiança ao aplicar a *Assemblage Theory* como lente ontológica, no contexto da IA generativa em publicidades. Ao conceptualizar a confiança como um sistema relacional emergente, resultante da interação ente marca, consumidor e tecnologia, a investigação ultrapassa a lacuna presente na literatura, que analisa os efeitos do *disclosure* de maneira unilateral e linear, ao demonstrar que a autenticidade, a transparência e a competência tecnológica podem ser simultaneamente estabilizadas ou fragilizadas. Neste sentido, este estudo contribui para a literatura ao evidenciar empiricamente o caráter paradoxal do *disclosure* de IA, na condição da IA como criadora do anúncio reforçou a transparência percebida, mas fragilizou a originalidade, uma dimensão da autenticidade percebida. Ademais, a investigação aprofunda a compreensão dos mecanismos que medeiam o *disclosure* de IA e a confiança. Embora a sinalização não tenha produzido um efeito direto na confiança, verificou-se um efeito indireto negativo por meio da autenticidade, reforçando a centralidade deste construto como mecanismo explicativo da confiança. O estudo também acrescenta nuance à literatura ao distinguir as tarefas realizadas pelo IA no *disclosure*, demonstrando que quando a inteligência artificial desempenha tarefas hedônicas (i.e., IA de Criação), os impactos sobre as propriedades emergentes foram mais significativos do que quando desempenha tarefas objetivas (i.e., IA de Sugestão). E por fim, clarificar que a familiaridade com a IA, por mais que tenha amplificado a transparência, não protegeu a autenticidade, nem se traduziu em maior confiança.

6.4 Contributos práticos

O presente estudo oferece importantes *insights* para a gestão, ao sugerir que a presença do *disclosure* de IA pode ser percebida pelos consumidores como uma prática ética e transparente por parte da marca. Contudo, quando esse *disclosure* não está alinhado com o seu posicionamento, pode reduzir a percepção de originalidade e, indiretamente, enfraquecer a confiança. Desta forma, ao realizarem o *disclosure*, as marcas devem salientar o papel da IA como ferramenta de apoio à equipa de marketing, capaz de amplificar a criatividade humana, e não como substituta da identidade criativa da marca (Baek et al., 2024; Rehman & Fareed, 2025; Wu et al., 2025). Uma vez que a autenticidade constitui um preditor central da confiança na marca, comunicar a centralidade do papel humano no processo criativo torna-se fundamental para que, perante o *disclosure*, os consumidores interpretem o uso da IA como sinal de sofisticação tecnológica, sem comprometer os ativos de confiança de longo prazo (Ali et al., 2025; Cillo & Rubera, 2025). Para tal, é necessário demonstrar o compromisso da marca com a inovação, comunicando os seus investimentos para as demais áreas da gestão, de modo que os consumidores percecionem a competência tecnológica da marca como um todo e reforcem a sua confiança nela (Khamitov et al., 2024).

6.5 Limitações do estudo e sugestões de investigação futura

Como limitações, destaca-se o uso de uma amostra não probabilística de conveniência e maioritariamente jovem, o que limita a generalização dos resultados do estudo. Além disto, a experiência recorreu a três marcas amplamente conhecidas pela amostra, pelo que opiniões prévias, familiaridade e experiências anteriores podem ter influenciado as respostas aos estímulos. Adicionalmente, o presente estudo investigou os efeitos do *disclosure* em três propriedades emergentes da confiança, não contemplando outras dimensões que também podem influenciar a confiança na marca. No que toca à investigação futura, estudos posteriores poderão testar diferentes categorias de produtos, marcas dentro de uma mesma categoria, marcas menos conhecidas ou até mesmo fictícias, de modo a compreender melhor as nuances dos efeitos do *disclosure* conforme estas variações. Também seria relevante, aplicar esta abordagem ontológica para investigar os efeitos do *disclosure* em outros comportamentos do consumidor, como a intenção de compra ou e-WOM, ou lealdade à marca. Por fim, recomenda-se recorrer a métodos qualitativos ou a experimentos longitudinais para compreender mais profundamente

como os consumidores interpretam o *disclosure* e se a exposição repetida a estes rótulos altera, ao longo do tempo, as percepções de autenticidade, transparência e confiança.

7. REFERÊNCIAS

- Ali, L., Ali, F., Abdalla, M. J., & Alotaibi, S. (2025). Beyond the hype: Evaluating the impact of generative AI on brand authenticity, image, and consumer behavior in the restaurant industry. *International Journal of Hospitality Management*, 131, 104318. <https://doi.org/10.1016/j.ijhm.2025.104318>
- Ammanath, B., Dutt, D., Perricos, C., & Sniderman, B. (2024, janeiro 23). *The state of generative AI in the enterprise*. Deloitte. <https://www.deloitte.com/ce/en/services/consulting/research/state-of-generative-ai-in-enterprise.html>
- Baek, T. H., Kim, J., & Kim, J. H. (2024). Effect of disclosing AI-generated content on prosocial advertising evaluation. *International Journal of Advertising*. <https://doi.org/10.1080/02650487.2024.2401319>
- Bloomberg. (2023, janeiro 6). *Generative AI to become a \$1.3 trillion market by 2032*. Bloomberg. <https://www.bloomberg.com/company/press/generative-ai-to-become-a-1-3-trillion-market-by-2032-research-finds/>
- Campbell, C., Plangger, K., Sands, S., & Kietzmann, J. (2022). Preparing for an era of deepfakes and AI-generated ads: A framework for understanding responses to manipulated advertising. *Journal of Advertising*, 51(1), 22–38. <https://doi.org/10.1080/00913367.2021.1909515>

- Canniford, R., & Bajde, D. (2015). *Assembling consumption: 1st Edition* (p. 1–18). Routledge.
https://www.researchgate.net/publication/305228092_ASSEMBLING_CONSUMPTION_INTRODUCTION#citations
- Chaudhuri, A., & Holbrook, M. B. (2001). The chain of effects from brand trust and brand affect to brand performance: The role of brand loyalty. *Journal of Marketing*, 65(2), 81–93. <https://doi.org/10.1509/jmkg.65.2.81.18255>
- Chen, Y., Wang, H., Rao Hill, S., & Li, B. (2024). Consumer attitudes toward AI-generated ads: Appeal types, self-efficacy and AI's social role. *Journal of Business Research*, 185, 114867. <https://doi.org/10.1016/j.jbusres.2024.114867>
- Cillo, P., & Rubera, G. (2025). Generative AI in innovation and marketing processes: A roadmap of research opportunities. *Journal of the Academy of Marketing Science*, 53(3), 684–701. <https://doi.org/10.1007/s11747-024-01044-7>
- Dang, J., & Liu, L. (2024). Extended artificial intelligence aversion: People deny humanness to artificial intelligence users. *Journal of Personality and Social Psychology*. <https://doi.org/10.1037/pspi0000480>
- DeLanda, M. (2006). *A new philosophy of society: Assemblage theory and social complexity*. Bloomsbury Academic. <https://doi.org/10.5040/9781350096769>
- Duffek, B., Eisingerich, A. B., Merlo, O., & Lee, G. (2025). Authenticity in influencer marketing: How can influencers and brands work together to build and maintain influencer authenticity? *Journal of Marketing*, 89(5), 21–46.
<https://doi.org/10.1177/00222429251319786>
- Feuerriegel, S., Hartmann, J., Janiesch, C., & Zschech, P. (2024). Generative AI. *Business & Information Systems Engineering*, 66(1), 111–126.
<https://doi.org/10.1007/s12599-023-00834-7>

- Glikson, E., & Woolley, A. W. (2020). Human trust in artificial intelligence: Review of empirical research. *Academy of Management Annals*, 14(2), 627–660.
<https://doi.org/10.5465/annals.2018.0057>
- Grewal, D., Saturnino, C. B., Davenport, T., & Guha, A. (2025). How generative AI Is shaping the future of marketing. *Journal of the Academy of Marketing Science*, 53(3), 702–722. <https://doi.org/10.1007/s11747-024-01064-3>
- Hartmann, J., Exner, Y., & Domdey, S. (2025). The power of generative marketing: Can generative AI create superhuman visual marketing content? *International Journal of Research in Marketing*, 42(1), 13–31.
<https://doi.org/10.1016/j.ijresmar.2024.09.002>
- Hayes, A. F., & Little, T. D. (2018). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach* (Second edition). The Guilford Press.
- Khamitov, M., Rajavi, K., Huang, D.-W., & Hong, Y. (2024). Consumer trust: Meta-analysis of 50 years of empirical research. *Journal of Consumer Research*, 51(1), 7–18. <https://doi.org/10.1093/jcr/ucad065>
- Kirkby, A., Baumgarth, C., & Henseler, J. (2023). To disclose or not disclose, is no longer the question – effect of AI-disclosed brand voice on brand authenticity and attitude. *Journal of Product & Brand Management*, 32(7), 1108–1122.
<https://doi.org/10.1108/JPBM-02-2022-3864>
- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *The Academy of Management Review*, 20(3), 709.
<https://doi.org/10.2307/258792>
- Mittelstadt, B. (2019). Principles alone cannot guarantee ethical AI. *Nature Machine Intelligence*, 1(11), 501–507. <https://doi.org/10.1038/s42256-019-0114-4>

- Pallant, J. (2020). *SPSS survival manual: A step by step guide to data analysis using IBM SPSS* (7º ed.). Routledge. <https://doi.org/10.4324/9781003117452>
- Portal, S., Abratt, R., & Bendixen, M. (2019). The role of brand authenticity in developing brand trust. *Journal of Strategic Marketing*, 27(8), 714–729. <https://doi.org/10.1080/0965254X.2018.1466828>
- Rajavi, K., Kushwaha, T., & Steenkamp, J.-B. E. M. (2019). In brands we trust? A multicategory, multicountry investigation of sensitivity of consumers' trust in brands to marketing-mix activities. *Journal of Consumer Research*, 46(4), 651–670. <https://doi.org/10.1093/jcr/ucz026>
- Rehman, S. U., & Fareed, M. (2025). Enhancing consumer trust in AI-created goods: The roles of transparency, simplicity, familiarity and authenticity in market integration. *Contemporary Economics*, 429. <https://doi.org/10.5709/ce.1897-9254.576>
- Sarstedt, M., & Mooi, E. (2019). *A concise guide to market research: The process, data, and methods using IBM SPSS statistics*. Springer Berlin Heidelberg. <https://doi.org/10.1007/978-3-662-56707-4>
- Saunders, M., Lewis, P., & Thornhill, A. (2023). *Research methods for business students* (Ninth edition). Pearson.
- Schnackenberg, A. K., Tomlinson, E., & Coen, C. (2021). The dimensional structure of transparency: A construct validation of transparency as disclosure, clarity, and accuracy in organizations. *Human Relations*, 74(10), 1628–1660. <https://doi.org/10.1177/0018726720933317>
- Shi, Y., & Jiang, Z. (2026). Consumer responses to AI disclosure labels: The role of novelty and authenticity. *Sage Open*, 16(1), 21582440261417793. <https://doi.org/10.1177/21582440261417793>

Statista. (2025, abril 15). *Artificial intelligence (AI) use in marketing*. Statista.

https://www.statista.com/topics/5017/ai-use-in-marketing/?srsltid=AfmBOopRTJ9IxL7LFfMzXPTt9Q5KIY-7_5NwloMnrLSVr51fa5FyUrYQ

Sundar, S. S., & Kim, J. (2019). Machine heuristic: When we trust computers more than humans with our personal information. *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*, 1–9.

<https://doi.org/10.1145/3290605.3300768>

The One Club For Creativity. (2023). *A.I. Ketchup*. The One Club For Creativity.

To, R. N., Wu, Y.-C., Kianian, P., & Zhang, Z. (2025). When AI doesn't sell prada: Why using AI-generated advertisements backfires for luxury brands. *Journal of Advertising Research*, 65(2), 202–236.

<https://doi.org/10.1080/00218499.2025.2454120>

Tokita, S. (2024). Content creation by generative AI and operator perception a focus on operator's profit-driven motivation. *2024 Portland International Conference on Management of Engineering and Technology (PICMET)*, 1–8.

<https://doi.org/10.23919/PICMET64035.2024.10653035>

Topsakal, Y. (2025). How familiarity, ease of use, usefulness, and trust influence the acceptance of generative artificial intelligence (AI)-assisted travel planning. *International Journal of Human–Computer Interaction*, 41(15), 9478–9491.

<https://doi.org/10.1080/10447318.2024.2426044>

Viglia, G., Zaefarian, G., & Ulqinaku, A. (2021). How to design good experiments in marketing: Types, examples, and methods. *Industrial Marketing Management*, 98, 193–206. <https://doi.org/10.1016/j.indmarman.2021.08.007>

- Wu, L., Dodoo, N. A., & Wen, T. J. (2025). Disclosing AI's involvement in advertising to consumers: A task-dependent perspective. *Journal of Advertising*, 54(1), 20–38. <https://doi.org/10.1080/00913367.2024.2309929>
- Zhu, Y., Zhang, J., Wu, J., & Liu, Y. (2022). AI is better when I'm sure: The influence of certainty of needs on consumers' acceptance of AI chatbots. *Journal of Business Research*, 150, 642–652. <https://doi.org/10.1016/j.jbusres.2022.06.044>

8. ANEXOS

Anexo A - Anúncios presentes nos questionários

 Move your body, move your mind asics 12 de novembro de 2025 ASICS (sem <i>disclosure</i>)	 Move your body, move your mind asics 12 de novembro de 2025 *Anúncio criado por IA Generativa ASICS (com <i>disclosure</i> de criação)	 Move your body, move your mind asics 12 de novembro de 2025 *Anúncio sugerido por IA Generativa ASICS (com <i>disclosure</i> de sugestão)
 ferrerorocherpt Áudio original TABLETES FERRERO ROCHER. DESFRUTE O SEU MOMENTO. Ferrero Rocher (sem <i>disclosure</i>)	 ferrerorocherpt Áudio original *Anúncio criado por IA Generativa Ferrero Rocher (com <i>disclosure</i> de criação)	 ferrerorocherpt Áudio original *Anúncio sugerido por IA Generativa Ferrero Rocher (com <i>disclosure</i> de sugestão)



Anexo B – Questionário online



Bem vindo(a)! Sou aluna do Mestrado em Marketing do ISEG – Lisbon School of Economics & Management. Este questionário foi desenvolvido no âmbito da minha dissertação de mestrado e gostaria de contar com seu contributo. Estou a desenvolver um estudo sobre as **percepções dos consumidores sobre as comunicações das marcas**.

A sua participação será completamente **confidencial e anónima**, em que os dados recolhidos serão de uso exclusivo para fins académicos.

O preenchimento deste questionário tem uma média 8 minutos. Não existem respostas certas ou erradas, apelo apenas que as responda com sinceridade.

O seu contributo é muito importante para a conclusão da minha dissertação.

Muito obrigada!
Giullia Scarpa



Aceita participar deste estudo?

☐ Sim

☐ Não

Tem mais de 18 anos?

☐ Sim

☐ Não

A Influência Do Uso De Inteligência Artificial Generativa Em Publicidades Na Confiança Da Marca

Costuma ver publicidades de marcas em seu *feed*?

☐ Sim

☐ Não

Costuma utilizar redes sociais (Instagram, Facebook, TikTok...)?

☐ Sim

☐ Não

"Recentemente, você percebeu que precisa de um novo tênis para exercitar-se. Durante seu almoço abriu sua rede social de preferência e enquanto rolava seu feed, você se deparou com o anúncio da ASICS™ sobre o lançamento do novo tênis. Você se interessou e examinou cuidadosamente a imagem e o texto, começando a formar a suas impressões."

Observe atentamente o anúncio abaixo. Em seguida responda as perguntas sobre a marca e a comunicação apresentada:



* Anúncio sugerido por IA Generativa

Confirma que visualizou a imagem anterior?

☐ Sim

☐ Não

Eu conheço esta marca.

☐ Sim

☐ Não

Eu já comprei um produto desta marca.

☐ Sim

☐ Não

Eu tenho familiaridade com esta marca.

☐ Sim

☐ Não

Antes de ver este anúncio, qual era sua opinião geral sobre a marca?

☐ Muito desfavorável

☐ Desfavorável

☐ Parcialmente desfavorável

☐ Nem Favorável nem desfavorável

☐ Parcialmente favorável

☐ Favorável

☐ Muito favorável

Por favor, indique o seu nível de concordância com cada uma das seguintes afirmações em relação a marca apresentada.

	Discordo Totalmente	Discordo	Discordo Parcialmente	Nem discordo nem concordo	Concordo Parcialmente	Concordo	Concordo Totalmente
1. Eu confio nesta marca.	☐	☐	☐	☐	☐	☐	☐
2. Esta marca é confiável	☐	☐	☐	☐	☐	☐	☐
3. Esta marca é honesta	☐	☐	☐	☐	☐	☐	☐
4. Esta marca é segura	☐	☐	☐	☐	☐	☐	☐

Havia um rótulo "Anúncio criado por IA Generativa" abaixo do anúncio.

☐ Sim

☐ Não

Estava claro para mim como o anúncio foi criado.

☐ Sim

☐ Não

Por favor, indique o seu nível de concordância com cada uma das seguintes afirmações.

Responder apenas com base no anúncio visualizado anteriormente.

	Discordo Totalmente	Discordo	Discordo Parcialmente	Nem discordo nem concordo	Concordo Parcialmente	Concordo	Conc Total
1. A marca tem história	☐	☐	☐	☐	☐	☐	☐
2. A marca é consistente ao longo do tempo	☐	☐	☐	☐	☐	☐	☐
3. A marca não permanece verdadeira	☐	☐	☐	☐	☐	☐	☐
4. A marca oferece continuidade	☐	☐	☐	☐	☐	☐	☐
5. A marca possui um propósito claro que segue	☐	☐	☐	☐	☐	☐	☐
6. A marca não retribui aos seus consumidores	☐	☐	☐	☐	☐	☐	☐
7. A marca não tem princípios morais	☐	☐	☐	☐	☐	☐	☐
8. A marca segue os seus princípios	☐	☐	☐	☐	☐	☐	☐
9. A marca se importa com seus consumidores	☐	☐	☐	☐	☐	☐	☐
10. A marca distingue-se claramente de outras	☐	☐	☐	☐	☐	☐	☐

11. Eu acho a marca única	☐	☐	☐	☐	☐	☐	☐
12. A marca não se destaca em relação às outras marcas	☐	☐	☐	☐	☐	☐	☐
13. A marca não vai me trair	☐	☐	☐	☐	☐	☐	☐
14. A marca é desonesta	☐	☐	☐	☐	☐	☐	☐
15. Minha experiência com a marca provou que ela não mantém as suas promessas	☐	☐	☐	☐	☐	☐	☐
16. As promessas da marca são creíveis	☐	☐	☐	☐	☐	☐	☐

Responder apenas com base no anúncio visualizado anteriormente.

	Discordo Totalmente	Discordo	Discordo Parcialmente	Nem discordo nem concordo	Concordo Parcialmente	Concordo	Conc Totalm
1. A informação que recebi deste anúncio abrange os aspectos que eu gostaria de conhecer	☐	☐	☐	☐	☐	☐	☐
2. A informação que recebi deste anúncio abrange os principais tópicos que eu gostaria de conhecer	☐	☐	☐	☐	☐	☐	☐
3. Este anúncio fornece a informação de que preciso sobre como ele foi criado	☐	☐	☐	☐	☐	☐	☐
4. Uma quantidade suficiente de informação foi apresentada neste anúncio	☐	☐	☐	☐	☐	☐	☐
5. A informação apresentada neste anúncio é fácil de compreender	☐	☐	☐	☐	☐	☐	☐
6. A informação deste anúncio é clara	☐	☐	☐	☐	☐	☐	☐
7. A informação apresentada neste anúncio é compreensível	☐	☐	☐	☐	☐	☐	☐

8. A linguagem utilizada neste anúncio é fácil de entender	☐	☐	☐	☐	☐	☐	☐
9. A informação apresentada neste anúncio parece ser verdadeira	☐	☐	☐	☐	☐	☐	☐
10. A informação apresentada neste anúncio parece correta	☐	☐	☐	☐	☐	☐	☐
11. A informação apresentada neste anúncio parece ser precisa	☐	☐	☐	☐	☐	☐	☐
12. A informação apresentada neste anúncio parece ser adequada	☐	☐	☐	☐	☐	☐	☐

Por favor, indique o seu nível de concordância com cada uma das seguintes afirmações.

Responder apenas com base no anúncio visualizado anteriormente.

	Discordo Totalmente	Discordo	Discordo Parcialmente	Nem discordo nem concordo	Concordo Parcialmente	Concordo	Concor Totalme
1. A marca é muito competente no seu setor	☐	☐	☐	☐	☐	☐	☐
2. A marca não é capaz de implementar as suas intenções	☐	☐	☐	☐	☐	☐	☒
3. A marca não é capaz e eficaz em empregar novas tecnologias para atingir seus objetivos	☐	☐	☐	☐	☐	☐	☐
4. A marca é eficaz	☐	☐	☐	☐	☐	☐	☐
5. A marca não é eficiente	☐	☐	☐	☐	☐	☐	☐

Para confirmar que você está prestando atenção, por favor clicar em **"Concordo"** para esta afirmação.

☐ Concordo

☐ Discordo

Qual marca estava presente no anúncio?

☐ Adidas

☐ ASICS

☐ Puma

☐ Nike

A Influência Do Uso De Inteligência Artificial Generativa Em Publicidades Na Confiança Da Marca

Por favor, indique o seu nível de concordância com cada uma das seguintes afirmações.

Responder apenas com base no anúncio visualizado anteriormente.

	Discordo Totalmente	Discordo	Discordo Parcialmente	Nem discordo nem concordo	Concordo Parcialmente	Concordo	Conc Totalm
1.Eu conheço muito sobre IA generativa	☐	☐	☐	☐	☐	☐	☐
2.Eu tenho familiaridade com IA generativa	☐	☐	☐	☐	☐	☐	☐
3.Eu tenho muito conhecimento sobre como a IA generativa opera	☐	☐	☐	☐	☐	☐	☐
4.Eu sei quando a IA generativa é utilizada	☐	☐	☐	☐	☐	☐	☐

Indique o seu sexo, por favor:

☐ Feminino

☐ Masculino

☐ Não binário

☐ Prefiro não dizer

Indique sua idade, por favor:

☐ 18 – 24 anos

☐ 25 – 34 anos

☐ 35 – 44 anos

☐ 45 – 54 anos

☐ 55 – 64 anos

☐ 65 anos ou mais

Indique a sua situação profissional, por favor:

☐ Estudante

☐ Trabalhador estudante

☐ Trabalhador por conta de outrem

☐ Trabalhador por conta própria

☐ Desempregado

☐ Reformado (aposentado)

Indique país reside atualmente, por favor:

☐ Brasil

☐ Chile

☐ Canadá

☐ Espanha

☐ Estados Unidos

☐ França

☐ Portugal

☐ Outro

Indique o seu rendimento individual mensal líquido, por favor:

☐ Não possuo renda

☐ Até R\$ 1.580

☐ R\$ 1.581 – R\$ 2.525

☐ R\$ 2.526– R\$ 10.885

☐ R\$ 10.886– R\$ 14.191

☐ Acima de R\$ 14.192

☐ Prefiro não responder

Anexo C - Escalas de medida

Tabela - Escalas de Medida			
Variável	Autor	Itens Originais	Itens Adaptados
Confiança na Marca	Chaudhuri e Holbrook (2001)	• CM1: I trust this brand • CM2: I rely on this brand • CM3: This is an honest brand • CM4: This brand is safe	• CM1: Eu confio nesta marca • CM2: Esta marca é confiável • CM3: Esta marca é honesta • CM4: Esta marca é segura
Autenticidade Percebida	Portal et al. (2019)	Continuity • APC1: A brand with history • APC2: The brand is consistent over time • APC3: The brand does not stay true to itself • APC4: The brand offers continuity • APC5: The brand has a clear concept that it pursues	Continuidade • APC1: A marca tem história • APC2: A marca é consistente ao longo do tempo • APC3: A marca não permanece verdadeira consigo mesma • APC4: A marca oferece continuidade • APC5: A marca possui um propósito claro que segue
		Integrity • API1: The brand does not give back to its consumers • API2: The brand has no moral principles • API3: The brand is true to a set of moral values • API4: The brand cares about its consumers	Integridade • API1: A marca não retribui aos seus consumidores • API2: A marca não tem princípios morais • API3: A marca segue os seus princípios • API4: A marca se importa com seus consumidores
		Originality • AP01: The brand clearly distinguishes itself from other brands • AP02: I think the brand is unique • AP03: The brand does not stand out from other brands	Originalidade • AP01: A marca distingue-se claramente de outras marcas • AP02: Eu acho a marca única • AP03: A marca não se destaca em relação às outras marcas
		Credibility • APC1: The brand will not betray me • APC2: The brand is dishonest • APC3: My experience with the brand has shown me that it does not keep its promises • APC4: The brand's promises are credible	Credibilidade • APC1: A marca não vai me trair • APC2: A marca é desonesta • APC3: Minha experiência com a marca provou que ela não mantém as suas promessas • APC4: As promessas da marca são credíveis
Transparência Percebida	Schnackenberg et al. (2020)	Disclosure • TPD1: The information I receive from (X) fully encompasses what I want to know about • TPD2: The information I receive from (X) covers all the topics I want to know about • TPD3: I have all the information I need from (X) • TPD4: A sufficient amount of information is presented by (X)	Disclosure • TPD1: A informação que recebi deste anúncio abrange os aspectos que eu gostaria de conhecer • TPD2: A informação que recebi deste anúncio cobre os principais tópicos que eu gostaria de conhecer • TPD3: Este anúncio fornece a informação de que preciso sobre como ele foi criado • TPD4: Uma quantidade suficiente de informação foi apresentada neste anúncio
		Clarity • TPC1: The information presented by (X) is understandable • TPC2: The information from (X) is clear • TPC3: The information from (X) is comprehensible • TPC4: The information from (X) is presented in a language I understand	Clareza • TPC1: A informação apresentada neste anúncio é fácil de compreender • TPC2: A informação deste anúncio é clara • TPC3: A informação apresentada neste anúncio é compreensível • TPC4: A linguagem utilizada neste anúncio é fácil de entender
		Accuracy • TPA1: The information from (X) appears to be true • TPA2: The information from (X) appears correct • TPA3: The information from (X) appears accurate • TPA4: The information from (X) appears right	Exatidão • TPA1: A informação apresentada neste anúncio parece ser verdadeira • TPA2: A informação apresentada neste anúncio parece correta • TPA3: A informação apresentada neste anúncio parece ser precisa • TPA4: A informação apresentada neste anúncio parece ser adequada

Anexo D - Amostra por subgrupo

		Disclosure			Total
		Sem disclosure	Com disclosure	Com disclosure	
			IA de Criação	IA de Sugestão	
Marca	ASICS	36	41	37	114
	Ferrero	35	40	44	119
	Rocher				
	Rexona	36	47	47	130
Total		107	128	128	363

Fonte: Elaboração própria

Anexo E - Perfil sociodemográfico da amostra

Tabela: Descrição sociodemográfica da amostra

Indicador	Opções de Resposta	%	N
Sexo (N= 363)	Feminino	55,10%	200
	Masculino	44,40%	161
	Não binário	0,6%	2
Idade (N= 363)	18-24 anos	68,90%	250
	25-34 anos	16,50%	60
	35-44 anos	9,60%	35
	45-54 anos	3,90%	14
	55-64 anos	0,8%	3
	65 anos ou mais	0,3%	1
Habilitações literárias (N= 363)	Inferior ao 12º ano	0	0
	12º ano completo	8,50%	31
	Licenciatura/ Bacharelado	56,20%	204
	Pós-graduação	11,30%	41
	Mestrado	24%	87
	Doutorado	0	0
Situação Profissional (N= 363)	Estudante	57,30%	208
	Trabalhador estudante	9,90%	36
	Trabalhador por conta de outrem	28,70%	104
	Trabalhador por conta própria	3%	11
	Reformado	0,3%	1
	Missing	0,8%	3
País de residência (N= 363)	Brasil	36,60%	133
	Chile	0,03%	1
	Espanha	1,10%	4
	Estados Unidos	1,10%	4
	França	0,06%	2
	Portugal	59,20%	215
	Outro	1,10%	4

Rendimento mensal líquido (EU e EUA) (N= 229)	Não aufero rendimentos	39,70%	144
	Menos de 500€	3,90%	14
	500€ - 1000€	4,10%	15
	1001€ - 1500€	11,30%	41
	1501€ - 2000€	1,70%	6
	2001€ - 2500€	1,40%	5
	Mais de 2500€	0,3%	1
	Prefiro não responder	0,8%	3
Rendimento mensal líquido (Brasil e Chile) (N= 134)	Não possuo renda	12,40%	45
	Até R\$1,580	0,03%	1
	R\$1581 - R\$2.525	2,20%	8
	R\$2.526 - R\$10.885	6,90%	25
	R\$10.886 - R\$14.191	8,50%	31
	Acima de R\$ 14.192	5,50%	20
	Prefiro não responder	1,10%	4

Fonte: Elaboração própria

Tabela - Conhecimento de Marca				
Pergunta	Marca	Opções de resposta	%	N
Q5 “Conhece a marca presente no anúncio?” (N= 363)	ASICS	Sim	100%	114
	Ferrero Rocher	Sim	100%	119
	Rexona	Sim	100%	130
Q6 “Tem familiaridade com a marca presente no anúncio?” (N= 363)	ASICS	Sim	98,20%	112
		Não	1,80%	2
	Ferrero Rocher	Sim	99,20%	118
		Não	0,8%	1
	Rexona	Sim	96,90%	126
		Não	3,10%	4
Q7 “Já comprou algum item da marca presente no anúncio?” (N= 363)	ASICS	Sim	50%	57
		Não	49,10%	56
		Missing	0,9%	1
	Ferrero Rocher	Sim	89,10%	106
		Não	10,90%	13
	Rexona	Sim	93,80%	122
Q8 “Antes de ver este anúncio, qual era a sua opinião sobre a marca?” (N= 363) Média/ Std. Deviation	ASICS	1= Muito Desfavorável - 7= Muito Favorável	100%	114
	Ferrero Rocher	1= Muito Desfavorável - 7= Muito Favorável	100%	119
	Rexona	1= Muito Desfavorável - 7= Muito Favorável	100%	130

Fonte: Elaboração própria

Anexo F - PCA e análise de fiabilidade

PCA e análise de fiabilidade interna dos índices globais

			ANÁLISE DE COMPONENTES PRINCIPAIS					ANÁLISE DE FIABILIDADE		
ÍNDICE	ITEM	N	KMO	Teste de esfericidade de Bartlett		Variância Total Explicada (% da Variância)	Comunalidades	Alfa de Cronbach	Estatísticas de Item Global	
				Aprox. X²	Sig.				Correlação de item total corrigida	Alfa de Cronbach se o item for excluído
Confiança na Marca	[CM1]	363	0,816	764,124	<0,001	72,818	0,794	0,865	0,733	0,824
	[CM2]						0,801		0,785	0,800
	[CM3]						0,594		0,623	0,880
	[CM4]						0,769		0,760	0,810
Autenticidade Percebida	[AP1]	363	0,891	2118,381	<0,001	61,852	0,546	0,871	0,491	0,865
	[AP2]						0,718		0,500	0,864
	[AP3]						0,498		0,412	0,870
	[AP4]						0,637		0,519	0,864
	[AP5]						0,637		0,573	0,863
	[AP6]						0,615		0,575	0,861
	[AP7]						0,577		0,585	0,861
	[AP8]						0,656		0,642	0,859
	[AP9]						0,655		0,590	0,861
	[AP10]						0,723		0,575	0,861
	[AP11]						0,750		0,527	0,863
	[AP12]						0,637		0,546	0,862
	[AP713]						0,548		0,272	0,877
	[AP714]						0,624		0,406	0,869
	[AP715]						0,499		0,526	0,863
	[AP16]						0,576		0,623	0,860
Transparência Percebida	[TP1]	363	0,848	2776,075	<0,001	71,606	0,794	0,871	0,334	0,877
	[TP2]						0,847		0,645	0,859
	[TP3]						0,595		0,653	0,057
	[TP4]						0,738		0,614	0,859
	[TP5]						0,716		0,663	0,854
	[TP6]						0,734		0,700	0,851
	[TP7]						0,800		0,415	0,883
	[TP8]						0,715		0,654	0,855
	[TP9]						0,488		0,638	0,858
	[TP10]						0,791		0,656	0,858
	[TP11]						0,763		0,684	0,859
	[TP12]						0,654		0,633	0,861
Competência Tecnológica	[CT1]	363	0,752	311,825	<0,001	46,909	0,44	0,708	0,441	0,675
	[CT2]						0,427		0,452	0,670
	[CT3]						0,476		0,488	0,651
	[CT4]						0,496		0,473	0,657
	[CT5]						0,508		0,506	0,642
Familiaridade IA	[FIA1]	363	0,836	831,99	<0,001	75,411	0,812	0,89	0,811	0,839
	[FIA2]						0,779		0,782	0,852
	[FIA3]						0,740		0,729	0,864
	[FIA4]						0,686		0,701	0,880

* Considerou-se um nível de significância de 5%

Fonte: Elaboração própria

PCA e análise de fiabilidade interna das dimensões da autenticidade e transparência percebida

			ANÁLISE DE FIABILIDADE				
Variável	Dimensão	Item	Eigenvalue >1	VARIMAX	Alfa de Cronbach	Estatísticas de Item por Dimensão	
						Correlação de item total corrigida	Alfa de Cronbach se o item for
Autenticidade Percebida	Benevolência	[AP6]	5,815	-0,635	0,805	0,609	0,764
		[AP7]		-0,575		0,585	0,769
		[AP8]		0,71		0,644	0,761
		[AP9]		0,762		0,655	0,755
		[AP13]		0,613		0,368	0,831
		[AP16]		0,632		0,603	0,767
	Continuidade	[AP1]	1,583	0,662	0,779	0,549	0,744
		[AP2]		0,784		0,647	0,692
		[AP4]		0,751		0,612	0,712
		[AP5]		0,645		0,535	0,751
	Originalidade	[AP10]	1,372	0,744	0,777	0,651	0,666
		[AP11]		0,820		0,654	0,652
		[AP12]		-0,636		0,545	0,773
	Credibilidade	[AP3]	1,127	0,666	0,558	0,359	0,494
		[AP714]		0,683		0,362	0,471
		[AP715]		0,574		0,402	0,412
Transparência Percebida	Disclosure	[TP1]	5,885	0,787	0,818	0,672	0,763
		[TP2]		0,847		0,754	0,723
		[TP3]		0,659		0,486	0,868
		[TP4]		0,798		0,722	0,734
	Clareza	[TP5]	1,554	0,789	0,885	0,757	0,852
		[TP6]		0,787		0,742	0,857
		[TP7]		0,826		0,812	0,832
		[TP8]		0,765		0,751	0,867
	Exatidão	[TP9]	1,153	0,664	0,789	0,405	0,883
		[TP10]		0,795		0,776	0,669
		[TP11]		0,741		0,720	0,681
		[TP12]		0,653		0,631	0,723

*Considerou-se um nível de significância de 5%

Fonte: Elaboração própria

Anexo G – Análise descritiva dos itens

Índice Globais	N	Mínimo	Máximo	Média	Desvio Padrão	Variância	Skewness	Kurtosis
							Statistics	Statistics
Confiança Marca	363	4	7	5,58	0,701	0,491	-0,208	-0,330
Autenticidade Percebida	363	1	7	5,21	0,611	0,373	-1,250	4,820
Transparência Percebida	363	2	7	5,17	0,777	0,604	-0,406	0,417
Competência Tecnológica	363	3	7	5,34	1,179	0,451	-0,794	1,866
Familiaridade com IA	363	1	7	4,63	0,671	1,390	-0,470	-0,570

Fonte: Elaboração própria

Anexo H – Teste da hipótese 1: pressupostos

Normalidade

Tipo de Disclosure	Kolmogorov-Smirnov			Shapiro-Wilk		
	Estatística	df1	Sig. (p)	Estatística	df2	Sig. (p)
Sem disclosure IA	0,192	107	< 0,001	0,927	107	< 0,001
Com disclosure IA Criação	0,149	128	< 0,001	0,951	128	< 0,001
Com disclosure IA Sugestão	0,145	128	< 0,001	0,957	128	< 0,001

Fonte: Elaboração própria

Levene

Critério de Cálculo	Estatística de Levene	df1	df2	Significância (p)
Baseado na Média	1,421	8	354	0,186
Baseado na Mediana	1,493	8	354	0,158
Baseado na Média Aparada	1,493	8	344,561	0,158
Baseado na Mediana e com df ajustado	1,439	8	354	0,179

Fonte: Elaboração própria

Anexo I - Teste da hipótese 1: two-way ANOVA

Estatísticas Descritivas da Confiança na Marca por Condição Experimental e Marca

Tipo de Marca	Tipo de Disclosure	N	Média (M)	Desvio-Padrão (DP)
ASICS	Sem disclosure IA	36	5,34	0,705
	Com disclosure IA Criação	41	5,49	0,685
	Com disclosure IA Sugestão	37	5,3	0,771
	Total	114	5,38	0,718
Ferrero Rocher	Sem disclosure IA	35	5,8	0,745
	Com disclosure IA Criação	40	5,52	0,784
	Com disclosure IA Sugestão	44	5,64	0,616
	Total	119	5,65	0,716
Rexona	Sem disclosure IA	36	5,83	0,57
	Com disclosure IA Criação	47	5,57	0,699
	Com disclosure IA Sugestão	47	5,71	0,608
	Total	130	5,7	0,636
Total Geral	Sem disclosure IA	107	5,66	0,708
	Com disclosure IA Criação	128	5,53	0,717
	Com disclosure IA Sugestão	128	5,57	0,678
	Total Global	363	5,58	0,701

Fonte: Elaboração própria

Two-Way Anova Confiança da Marca

Fonte de Variabilidade	Soma de Quadrados (SS)	Graus de Liberdade (df)	Quadrado Médio (MS)	F	Significância (p)	Tamanho do Efeito (η^2)
Modelo Corrigido	10,329	8	1,291	2,728	0,006	0,058
Tipo de Marca	7,357	2	3,679	7,773	< 0,001	0,042
Tipo de Disclosure	1,076	2	0,538	1,137	0,322	0,006
Tipo_Marca*Tipo_Disclosure	2,471	4	0,618	1,305	0,268	0,015
Erro	167,523	354	0,473			
Total Corrigido	177,853	362				

Fonte: Elaboração própria

Post-Hoc Bonferroni Tipo_Marca

Tipo de Marca	Tipo de Marca	Diferença de Médias (I-J)	Erro Padrão	Significância (p-value)
ASICS	Ferrero Rocher	-0,276*	0,09	0,007
	Rexona	-0,329*	0,089	< 0,001
Ferrero Rocher	ASICS	0,276*	0,09	0,007
	Rexona	-0,053	0,088	1
Rexona	ASICS	0,329*	0,089	< 0,001
	Ferrero Rocher	0,053	0,088	1

Fonte: Elaboração própria

Anexo J – Teste hipótese 2: pressupostos

Normalidade

Tipo de disclosure	n	Kolmogorov-Smirnov <i>D</i>	<i>p</i>	Shapiro-Wilk <i>W</i>	<i>p</i>
Sem disclosure IA	107	.159	< .001	.911	< .001
Com disclosure IA — Criação	128	.112	< .001	.900	< .001
Com disclosure IA — Sugestão	128	.177	< .001	.931	< .001

Variável: Autenticidade

Fonte: Elaboração própria

Levene

Critério	<i>F</i> de Levene	df1	df2	<i>p</i>
Baseado na média	0.607	8	354	.772
Baseado na mediana	0.688	8	354	.702
Baseado na mediana com df ajustados	0.688	8	308.335	.702
Baseado na média aparada	0.638	8	354	.745

Variável: Autenticidade

Fonte: Elaboração própria

Anexo K - Teste hipótese 2: two-way ANOVA

Nível Macro

Estatísticas Descritivas da Autenticidade por Condição e Tipo de Marca

Tipo de Marca	Tipo de Disclosure	<i>N</i>	<i>M</i>	<i>DP</i>
ASICS	Sem disclosure IA	36	5,19	0,455
	Com disclosure IA Criação	41	5,13	0,514
	Com disclosure IA Sugestão	37	5,12	0,685
	Total	114	5,15	0,555
Ferrero Rocher	Sem disclosure IA	35	5,39	0,606
	Com disclosure IA Criação	40	5,1	0,786
	Com disclosure IA Sugestão	44	5,22	0,567
	Total	119	5,23	0,664
Rexona	Sem disclosure IA	36	5,31	0,573
	Com disclosure IA Criação	47	5,14	0,613
	Com disclosure IA Sugestão	47	5,31	0,622
	Total	130	5,25	0,606
Total Geral	Sem disclosure IA	107	5,3	0,549
	Com disclosure IA Criação	128	5,13	0,64
	Com disclosure IA Sugestão	128	5,23	0,623
	Total geral	363	5,21	0,611

Fonte: Elaboração própria

Two-way ANOVA Autenticidade Percebida

Fonte	Soma dos Quadrados	df	MS	<i>F</i>	<i>p</i>	η^2p
Modelo corrigido	3.241	8	0.405	1.089	.370	.024
Intercepto	9.754.480	1	9.754.480	26.222.018	< .001	.987
Tipo de Marca	0.807	2	0.403	1.084	.339	.006
Tipo de Disclosure	1.712	2	0.856	2.301	.102	.013
Tipo_Marca * Tipo_Disclosure	0.770	4	0.192	0.517	.723	.006
Erro	131.687	354	0.372	—	—	—
Total corrigido	134.928	362	—	—	—	—

Fonte: Elaboração própria

Nível Micro

Levene

Variável dependente	Critério	F de Levene	df1	df2	p
Benevolência	Baseado na média	0.948	8	354	.477
Continuidade	Baseado na média	0.707	8	354	.685
Originalidade	Baseado na média	0.710	8	354	.682

Fonte: Elaboração própria

Teste M de Box para igualdade das matrizes de covariância

M de Box	F	df1	df2	p
86,442	1,745	48	169,176,810	.001

Fonte: Elaboração própria

Testes Multivariados MANOVA

Efeito	Estatística Multivariada	Valor	F	df	df do Erro	Significância (p)	Tamanho do Efeito (η^2)
Tipo de Marca	Traço de Pillai	0,083	5,066	6	706	< 0,001	0,041
	Lambda de Wilks	0,918	5,136	6	704	< 0,001	0,042
Tipo de Disclosure	Traço de Pillai	0,039	2,342	6	706	0,03	0,02
	Lambda de Wilks	0,961	2,35	6	704	0,03	0,02
Tipo_Marca*Tipo_Disclosure	Traço de Pillai	0,025	0,758	12	1062	0,695	0,008
	Lambda de Wilks	0,975	0,757	12	931,596	0,696	0,009

Fonte: Elaboração própria

Resultados dos Testes dos Efeitos Entre-Sujeitos (Univariados) para cada Subdimensão da Autenticidade Percebida

Fonte	Variável Dependente	Soma de Quadrados (SS)	df	Quadrado Médio (MS)	F	Significância (p)	η^2
Tipo de Disclosure	Benevolência	0,845	2	0,423	0,81	0,446	0,005
	Continuidade	1,29	2	0,645	1,376	0,254	0,008
	Originalidade	10,594	2	5,297	5,86	0,003	0,032
Tipo de Marca	Benevolência	1,774	2	0,887	1,70	0,184	0,01
	Continuidade	2,143	2	1,071	2,285	0,103	0,013
	Originalidade	9,942	2	4,971	5,50	0,004	0,03
Erro	Benevolência	184,69	354	0,522			
	Continuidade	165,945	354	0,469			
	Originalidade	319,963	354	0,904			

Fonte: Elaboração própria

Estimativas das médias por Tipo de Disclosure

Variável dependente	Marca	M	EP	IC 95% [LL, LS]
Originalidade	Sem disclosure IA	5,311	0,092	[5,130; 5,492]
	Com disclosure IA Criação	4,887	0,084	[4,722; 5,053]
	Com disclosure IA Sugestão	5,036	0,084	[4,870; 5,203]

Fonte: Elaboração Própria

Comparações Múltiplas Post-Hoc (Bonferroni) para o Efeito do Tipo de Disclosure nas Subdimensões da Autenticidade

Variável Dependente	(I) Tipo de Disclosure	(J) Tipo de Disclosure	Diferença de Médias (I-J)	Erro Padrão (EP)	Significância (p)	Intervalo de Confiança 95%
Originalidade	Sem disclosure IA	Com disclosure IA Criação	0,424*	0,125	0,002	[0,124; 0,724]
		Com disclosure IA Sugestão	0,274	0,125	0,086	[-0,026; 0,575]
	Com disclosure IA Criação	Com disclosure IA Sugestão	-0,149	0,119	0,635	[-0,436; 0,138]

Fonte: Elaboração própria

Estimativas das médias por Tipo de Marca

Variável dependente	Marca	M	EP	IC 95% [LI, LS]
Originalidade	ASICS	4,892	0,089	[4,717; 5,068]
	Ferrero Rocher	5,301	0,088	[5,129; 5,473]
	Rexona	5,041	0,084	[4,876; 5,207]

Fonte: Elaboração Própria

Comparações Múltiplas Post-Hoc (Bonferroni) para o Efeito do Tipo de marca nas Subdimensões da Autenticidade

Variável Dependente	(I) Tipo de Disclosure	(J) Tipo de Disclosure	Diferença de Médias (I-J)	Erro Padrão (EP)	Significância (p)	Intervalo de Confiança 95%
Originalidade	ASICS	Ferrero Rocher	-0,408	0,125	0,004	[-0,709; -0,108]
		Rexona	-0,149	0,123	0,675	[-0,444; 0,146]
	Ferrero Rocher	ASICS	0,408	0,125	0,004	[0,108; 0,709]
		Rexona	0,26	0,121	0,099	[-0,032; 0,551]
	Rexona	ASICS	0,149	0,123	0,675	[-0,146; 0,444]
		Ferrero Rocher	-0,26	0,121	0,099	[-0,551; 0,032]

Fonte: Elaboração própria

Anexo L- Regressão linear múltipla – H3, H5 e H7

Nível Macro

Matriz de Correlações de Pearson entre as Variáveis Globais (N = 363)

Variável	1	2	3	4
1. Confiança na Marca Global	—			
2. Competência Tecnológica Global	0,417**	—		
3. Transparência Percebida Global	0,249**	0,407**	—	
4. Autenticidade Percebida Global	0,464**	0,675**	0,371**	—

Nota. A correlação é significativa ao nível 0,01 (2 extremidades).

Fonte: Elaboração própria

Resumo do Modelo de Regressão Linear Múltipla e Diagnósticos de Independência para a Confiança na Marca

Modelo	Coefficiente de Correlação (R)	Coefficiente de Determinação (R2)	R2 Ajustado	Erro Padrão do Estimador	Estatística de Durbin-Watson
1	0,488	0,238	0,231	0,614	1,749

Nota: Preditores: Competência Tecnológica, Transparência, Autenticidade. Variável Dependente: Confiança na Marca

Fonte: Elaboração própria

Análise de Variância (ANOVA) de Significância Global do Modelo

Modelo	Fonte de Variação	Soma de Quadrados	Graus de Liberdade (gl)	Quadrado Médio	Estatística F	Significância (p)
1	Regressão	42,302	3	14,101	37,345	< 0,001
	Residual	135,551	359	0,378		
	Total	177,853	362			

Nota: Preditores: Competência Tecnológica, Transparência, Autenticidade. Variável Dependente: Confiança na Marca

Fonte: Elaboração própria

Coefficientes de Regressão Múltipla e Indicadores de Colinearidade para a Confiança na Marca

Variável Preditora	Coefficientes Não Padronizados B	Erro Padrão de B	Coefficientes Padronizados β	Estatística t	Significância (p)	Estatísticas de Colinearidade Tolerance	Estatísticas de Colinearidade VIF
(Constante)	2,396	0,311	—	7,702	< 0,001	—	—
Autenticidade	0,374	0,072	0,326	5,157	< 0,001	0,533	1,877
Transparência	0,052	0,046	0,057	1,122	0,263	0,817	1,224
Competência	0,181	0,067	0,174	2,709	0,007	0,516	1,94

Nota. Variável Dependente: Confiança na Marca

Fonte: Elaboração própria

Nível Micro

Matriz de Correlações de Pearson entre a Confiança na Marca e as Subdimensões dos Preditores (N = 363)

Variável	1	2	3	4	5	6	7	8
1. Confiança da Marca	—							
2. Competência Tecnológica	0,417**	—						
3. TP_Disclosure	0,161**	0,220**	—					
4. TP_Clareza	0,203**	0,379**	0,521**	—				
5. TP_Exatidão	0,276**	0,480**	0,423**	0,612**	—			
6. AU_Benevolência	0,306**	0,519**	0,151**	0,185**	0,382**	—		
7. AU_Continuidade	0,454**	0,575**	0,179**	0,286**	0,456**	0,461**	—	
8. AU_Originalidade	0,319**	0,465**	0,154**	0,130*	0,320**	0,485**	0,464**	—

Nota. *p < 0,05; **p < 0,01.

Fonte: Elaboração própria

Resumo do Modelo de Regressão Linear Múltipla Micro para as Subdimensões (N = 363)

Modelo	Coefficiente de Correlação (R)	Coefficiente de Determinação (R2)	R2 Ajustado	Erro Padrão do Estimador	Estatística de Durbin-Watson
1	0,501	0,251	0,236	0,612	1,1803

Nota: Preditores: Autenticidade (Originalidade, Benevolência e Continuidade); Transparência (Disclosure, Clareza e Exatidão); Competência Tecnológica. Variável Dependente: Confiança na Marca

Fonte: Elaboração própria

Análise de Variância (ANOVA) de Significância Global do Modelo Micro

Modelo	Fonte de Variação	Soma de Quadrados	Graus de Liberdade (df)	Quadrado Médio	Estatística F	Significância (p)
1	Regressão	46,681	7	6,383	17,016	< 0,001
	Residual	133,171	355	0,37		
	Total	177,853	362			

Nota: Preditores: Autenticidade (Originalidade, Benevolência e Continuidade); Transparência (Disclosure, Clareza e Exatidão); Competência Tecnológica. Variável Dependente: Confiança na Marca

Fonte: Elaboração própria

Coefficientes da Regressão Múltipla Micro e Indicadores de Colinearidade

Variável Preditora	Coefficiente B	Beta (β)	t	Sig. (p)	VIF
Competência Tecnológica	0,196	0,188	2,95	0,003	1,924
AU_Continuidade	0,294	0,288	4,761	< 0,001	1,732
AU_Benevolência	0,033	0,034	0,581	0,561	1,6
AU_Originalidade	0,055	0,076	1,352	0,177	1,517
TP_Disclosure	0,026	0,048	0,88	0,38	1,416
TP_Clareza	0,017	0,017	0,263	0,793	1,94
TP_Exatidão	-0,011	-0,013	-0,205	0,838	2,03

Nota. Variável Dependente: Confiança na Marca

Fonte: Elaboração própria

Anexo M – Teste hipótese 4: pressupostos

Normalidade						
Tipo de disclosure	Kolmogorov-Smirnov	df	p	Shapiro-Wilk	df	p
Sem disclosure IA	.095	107	.018	.930	107	< .001
Com disclosure IA Criação	.169	128	< .001	.895	128	< .001
Com disclosure IA Sugestão	.134	128	< .001	.936	128	< .001
Fonte: Elaboração própria						
Levene						
Crítério do teste	Estatística de Levene	df1	df2	p		
Baseado na média	2.318	8	354	.020		
Baseado na mediana	1.873	8	354	.063		
Baseado na mediana com gl ajustados	1.873	8	329.936	.063		
Baseado na média aparada	2.211	8	354	.026		
Fonte: Elaboração própria						

Anexo N - Teste hipótese 4: two-way ANOVA

Macro

Estatísticas Descritivas para a Transparência Global em Função do Tipo de Marca e do Tipo de Disclosure

Tipo de Marca	Tipo de Disclosure	Média (M)	Desvio-Padrão (DP)	N
ASICS	Sem disclosure IA	4,88	0,587	36
	Com disclosure IA Criação	5,59	0,642	41
	Com disclosure IA Sugestão	4,94	0,787	37
	Total	5,16	0,746	114
Ferrero Rocher	Sem disclosure IA	5,03	0,882	35
	Com disclosure IA Criação	5,42	0,586	40
	Com disclosure IA Sugestão	4,92	0,854	44
	Total	5,12	0,807	119
Rexona	Sem disclosure IA	5,15	0,706	36
	Com disclosure IA Criação	5,56	0,722	47
	Com disclosure IA Sugestão	4,97	0,782	47
	Total	5,23	0,778	130
Total Geral	Sem disclosure IA	5,02	0,735	107
	Com disclosure IA Criação	5,52	0,655	128
	Com disclosure IA Sugestão	4,95	0,803	128
	Total	5,17	0,777	363

Fonte: Elaboração própria

Resultados da ANOVA a Dois Fatores (Tests of Between-Subjects Effects) para a Transparência Global

Modelo Corrigido	26,978	8	3,372	6,232	< 0,001	0,123
Intercepto	9558,407	1	9558,407	17664,88	< 0,001	0,98
Tipo_Marca	0,708	2	0,354	0,654	0,521	0,004
Tipo_Disclosure	24,661	2	12,331	22,788	< 0,001	0,114
Tipo_Marca*Tipo_Disclosure	1,358	4	0,339	0,627	0,643	0,007
Erro	191,548	354	0,541			
Total	9926,701	363				
Total Corrigido	218,526	362				

Fonte: Elaboração própria

Comparações Múltiplas Post-Hoc (Bonferroni) para o Efeito Principal do Tipo de Disclosure na Transparência Global

(I) Tipo de Disclosure	(J) Tipo de Disclosure	Diferença de Médias (I-J)	Erro Padrão (EP)	Significância p	Intervalo de Confiança 95%
Sem disclosure IA	Com disclosure IA Criação	-0,503*	0,096	< 0,001	[-0,735; -0,271]
	Com disclosure IA Sugestão	0,075	0,097	1	[-0,157; 0,307]
	Sem disclosure IA	0,503*	0,096	< 0,001	[0,271; 0,735]
Com disclosure IA Criação	Com disclosure IA Sugestão	0,578*	0,092	< 0,001	[0,356; 0,800]
	Sem disclosure IA	-0,075	0,097	1	[-0,307; 0,157]
	Com disclosure IA Criação	-0,578*	0,092	< 0,001	[-0,800; -0,356]

Fonte: Elaboração própria

Micro

Teste de pressuposto multivariado para as subdimensões da Transparência Percebida

Box's M	F	df1	df2	Sig. <i>p</i>
918.560	1.864	48	169.176.810	<.001

Nota: A significância do Teste de Box indica violação da homogeneidade das matrizes de covariância. Por esse motivo, privilegiou-se a interpretação do Traço de Pillai nos testes multivariados.

Fonte: Elaboração própria.

Resultados dos testes multivariados da MANOVA para as subdimensões da Transparência Percebida

Efeito	Estatística multivariada	Valor	F	df da hipótese	df do erro	p	η^2
Tipo de Marca	Traço de Pillai	0,011	0,646	6	706	0,693	0,005
Tipo de Disclosure	Traço de Pillai	0,304	21,068	6	706	< 0,001	0,152
Tipo de Marca $\times$ Tipo	Traço de Pillai	0,012	0,356	12	1062	0,978	0,004

Nota. Variáveis dependentes: TP_Disclosure, TP_Clareza e TP_Exatidão.

Fonte: Elaboração própria.

Efeitos univariados do Tipo de Disclosure nas subdimensões da Transparência Percebida

Variável dependente	Soma de Quadrados	df	Quadrado Médio	F	p	η^2
TP_Disclosure	144,208	2	72,104	55,582	< 0,001	0,239
TP_Clareza	8,647	2	4,324	9,049	< 0,001	0,049
TP_Exatidão	0,139	2	0,07	0,1	0,904	0,001

Nota. Apenas o efeito principal do Tipo de Disclosure é apresentado, por ser o efeito relevante para a H4.

Fonte: Elaboração própria.

Comparações múltiplas post-hoc de Bonferroni para o efeito do Tipo de Disclosure nas subdimensões da Transparência Percebida

Variável dependente	Comparação	Diferença de médias (I-J)	EP	p	IC 95%
TP_Disclosure	Sem disclosure IA vs. IA Criação	-1,189	0,149	< 0,001	[-1,549; -0,830]
	Sem disclosure IA vs. IA Sugestão	0,222	0,15	0,418	[-0,138; 0,581]
	IA Criação vs. IA Sugestão	1,411	0,143	< 0,001	[1,067; 1,755]
TP_Clareza	Sem disclosure IA vs. IA Criação	-0,346	0,091	< 0,001	[-0,565; -0,128]
	Sem disclosure IA vs. IA Sugestão	-0,046	0,091	1	[-0,264; 0,172]
	IA Criação vs. IA Sugestão	0,301	0,087	0,002	[0,092; 0,509]
TP_Exatidão	Sem disclosure IA vs. IA Criação	0,028	0,109	1	[-0,235; 0,290]
	Sem disclosure IA vs. IA Sugestão	0,049	0,109	1	[-0,214; 0,312]
	IA Criação vs. IA Sugestão	0,021	0,104	1	[-0,230; 0,273]

Nota. Ajustamento para comparações múltiplas: Bonferroni. As diferenças de médias baseiam-se nas médias marginais estimadas.

Fonte: Elaboração própria.

Anexo O – Teste hipótese 6: pressupostos

Normalidade

Tipo de Marca	Tipo de Disclosure	Kolmogorov-Smirnov			Shapiro-Wilk		
		Estatística	df	Sig. (p)	Estatística	df	Sig. (p)
ASICS	Sem disclosure IA	0,227	36	< 0,001	0,879	36	0,001
	Com disclosure IA Criação	0,221	41	< 0,001	0,878	41	< 0,001
	Com disclosure IA Sugestão	0,266	37	< 0,001	0,838	37	< 0,001
Ferrero Roche	Sem disclosure IA	0,163	35	0,02	0,919	35	0,012
	Com disclosure IA Criação	0,166	40	0,007	0,936	40	0,025
	Com disclosure IA Sugestão	0,219	44	< 0,001	0,866	44	< 0,001
Rexona	Sem disclosure IA	0,25	36	< 0,001	0,859	36	< 0,001
	Com disclosure IA Criação	0,23	47	< 0,001	0,874	47	< 0,001
	Com disclosure IA Sugestão	0,207	47	< 0,001	0,883	47	< 0,001

Fonte: Elaboração própria

Levene

Teste	Estatística de Levene	df1	df2	p
Baseado na média	0,35	8	354	0,945
Baseado na mediana	0,253	8	354	0,98
Baseado na média aparada	0,311	8	354	0,962

Fonte: Elaboração própria.

Anexo P – Tese hipótese 6: two-way ANOVA

Estatísticas descritivas da Competência Tecnológica Global por tipo de marca e tipo de disclosure

Tipo de Marca	Tipo de Disclosure	M	DP	N
ASICS	Sem disclosure IA	5,33	0,599	36
	Com disclosure IA Criação	5,43	0,594	41
	Com disclosure IA Sugestão	5,42	0,747	37
	Total	5,39	0,645	114
Ferrero Rocher	Sem disclosure IA	5,34	0,706	35
	Com disclosure IA Criação	5,11	0,77	40
	Com disclosure IA Sugestão	5,34	0,677	44
	Total	5,26	0,721	119
Rexona	Sem disclosure IA	5,37	0,57	36
	Com disclosure IA Criação	5,33	0,682	47
	Com disclosure IA Sugestão	5,46	0,667	47
	Total	5,39	0,645	130
Total Geral	Sem disclosure IA	5,35	0,621	107
	Com disclosure IA Criação	5,29	0,692	128
	Com disclosure IA Sugestão	5,4	0,691	128
	Total	5,35	0,671	363

Fonte: Elaboração própria.

Resultados da Two-Way ANOVA I para a Competência Tecnológica Global

Fonte de variação	Soma de Quadrados	df	Quadrado Médio	F	p	η^2
Modelo corrigido	3,405	8	0,426	0,943	0,481	0,021
Tipo de Marca	1,24	2	0,62	1,374	0,255	0,008
Tipo de Disclosure	0,867	2	0,434	0,961	0,384	0,005
Tipo de Marca * Tipo de Disclosure	1,214	4	0,304	0,672	0,612	0,008
Erro	159,801	354	0,451	—	—	—
Total corrigido	163,206	362	—	—	—	—

Nota. Variável dependente: Competência Tecnológica Global. $R^2 = 0,021$; R^2 ajustado = -0,001.

Fonte: Elaboração própria.

Anexo Q – Teste hipótese H8a: PROCESS modelo 1

Macro

Análise de Moderação da Familiaridade com a IA na Relação entre o Tipo de Disclosure e a Autenticidade Percebida (Macro)

Variável / Efeito	Coefficiente (B)	Erro Padrão (SE)	Estatística t	Significância (p)	IC 95% Inferior (LLCI)	IC 95% Superior (ULCI)
Constante	5,3187	0,0754	70,5802	< 0,001	5,1705	5,467
IA Criação vs. Sem Disclosure (X1)	-0,1497	0,0811	-1,8461	0,0657	-0,3092	0,0098
IA Sugestão vs. Sem Disclosure (X2)	-0,0625	0,0818	-0,7639	0,4454	-0,2234	0,0984
Familiaridade com a IA (W)	-0,0547	0,0456	-1,1988	0,2314	-0,1443	0,035
Interação * Familiaridade (Int_1)	-0,0053	0,0671	-0,0788	0,9372	-0,1268	0,1162
Interação X2 * Familiaridade (Int_2)	0,0846	0,067	1,2618	0,2079	-0,0473	0,2164
Covariável: ASICS (M1)	-0,1018	0,0782	-1,3014	0,194	-0,2556	0,052
Covariável: Ferrero (M2)	-0,0168	0,0774	-0,2172	0,8281	-0,1691	0,1354
Efeito de Interação Global (X x W)			F(2, 355) = 1,0780\$	0,3414	R2= 0,0059\$	

Nota. N=363. Variável de Referência para o Disclosure: *Sem Disclosure* (Grupo 1). Variável de Referência para a Marca: *Rexona* (Grupo 3). Variável Dependente (Y): Autenticidade Percebida Global (AU).

Fonte: Elaboração Própria

Micro

Dimensão Benevolência

Análise de Moderação da Familiaridade com a IA na Relação entre o Tipo de Disclosure e a Benevolência Percebida (Micro)

Variável / Efeito	Coefficiente (B)	Erro Padrão (SE)	Estatística t	Significância (p)	IC 95% Inferior (LLCI)	IC 95% Superior (ULCI)
Constante	5,3092	0,0891	59,5918	< 0,001	5,134	5,4844
IA Criação vs. Sem Disclosure (X1)	-0,0964	0,0959	-1,0060	0,3151	-0,2850	0,0921
IA Sugestão vs. Sem Disclosure (X2)	-0,0285	0,0967	-0,2948	0,7683	-0,2187	0,1617
Familiaridade com a IA (W)	-0,0035	0,0539	-0,0645	0,9486	-0,1095	1,0255
Interação * Familiaridade (Int_1)	-0,0375	0,0794	-0,4722	0,6371	-0,1936	0,1186
Interação X2 * Familiaridade (Int_2)	0,1292	0,0793	1,6307	0,1039	-0,0266	0,2851
Covariável: ASICS (M1)	-0,0115	0,0925	-0,1249	0,9007	-0,1934	0,1703
Covariável: Ferrero (M2)	-0,1702	0,0915	-1,8601	0,0637	-0,3502	0,0098
Efeito de Interação Global (X x W)			F(2, 355) = 2,2936	0,1024	R2 = 0,0125	

Nota. N=363. Variável de Referência para o Disclosure: *Sem Disclosure* (Grupo 1). Variável de Referência para a Marca: *Rexona* (Grupo 3). Variável Dependente (Y): Benevolência da Autenticidade (AP_DBen).

Fonte: Elaboração Própria

Dimensão Continuidade

Análise de Moderação da Familiaridade com a IA na Relação entre o Tipo de Disclosure e a Continuidade Percebida (Micro)

Variável / Efeito	Coefficiente (B)	Erro Padrão (SE)	Estatística t	Significância (p)	IC 95% Inferior (LLCI)	IC 95% Superior (ULCI)
Constante	5,5258	0,0841	65,7084	< 0,001	5,3604	5,6912
IA Criação vs. Sem Disclosure (X1)	-1,1017	0,0905	-1,1240	0,2618	-0,2797	0,0763
IA Sugestão vs. Sem Disclosure (X2)	-0,0832	0,0913	-0,9110	0,3629	-0,2627	0,0964
Familiaridade com a IA (W)	-0,0897	0,0509	-1,7628	0,0788	-0,1898	0,0104
Interação * Familiaridade (Int_1)	-0,0311	0,0749	-0,4155	0,6781	-0,1785	1,1162
Interação X2 * Familiaridade (Int_2)	0,0657	0,0748	0,8778	0,3806	-0,0815	0,2128
Covariável: ASICS (M1)	-0,1443	0,0873	-1,6536	0,0991	-0,3160	0,0273
Covariável: Ferrero (M2)	0,0306	0,0864	0,3545	0,7232	-0,1393	0,2005
Efeito de Interação Global (X x W)			F(2, 355) = 0,8174	0,4424	R2 = 0,0044	

Nota. N=363. Variável de Referência para o Disclosure: *Sem Disclosure* (Grupo 1). Variável de Referência para a Marca: *Rexona* (Grupo 3). Variável Dependente (Y): Continuidade da Autenticidade (AP_DCCont).

Fonte: Elaboração Própria

Dimensão Originalidade

Análise de Moderação da Familiaridade com a IA na Relação entre o Tipo de Disclosure e a Originalidade Percebida (Micro)

Variável / Efeito	Coefficiente (B)	Erro Padrão (SE)	Estatística t	Significância (p)	IC 95% Inferior (LLCI)	IC 95% Superior (ULCI)
Constante	5,269	0,1177	44,7768	< 0,001	5,0375	5,5504
IA Criação vs. Sem Disclosure (X1)	-0,4269	0,1266	-3,3714	0,0008	-0,6759	-0,1779
IA Sugestão vs. Sem Disclosure (X2)	-0,02563	0,1277	-2,0068	0,0455	-0,5075	-0,0051
Familiaridade com a IA (W)	-0,0175	0,0712	-0,2456	0,8061	-0,1575	0,1225
Interação * Familiaridade (Int_1)	0,1223	0,1049	1,1662	0,2443	-0,0839	0,3285
Interação X2 * Familiaridade (Int_2)	-0,0066	0,1047	-0,0632	0,9496	-0,2125	0,1992
Covariável: ASICS (M1)	-0,1423	0,1221	-1,1650	0,2448	-0,3824	0,0979
Covariável: Ferrero (M2)	0,2496	0,1209	2,0652	0,0396	0,0119	0,4873
Efeito de Interação Global (X x W)			F(2, 355) = 0,9134	0,4021	R2 = 0,0048	

Nota. N=363. Variável de Referência para o Disclosure: *Sem Disclosure* (Grupo 1). Variável de Referência para a Marca: *Rexona* (Grupo 3). Variável Dependente (Y): Originalidade da Autenticidade (AP_DOori).

Fonte: Elaboração Própria

Anexo R – Teste hipótese H8b: PROCESS model 1

Macro

Análise de Moderação da Familiaridade com a IA na Relação entre o Tipo de Disclosure e a Transparência Percebida Global (Macro)

Variável / Efeito	Coefficiente (B)	Erro Padrão (SE)	Estatística t	Significância (p)	IC 95% Inferior	IC 95% Superior
Constante	5,0309	0,0883	57,0043	< 0,001	4,8573	5,2044
IA Criação vs. Sem Disclosure (X1)	0,537	0,095	5,6544	< 0,001	0,3502	0,7238
IA Sugestão vs. Sem Disclosure (X2)	0,0168	0,0958	0,1757	0,8606	-0,1716	0,2052
Familiaridade com a IA (W)	-0,1101	0,0534	-2,0622	0,0399	-0,2152	-0,0051
Interação * Familiaridade (Int_1)	0,1436	0,0786	1,8259	0,0687	-0,0111	0,2983
Interação X2 * Familiaridade (Int_2)	-0,1525	0,0785	-1,9419	0,0529	-0,3068	0,0019
Covariável: ASICS (M1)	-0,0617	0,0916	-0,6740	0,5007	-0,2419	0,1184
Covariável: Ferrero (M2)	-0,0805	0,0907	-0,8876	0,3753	-0,2588	0,0978
Efeito de Interação Global (X x W)			F(2, 355) = 6,6305	0,0015	R2 = 0,0308	

Níveis de Moderação da Familiaridade com a IA na relação entre o Tipo de Disclosure e a Transparência Percebida Global

Níveis da Familiaridade com a IA (W)	Efeito (B) em X1	Erro Padrão (SE)	Estatística t	Significância (p)	IC 95% Inferior	IC 95% Superior
Familiaridade Baixa (-1,1350)	0,374	0,1237	3,0237	0,0027	0,1307	0,6173
Familiaridade Média (0,1150)	0,5535	0,0963	5,7483	< 0,001	0,3641	0,7429
Familiaridade Alta (1,1150)	0,6971	0,1355	5,144	< 0,001	0,4306	0,9636

Nota. N=363. Variável de Referência para o Disclosure: *Sem Disclosure* (Grupo 1). Variável de Referência para a Marca: *Rexona* (Grupo 3). Variável Dependente (Y): Transparência Percebida Global (TP_Globa).

Micro

Dimensão Disclosure

Análise de Moderação da Familiaridade com a IA na Relação entre o Tipo de Disclosure e a Dimensão de Divulgação da Transparência (Micro)

Variável / Efeito	Coefficiente (B)	Erro Padrão (SE)	Estatística t	Significância (p)	IC 95% Inferior	IC 95% Superior
Constante	4,1177	0,1336	30,8151	< 0,001	3,8549	4,3805
IA Criação vs. Sem Disclosure (X1)	1,2567	0,1438	8,7401	< 0,001	0,974	1,5395
IA Sugestão vs. Sem Disclosure (X2)	-0,0385	0,1451	-0,2657	0,7906	-0,3238	0,2467
Familiaridade com a IA (W)	-0,2349	0,0809	-2,9051	0,0039	-0,3939	-0,0759
Interação * Familiaridade (Int_1)	0,402	0,1191	3,3766	0,0008	0,1679	0,6362
Interação X2 * Familiaridade (Int_2)	-0,2530	0,1189	-2,1285	0,034	-0,4868	-0,0192
Covariável: ASICS (M1)	-0,0336	0,1387	-0,2420	0,8089	-0,3063	0,2392
Covariável: Ferrero (M2)	-0,0367	0,1373	-0,2673	0,7894	-0,3066	0,2333
Efeito de Interação Global (X x W)			F(2, 355) = 14,4152	< 0,001	R2 = 0,0551	

Níveis de Moderação da Familiaridade com a IA na relação entre o Tipo de Disclosure e a Dimensão da Transparência Percebida Disclosure

Níveis da Familiaridade com a IA (W)	Efeito (B) em X1	Significância (p)	Efeito (B) em X2	Significância (p)
Familiaridade Baixa (\$-1,1350\$)	0,8004	< 0,001	0,2486	0,2035
Familiaridade Média (\$0,1150\$)	1,303	< 0,001	-0,0676	0,6437
Familiaridade Alta (\$1,1150\$)	1,705	< 0,001	-0,3206	0,1086

Marca: *Rexona* (Grupo 3). Variável Dependente (Y): Dimensão Divulgação da Transparência (TP_DDsc).

Fonte: Elaboração Própria

Dimensão Clareza

Análise de Moderação da Familiaridade com a IA na Relação entre o Tipo de Disclosure e a Dimensão de Clareza da Transparência (Micro)

Variável / Efeito	Coefficiente (B)	Erro Padrão (SE)	Estatística t	Significância (p)	IC 95% Inferior	IC 95% Superior
Constante	5,6033	0,0852	65,7734	< 0,001	5,4357	5,7708
IA Criação vs. Sem Disclosure (X1)	0,3442	0,0917	3,7546	0,0002	0,1639	0,5245
IA Sugestão vs. Sem Disclosure (X2)	0,0644	0,0925	0,6969	0,4863	-0,1174	0,2463
Familiaridade com a IA (W)	0,007	0,0759	0,0928	0,9261	-0,1422	0,1563
Interação * Familiaridade (Int_1)	0,0071	0,0515	0,1375	0,8907	-0,0943	0,1085
Interação X2 * Familiaridade (Int_2)	-0,1170	0,0758	-1,5444	0,1234	-0,2661	0,032
Covariável: ASICS (M1)	-0,0699	0,0884	-0,7903	0,4299	-0,2437	0,104
Covariável: Ferrero (M2)	-0,1440	0,0875	-1,6456	0,1007	-0,3161	0,0281
Efeito de Interação Global (X x W)			F(2, 355) = 1,6110	0,2011	R2 = 0,0085	

Nota. N=363. Variável de Referência para o Disclosure: *Sem Disclosure* (Grupo 1). Variável de Referência para a Marca: *Rexona* (Grupo 3). Variável

Fonte: Elaboração Própria

Dimensão Exatidão

Análise de Moderação da Familiaridade com a IA na Relação entre o Tipo de Disclosure e a Dimensão de Exatidão da Transparência (Micro)

Variável / Efeito	Coefficiente (B)	Erro Padrão (SE)	Estatística t	Significância (p)	IC 95% Inferior	IC 95% Superior
Constante	5,3716	0,1015	52,9389	< 0,001	5,172	5,5711
IA Criação vs. Sem Disclosure (X1)	0,01	0,1092	0,0917	0,927	-0,2047	0,2247
IA Sugestão vs. Sem Disclosure (X2)	0,0246	0,1101	0,2232	0,8235	-0,1920	0,2412
Familiaridade com a IA (W)	-0,1026	0,0614	-1,6705	0,0957	-0,2233	0,0182
Interação * Familiaridade (Int_1)	0,0217	0,0904	0,2399	0,8106	-0,1561	0,1995
Interação X2 * Familiaridade (Int_2)	-0,0873	0,0903	-0,9674	0,334	-0,2648	0,0902
Covariável: ASICS (M1)	-0,0818	0,1053	-0,7765	0,438	-0,2889	0,1253
Covariável: Ferrero (M2)	-0,0607	0,1042	-0,5825	0,5606	-0,2657	0,1443
Efeito de Interação Global (X x W)			F(2, 355) = 0,7708	0,4634	R² = 0,0042	

Nota. N=363. Variável de Referência para o Disclosure: *Sem Disclosure* (Grupo 1). Variável de Referência para a Marca: *Rexona* (Grupo 3). Variável Dependente (Y): Dimensão Exatidão da Transparência (TP_DExat).

Fonte: Elaboração Própria

Anexo S - Teste hipótese H8c: PROCESS model 1

Competência Tecnológica

Análise de Moderação da Familiaridade com a IA na Relação entre o Tipo de Disclosure e a Competência Tecnológica Global (Macro)

Variável / Efeito	Coefficiente (B)	Erro Padrão (SE)	Estatística t	Significância (p)	IC 95% Inferior (LLCI)	IC 95% Superior (ULCI)
Constante	5,3801	0,0834	64,5246	< 0,001	5,2162	5,5441
IA Criação vs. Sem Disclosure (X1)	-0,0534	0,0897	-0,5953	0,552	-0,2299	0,123
IA Sugestão vs. Sem Disclosure (X2)	0,0681	0,0905	0,752	0,4525	-0,1099	0,2461
Familiaridade com a IA (W)	-0,0095	0,0505	-0,1883	0,8507	-1,0870	0,0897
Interação * Familiaridade (Int_1)	-0,0294	0,0743	-0,3960	0,6924	-0,1755	0,1167
Interação X2 * Familiaridade (Int_2)	-0,0156	0,0742	-0,2104	0,8335	-0,1615	0,1303
Covariável: ASICS (M1)	0,011	0,0865	0,1275	0,8986	-0,1591	0,1812
Covariável: Ferrero (M2)	-0,1201	0,0857	-1,4020	0,1618	-0,2885	0,0484
Efeito de Interação Global (X x W)			F(2, 355) = 0,0788	0,9242	R² = 0,0004	

Nota. N=363. Variável de Referência para o Disclosure: *Sem Disclosure* (Grupo 1). Variável de Referência para a Marca: *Rexona* (Grupo 3). Variável Dependente (Y): Competência Tecnológica Global (CP_G).

Fonte: Elaboração Própria

Anexo T – Mediação paralela: PROCESS modelo 4

Efeitos das condições experimentais nas propriedade emergentes (Caminho A)

Antecedentes	Condição	Coefficiente (B)	Erro Padrão (SE)	Estatística t	Significância (p)
Autenticidade	IA Criação	-0,1708	0,0797	-2,1436	0,0327
	IA Sugestão	-0,0739	0,0797	-0,9271	0,3545
	ASICS	-0,1076	0,0781	-1,377	0,1693
	Ferrero Rocher	-0,0215	0,0772	-0,278	0,7812
Transparência	IA Criação	0,5027	0,0962	5,2259	< 0,001
	IA Sugestão	-0,0762	0,0962	-0,7919	0,429
	ASICS	-0,073	0,0943	-0,8094	0,4188
	Ferrero Rocher	-0,0942	0,0932	-1,0111	0,312
Competência Tecnológica	IA Criação	-0,0587	0,0879	-0,6684	0,5043
	IA Sugestão	0,0595	0,0879	0,6772	0,4987
	ASICS	0,0089	0,0861	0,1036	0,9176
	Ferrero Rocher	-0,1259	0,0851	-1,4799	0,1398

Nota: Condição "Sem disclosure" utilizada como referência; "Rexona" utilizada como covariável de referência.

Fonte: Elaboração própria.

Resumo do Modelo da Confiança da Marca (Modelo 4)

R	R ²	MSE	F	df1	df2	p
0,5225	0,273	0,3642	19,0446	7	355	< 0,001

Nota: Variável Dependente: Confiança na Marca; Variável Independente: Tipo_Disclosure; Mediadores: Autenticidade, Transparência e Competência Tecnológica; Covariáveis: Tipo_Marca

Fonte: Elaboração própria.

Modelo de Mediação Paralela para Confiança da Marca (Caminho B)

Preditor	B	EP	t	p	IC 95% inferior	IC 95% superior
Constante	2,5903	0,3162	8,1927	< 0,001	1,9685	3,2121
IA Criação	-0,0932	0,0846	-1,1022	0,2711	-0,2596	0,0731
IA Sugestão	-0,0871	0,0797	-1,094	0,2747	-0,2438	0,0695
Autenticidade	0,3213	0,0731	4,3948	< 0,001	0,1775	0,4651
Transparência	0,0591	0,0497	1,1901	0,2348	-0,0386	0,1569
Competência Tecnológica	0,2181	0,0672	3,2433	0,0013	0,0858	0,3503
ASICS	-0,2815	0,078	-3,6105	< 0,001	-0,4348	-0,1282
Ferrero Rocher	-0,0117	0,077	-0,1515	0,8797	-0,163	0,1397

Nota: Variável Dependente: Confiança na Marca; Variável Independente: Tipo_Disclosure; Mediadores: Autenticidade, Transparência e Competência Tecnológica; Covariáveis: Tipo_Marca; Condição "Sem disclosure" utilizada como referência; "Rexona" utilizada como covariável de referência.

Fonte: Elaboração própria.

Efeitos Indiretos do *disclosure* na Confiança da Marca

Mediador	Comparação	Efeito	BootSE	BootLLCI	BootULCI
Autenticidade	IA Criação	-0,0549	0,0293	-0,1201	-0,0039
	IA Sugestão	-0,0238	0,0274	-0,0824	0,0301
Transparência	IA Criação	0,0297	0,0301	-0,0279	0,094
	IA Sugestão	-0,0045	0,0096	-0,0295	0,0101
Competência Tecnológica	IA Criação	-0,0128	0,0203	-0,0573	0,0255
	IA Sugestão	0,013	0,0207	-0,0223	0,0599

Nota: Variável Dependente: Confiança na Marca; Variável Independente: Tipo_Disclosure; Mediadores: Autenticidade, Transparência e Competência Tecnológica; Covariáveis: Tipo_Marca; Condição "Sem disclosure" utilizada como referência; "Rexona" utilizada como covariável de referência.

Fonte: Elaboração própria.

Anexo U- Moderação mediada: PROCESS modelo 7**Interação entre o *disclosure* e a Familiaridade com IA por dimensão**

Dimensão	R ²	F	df1	df2	p
Autenticidade	0,0059	1,078	2	355	0,3414
Transparência	0,0308	6,6305	2	355	0,0015
Competência tecnológica	0,0004	0,0788	2	355	0,9242

Fonte: Elaboração Própria**Moderação da Familiaridade com a IA na relação entre o Tipo de Disclosure e a Transparência**

Tipo de Disclosure	Familiaridade com IA	Efeito	EP	t	p	IC 95% inferior	IC 95% superior
IA Criação	Familiaridade Baixa (-1,1350)	0,374	0,1237	3,0237	0,0027	0,1307	0,6173
	Familiaridade Média (0,1150)	0,5535	0,0963	5,7483	< 0,001	0,3641	0,7429
	Familiaridade Alta (1,1150)	0,6971	0,1355	5,144	< 0,001	0,4306	0,9636
IA Sugestão	Familiaridade Baixa (-1,1350)	0,1899	0,1289	1,473	0,1416	-0,0636	0,4434
	Familiaridade Média (0,1150)	-0,0007	0,0965	-0,0073	0,9942	-0,1905	0,1891
	Familiaridade Alta (1,1150)	-0,1532	0,1317	-1,1633	0,2455	-0,4121	0,1058

Fonte: Elaboração própria.**Resumo do Modelo da Confiança da Marca (Modelo 7)**

R	R ²	MSE	F	df1	df2	p
0,5225	0,273	0,3642	19,0446	7	355	< 0,001

Modelo de Moderação Mediada para Confiança da Marca

Preditor	B	EP	t	p	IC 95% inferior	IC 95% superior
Constante	2,5903	0,3162	8,1927	< 0,001	1,9685	3,2121
IA Criação	-0,0932	0,0846	-1,1022	0,2711	-0,2596	0,0731
IA Sugestão	-0,0871	0,0797	-1,094	0,2747	-0,2438	0,0695
Autenticidade	0,3213	0,0731	4,3948	< 0,001	0,1775	0,4651
Transparência	0,0591	0,0497	1,1901	0,2348	-0,0386	0,1569
Competência Tecnológica	0,2181	0,0672	3,2433	0,0013	0,0858	0,3503
ASICS	-0,2815	0,078	-3,6105	< 0,001	-0,4348	-0,1282
Ferrero Rocher	-0,0117	0,077	-0,1515	0,8797	-0,163	0,1397

Nota: Variável Dependente: Confiança na Marca; Variável Independente: Tipo_Disclosure; Mediadores: Autenticidade, Transparência e Competência Tecnológica; Covariáveis: Tipo_Marca; Condição "Sem disclosure" utilizada como referência; "Rexona" utilizada como covariável de referência.

Fonte: Elaboração própria.

Tabela: Efeitos Indiretos Condicionais e Índices de Mediação Moderada (Modelo 7)

Variável Mediadora	Nível de Familiaridade (FIA)	Efeito Indireto	Boot SE	Boot LLCI	Boot ULCI
IA Criação -> Confiança na Marca					
Autenticidade	Baixa (-1,1350)	-0,0462	0,044	-0,1485	0,0215
	Média (0,1150)	-0,0483	0,0298	-0,1112	0,004
	Alta (1,1150)	-0,0500	0,0445	-0,1371	0,0464
Transparência	Baixa (-1,1350)	0,0221	0,0265	-0,0179	0,0881
	Média (0,1150)	0,0327	0,0328	-0,0303	0,1005
	Alta (1,1150)	0,0412	0,0411	-0,0397	0,1266
Competência	Baixa (-1,1350)	-0,0044	0,0269	-0,0622	0,0475
	Média (0,1150)	-0,0124	0,0218	-0,0599	0,0281
	Alta (1,1150)	-0,0188	0,0345	-0,0931	0,0469
IA Sugestão -> Confiança na Marca					
Autenticidade	Baixa (-1,1350)	-0,0509	0,0462	-0,1558	0,0275
	Média (0,1150)	-0,0170	0,0294	-0,0819	0,0381
	Alta (1,1150)	-0,0102	0,0412	-0,0731	0,1002
Transparência	Baixa (-1,1350)	0,0112	0,0177	-0,0139	0,056
	Média (0,1150)	0	0,0081	-0,0191	0,0162
	Alta (1,1150)	-0,0091	0,0154	-0,0499	0,0113
Competência	Baixa (-1,1350)	0,0187	0,0357	-0,0467	0,0977
	Média (0,1150)	0,0145	0,0221	-0,0248	0,0641
	Alta (1,1150)	0,011	0,0302	-0,0446	0,0781

Fonte: Elaboração própria

Índices de Mediação Moderada

Tipo de Disclosure	Variáveis Mediadoras	Índice	Boot SE	Boot LLCI	Boot ULCI
IA Criação	Autenticidade	-0,0017	0,0293	-0,050	0,0679
	Transparência	0,0085	0,0118	-0,0136	0,0356
	Competência	-0,0064	0,0202	-0,0477	0,0346
IA Sugestão	Autenticidade	0,272	0,0286	-0,0232	0,0912
	Transparência	-0,09	0,0128	-0,0424	0,0075
	Competência	-0,034	0,217	-0,0487	0,0395

Fonte: Elaboração própria