

**UNIVERSIDADE FEDERAL DE SANTA CATARINA
CENTRO SOCIOECONÔMICO
DEPARTAMENTO DE CIÊNCIAS DA ADMINISTRAÇÃO**

Bianca Maria dos Santos Gonzaga

**Co-criação entre gestores e Inteligência Artificial Generativa: a influência da
confiança na IA em perspectiva intercultural.**

Florianópolis

2026

Bianca Maria dos Santos Gonzaga

Co-criação entre gestores e Inteligência Artificial Generativa: a influência da confiança na IA em perspectiva intercultural.

Trabalho de Conclusão de Curso submetido ao curso de Administração do Centro Socioeconômico da Universidade Federal de Santa Catarina como requisito parcial para a obtenção do título de Bacharel(a) em Administração.

Orientador(a): Prof.(a) Janaína Gularte Cardoso, Dr.^a

Florianópolis

2026

Gonzaga, Bianca Maria dos Santos

Co-criação entre gestores e Inteligência Artificial Generativa: a influência da confiança na IA em perspectiva intercultural. / Bianca Maria dos Santos Gonzaga ; orientadora, Janaína Gularte Cardoso, 2026.

104 p.

Trabalho de Conclusão de Curso (graduação) - Universidade Federal de Santa Catarina, Centro Socioeconômico, Graduação em Administração, Florianópolis, 2026.

Inclui referências.

1. Administração. 2. Inteligência Artificial Generativa. 3. confiança na IA. 4. humano-IA. 5. comparação Brasil-China. I. Cardoso, Janaína Gularte. II. Universidade Federal de Santa Catarina. Graduação em Administração. III. Título.

Bianca Maria dos Santos Gonzaga

Co-criação entre gestores e Inteligência Artificial Generativa: a influência da confiança na IA em perspectiva intercultural.

Este Trabalho de Conclusão de Curso foi julgado adequado para obtenção do título de Bacharel e aprovado em sua forma final pelo Curso de Administração.

Florianópolis, 18 de Agosto de 2026.

Banca examinadora

Prof.(a) Janaína Gularte Cardoso, Dr.^a
Orientadora
Universidade Federal de Santa Catarina

Prof. Leandro Dorneles dos Santos, Dr.
Avaliador
Universidade Federal de Santa Catarina

Prof. Alexandre Moraes Ramos, Dr.
Avaliador
Universidade Federal de Santa Catarina

Dedico este trabalho
à minha família, por
todo o amor, apoio e
incentivo ao longo da
minha trajetória.

AGRADECIMENTOS

É curioso pensar que uma graduação começa, oficialmente, dentro de uma sala de aula, mas algumas das experiências que mais definiram a minha trajetória aconteceram muito além dela.

À Universidade Federal de Santa Catarina, agradeço por ter sido o ponto de partida de caminhos que, quando entrei na universidade, eu nem sabia que existiam. Foi a partir da UFSC que pude conhecer outros lugares, atravessar fronteiras, voltar com novas perguntas e entender que uma formação acadêmica também acontece nos encontros, nas experiências e nas possibilidades que surgem pelo caminho. Mais do que o espaço onde construí minha graduação, a universidade se tornou parte de uma trajetória que mudou a forma como eu observo o mundo e, principalmente, o quanto ainda quero conhecê-lo.

À Fudan University, agradeço pela oportunidade de viver parte dessa trajetória na China. Estar em um contexto tão diferente do meu me permitiu não apenas desenvolver uma etapa importante desta pesquisa, mas também enxergar de perto outras formas de pensar, aprender e interpretar o mundo. Talvez essa tenha sido uma das maiores constantes desses anos: perceber que quanto mais lugares conheço e quanto mais perspectivas encontro, mais perguntas aparecem. E considero isso uma coisa boa.

À minha orientadora, professora Janaína, agradeço pela confiança nas minhas ideias, pelo apoio durante a construção deste trabalho e, especialmente, pela liberdade de transformar interesses e inquietações em pesquisa. Aos professores que fizeram parte da minha formação, dentro e fora da sala de aula, agradeço pelas conversas, provocações e referências que ajudaram a ampliar minha forma de pensar. Nem todo aprendizado aparece diretamente nas páginas de um trabalho acadêmico, mas isso não significa que ele não esteja presente.

Aos participantes desta pesquisa, no Brasil e na China, agradeço pela disponibilidade e pela confiança em compartilhar suas experiências comigo. Pesquisa também é construída a partir da disposição de alguém em parar, refletir e dividir aquilo que sabe. Este trabalho existe porque diferentes pessoas aceitaram fazer exatamente isso.

Aos amigos e colegas que encontrei durante esses anos, fica meu carinho por tudo o que existiu: as conversas, as viagens, os encontros inesperados, os aprendizados e os momentos que acabaram se tornando parte da trajetória tanto quanto qualquer

disciplina. Muitas oportunidades foram importantes pelo que me permitiram fazer. Outras se tornaram importantes, principalmente, pelas pessoas que colocaram no meu caminho.

À minha família, agradeço pelo amor, pelo apoio e por uma segurança que aprendi a valorizar ainda mais conforme fui cada vez mais longe: a de saber que sempre existe um lugar para onde voltar. Cada experiência que pude viver carrega também um pouco da confiança que vocês depositaram em mim muito antes de eu saber exatamente onde ela poderia me levar.

Acima de tudo, agradeço a Deus pela saúde, pelas oportunidades e pela possibilidade de continuar me movimentando, aprendendo e encontrando pessoas capazes de ampliar a minha visão de mundo. Sou especialmente grata pelas portas que se abriram antes mesmo de eu saber procurá-las, e por ter tido condições, coragem e curiosidade suficientes para atravessá-las.

Por fim, agradeço à universidade pública e a todos que ainda acreditam no conhecimento como uma forma concreta de transformação. Existe algo muito particular em começar uma trajetória sem conseguir prever aonde ela pode levar e, alguns anos depois, olhar para trás e perceber quantos caminhos nasceram daquele primeiro passo.

Este trabalho encerra uma etapa, mas talvez o que eu mais leve dela não seja uma conclusão. Levo a certeza de que continuar curiosa, aberta a outras perspectivas e disposta a atravessar caminhos que ainda não conheço foi o que tornou essa trajetória possível até aqui.

E espero que continue sendo.

“君子和而不同，小人同而不和。”

——孔子《论语·子路》

RESUMO

A expansão da Inteligência Artificial Generativa (IAG) tem deslocado o debate organizacional da adoção tecnológica para as formas de interação entre gestores e sistemas generativos. Esta pesquisa analisa como a confiança influencia a cocriação humano-IAG em rotinas gerenciais, comparando os contextos brasileiro e chinês. Adotou-se abordagem qualitativa, interpretativa, exploratória e comparativa, com entrevistas semiestruturadas realizadas com dez gestores, cinco em cada país. A análise temática utilizou procedimentos de codificação inspirados na Grounded Theory, com apoio de IA e validação independente pela pesquisadora. Os resultados indicam que a confiança é construída e recalibrada ao longo do uso, variando conforme a criticidade da tarefa, a verificabilidade dos outputs e as condições organizacionais e institucionais. Essa confiança regula o grau de delegação, supervisão e refinamento na interação com a IAG. A comparação revela ampla convergência na preservação do julgamento humano, enquanto, nos casos chineses, segurança de dados, compliance e uso de ferramentas aprovadas aparecem de forma mais recorrente. Conclui-se que a confiança atua como mecanismo de calibração da participação da IAG no trabalho gerencial, estruturando processos de cocriação baseados em delegação seletiva, validação humana e possibilidade de retomada manual das tarefas. O estudo contribui para a compreensão da confiança e da cocriação humano-IAG em contextos do Sul Global.

Palavras-chave: Inteligência Artificial Generativa; confiança na IA; co-criação humano-IA; gestores; comparação Brasil-China.

ABSTRACT

The expansion of Generative Artificial Intelligence (GenAI) has shifted the organizational debate from technology adoption toward the forms of interaction between managers and generative systems. This study analyzes how trust influences human–GenAI co-creation in managerial routines by comparing the Brazilian and Chinese contexts. A qualitative, interpretive, exploratory, and comparative approach was adopted, based on semi-structured interviews with ten managers, five in each country. Thematic analysis employed coding procedures inspired by Grounded Theory, with AI support and independent validation by the researcher. The findings indicate that trust is constructed and recalibrated through use, varying according to task criticality, output verifiability, and organizational and institutional conditions. This trust regulates the degree of delegation, supervision, and refinement in interactions with GenAI. The comparison reveals broad convergence in the preservation of human judgment, while data security, compliance, and the use of approved tools appear more recurrently in the Chinese cases. The study concludes that trust acts as a mechanism for calibrating GenAI participation in managerial work, structuring co-creation processes based on selective delegation, human validation, and the possibility of manually resuming tasks. The study contributes to the understanding of trust and human–GenAI co-creation in Global South contexts.

Keywords: Generative Artificial Intelligence; trust in AI; human–AI co-creation; managers; Brazil–China comparison.

LISTA DE FIGURAS

Figura 1 - Articulação entre o problema de pesquisa, o modelo analítico e o percurso metodológico.....	43
Figura 2 - Trajetórias e mecanismos de articulação entre incorporação da IAG, confiança e cocriação humano-IA nos contextos brasileiro e chinês.....	81

LISTA DE QUADROS

Quadro 1 - Articulação entre problema de pesquisa, objetivos específicos e dimensões analíticas.....	42
Quadro 2 - Articulação entre objetivos específicos, roteiro de entrevistas e o processo de análise dos dados.....	46
Quadro 3 - Evolução das categorias e verificação da saturação teórica nos contextos brasileiro e chinês.....	50
Quadro 4 - Evidências processuais de saturação nas entrevistas finais.....	51
Quadro 5 - Critérios de rigor qualitativo e procedimentos adotados.....	60
Quadro 6 - Concordância entre as classificações da pesquisadora e do assistente de IA.....	62
Quadro 7 - Visão geral da amostra.....	63
Quadro 8 - Caracterização dos participantes.....	64
Quadro 9 - Perfil de utilização da IA.....	66
Quadro 10 - Síntese dos contextos analisados.....	67
Quadro 11 - Categorias emergentes da incorporação da IAG nas rotinas gerenciais.....	68
Quadro 12 - Comparação da construção, percepção e mobilização da confiança na IAG entre gestores brasileiros e chineses.....	70
Quadro 13 - Comparação dos processos de interação e co-criação entre gestores e IAG.....	74
Quadro 14 - Etapas do processo de co-criação humano-IA identificadas na pesquisa.....	75
Quadro 15 - Síntese comparativa dos resultados entre Brasil e China.....	77
Quadro 16 - Implicações gerenciais decorrentes dos resultados da pesquisa.....	83
Quadro 17 - Limitações da pesquisa e agenda de estudos futuros.....	86
Quadro 18 - Estrutura do Roteiro de entrevistas.....	98
Quadro 19 - Roteiro de entrevista semiestruturada.....	99
Quadro 20 - Declaração de uso de Inteligência Artificial Generativa.....	101

LISTA DE PROMPTS

Prompt 1 – System prompt para assistente de IA em codificação aberta (modelo relacional CTX/TRUST/COCRI/LINK).....	57
--	----

LISTA DE ABREVIATURAS E SIGLAS

AI Artificial Intelligence (Inteligência Artificial)

IA Inteligência Artificial

IAG Inteligência Artificial Generativa

TCLE Termo de Consentimento Livre e Esclarecido

SUMÁRIO

1 INTRODUÇÃO	16
1.1 PROBLEMA DE PESQUISA	18
1.2 OBJETIVOS	19
1.3 JUSTIFICATIVA	19
2 FUNDAMENTAÇÃO TEÓRICA	22
2.1 INTELIGÊNCIA ARTIFICIAL GENERATIVA NO CONTEXTO DO TRABALHO	23
2.2 CO-CRIAÇÃO ENTRE HUMANOS E INTELIGÊNCIA ARTIFICIAL GENERATIVA	26
2.3 CONFIANÇA NA INTELIGÊNCIA ARTIFICIAL GENERATIVA	29
2.4 DIMENSÃO INTERCULTURAL DA CONFIANÇA NA INTELIGÊNCIA ARTIFICIAL (BRASIL VS. CHINA)	33
2.5 SÍNTESE TEÓRICA E CONSTRUÇÃO DO MODELO ANALÍTICO	37
3 PROCEDIMENTOS METODOLÓGICOS	40
3.1 NATUREZA DA PESQUISA	41
3.2 ESTRATÉGIA DE PESQUISA E ARTICULAÇÃO COM O MODELO ANALÍTICO	43
3.3 COLETA DE DADOS	44
3.3.1 CRITÉRIOS DE SELEÇÃO DOS PARTICIPANTES	46
3.3.2 INSTRUMENTO DE COLETA	47
3.3.3 PROCEDIMENTOS PARA COLETA DE DADOS	47
3.3.4 SATURAÇÃO TEÓRICA E SUFICIÊNCIA DO MATERIAL EMPÍRICO	48
3.4 PRÉ-TESTE DO INSTRUMENTO DE COLETA	53
3.5 ANÁLISE DE DADOS	54
3.6 USO DE INTELIGÊNCIA ARTIFICIAL COMO SUPORTE METODOLÓGICO	56
3.7 CONSIDERAÇÕES SOBRE RIGOR METODOLÓGICO	59
4 RESULTADOS E DISCUSSÃO	62
4.1 AVALIAÇÃO DA CONCORDÂNCIA DO PROCEDIMENTO DE CODIFICAÇÃO ASSISTIDA POR IA	62
4.2 PERFIL DOS PARTICIPANTES E CARACTERIZAÇÃO DOS CONTEXTOS	63
4.2.1 CARACTERIZAÇÃO DOS PARTICIPANTES	64
4.2.2 PERFIL DE UTILIZAÇÃO DA IA GENERATIVA	65
4.2.3 CARACTERIZAÇÃO DE CONTEXTO	66
4.2.4 SÍNTESE COMPARATIVA	67
4.3 INCORPORAÇÃO DA IAG NAS ROTINAS GERENCIAIS (OBJETIVO ESPECÍFICO A)	68
4.4 CONSTRUÇÃO, PERCEPÇÃO E MOBILIZAÇÃO DA CONFIANÇA NA IA (OBJETIVO ESPECÍFICO B)	70
4.5 PROCESSOS DE INTERAÇÃO E CO-CRIAÇÃO HUMANO-IA (OBJETIVO ESPECÍFICO C)	73
4.6 ANÁLISE COMPARATIVA E DISCUSSÃO INTEGRADA BRASIL-CHINA (OBJETIVO ESPECÍFICO D)	76
5 CONSIDERAÇÕES FINAIS	82

1 INTRODUÇÃO

A expansão da Inteligência Artificial Generativa (IAG) insere-se na transformação digital do trabalho e das organizações. Brynjolfsson e McAfee (2014) argumentam que avanços tecnológicos têm ampliado a capacidade de sistemas computacionais realizarem tarefas anteriormente associadas a capacidades humanas complexas. A IAG aprofunda esse movimento ao viabilizar atuação direta sobre atividades linguísticas, interpretativas e criativas, enquanto Davenport e Ronanki (2018) demonstram sua incorporação crescente a processos organizacionais de comunicação, análise e tomada de decisão. O fenômeno deixa de restringir-se à substituição de tarefas rotineiras e passa a incidir sobre o núcleo cognitivo do trabalho gerencial.

Consoante a isso, Brynjolfsson, Li e Raymond (2025) identificam ganhos de produtividade decorrentes do uso de sistemas generativos, com efeitos mais pronunciados entre trabalhadores menos experientes, ao passo que Johri, Schleiss e Ranade (2025) evidenciam que a apropriação da IAG varia conforme função, expertise e contexto organizacional. Esses achados deslocam a análise da tecnologia em si para as práticas de uso, nas quais a atuação gerencial se define pela mediação entre julgamento humano e respostas geradas por sistemas.

Nesse cenário, compreender a IAG apenas como automação mostra-se insuficiente. Como sugerem Johri, Schleiss e Ranade (2025), a interação com a tecnologia é heterogênea e situada. Em práticas profissionais, gestores formulam comandos, contextualizam, interpretam e validam respostas, configurando uma dinâmica de co-criação humano-IA, na qual os resultados emergem da interação contínua entre humano e sistema. Embora o conceito de co-criação derive de Prahalad e Ramaswamy (2000; 2004), que o definem como processo de geração conjunta de valor mediante a interação ativa entre agentes, deslocando a criação de valor da produção unilateral para a experiência interativa, este estudo o mobiliza como lente analítica para compreender práticas de trabalho mediadas por IAG. Woolley (2025) reforça essa perspectiva ao caracterizar a IAG como tecnologia de coordenação que amplia capacidades cognitivas no trabalho colaborativo.

A confiança na IA emerge, nesse contexto, como dimensão central. Cvetkovic et al. (2026) demonstram que a confiança depende de fatores humanos e contextuais, influenciando diretamente adoção e uso, enquanto Daly, Wiewiora e Hearn (2025) a concebem como processo dinâmico associado à percepção de confiabilidade e às

condições de uso. Assim, a confiança atua como mecanismo que media a co-criação, condicionando a interpretação, a validação e a incorporação das respostas geradas.

Essa relação, contudo, não se manifesta de forma homogênea. Ikkatai et al. (2024) mostram que percepções e preocupações em relação à IA variam entre países, e Cvetkovic et al. (2026) evidenciam que seus determinantes são contextuais. Tais achados permitem compreender a confiança como fenômeno situado, condicionado por fatores institucionais, culturais e tecnológicos.

A escolha da China como contexto comparativo fundamenta-se em sua posição como polo global de desenvolvimento e regulação da IA. Zou e Zhang (2025) descrevem forte atuação estatal e implementação de regulações específicas para sistemas generativos, enquanto Wang et al. (2025) destacam a centralidade do Estado e a articulação entre políticas públicas e desenvolvimento tecnológico. Esse arranjo institucional distinto sugere que a relação entre tecnologia e sociedade assume configurações específicas, sendo plausível inferir, à luz de Cvetkovic et al. (2026), que tais diferenças influenciam a construção da confiança e, por consequência, as formas de uso da IAG. Em contraste, no Brasil, a heterogeneidade na adoção da IA, associada a variações nas capacidades organizacionais e tecnológicas, indica que a confiança tende a ser construída de forma mais situada e dependente da experiência de uso (Ribeiro; Segatto, 2025).

A comparação entre os dois países ganha densidade adicional quando se considera que Brasil e China integram o mesmo arranjo multilateral de cooperação em inteligência artificial. O relatório do *Center for Global AI Innovative Governance*, da Universidade Fudan, documenta que os países do BRICS elevaram a IA a tema prioritário de suas agendas, com iniciativas como a Declaração de Kazan, de 2024, e o Centro China-BRICS de Desenvolvimento e Cooperação em Inteligência Artificial, ao mesmo tempo em que preservam estratégias nacionais e modelos regulatórios distintos (Jiang et al., 2026). Nesse quadro, Brasil e China constituem polos institucionais contrastantes de um mesmo espaço de coordenação, o que reforça a pertinência analítica da comparação: examina-se como dois países expostos a agendas convergentes de desenvolvimento tecnológico produzem condições institucionais distintas para a construção da confiança na IA e, por consequência, para a co-criação entre gestores e sistemas generativos.

No contexto brasileiro, Mangoni (2025) demonstra que a co-criação humano-IA se estrutura a partir da validação crítica das respostas da IA, dependente do repertório e

do senso crítico dos gestores. A adoção da tecnologia no país ocorre de forma heterogênea entre organizações, associada a diferentes níveis de capacidade em tecnologia da informação, infraestrutura digital e institucionalização de práticas tecnológicas (Ribeiro; Segatto, 2025). Contudo, por limitar-se ao contexto nacional, aquele estudo não contempla variações interculturais do fenômeno.

Mangoni (2025) investigou como gestores utilizam a Inteligência Artificial Generativa conversacional em suas rotinas profissionais e quais padrões de interação e cocriação emergem dessas práticas. Seus resultados identificam a cocriação humano-IA como fenômeno central e evidenciam que sua realização envolve processos de refinamento iterativo, validação crítica e mobilização do repertório profissional do gestor. A confiança também aparece no modelo empírico do estudo entre as condições que intervêm na relação com a tecnologia. Contudo, não constitui sua questão teórica central nem é examinada sistematicamente como processo de formação, erosão e recalibração ao longo das interações.

Permanece, portanto, em aberto compreender como a confiança é construída e mobilizada em episódios concretos de interação com sistemas generativos, como experiências de acerto e falha modificam as decisões de aceitar, verificar, transformar ou rejeitar seus *outputs* e em que medida esses mecanismos se configuram de forma distinta em diferentes contextos. Esta pesquisa avança essa agenda ao colocar a confiança no centro da explicação processual da cocriação humano-IA e ao examinar comparativamente gestores no Brasil e na China, considerando de forma analiticamente distinta as condições culturais, institucionais, organizacionais e experienciais que podem participar da construção dessa relação.

A partir dessa lacuna, este estudo posiciona-se como continuidade analítica de Mangoni (2025), mantendo a co-criação humano-IA como fenômeno central e incorporando a dimensão paralela entre contextos. Especificamente, desloca o foco da adoção tecnológica para a mediação da confiança na IA e examina como essa confiança influencia a co-criação entre gestores e sistemas generativos em dois contextos distintos, Brasil e China, integrando co-criação, confiança e governança em perspectiva comparativa.

1.1 PROBLEMA DE PESQUISA

Como a confiança na IA, em perspectiva intercultural, influencia a co-criação entre gestores e Inteligência Artificial Generativa no Brasil e na China?

1.2 OBJETIVOS

1.2.1 Objetivo geral:

- Analisar como a confiança na IA influencia a co-criação entre gestores e Inteligência Artificial Generativa em rotinas de trabalho, considerando a comparação entre os contextos brasileiro e chinês.

1.2.2 Objetivos específicos:

- a) Mapear a incorporação da IAG nas rotinas de gestores nos contextos brasileiro e chinês.
- b) Examinar como a confiança na IA é construída, percebida e mobilizada nessas práticas.
- c) Caracterizar os processos de interação e co-criação entre gestores e sistemas de IAG.
- d) Comparar similaridades e diferenças entre Brasil e China quanto à confiança e a seus efeitos na co-criação.

1.3 JUSTIFICATIVA

Do ponto de vista teórico, esta pesquisa insere-se na literatura sobre os efeitos da Inteligência Artificial Generativa nas práticas organizacionais. Brynjolfsson e McAfee (2014) discutem a reconfiguração do trabalho decorrente do avanço tecnológico, enquanto Davenport e Ronanki (2018) analisam a incorporação da IA aos processos organizacionais. Mais recentemente, Brynjolfsson, Li e Raymond (2025) identificam efeitos sobre a produtividade, e Johri, Schleiss e Ranade (2025) demonstram que a apropriação da IAG varia conforme as características dos usuários, das atividades e dos contextos de uso. Este estudo avança nessa discussão ao examinar a confiança como dimensão central das práticas de co-criação entre gestores e sistemas generativos.

Essa relação não é compreendida como simples adoção ou substituição tecnológica. A perspectiva da *Sociomateriality* permite analisar tecnologia e prática como elementos constituídos conjuntamente nas atividades organizacionais

(ORLIKOWSKI, 2007). De forma complementar, Leonardi (2011) demonstra que os efeitos das tecnologias emergem da interação entre suas propriedades materiais e as ações dos indivíduos. Sob essa perspectiva, a incorporação da IAG envolve um processo contínuo no qual gestores adaptam suas práticas às possibilidades percebidas da tecnologia e, simultaneamente, definem como suas capacidades foram integradas ao trabalho. A confiança torna-se relevante nesse processo porque participa das decisões de validar, modificar, incorporar ou rejeitar as contribuições produzidas pela IA.

Embora a confiança em IA constitua um campo de pesquisa amplo, sua produção apresenta assimetrias importantes. Benk et al. (2025), ao examinarem estudos publicados entre 2000 e 2023, identificaram 1.156 trabalhos empíricos centrais e uma base de 36.306 referências citadas; entretanto, 80,5% dos artigos do corpus principal originam-se de países do Norte Global. Dang e Li (2026), em revisão sistemática de 562 estudos empíricos quantitativos, também identificam concentração geográfica da produção e menor representação da América do Sul. Além disso, os autores caracterizam a investigação intercultural sobre confiança em IA como fragmentada e apontam a necessidade de compreender sua formação dinâmica ao longo das interações, para além de abordagens centradas predominantemente na mensuração de antecedentes, consequências e níveis de confiança.

É nessa interseção que se situa a lacuna desta pesquisa. Na estratégia de busca realizada para este estudo, não foram identificados estudos *peer-reviewed* diretamente dedicados à comparação entre Brasil e China sobre como gestores constroem e mobilizam confiança em processos de cocriação com IAG. Esse resultado não permite afirmar a inexistência de pesquisas sobre o tema, mas evidencia a limitada presença desse recorte na literatura recuperada. A investigação proposta concentra-se, assim, nos processos pelos quais experiências de interação podem fortalecer, reduzir ou recalibrar a confiança e em como essas mudanças se relacionam às práticas de validação, delegação e incorporação das contribuições da IAG.

Do ponto de vista prático, a crescente incorporação da IAG ao trabalho gerencial torna relevante compreender não apenas sua adoção, mas as condições sob as quais os gestores efetivamente delegam atividades, utilizam seus resultados e mantêm mecanismos de supervisão. Já em 2024, o *Work Trend Index*, elaborado pela Microsoft e pelo LinkedIn a partir de uma pesquisa com 31 mil profissionais em 31 países, identificou que 75% dos trabalhadores do conhecimento já utilizavam IA no trabalho, enquanto 79% dos líderes consideravam sua adoção necessária para manter a

competitividade organizacional (MICROSOFT; LINKEDIN, 2024). Ao mesmo tempo, 60% dos líderes indicaram ausência de visão e planejamento organizacional adequados para sua implementação, evidenciando uma distância entre adoção e institucionalização da tecnologia. Esse cenário amplia a necessidade de compreender os critérios empregados pelos profissionais para decidir quando utilizar, verificar, modificar ou rejeitar respostas produzidas por sistemas generativos. No ano seguinte, Mangoni (2025) identificou a relevância do julgamento gerencial na validação dos resultados produzidos pela IA, aspecto que dialoga com a discussão de Woolley (2025) sobre o potencial dessas tecnologias como mecanismos de coordenação cognitiva. Nesse sentido, compreender como a confiança é construída e mobilizada pode oferecer subsídios para a definição de práticas de uso nas quais contribuição algorítmica, supervisão e julgamento humano sejam articulados conforme as características e a criticidade das atividades.

Do ponto de vista social e gerencial, a expansão da IAG sobre atividades cognitivas amplia os efeitos da tecnologia para além de ganhos de eficiência, alcançando questões relacionadas à autonomia profissional, à responsabilização pelas decisões, à qualificação dos trabalhadores e à distribuição do julgamento entre pessoas e sistemas. Estudo da Organização Internacional do Trabalho e do *National Research Institute* da Polônia estima que um em cada quatro trabalhadores no mundo está em ocupações com algum grau de exposição à IAG e ressalta que, diante da permanência da contribuição humana em grande parte dessas atividades, a transformação do trabalho constitui resultado mais provável do que sua substituição integral (GMYREK et al., 2025). Nesse contexto, compreender como gestores constroem e recalibram sua confiança na tecnologia torna-se relevante para analisar como se redefinem os limites entre delegação algorítmica, supervisão e responsabilidade humana. Evidências de que os benefícios da IAG variam conforme experiência, função e capacidade de interação com a tecnologia reforçam que seus efeitos não são uniformes (BRYNJOLFSSON; LI; RAYMOND, 2025; JOHRI; SCHLEISS; RANADE, 2025). Investigar a confiança nesse contexto permite compreender um mecanismo pelo qual gestores estabelecem diferentes níveis de delegação, controle e validação, contribuindo para o debate sobre formas de integração da IAG ao trabalho que preservem a responsabilidade e o julgamento humano.

No eixo institucional e global, a comparação entre Brasil e China permite examinar esses processos em contextos institucionais e tecnológicos distintos, sem

pressupor que diferenças nacionais correspondam diretamente a diferenças culturais. Ikkatai et al. (2024) e Cvetkovic et al. (2026) demonstram que percepções e confiança em IA variam entre contextos, enquanto Zou e Zhang (2025) e Wang et al. (2025) destacam características específicas da governança e do desenvolvimento da IA na China. Nesse sentido, a comparação busca investigar empiricamente como condições institucionais, organizacionais e socioculturais são mobilizadas na construção da confiança e nas práticas de co-criação, em vez de assumir previamente a direção de sua influência.

Essa discussão também se conecta ao debate sobre a apropriação da IA no Sul Global. Jiang et al. (2026) argumentam que, nas economias do BRICS, a conversão das capacidades da IA em benefícios de desenvolvimento depende das condições institucionais em que sua adoção ocorre. Ao investigar gestores no Brasil e na China, esta pesquisa desloca essa discussão para o nível das práticas organizacionais, examinando como condições contextuais se manifestam nas interações cotidianas entre profissionais e sistemas generativos. Brasil e China participam de um mesmo espaço de cooperação multilateral em inteligência artificial, ao mesmo tempo em que apresentam configurações institucionais, regulatórias e tecnológicas distintas, o que permite articular o debate global sobre desenvolvimento e governança da IA às práticas organizacionais observadas nos dois contextos.

Diante dessas dimensões, esta pesquisa busca contribuir para preencher simultaneamente uma lacuna teórica, empírica e contextual. Ao colocar a confiança no centro da explicação da cocriação humano-IA, o estudo procura ampliar a compreensão de sua formação dinâmica nas práticas gerenciais; ao examinar os processos de delegação, validação e incorporação das respostas da IA, busca produzir conhecimento aplicável à gestão de sua participação no trabalho; e, ao comparar Brasil e China, amplia a presença de contextos do Sul Global em uma literatura ainda geograficamente concentrada no Norte Global (BENK et al., 2025; DANG; LI, 2026). Dessa forma, a pesquisa articula contribuições teóricas, práticas, sociais, gerenciais e institucionais para compreender não apenas se gestores confiam na IA, mas como essa confiança é construída, recalibrada e mobilizada na definição dos limites da participação da tecnologia no trabalho gerencial.

2 FUNDAMENTAÇÃO TEÓRICA

A fundamentação teórica deste estudo estrutura o arcabouço conceitual para analisar a co-criação entre gestores e IAG no contexto organizacional. Partindo da premissa de que a IAG não se limita à adoção tecnológica, mas reconfigura práticas cognitivas e decisórias, a seção organiza-se em torno da inserção da IA no trabalho, da co-criação humano-IA, da confiança como variável central e de sua dimensão intercultural.

Em continuidade analítica ao estudo de Mangoni (2025), que evidencia o uso da IAG como suporte à comunicação, ao aprendizado e à decisão, este trabalho desloca o foco da adoção para a dinâmica relacional entre humano e tecnologia. A co-criação é compreendida como processo iterativo de interação contínua, no qual gestores e sistemas de IA co-produzem soluções mediante ciclos de geração, avaliação e refinamento.

Nesse contexto, a confiança na IA emerge como elemento estruturante ao mediar decisões de uso, validação e incorporação dos *outputs*, ao mesmo tempo em que se reconhece seu caráter contextual, condicionado por fatores institucionais, culturais e tecnológicos. A dimensão intercultural, operacionalizada pela comparação entre Brasil e China, é incorporada como eixo analítico para examinar como diferentes arranjos de governança moldam a confiança e, consequentemente, a co-criação humano-IA. A seção desenvolve-se de forma progressiva, articulando IAG no trabalho, co-criação, confiança e contexto intercultural, com vistas à construção do modelo analítico que relaciona contexto, confiança e co-criação.

2.1 INTELIGÊNCIA ARTIFICIAL GENERATIVA NO CONTEXTO DO TRABALHO

A inserção da IAG no trabalho pode ser compreendida como desdobramento da transformação digital das organizações, na medida em que sistemas computacionais passam a participar de atividades antes dependentes de capacidades humanas complexas. Brynjolfsson e McAfee (2014) sustentam que o avanço tecnológico reconfigura o conteúdo do trabalho ao deslocar para sistemas digitais tarefas cognitivas antes executadas por pessoas. Em convergência, Davenport e Ronanki (2018) mostram que a IA já vinha sendo incorporada a processos empresariais voltados à automação, à geração de insights e à interação com usuários. No caso da IAG, esse movimento adquire novo alcance, porque a tecnologia deixa de atuar apenas sobre tarefas rigidamente estruturadas e passa a operar diretamente sobre linguagem, interpretação e

produção de conteúdo, ampliando sua incidência sobre atividades analíticas e comunicacionais do trabalho organizacional.

Essa mudança de perspectiva é particularmente relevante, uma vez que desloca o debate de uma concepção estrita de automação para uma discussão sobre reestruturação do trabalho. Brynjolfsson, Li e Raymond (2025) mostram que a adoção de um sistema generativo em larga escala elevou a produtividade de agentes de suporte, com efeitos mais intensos entre trabalhadores menos experientes, o que sugere que a tecnologia não apenas executa tarefas, mas redistribui conhecimento e acelera curvas de aprendizagem. Os autores ressaltam que ferramentas generativas diferem de formas anteriores de computação por conseguirem atuar em tarefas não rotineiras apoiadas em conhecimento tácito, incluindo comunicação profissional, análise e formulação textual, ampliando o escopo de complementaridade entre humanos e sistemas. Assim, a IAG não deve ser tratada apenas como instrumento de substituição, mas como tecnologia capaz de alterar a organização do trabalho, os critérios de desempenho e a distribuição de capacidades no interior das firmas.

No campo da Administração, essa mudança é especialmente expressiva porque incide sobre atividades centrais da prática gerencial. A incorporação da IAG aparece vinculada a planejamento, comunicação, organização do trabalho e suporte à decisão, o que desloca a atuação do gestor de uma posição exclusivamente executiva ou supervisora para uma função de mediação entre julgamento humano e respostas produzidas por sistemas generativos. Essa formulação é compatível com os achados de Mangoni (2025), cujo estudo identifica que gestores recorrem à IAG para agilizar tarefas operacionais, apoiar o próprio aprendizado, qualificar a comunicação e subsidiar decisões de caráter estratégico, ao mesmo tempo em que dependem de repertório, experiência e senso crítico para validar os resultados produzidos pela tecnologia. Mangoni (2025) evidencia empiricamente que o uso gerencial da IA não se reduz ao consumo passivo de respostas, mas envolve apropriação situada, validação e avaliação crítica.

Essa característica impede que o impacto da IAG seja interpretado de forma homogênea. Johri, Schleiss e Ranade (2025) mostram que o uso de IAG varia amplamente entre funções, níveis de domínio técnico e tipos de atividade, havendo desde usos mais superficiais, com ferramentas prontas, até integrações mais sofisticadas em fluxos de trabalho complexos. Os autores sustentam uma perspectiva de *augmentation*, segundo a qual o interesse analítico central não está na automação plena,

mas em como a tecnologia amplia experiências humanas e resultados do trabalho em contextos concretos. A partir dessa evidência, é plausível inferir que os efeitos da IAG não decorrem automaticamente de sua disponibilidade técnica, mas das formas pelas quais usuários e organizações a incorporam, adaptam e legitimam no cotidiano profissional.

Essa incorporação pode ser compreendida diante da *Affordance Theory*, segundo a qual os efeitos de uma tecnologia não decorrem exclusivamente de suas características técnicas, mas das possibilidades de ação que os usuários percebem e efetivamente mobilizam em contextos específicos. Nessa ótica, uma mesma tecnologia pode assumir funções distintas conforme as capacidades, objetivos, experiências e práticas dos indivíduos que a utilizam, de modo que sua incorporação ao trabalho resulta da interação entre propriedades tecnológicas e interpretações humanas.

No contexto da IAG, essa perspectiva permite compreender por que gestores utilizam sistemas semelhantes para finalidades distintas, como apoio à comunicação, geração de ideias, aprendizagem, análise de informação ou suporte à decisão. Assim, a incorporação da IAG não depende apenas da disponibilidade da tecnologia, mas das *affordances* percebidas pelos gestores, isto é, das possibilidades de uso que conseguem identificar, experimentar e integrar às suas rotinas de trabalho. Essa abordagem contribui para compreender que diferenças na incorporação da IA podem decorrer não apenas de fatores tecnológicos, mas também da forma como indivíduos e organizações percebem e exploram suas potencialidades.

Sob essa perspectiva, a IAG pode ser definida, para fins deste estudo, como tecnologia sociotécnica de uso organizacional que reconfigura práticas de trabalho ao combinar geração de conteúdo, apoio cognitivo e mediação interacional. Essa formulação não corresponde a uma definição literal presente em uma única fonte, mas constitui inferência analítica construída a partir de três elementos articulados: a participação crescente da IA em processos organizacionais (Davenport; Ronanki, 2018), sua capacidade de elevar produtividade e alterar a distribuição de expertise no trabalho (Brynjolfsson; Li; Raymond, 2025) e sua apropriação heterogênea em contextos profissionais distintos, em chave de *augmentation* (Johri; Schleiss; Ranade, 2025).

É precisamente nesse ponto que a continuidade analítica com Mangoni (2025) se torna teoricamente produtiva. Se o estudo anterior demonstrou que, no contexto brasileiro, a IAG já se insere na rotina de gestores como apoio à comunicação, ao aprendizado e à decisão, e que seu uso exige validação crítica, o presente trabalho

avança ao tratar essa inserção não apenas como adoção tecnológica, mas como base para compreender relações de co-criação mediadas pela confiança. Esta seção delimita, portanto, o primeiro eixo da fundamentação: a IAG interessa ao campo da Administração não por representar apenas nova ferramenta digital, mas por alterar a gramática do trabalho gerencial, criando condições para formas mais complexas de interação entre gestores e sistemas generativos. Com base nesse enquadramento, a seção seguinte desenvolve a co-criação entre humanos e IA como lente analítica do fenômeno.

2.2 CO-CRIAÇÃO ENTRE HUMANOS E INTELIGÊNCIA ARTIFICIAL GENERATIVA

A noção de co-criação, conforme desenvolvida por Prahalad e Ramaswamy (2004), refere-se à geração conjunta de valor por meio da interação ativa entre múltiplos agentes. Nessa concepção, o valor não é produzido de forma unilateral, mas emerge de processos interativos caracterizados por diálogo, transparência e participação, o que desloca o foco da produção para a interação e estabelece a co-criação como fenômeno relacional e processual. Ao transpor esse conceito para o contexto da IAG, contudo, torna-se necessário ampliar sua fundamentação, uma vez que a formulação original está centrada predominantemente na interação entre atores humanos e organizacionais, enquanto a co-criação humano-IA envolve também a integração de capacidades tecnológicas nos processos de geração de valor.

Essa compreensão foi ampliada pela *Service-Dominant Logic* (S-D Logic), proposta por Vargo e Lusch (2004), segundo a qual a criação de valor ocorre por meio da integração de recursos mobilizados pelos diferentes atores envolvidos em determinado contexto. Em desenvolvimentos posteriores, essa perspectiva passou a enfatizar que a cocriação ocorre em ecossistemas de serviço estruturados por relações, instituições e processos contínuos de integração de recursos (VARGO; LUSCH, 2016). Essa abordagem amplia a formulação clássica de cocriação ao deslocar a análise da interação específica entre organização e consumidor para configurações mais amplas nas quais diferentes recursos e capacidades são articulados na geração de valor.

No contexto digital, essa perspectiva torna-se particularmente relevante porque tecnologias passam a integrar os recursos mobilizados nos processos de criação de valor. Lusch e Nambisan (2015) mostram que a inovação em ecossistemas digitais depende da recombinação contínua de recursos distribuídos entre diferentes atores e

elementos tecnológicos. Assim, a análise da cocriação humano-IA pode ser situada em uma trajetória conceitual que parte da interação como fundamento da criação conjunta de valor, avança para a integração de recursos proposta pela *S-D Logic* e alcança contextos digitais nos quais capacidades humanas e tecnológicas participam conjuntamente da produção dos resultados.

A inteligência artificial, especialmente em sua vertente generativa, apresenta características que a diferenciam de tecnologias tradicionais. Trata-se de sistemas capazes de aprender com dados, reconhecer padrões e gerar respostas que simulam competências cognitivas humanas, operando dentro de parâmetros definidos por seus desenvolvedores. Essa capacidade de geração e adaptação introduz uma mudança qualitativa na interação, uma vez que a IA deixa de funcionar apenas como suporte passivo e passa a contribuir diretamente para a geração e o refinamento dos resultados produzidos ao longo do processo. Nesse sentido, a interação entre humano e IA assume contornos de coprodução cognitiva, cuja caracterização é desenvolvida ao longo desta seção.

Na interação com sistemas generativos, essa integração de recursos assume caráter iterativo, estruturando-se por ciclos nos quais o humano formula entradas, a IA gera saídas e o humano avalia, ajusta e redefine os parâmetros da interação. Estudos empíricos indicam que o uso de IAG no trabalho ocorre de maneira experimental e progressiva, com usuários refinando suas instruções ao longo do tempo para melhorar a qualidade dos resultados (Mangoni, 2025). Essa dinâmica evidencia que a co-criação não é evento pontual, mas fluxo contínuo de ajustes sucessivos, revelando que a qualidade do resultado depende não apenas da capacidade técnica da IA, mas da habilidade do usuário em interpretar, questionar e orientar o sistema.

É no decorrer dessas interações que a confiança na IA se torna relevante para o processo de cocriação. A interação iterativa pressupõe decisões constantes por parte do usuário, como aceitar, modificar ou rejeitar as respostas geradas, decisões mediadas pelo nível de confiança atribuído ao sistema. Yang e Wibowo (2022) sustentam que a confiança em IA influencia diretamente sua aceitação e uso, enquanto Daly, Wiewiora e Hearn (2025) a caracterizam como processo dinâmico associado às condições de uso e à experiência organizacional. Além disso, a confiança pode aumentar ou diminuir em função do desempenho percebido da IA ao longo das interações, o que a conecta diretamente à lógica iterativa da co-criação.

A relação entre humano e IA pode ser interpretada sob a lógica da inteligência híbrida, na qual capacidades humanas e computacionais são combinadas de forma complementar. Brynjolfsson, Li e Raymond (2025) mostram que sistemas de IAG podem ampliar o desempenho de trabalhadores ao disponibilizar padrões e soluções derivados de grandes volumes de dados, especialmente para usuários menos experientes. Contudo, esses ganhos não são homogêneos e dependem da capacidade do usuário de interagir criticamente com o sistema, o que indica que a co-criação não elimina o papel humano, mas o remodela, deslocando-o para funções de direção, interpretação e validação.

Sob essa ótica, a IA pode ser compreendida como parceiro cognitivo complementar, cuja atuação depende da mediação humana. Conforme Mangoni (2025), gestores utilizam a IA para apoiar processos de ideação, análise e decisão, mas mantêm a responsabilidade pela validação e pelo direcionamento das respostas. Esse achado reforça que a co-criação humano-IA é marcada por assimetria funcional: enquanto a IA contribui com capacidade de processamento e geração, o humano permanece responsável pela atribuição de sentido e pela avaliação crítica. A efetividade dessa relação está, portanto, diretamente condicionada ao grau de confiança do usuário na tecnologia, que orienta o nível de delegação cognitiva atribuído à IA.

Ao mesmo tempo, essa confiança não se estabelece de forma universal. Ikkatai et al. (2024) mostram que percepções, atitudes e níveis de preocupação em relação à IA variam entre países, enquanto Cvetkovic et al. (2026) evidenciam que a confiança é influenciada por fatores culturais, institucionais e pelos níveis de familiaridade com a tecnologia, associando autoeficácia em IA e atitudes positivas à confiança, sem homogeneidade entre diferentes contextos nacionais. Isso sugere que a co-criação humano-IA deve ser compreendida como fenômeno situado, no qual a forma e a intensidade da interação são mediadas por condições contextuais que influenciam a confiança na tecnologia.

Dessa forma, a cocriação entre humanos e IAG pode ser definida como um processo iterativo de integração de recursos e coprodução cognitiva, estruturado por interações contínuas entre humanos e sistemas generativos, no qual a geração de valor decorre da articulação entre capacidades humanas e computacionais. Nesse processo, a IAG contribui para a geração e o refinamento de conteúdos e alternativas, enquanto o gestor mobiliza conhecimento contextual, julgamento e capacidade de validação. Esse processo é orientado pelo julgamento humano e condicionado pelo nível de confiança

na IA, que influencia diretamente as decisões de uso, validação e incorporação das respostas geradas. Ao reconhecer a natureza dinâmica e contextual dessa confiança, abre-se espaço para sua análise como variável central, especialmente em perspectiva intercultural, como proposto neste estudo.

2.3 CONFIANÇA NA INTELIGÊNCIA ARTIFICIAL GENERATIVA

A confiança constitui construto central nas ciências sociais, compreendida como a disposição de um agente em aceitar vulnerabilidade com base em expectativas positivas sobre o comportamento de outro (RHEU et al., 2021). No contexto tecnológico, essa definição é reconfigurada, uma vez que o objeto da confiança deixa de ser exclusivamente humano e passa a incluir sistemas artificiais. Nessa perspectiva, a confiança em tecnologia refere-se à crença de que um sistema irá operar de forma confiável, previsível e alinhada aos interesses do usuário, reduzindo incertezas associadas ao seu uso (YANG; WIBOWO, 2022). Essa dimensão torna-se particularmente relevante no caso da inteligência artificial, cuja adoção envolve níveis elevados de complexidade, autonomia e opacidade.

A especificidade da confiança em sistemas de inteligência artificial decorre, sobretudo, de sua limitada transparência operacional. Sistemas baseados em IA operam por meio de modelos complexos, frequentemente inacessíveis à compreensão direta do usuário, o que dificulta a interpretação dos processos que levam à geração de respostas. Essa característica compromete a previsibilidade do sistema e desloca a base da confiança da compreensão para a experiência. Nesse sentido, a consistência de desempenho e a qualidade das respostas passam a atuar como elementos centrais para a formação da confiança, uma vez que permitem ao usuário construir expectativas sobre o comportamento futuro da tecnologia (NORDHEIM et al., 2019; SHERIDAN, 2019; YANG; WIBOWO, 2022).

A confiança em IA apresenta natureza multidimensional, sendo estruturada a partir de diferentes critérios de avaliação. Um primeiro conjunto de abordagens a define com base em estágios progressivos de formação, envolvendo previsibilidade, dependabilidade e fé, nos quais a repetição consistente de desempenho favorece a consolidação da confiança (RHEU et al., 2021; SHERIDAN, 2019). Em paralelo, outra vertente conceitual enfatiza dimensões como habilidade, benevolência e integridade, indicando que os usuários avaliam simultaneamente a competência técnica do sistema, o

alinhamento de seus resultados com os interesses do usuário e sua conformidade com normas éticas (THIEBES et al., 2021; YANG; WIBOWO, 2022). Essas abordagens evidenciam que a confiança em IA não se restringe à performance técnica, incorporando julgamentos normativos e expectativas sociais.

Além disso, a confiança pode ser analisada a partir de diferentes objetos de referência, distinguindo-se entre confiança na tecnologia e confiança na organização que a desenvolve ou implementa. A confiança na tecnologia envolve percepções sobre desempenho, segurança e qualidade do sistema, enquanto a confiança na organização está associada à reputação, à responsabilidade no uso de dados e à conformidade com normas sociais (SÖLLNER et al., 2016; THIEBES et al., 2021). Essa distinção é relevante porque a aceitação de sistemas de IA depende não apenas de suas características técnicas, mas também das condições institucionais que sustentam sua operação.

A formação da confiança em IA também pode ser compreendida a partir de suas dimensões cognitivas, afetivas e comportamentais. A dimensão cognitiva relaciona-se à avaliação racional do sistema, considerando aspectos como desempenho, utilidade e risco percebido. A dimensão afetiva envolve predisposições emocionais e atitudes em relação à tecnologia, independentemente de experiência direta. A dimensão comportamental decorre de interações repetidas com o sistema ao longo do tempo, sendo diretamente influenciada pela experiência acumulada do usuário (GLIKSON; WOOLLEY, 2020; YANG; WIBOWO, 2022). Essa estrutura evidencia que a confiança é construída simultaneamente por avaliação racional, disposição psicológica e aprendizagem prática.

A natureza dinâmica da confiança decorre de sua dependência da experiência. Interações bem-sucedidas tendem a reforçar a confiança, ao passo que falhas, inconsistências ou resultados inesperados contribuem para sua erosão. Esse processo implica que a confiança não é condição estática, mas fenômeno que se ajusta continuamente à medida que o usuário interage com a tecnologia (YANG; WIBOWO, 2022). Tal característica torna a confiança particularmente relevante em contextos de uso contínuo e iterativo, como ocorre na interação com sistemas de IAG.

Nesse contexto, a confiança assume papel estruturante na co-criação entre humanos e IAG. A interação humano-IA é mediada por decisões constantes de aceitação, modificação ou rejeição das respostas geradas, decisões que dependem diretamente do nível de confiança atribuído ao sistema, o qual orienta o grau de

delegação cognitiva e a intensidade da intervenção humana. Níveis mais elevados de confiança tendem a ampliar a incorporação das respostas da IA no processo de trabalho, enquanto níveis reduzidos levam a maior verificação, ajuste e controle por parte do usuário. Dessa forma, a confiança não atua apenas como antecedente da adoção, mas como mecanismo que regula a forma e a profundidade da co-criação humano-IA.

A relação entre confiança e julgamento humano introduz tensão central nesse processo. A confiança é necessária para viabilizar o uso da IA, mas sua distribuição inadequada pode comprometer a qualidade das decisões. A atribuição excessiva de confiança pode conduzir à aceitação acrítica das respostas geradas, enquanto a desconfiança excessiva pode limitar o aproveitamento das capacidades da tecnologia. Nesse sentido, a efetividade da co-criação depende da capacidade do usuário de calibrar sua confiança, mantendo o julgamento crítico mesmo diante de sistemas com alto desempenho. No contexto gerencial, essa calibração é particularmente relevante, uma vez que decisões apoiadas por IA podem envolver consequências estratégicas e organizacionais (Mangoni, 2025).

Por fim, a confiança em inteligência artificial deve ser compreendida como fenômeno contextual, influenciado por fatores sociais, culturais e institucionais. Dimensões culturais, como individualismo, aversão à incerteza e orientação de longo prazo, afetam a disposição dos indivíduos em confiar em tecnologias, influenciando tanto a formação quanto a manifestação da confiança (HALLIKAINEN; LAUKKANEN, 2018; ROBINSON, 2020). Além disso, Cvetkovic et al. (2026) indicam que atitudes, níveis de familiaridade e experiências prévias com a tecnologia influenciam a confiança de maneira diferenciada entre países. Esse conjunto de fatores indica que a confiança na IA não se configura como construto universal, mas como fenômeno situado, cuja compreensão exige considerar as especificidades culturais e institucionais dos contextos de uso.

Além de seu caráter dinâmico, a confiança na Inteligência Artificial pode ser compreendida à luz da Paradox Theory, segundo a qual organizações e gestores convivem simultaneamente com demandas aparentemente contraditórias, mas que não podem ser eliminadas, apenas administradas (Lewis, 2000; Smith; Lewis, 2011; Schad et al., 2016). Nessa perspectiva, a interação entre gestores e sistemas de IAG não envolve escolhas dicotômicas entre confiar ou não confiar na tecnologia, mas a necessidade de equilibrar tensões permanentes inerentes ao seu uso. A confiança

torna-se, assim, um processo de calibração contínua, no qual diferentes exigências coexistem e precisam ser conciliadas no cotidiano organizacional.

No contexto da co-criação entre gestores e a IAG, essas tensões manifestam-se em diferentes dimensões. A busca por maior produtividade convive com o risco de dependência excessiva da tecnologia; a automação amplia a capacidade de processamento de informações, mas preserva a necessidade de julgamento humano; a rapidez na geração de respostas exige processos de validação e verificação; e a crescente autonomia dos sistemas de IA amplia a responsabilidade do gestor sobre as decisões tomadas com seu apoio. Diante desse enquadramento, a efetividade da co-criação não decorre da eliminação desses paradoxos, mas da capacidade do gestor de administrá-los por meio de níveis adequados de confiança, preservando o equilíbrio entre delegação cognitiva e controle crítico. Assim, a confiança deixa de representar apenas um antecedente da adoção tecnológica e passa a constituir um mecanismo de gestão das tensões inerentes à interação entre humanos e sistemas de IAG.

Neste cenário, a confiança na IAG pode ser definida como construto multidimensional, dinâmico e contextual, que expressa a expectativa de que o sistema atuará de forma competente, previsível e alinhada aos interesses do usuário. No âmbito da co-criação humano-IA, a confiança opera como variável central ao mediar uso, validação e incorporação das respostas geradas, influenciando diretamente o nível de colaboração entre humano e tecnologia e, por consequência, a legitimidade e a aplicabilidade de conteúdos suportados por IA em contextos organizacionais relevantes.

Ao reconhecer seu caráter contextual, torna-se possível avançar na análise de como diferentes ambientes culturais, como Brasil e China, moldam a construção e a operacionalização da confiança na IA, aspecto que orienta o desenvolvimento empírico deste estudo.

Embora parte relevante da literatura sobre confiança em IA privilegie abordagens quantitativas voltadas à mensuração de seus antecedentes, níveis e consequências, esse tipo de desenho oferece alcance limitado para compreender como a confiança é construída, mobilizada e revisada ao longo de experiências concretas de interação. Nesta pesquisa, o interesse não está em mensurar quanto gestores confiam na IAG, mas em compreender os processos interpretativos pelos quais avaliam respostas, atribuem significado a erros e acertos e decidem validar, modificar, incorporar ou rejeitar os resultados produzidos pelo sistema. A abordagem qualitativa mostra-se, portanto, epistemologicamente coerente com a natureza dinâmica e contextual do

fenômeno, ao permitir reconstruir essas experiências e identificar como critérios de confiança são produzidos e mobilizados nas práticas de uso. Em perspectiva intercultural, essa escolha também possibilita examinar como fatores culturais, institucionais, organizacionais e experienciais participam desses processos nos contextos brasileiro e chinês, sem reduzir a comparação à mensuração de diferenças agregadas entre países (DANG; LI, 2026).

2.4 DIMENSÃO INTERCULTURAL DA CONFIANÇA NA INTELIGÊNCIA ARTIFICIAL (BRASIL VS. CHINA)

A confiança na inteligência artificial deve ser compreendida como construto situado, cuja formação depende das condições institucionais, organizacionais e culturais que estruturam a interação entre humanos e tecnologia. No contexto organizacional, a confiança atua como mecanismo que viabiliza a adoção e o uso da IA, influenciando a disposição dos indivíduos em delegar tarefas, incorporar recomendações e interagir com sistemas automatizados (DALY; WIEWIORA; HEARN, 2025). Essa relação indica que a confiança não se limita a uma percepção individual, mas é moldada por arranjos institucionais que condicionam expectativas, percepções de risco e critérios de validação da tecnologia.

A dimensão cultural pode ser examinada por meio de valores e normas associados à autoridade, à autonomia, às relações coletivas e à incerteza. Hofstede (2001) propõe dimensões como distância de poder, individualismo/coletivismo e aversão à incerteza, enquanto o GLOBE distingue, entre outras, distância de poder, aversão à incerteza, coletivismo institucional e coletivismo intragrupo (HOUSE et al., 2004). Essas dimensões permitem analisar como diferentes contextos estruturam percepções de legitimidade, risco e previsibilidade, aspectos relacionados à formação da confiança.

Sua influência, contudo, não é determinística. Doney, Cannon e Mullen (1998) argumentam que valores e normas sociais podem afetar os processos de construção da confiança: a distância de poder pode influenciar o peso atribuído à autoridade; orientações individualistas ou coletivistas, a relevância da experiência individual, das normas do grupo e da validação social; e a relação com a incerteza, a busca por previsibilidade, controle e verificação. Neste estudo, essas dimensões são utilizadas como lentes analíticas, sem atribuir características uniformes a gestores brasileiros ou

chineses e considerando variações organizacionais, profissionais e individuais em cada contexto nacional (LEIDNER; KAYWORTH, 2006).

A governança da inteligência artificial constitui um dos principais vetores dessa configuração. A forma como os Estados estruturam regulação, incentivos e mecanismos de controle define o ambiente no qual a IA é desenvolvida e utilizada, influenciando diretamente a percepção de legitimidade e previsibilidade associada à tecnologia. No caso chinês, a governança da IA é caracterizada por abordagem *state-led*, na qual o desenvolvimento tecnológico é articulado a objetivos estratégicos nacionais, como crescimento econômico, segurança e estabilidade social (ZOU; ZHANG, 2025). Esse arranjo institucional não apenas orienta o desenvolvimento da tecnologia, mas também estabelece parâmetros normativos que estruturam a forma como a IA é percebida e utilizada nas organizações.

A regulação da IA na China opera por meio de instrumentos juridicamente vinculantes e de lógica adaptativa, baseada na introdução contínua e iterativa de normas específicas para diferentes aplicações tecnológicas. Na análise de Zou e Zhang (2025), esse modelo regulatório apresenta caráter reativo e incremental, no qual diferentes instrumentos são progressivamente formulados, ajustados e substituídos em resposta à evolução tecnológica e aos riscos emergentes, resultando em arcabouço normativo cada vez mais detalhado e orientado a aplicações específicas. No caso da IAG, as *Interim Measures for the Management of Generative AI Services* estabelecem obrigações explícitas para provedores, incluindo requisitos relativos à legalidade e qualidade dos dados de treinamento, à proteção de informações pessoais, à identificação de usuários, à garantia de segurança e estabilidade dos sistemas e ao monitoramento e reporte de conteúdos ilícitos (ZOU; ZHANG, 2025).

Essa ênfase na conformidade normativa e no controle de conteúdo evidencia a articulação entre governança tecnológica e objetivos mais amplos de regulação social. Nesse sentido, a combinação entre normatização formal, elevado grau de especificidade regulatória e adaptação contínua configura um modelo que busca simultaneamente promover o desenvolvimento tecnológico e mitigar riscos, ao mesmo tempo em que delimita responsabilidades ao longo da cadeia de valor e reduz ambiguidades regulatórias associadas ao uso da IA.

Esse ambiente institucional constitui uma das condições contextuais potencialmente relevantes para a forma como gestores interpretam e utilizam sistemas de IA. A presença de parâmetros regulatórios explícitos e de uma estratégia nacional

articulada pode constituir fonte de legitimidade e previsibilidade percebida em relação à tecnologia, na medida em que normas e estruturas institucionais participam da formação das expectativas sobre sua utilização (BUNDUCHI; SITAR-TĂUT; MICAN, 2025). Entretanto, a forma como esses elementos são interpretados e mobilizados pelos gestores permanece uma questão empírica, não sendo possível inferir, a priori, que maior estruturação regulatória corresponda necessariamente a níveis mais elevados de confiança na IA.

No nível gerencial, essa estrutura influencia a forma como a IA é incorporada aos processos decisórios. Sistemas de IA operam como *advisors* que produzem recomendações baseadas em análise de dados, deslocando o papel do gestor da geração de alternativas para a avaliação e seleção de soluções (KEDING; MEISSNER, 2021). A confiança atribuída a esses sistemas afeta diretamente o grau de adesão às recomendações apresentadas, uma vez que níveis mais elevados de confiança aumentam a probabilidade de incorporação das sugestões da IA nas decisões estratégicas, processo que pode ampliar a delegação cognitiva e reduzir o nível de intervenção humana na validação das respostas.

Ademais, Keding e Meissner (2021) evidenciam que níveis elevados de confiança em sistemas de IA podem gerar padrões de *overreliance*, nos quais gestores passam a atribuir maior peso às recomendações algorítmicas em detrimento de sua própria avaliação crítica. Esse fenômeno revela que a confiança não apenas viabiliza o uso da IA, mas também reconfigura o equilíbrio entre julgamento humano e automação, influenciando a forma como decisões são construídas no contexto organizacional.

Em contraste, o contexto brasileiro apresenta arranjo institucional de governança da inteligência artificial ainda em consolidação, marcado por fragmentação regulatória, baixa coordenação entre políticas e atores e limitações na capacidade de enforcement, o que reduz a efetividade das diretrizes normativas e impõe desafios à sua implementação, especialmente diante das restrições institucionais e dos custos associados à conformidade (SERRA; SCAFUTO; SERRA VIEIRA, 2025). Nesse cenário, a incorporação da tecnologia ocorre em ambiente de menor previsibilidade institucional, no qual a distância entre formulação normativa e capacidade de execução limita a efetividade prática da regulação.

Essa caracterização não implica ausência de institucionalização formal. Jiang et al. (2026) registram que o Brasil dispõe de instrumentos regulatórios consolidados na esfera de proteção de dados, como a Lei Geral de Proteção de Dados e a atuação da

Autoridade Nacional de Proteção de Dados, que passou a exigir o reporte de incidentes de segurança envolvendo algoritmos, além de política nacional dedicada, o Plano Brasileiro de Inteligência Artificial 2024-2028, com investimentos públicos expressivos em infraestrutura. A distinção relevante para este estudo não reside, portanto, na existência ou inexistência de normas, mas na distância entre a formalização regulatória e a capacidade de implementação e coordenação apontada por Serra, Scafuto e Serra Vieira (2025). É essa defasagem que desloca a construção da legitimidade da IA para o nível organizacional e reforça o argumento de que, no Brasil, a confiança tende a se formar de maneira situada e dependente da experiência de uso.

Nesse contexto, capacidades internas, arranjos de governança e infraestrutura tecnológica constituem condições relevantes para a incorporação organizacional da IA e podem participar da construção da confiança na tecnologia (RIBEIRO; SEGATTO, 2025). Entretanto, o peso atribuído a essas condições, em comparação com fatores culturais, institucionais e relacionados à experiência individual, deverá ser examinado empiricamente. Esse deslocamento institucional implica que critérios de validação, uso e controle da IA não são definidos de forma homogênea, mas emergem de práticas situadas, nas quais experiências acumuladas, avaliações de risco e competências técnicas condicionam a forma de incorporação da tecnologia, reforçando a heterogeneidade nos padrões de adoção e de confiança do usuário.

Além dos arranjos regulatórios, os dois contextos diferem quanto à configuração dos ecossistemas tecnológicos nos quais a confiança é exercida. Jiang et al. (2026) mostram que os países do BRICS ocupam posições distintas em relação à dependência de tecnologias ocidentais: o Brasil é classificado entre os países que buscam alternativas multilaterais para reduzir a dependência de provedores estrangeiros de computação em nuvem e de modelos, enquanto a China, alvo prioritário dos controles de exportação norte-americanos, orienta-se à construção de um ecossistema tecnológico próprio. Para os fins deste estudo, é plausível inferir que essa assimetria estrutural repercute sobre o objeto da confiança discutido na seção 2.3: gestores brasileiros interagem majoritariamente com sistemas fornecidos por corporações estrangeiras, cuja governança escapa ao ambiente institucional nacional, ao passo que gestores chineses tendem a utilizar sistemas de provedores domésticos submetidos à regulação estatal descrita por Zou e Zhang (2025). Nessas condições, a distinção entre confiança na tecnologia e confiança na organização provedora (SÖLLNER et al., 2016; THIEBES et

al., 2021) assume contornos específicos em cada contexto, constituindo dimensão adicional a ser examinada empiricamente na comparação entre os dois países.

A articulação entre essas dimensões permite compreender que diferenças entre Brasil e China podem incidir não apenas sobre a intensidade da confiança, mas sobre os mecanismos pelos quais ela é construída e mobilizada. Valores culturais podem influenciar a interpretação de autoridade, autonomia e incerteza; instituições e estruturas regulatórias podem fornecer referências de legitimidade e previsibilidade; práticas organizacionais podem estabelecer limites e critérios para o uso da tecnologia; e experiências individuais podem alterar progressivamente as expectativas sobre seu desempenho. A interação entre essas dimensões pode, por sua vez, refletir-se nas práticas de validação, delegação, ajuste e incorporação das respostas geradas pela IAG. A forma assumida por essas relações em cada contexto permanece aberta à investigação empírica.

Dessa forma, a comparação entre Brasil e China permite analisar como diferentes configurações culturais, institucionais, organizacionais e experienciais participam da construção e da mobilização da confiança na IAG e, conseqüentemente, dos processos de cocriação humano-IA. A dimensão intercultural não se reduz, portanto, à comparação entre regimes de governança, mas incorpora valores e normas que podem estruturar diferentes expectativas sobre autoridade, autonomia, coletividade e incerteza. Nesse enquadramento, o estudo busca identificar empiricamente quais dessas condições adquirem relevância nos contextos analisados e como se relacionam às práticas de validação, delegação e incorporação das contribuições da IAG.

2.5 SÍNTESE TEÓRICA E CONSTRUÇÃO DO MODELO ANALÍTICO

A fundamentação teórica desenvolvida nas seções anteriores permite estruturar o fenômeno investigado a partir da interação entre gestores e IAG no contexto organizacional. Conforme definido na seção 2.2, a co-criação humano-IA constitui processo iterativo de coprodução cognitiva entre gestores e sistemas de IA, estruturado por ciclos de formulação, geração e refinamento de *outputs*. Nesse processo, a IA atua como sistema de suporte à geração de alternativas e à análise de informações, enquanto o gestor mantém o papel de direção, interpretação e validação dos resultados, configurando relação de complementaridade funcional (PRAHALAD; RAMASWAMY, 2004; MANGONI, 2025; KEDING; MEISSNER, 2021).

A incorporação da IA nos processos de trabalho desloca o papel do gestor da produção direta de soluções para a avaliação e seleção de alternativas geradas com suporte tecnológico, especialmente em contextos decisórios complexos (KEDING; MEISSNER, 2021). Esse deslocamento evidencia que a co-criação não se limita ao uso instrumental da tecnologia, mas envolve a integração entre capacidades humanas e computacionais em processos contínuos de interação. Nesse contexto, a efetividade da co-criação depende não apenas das capacidades técnicas da IA, mas da forma como o gestor interage com o sistema, particularmente no que se refere à aceitação, modificação e validação das respostas geradas.

A confiança na inteligência artificial emerge, nesse cenário, como variável central para compreender a dinâmica dessa interação. Conforme discutido, a confiança constitui construto multidimensional e dinâmico que influencia a disposição dos indivíduos em utilizar sistemas de IA e incorporar suas recomendações (YANG; WIBOWO, 2022; DALY; WIEWIORA; HEARN, 2025). No contexto da co-criação, a confiança atua como mecanismo mediador que regula o comportamento do gestor ao longo do processo interativo, influenciando diretamente o grau de aceitação dos *outputs*, a intensidade da validação e o nível de delegação cognitiva atribuído à IA.

Essa mediação torna-se particularmente evidente no nível gerencial. A confiança orienta decisões relativas ao uso da IA, determinando se o gestor adota postura de maior dependência das recomendações algorítmicas ou mantém maior controle e intervenção sobre os resultados. Níveis mais elevados de confiança tendem a ampliar a incorporação das respostas da IA nos processos decisórios, enquanto níveis mais reduzidos implicam maior verificação, ajuste e questionamento das saídas do sistema. Ao mesmo tempo, a literatura evidencia que a confiança excessiva pode levar a padrões de *overreliance*, nos quais gestores atribuem peso desproporcional às recomendações da IA, reduzindo o exercício do julgamento crítico (KEDING; MEISSNER, 2021). Dessa forma, a confiança não apenas viabiliza a co-criação, mas redefine o equilíbrio entre autonomia da tecnologia e controle humano.

A formação da confiança, por sua vez, não ocorre de maneira homogênea, sendo condicionada por fatores contextuais que moldam a relação entre humanos e tecnologia. A dimensão intercultural introduzida neste estudo evidencia que a confiança em IA é influenciada por arranjos institucionais, estruturas de governança e ambientes tecnológicos específicos. No caso chinês, a presença de governança *state-led*, caracterizada por forte coordenação estatal e regulação ativa da inteligência artificial,

contribui para a construção de ambiente de maior previsibilidade e legitimidade da tecnologia (ZOU; ZHANG, 2025). Esse contexto pode influenciar a construção da confiança ao oferecer referências institucionais, regulatórias e tecnológicas para o uso da IAG. A direção dessa influência, contudo, não é assumida a priori: maior formalização pode tanto ampliar percepções de previsibilidade quanto intensificar cautelas, restrições e exigências de validação, questão que permanece aberta à investigação empírica.

Em contraste, o contexto brasileiro apresenta organização institucional menos centralizada, marcada por fragmentação regulatória, limitações de coordenação e restrições na capacidade de implementação das diretrizes normativas (SERRA; SCAFUTO; SERRA VIEIRA, 2025), na qual a confiança na IA tende a ser construída de forma mais distribuída, dependendo de experiências organizacionais, percepções individuais, capacidades tecnológicas e características específicas das aplicações (RIBEIRO; SEGATTO, 2025). Essa composição implica que a confiança, enquanto mecanismo mediador, assume formas distintas entre os contextos, influenciando a maneira como gestores interagem com sistemas de IA em suas práticas de trabalho.

A articulação entre esses elementos permite explicitar o modelo analítico proposto neste estudo, no qual a co-criação entre gestores e IAG constitui o fenômeno central, estruturado por interações nas quais o gestor exerce papel ativo de interpretação, validação e direcionamento. Nesse processo, a confiança na IA atua como mecanismo mediador que regula o comportamento gerencial, influenciando o uso da tecnologia, o grau de aceitação dos *outputs*, a intensidade da validação e o nível de delegação cognitiva. Adicionalmente, propõe-se que essa confiança é condicionada por fatores contextuais de natureza intercultural, especialmente aqueles relacionados à governança da IA, às estruturas institucionais e ao ambiente tecnológico, de modo que o contexto institucional e cultural molda a confiança, e esta regula a forma como gestores utilizam, validam e incorporam a IA em seus processos de trabalho.

No nível gerencial, esse modelo implica que diferenças contextuais entre Brasil e China produzem variações no comportamento dos gestores, especialmente no que se refere à delegação de decisões, à validação de *outputs* e à intensidade da interação com o sistema. Assim, a co-criação humano-IA não decorre apenas das capacidades tecnológicas disponíveis, mas da forma como a confiança é construída e operacionalizada em contextos específicos. Em síntese, o estudo propõe um modelo analítico no qual: (i) o contexto institucional e cultural atua como condição estruturante,

ao moldar percepções, expectativas e formas de regulação da tecnologia; (ii) a confiança na IA atua como mecanismo mediador, ao regular o comportamento gerencial, a intensidade da validação e o grau de delegação cognitiva; e (iii) a co-criação humano-IA constitui o resultado observável dessa mediação, manifestando-se em práticas concretas de interação, produção de conteúdo e apoio à decisão. A contribuição teórica deste estudo reside na explicitação de que a co-criação entre gestores e IAG resulta da confiança construída em contextos institucionais distintos, orientando a análise empírica ao examinar como esses contextos influenciam, por meio da confiança, a interação entre gestores e sistemas generativos.

A consolidação desse modelo analítico proposto permite compreender que a co-criação entre gestores e IAG transcende a adoção de uma tecnologia específica, configurando-se como processo contínuo de desenvolvimento de competências individuais, organizacionais e relacionais. Nesse contexto, este estudo compreende a competência estratégica híbrida como a capacidade de integrar, de forma articulada, competências humanas e recursos proporcionados pela IA nos processos de trabalho e de decisão. Em nível individual, essa competência pressupõe o desenvolvimento da literacia em Inteligência Artificial (*AI Literacy*), permitindo aos gestores compreender criticamente as potencialidades, limitações e implicações do uso da IAG. Em nível organizacional, apoia-se nas capacidades dinâmicas (*Dynamic Capabilities*), que possibilitam adaptar continuamente processos, rotinas e práticas diante das transformações tecnológicas (TEECE; PISANO; SHUEN, 1997; TEECE, 2007). No nível relacional, manifesta-se por meio de formas colaborativas de interação entre humanos e sistemas inteligentes, conforme discutido pela literatura *Human-AI Teaming*, na qual a geração de valor decorre da complementaridade entre capacidades humanas e computacionais (DELLERMANN et al., 2019; JARRAHI, 2018; NG et al., 2021). A partir desse entendimento, a confiança deixa de representar apenas um antecedente de adoção tecnológica e passa a atuar como mecanismo que viabiliza a construção dessa competência estratégica híbrida, ao regular a intensidade da colaboração, da delegação cognitiva e do julgamento crítico exercido pelos gestores em processos de co-criação com a Inteligência Artificial Generativa.

3 PROCEDIMENTOS METODOLÓGICOS

O presente estudo investiga a co-criação entre gestores e Inteligência Artificial Generativa (IAG), com foco na influência da confiança na IA em perspectiva intercultural, considerando comparativamente os contextos do Brasil e da China. A construção metodológica foi delineada de modo a assegurar coerência com a base teórica do trabalho, ao mesmo tempo em que explicita sua contribuição original na esteira do estudo de Mangoni (2025), que constitui sua referência empírica e analítica inicial.

Este estudo adapta a abordagem metodológica proposta por Mangoni (2025), reorientando-a para a análise da confiança na IA e de sua variação em contextos interculturais. Tal adaptação implica não apenas ajustes operacionais, mas redefinição do foco analítico, deslocando a centralidade do uso da IA para a compreensão da confiança como mecanismo mediador da co-criação humano-IA, bem como para a incorporação explícita da dimensão institucional e cultural na explicação do fenômeno.

3.1 NATUREZA DA PESQUISA

A pesquisa caracteriza-se como qualitativa, de natureza interpretativa e com delineamento exploratório. Essa escolha decorre da necessidade de compreender fenômeno relacional, dinâmico e ainda em consolidação teórica, no qual se articulam práticas gerenciais, percepções subjetivas e condicionantes contextuais. Conforme Creswell, Lopes e Silva (2021), o delineamento qualitativo mostra-se apropriado quando a investigação se volta à interpretação dos sentidos que os sujeitos conferem à própria experiência, sem dissociá-los do ambiente em que se produzem.

A abordagem qualitativa permite apreender os significados atribuídos pelos gestores às suas interações com sistemas de IAG, particularmente no que se refere à construção, à manutenção e à revisão da confiança na tecnologia. Essa dimensão é central, uma vez que a confiança não se configura como atributo estático, mas como processo situado, influenciado por experiências de uso, expectativas de desempenho e condições institucionais. Kaplan et al. (2023) conceituam a confiança em IA como fenômeno relacional, estruturado pela interação entre atributos do usuário, características do sistema e o contexto de uso, indicando que sua formação não pode ser compreendida de forma isolada, mas como resultado de configuração sociotécnica específica.

O caráter exploratório justifica-se pela limitação de estudos empíricos que articulem, de forma integrada, confiança na IA, co-criação humano-IA e variações interculturais em contextos organizacionais. Embora a literatura avance na discussão da adoção tecnológica, ainda são incipientes as análises que capturam a dimensão prática e contextual dessas interações sob perspectiva comparativa (Gil, 2002).

Outrossim, o estudo assume caráter comparativo, ao analisar dois contextos nacionais distintos. Essa escolha ancora-se na compreensão de que a confiança na IA é condicionada por estruturas institucionais, regimes regulatórios e padrões culturais que influenciam tanto as percepções quanto os comportamentos dos usuários. Diferenças entre países na forma de governança da IA e nas prioridades políticas associadas ao seu desenvolvimento evidenciam a relevância dessa dimensão para a análise do fenômeno.

Quadro 1 - Articulação entre problema de pesquisa, objetivos específicos e dimensões analíticas

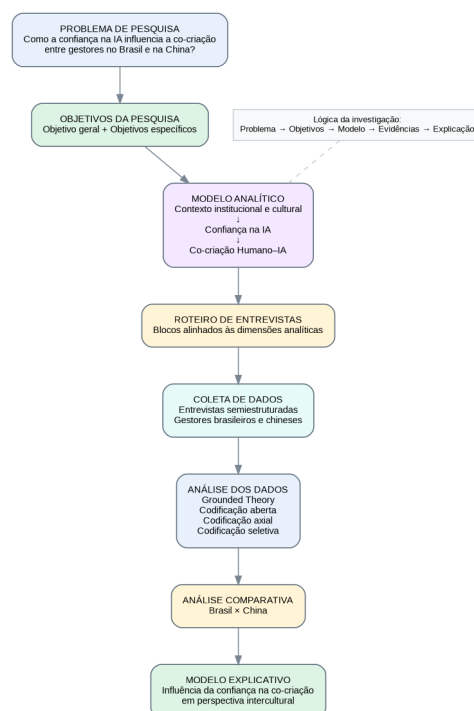
Problemas/ Objetivos	Dimensão analítica	Evidência buscada
Como os gestores incorporam a IAG?	Incorporação da IA às rotinas	Práticas de uso, tipos de tarefas, frequência e finalidade
Como a confiança é construída?	Confiança na IA	Critérios de confiança, validação, percepção de riscos, revisão da confiança
Como ocorre a co-criação?	Interação humano-IA	Estratégias de interação, refinamento, colaboração e delegação cognitiva
Como Brasil e China diferem?	Contexto institucional e cultural	Similaridades, diferenças e influência do contexto sobre confiança e co-criação.

Fonte: Elaborado pela Autora (2026)

O Quadro 1 sistematiza a articulação entre o problema de pesquisa, os objetivos específicos, as dimensões analíticas e as evidências empíricas buscadas na coleta e na análise dos dados.

A figura 1 apresenta, de forma sintética, o percurso metodológico da pesquisa, evidenciando a articulação entre o problema de pesquisa, os objetivos, o modelo analítico, a coleta e a análise dos dados. Essa representação permite visualizar a sequência lógica que orienta a investigação, desde a definição do fenômeno estudado até a construção do modelo explicativo proposto.

Figura 1 - Articulação entre o problema de pesquisa, o modelo analítico e o percurso metodológico.



Fonte: Elaborado pela Autora (2026), com base em Creswell e Creswell (2021), Strauss e Corbin (2008) e Mangoni (2025)

3.2 ESTRATÉGIA DE PESQUISA E ARTICULAÇÃO COM O MODELO ANALÍTICO

A estratégia de pesquisa fundamenta-se em abordagem qualitativa comparativa, orientada à construção de modelo analítico explicativo da relação entre contexto, confiança e co-criação. Como desdobramento ao estudo de Mangoni (2025), que investigou a co-criação humano-IA no contexto brasileiro, este trabalho amplia o escopo analítico ao incorporar a confiança na IA como dimensão central e ao introduzir a comparação com o contexto chinês.

Diferentemente da estrutura analítica baseada no paradigma condicional adotada por Mangoni (2025), o presente estudo propõe modelo analítico próprio, organizado a partir de lógica relacional que articula três dimensões centrais: contexto institucional e cultural, confiança na IA e co-criação humano-IA. Nesse modelo, o contexto atua como condição estruturante que molda percepções, expectativas e formas de regulação da tecnologia; a confiança emerge como mecanismo mediador que orienta o

comportamento gerencial e define os limites e as possibilidades da interação com a IA; e a co-criação constitui o resultado observável dessas interações, manifestando-se em práticas concretas de trabalho, produção de conteúdo e apoio à decisão.

A incorporação dessa lógica permite avançar em relação ao estudo de Mangoni (2025) ao introduzir camada explicativa adicional, centrada na mediação da confiança e na influência do contexto. Além disso, possibilita a comparação analítica entre os achados do estudo anterior, realizado no Brasil, e os dados produzidos no contexto chinês, evidenciando convergências, divergências e especificidades na forma como gestores constroem e mobilizam a confiança na IA em diferentes ambientes institucionais.

A sistematização apresentada no Quadro 1 orienta a construção do roteiro de entrevistas e assegura que cada objetivo específico seja contemplado de forma consistente ao longo do percurso metodológico.

3.3 COLETA DE DADOS

A coleta de dados foi realizada por meio de entrevistas semiestruturadas com gestores que utilizam ferramentas de inteligência artificial generativa em suas atividades profissionais. A escolha dessa técnica justifica-se por sua adequação à exploração de percepções, interpretações e práticas, permitindo simultaneamente profundidade analítica e comparabilidade entre os contextos investigados (Gil, 2002).

Os participantes foram selecionados de forma intencional, considerando sua atuação em funções gerenciais e o uso recorrente de ferramentas de IAG no contexto de trabalho. Em alinhamento aos critérios adotados por Mangoni (2025), estabeleceu-se como condição de inclusão o exercício de cargo de gestão e o uso de ferramentas de IAG por período mínimo de seis meses, de modo a assegurar familiaridade suficiente com a tecnologia, permitindo que os relatos refletissem rotinas de uso já incorporadas ao trabalho gerencial, e não reações preliminares de adoção recente. A definição desses critérios buscou assegurar que os participantes dispusessem de repertório de experiência capaz de sustentar reflexão consistente sobre a interação com a tecnologia e sobre a maneira como a confiança se forma e se revisa em situações concretas de uso.

A seleção buscou abranger gestores de áreas funcionais distintas, entre elas marketing, finanças, recursos humanos, operações e estratégia, de forma a diversificar os pontos de vista organizacionais representados na amostra (Creswell; Lopes; Silva,

2021). Dado o caráter comparativo do estudo, a amostra foi estruturada em dois subconjuntos correspondentes aos contextos brasileiro e chinês, aos quais foram aplicados critérios equivalentes de inclusão e coleta, condição necessária à comparabilidade dos achados. Ambos os conjuntos de entrevistas constituem material empírico original desta pesquisa. Os resultados de Mangoni (2025) são utilizados como referência analítica para discutir convergências, ampliações e divergências em relação aos achados produzidos no contexto brasileiro.

A definição do número de participantes ancorou-se no critério de saturação teórica formulado por Strauss e Corbin (2008), segundo o qual a incorporação de novos casos é interrompida quando deixa de gerar categorias inéditas ou de alterar de forma substantiva as já delineadas. Esse parâmetro sustenta a densidade interpretativa exigida pela construção do modelo analítico proposto e foi observado separadamente em cada contexto nacional, com vistas a preservar a especificidade dos padrões que emergem em cada ambiente institucional.

O roteiro de entrevistas foi estruturado com base nas dimensões do modelo analítico, contemplando aspectos relacionados à percepção de confiabilidade da IA, à avaliação de riscos, às estratégias de validação das respostas, às formas de uso em atividades gerenciais e à influência do contexto organizacional e cultural. Essa estrutura permite não apenas compreender as práticas individuais, mas também identificar padrões e variações entre os contextos nacionais analisados. A realização das entrevistas ocorreu em ambiente remoto, com agendamento prévio e registro em áudio autorizado pelos participantes mediante Termo de Consentimento Livre e Esclarecido (TCLE); o conteúdo foi integralmente transcrito para subsidiar a fase analítica.

Com o objetivo de assegurar a coerência entre o problema de pesquisa, os objetivos específicos e a produção dos dados empíricos, o roteiro de entrevistas foi estruturado em blocos temáticos alinhados às dimensões analíticas do estudo. Essa organização permitiu que cada objetivo específico fosse contemplado por questões capazes de gerar evidências relevantes para a análise qualitativa, preservando, ao mesmo tempo, a flexibilidade inerente às entrevistas semiestruturadas. O Quadro 2 apresenta a relação entre os objetivos específicos da pesquisa, os blocos do roteiro de entrevistas, o tipo de evidência e o correspondente procedimento de codificação utilizado na análise.

Quadro 2 - Articulação entre objetivos específicos, roteiro de entrevistas e o processo de análise dos dados

Objetivo específico	Bloco temático da entrevista	Evidências empíricas buscadas	Procedimento de análise
Mapear a incorporação da IAG	Uso da IA na rotina	Relatos sobre tarefas, frequência e finalidade	Codificação aberta
Examinar a confiança	Construção da confiança	Critérios de confiança, validação e risco	Codificação aberta e axial
Caracterizar a co-criação	Interação gestor-IA	Estratégias de colaboração, refinamento e decisão	Codificação axial
Comparar Brasil x China	Contexto Institucional	Similaridades e diferenças	Codificação seletiva

Fonte: Elaborado pela Autora (2026)

3.3.1 CRITÉRIOS DE SELEÇÃO DOS PARTICIPANTES

Os participantes dessa pesquisa foram gestores que utilizam IAG em suas rotinas de trabalho e que exerçam atividades de natureza decisória, estratégica ou gerencial em organizações situadas no Brasil e na China. A seleção ocorreu por amostragem intencional (*purposeful sampling*), de abrangência internacional, buscando participantes que possuam experiência efetiva na utilização da IAG em contextos organizacionais e que possam contribuir com informações relevantes para a compreensão do fenômeno investigado. Considerando o caráter comparativo da pesquisa, buscou-se contemplar gestores de diferentes setores econômicos, preservando como critério comum o uso recorrente da Inteligência Artificial Generativa em suas atividades profissionais.

Definiu-se como parâmetro inicial de planejamento a realização de cinco entrevistas em cada contexto nacional, quantitativo passível de ampliação ou redução conforme a suficiência do material empírico para responder ao problema de pesquisa e sustentar a comparação entre os contextos analisados.

No contexto chinês, o recrutamento foi realizado por meio de redes profissionais previamente estabelecidas pela pesquisadora, incluindo contatos vinculados à sua atuação em uma organização de origem chinesa, além de contatos acadêmicos e profissionais acessíveis durante o período de estudos na *Fudan University*, em Shanghai, China. Foram convidados profissionais que atendiam aos critérios de inclusão da pesquisa, incluindo docentes que também exerciam funções de gestão em

organizações. Os convites foram realizados individualmente, com registro do número de pessoas contatadas, das respostas recebidas e das entrevistas realizadas. A seleção buscou incluir participantes de diferentes setores econômicos. A composição final da amostra, assim como eventuais limitações de acesso e diversidade, foi apresentada no estudo.

3.3.2 INSTRUMENTO DE COLETA

A coleta de dados foi realizada por meio de entrevistas semiestruturadas, por possibilitarem aprofundar a compreensão das experiências, percepções e práticas dos participantes, preservando simultaneamente um roteiro comum que assegura comparabilidade entre os contextos brasileiro e chinês. O roteiro de entrevistas foi elaborado a partir do problema de pesquisa, dos objetivos específicos e do modelo analítico desenvolvido na fundamentação teórica, contemplando blocos temáticos relacionados à incorporação da Inteligência Artificial Generativa nas rotinas gerenciais, à construção da confiança, aos processos de cocriação entre gestores e sistemas de IAG e às influências do contexto organizacional e cultural.

A natureza semiestruturada das entrevistas permitiu explorar aspectos emergentes ao longo da interação com os participantes, preservando abertura para elementos não previstos inicialmente no modelo analítico. O roteiro completo encontra-se no Apêndice C.

3.3.3 PROCEDIMENTOS PARA COLETA DE DADOS

As entrevistas foram realizadas por videoconferência, modalidade adotada em razão da distribuição geográfica dos participantes nos dois contextos nacionais. Mediante Termo de Consentimento Livre e Esclarecido (TCLE), cujo modelo consta do Apêndice B, todas as entrevistas foram gravadas em áudio para assegurar fidelidade ao conteúdo produzido e posteriormente transcritas integralmente para análise.

Após a transcrição, o material foi submetido à revisão pela pesquisadora, com consulta aos registros originais de áudio nos trechos em que a transcrição apresentava ambiguidades, palavras incompletas ou inconsistências semânticas. Essa etapa teve como finalidade preservar o sentido das falas antes do processo de codificação. As

correções restringiram-se à fidelidade da transcrição, sem alteração substantiva das respostas dos participantes.

A coleta foi acompanhada pela análise preliminar do material produzido, o que possibilitou o aprofundamento de aspectos que se revelaram relevantes ao longo da pesquisa. Os ajustes realizados no roteiro foram registrados e justificados, preservando-se o núcleo comum de questões necessário à comparação entre os contextos brasileiro e chinês.

Para preservar o anonimato, os participantes foram identificados exclusivamente por códigos alfanuméricos, utilizando-se BR01 a BR05 para o contexto brasileiro e CN01 a CN05 para o contexto chinês. Nomes pessoais, organizações, unidades administrativas, cargos específicos, localidades, projetos e demais elementos que permitissem a identificação direta ou indireta foram removidos ou generalizados. Informações profissionais foram apresentadas em categorias amplas e faixas de experiência, evitando combinações de características que possibilitassem a reidentificação dos participantes. Excertos potencialmente identificáveis foram apresentados por meio de paráfrases analíticas, preservando-se seu conteúdo substantivo.

3.3.4 SATURAÇÃO TEÓRICA E SUFICIÊNCIA DO MATERIAL EMPÍRICO

A definição do ponto de encerramento da coleta foi orientada pelo critério de saturação teórica, entendido como o momento em que a incorporação de novos casos deixa de produzir categorias analíticas inéditas ou de modificar substantivamente as propriedades e relações das categorias já construídas (STRAUSS; CORBIN, 2008). A previsão inicial de aproximadamente cinco participantes por contexto constituiu referência de planejamento da coleta, e não definição prévia de suficiência amostral. Foram realizadas cinco entrevistas com gestores em cada contexto, brasileiro e chinês, e a decisão de encerramento foi estabelecida a partir da evolução da própria análise e da suficiência do material empírico para responder ao problema de pesquisa e sustentar a comparação entre os contextos.

A avaliação da saturação foi conduzida de forma processual, mediante análise comparativa das entrevistas e acompanhamento da emergência, recorrência e estabilização das categorias em cada contexto nacional. Considerou-se como evidência de saturação o momento em que novos casos deixaram de produzir categorias analíticas

inéditas, de modificar substantivamente as propriedades das categorias já constituídas ou de alterar as relações centrais entre contexto, confiança e co-criação. A presença de variações entre participantes não foi interpretada, por si só, como ausência de saturação, desde que essas variações pudessem ser incorporadas às categorias existentes sem necessidade de reformulação da estrutura analítica.

No contexto chinês, a estabilização das categorias tornou-se particularmente observável nas entrevistas CN04 e CN05. CN04 aprofundou propriedades já associadas às categorias de incorporação seletiva da IAG, critérios de confiança e influência do contexto institucional, especialmente quanto à segurança de dados, aos limites de uso e às condições organizacionais. CN05 reiterou mecanismos previamente identificados, particularmente a validação dos outputs, a delegação seletiva conforme a criticidade e a verificabilidade da tarefa, o refinamento iterativo e a preservação do julgamento humano como instância final de validação e decisão. Nenhuma das duas entrevistas resultou na criação de nova categoria estruturante ou de relação capaz de alterar a estrutura analítica já constituída ($n = 0$). Os casos acrescentaram propriedades, exemplos e variações internas, contribuindo para a densificação das categorias existentes e fornecendo evidência da estabilização alcançada no processo de análise.

No contexto brasileiro, a mesma verificação foi realizada nas entrevistas finais. BR04 introduziu uma variação interna relevante ao apresentar uso menos sofisticado da IAG e ausência de influência direta da ferramenta sobre determinadas decisões, sem, contudo, exigir a criação de nova categoria. BR05 reiterou padrões previamente observados, especialmente o uso da IAG como ponto de partida, a validação dos outputs, a delegação seletiva e a retomada manual da tarefa quando necessário. Dessa forma, as variações observadas contribuíram para ampliar a densidade interpretativa das categorias, mas não modificaram sua estrutura.

Para tornar observável o processo de saturação, foi acompanhada a evolução das categorias ao longo da sequência de entrevistas em cada contexto. Após a codificação de cada novo caso, os códigos e propriedades identificados foram comparados com aqueles provenientes das entrevistas anteriores, registrando-se a emergência de novas categorias, o aprofundamento de categorias existentes e eventuais alterações nas relações analíticas. A saturação não foi inferida exclusivamente da recorrência dos temas, mas da progressiva redução da emergência de novas categorias e da capacidade da estrutura analítica existente de incorporar as variações apresentadas pelos novos casos.

Quadro 3 - Evolução das categorias e verificação da saturação teórica nos contextos brasileiro e chinês

Entrevista	Categorias/ mecanismos emergentes	Novas categorias		Aprofundamento s/ variações
BR01	Verificabilidade; contextualização dos prompts; validação; julgamento humano; incorporação seletiva	Formação das categorias de referência	Uso progressivamente contextualizado ; validação de outputs e delimitação da participação da IAG	Formação inicial
BR02	Confiança experiencial; validação independente; delegação seletiva; refinamento; julgamento humano	Expansão das categorias iniciais	Ampliação das formas de validação e dos limites de delegação conforme tarefa e risco	Expansão
BR03	Refinamento iterativo; uso mais integrado da IAG; julgamento humano; delegação conforme criticidade e verificabilidade	Predomínio de densificação	Maior sofisticação da interação e explicitação dos critérios de supervisão	Densificação
BR04	Incorporação da IAG; apoio à decisão; julgamento humano	Nenhuma nova categoria estruturante	Uso menos sofisticado e ausência de influência direta da IAG sobre determinadas decisões	Estabilização com variação interna
BR05	IAG como ponto de partida; validação; delegação seletiva; retomada manual	Nenhuma nova categoria estruturante	Reiteração de mecanismos existentes e possibilidade de retirada da IAG	Confirmação da estabilização
CN01	Confiança experiencial; verificabilidade dos outputs; contextualização dos prompts; preservação do julgamento humano; uso de ferramentas aprovadas para dados sensíveis	Formação das categorias de referência	Curva de aprendizagem; aumento progressivo de contexto; aceitação condicionada à expertise	Formação inicial
CN02	Automação e agentes; integração ao fluxo de trabalho; sistemas	Expansão das categorias iniciais	Maior sofisticação do uso e	Expansão

	internos/protegidos; contextualização e restrições no uso		aprofundament o dos limites institucionais	
CN03	Validação factual; verificabilidade; julgamento humano sobre estratégia/compliance; refinamento e possibilidade de abandono; uso de plataformas internas para dados confidenciais	Predomínio de densificação	Maior rigor na validação; decomposição após falhas; explicitação dos limites da delegação	Densificação
CN04	Incorporação seletiva; confiança; contexto institucional	Nenhuma nova categoria estruturante	Segurança de dados; limites de uso; condições organizacionais	Estabilização
CN05	Validação; delegação seletiva; refinamento; julgamento humano	Nenhuma nova categoria estruturante	Novos exemplos e variações internas	Confirmação da estabilização

Fonte: Elaborado pela Autora com base na análise comparativa e codificação das entrevistas (2026).

Para tornar observável o processo de saturação, foi acompanhada a evolução das categorias ao longo da sequência de entrevistas em cada contexto. Após a codificação de cada novo caso, os códigos e propriedades identificados foram comparados com aqueles provenientes das entrevistas anteriores, registrando-se a emergência de novas categorias, o aprofundamento de categorias existentes e eventuais alterações nas relações analíticas. A saturação não foi inferida exclusivamente da recorrência dos temas, mas da progressiva redução da emergência de novas categorias e da capacidade da estrutura analítica existente de incorporar as variações apresentadas pelos novos casos.

Quadro 4 - Evidências processuais de saturação nas entrevistas finais

Caso	Nova categoria estruturante	Evidências acrescentadas/ reiteradas	Efeito sobre a estrutura analítica
CN04	0	Limites de uso; segurança de dados; condições institucionais	Aprofundamento de categorias existentes
CN05	0	Verificação dos outputs; delegação seletiva; refinamento; preservação do julgamento humano	Reiteração e densificação das categorias existentes

BR04	0	Menor sofisticação de uso; menor influência da IAG sobre decisões	Variação interna sem nova categoria
BR05	0	IAG como ponto de partida; validação; delegação seletiva; retomada manual	Reiteração e densificação das categorias existentes

Fonte: Elaborado pela Autora com base na análise comparativa e codificação das entrevistas (2026).

Como demonstra o Quadro 4, a contribuição incremental das entrevistas chinesas diminuiu ao longo do processo analítico. Enquanto os casos iniciais participaram da formação e do desenvolvimento das categorias, CN04 não introduziu nova categoria estruturante, embora tenha aprofundado propriedades relacionadas à segurança de dados, aos limites de uso e às condições institucionais. CN05 igualmente não produziu nova categoria, reiterando mecanismos de validação dos outputs, delegação seletiva, refinamento iterativo e preservação do julgamento humano. A ocorrência consecutiva de dois casos sem emergência de novas categorias estruturantes, acompanhada da recorrência e densificação das categorias existentes, constituiu a evidência utilizada para considerar estabilizada a estrutura analítica nesse contexto.

A estabilização observada nas entrevistas finais não pressupõe a repetição integral das categorias por todos os participantes, mas a ausência de evidências que exigissem a criação de novas categorias, a reformulação de suas propriedades centrais ou a alteração das relações analíticas posteriormente consolidadas nos Quadros 11 a 15. Esses quadros constituem a estrutura analítica resultante do processo de categorização, enquanto a saturação se refere ao processo de verificação de que a incorporação de novos casos deixou de modificar substantivamente essa estrutura.

A recorrência dos mecanismos entre os casos forneceu evidência adicional dessa estabilização. Conforme sintetizado no Quadro 14 a formulação e delimitação da tarefa e a geração de respostas pela IAG foram identificadas nos dez casos analisados (BR01–BR05 e CN01–CN05), assim como práticas de avaliação crítica associadas a fontes e verificabilidade. O refinamento iterativo também se mostrou recorrente nos dois corpora. Nos Quadros 11 a 13 e 15, esses mecanismos reaparecem associados à incorporação seletiva da IAG, à validação dos outputs, à calibração da delegação conforme criticidade e verificabilidade e à manutenção do julgamento humano como instância final de decisão.

A decisão de encerrar a coleta apoiou-se, portanto, na convergência entre a ausência de novas categorias estruturantes nas entrevistas finais, a recorrência dos mecanismos centrais entre os casos e a capacidade das categorias existentes de incorporar variações sem alteração das relações analíticas construídas. A saturação foi considerada alcançada quando os casos adicionais passaram a densificar e exemplificar categorias existentes, sem gerar novas categorias estruturantes ou relações capazes de modificar a interpretação do fenômeno. A suficiência do material empírico foi, assim, estabelecida a partir da evolução da própria análise, e não exclusivamente do número inicialmente previsto de participantes.

3.4 PRÉ-TESTE DO INSTRUMENTO DE COLETA

Antes da aplicação definitiva, o roteiro passou por etapa de validação preliminar, voltada a examinar a inteligibilidade das perguntas, sua aderência aos objetivos da pesquisa e o tempo demandado por sua aplicação, além de detectar formulações ambíguas passíveis de correção (Gil, 2002). Seguindo a orientação metodológica de Mangoni (2025), essa validação envolveu profissionais que satisfaziam as condições de inclusão do estudo, porém sem compor o conjunto de casos efetivamente analisados, o que preservou a separação entre a etapa de calibragem do instrumento e a etapa de geração dos dados. Os registros dessa fase não integraram o material submetido à análise, tendo servido unicamente ao aperfeiçoamento do roteiro.

Dado o caráter comparativo do estudo, o instrumento também foi adaptado linguisticamente para sua aplicação no contexto chinês. O roteiro originalmente elaborado em português foi convertido para o inglês, idioma utilizado na realização das entrevistas com os participantes do contexto chinês. A adaptação priorizou equivalência conceitual, e não tradução estritamente literal, especialmente nos termos associados à confiança na IA, validação de respostas, delegação cognitiva e co-criação. O núcleo temático, a sequência dos blocos e os objetivos analíticos das perguntas foram preservados nas duas versões do instrumento, buscando reduzir a possibilidade de que diferenças de formulação fossem confundidas com diferenças substantivas entre os contextos.

As entrevistas do contexto chinês foram analisadas em seu registro original em inglês. Para apresentação e integração analítica em português, os conteúdos foram traduzidos pela pesquisadora, preservando-se o sentido substantivo das respostas.

Eventuais ambiguidades linguísticas foram tratadas com cautela durante a codificação e consideradas como limitação da comparação intercultural.

O pré-teste foi organizado em duas frentes. Na primeira, de natureza linguística e conceitual, avaliou-se a compreensibilidade dos termos e a correspondência entre a formulação das perguntas e as dimensões do modelo analítico, quais sejam, contexto institucional e cultural, confiança na IA e co-criação humano-IA. Na segunda, de natureza operacional, verificou-se o tempo médio de aplicação, a sequência lógica dos blocos temáticos e a capacidade das questões de eliciar relatos sobre práticas concretas de interação com sistemas generativos.

Os resultados dessa etapa motivaram ajustes no roteiro, entre os quais o rearranjo de perguntas de baixa inteligibilidade, a eliminação de sobreposições temáticas e a reordenação dos blocos, com o propósito de tornar a condução mais fluida e de elevar a riqueza dos relatos obtidos. Tais ajustes reforçaram a consistência interna do instrumento e uniformizaram as condições de coleta entre os dois contextos, requisito de particular importância em desenhos comparativos, nos quais variações não controladas no próprio instrumento tenderiam a se confundir com diferenças reais entre os casos.

3.5 ANÁLISE DE DADOS

A análise dos dados foi realizada por meio de análise temática, utilizando procedimentos de codificação e comparação inspirados na Grounded Theory, conforme sistematizada por Strauss e Corbin (2008). Essa escolha permitiu identificar e organizar padrões no material empírico sem pressupor a aplicação integral dos procedimentos da Grounded Theory, especialmente aqueles relacionados à simultaneidade sistemática entre coleta e análise e à amostragem teórica. Essa opção justifica-se pela pertinência do método ao exame de processos sociais de contornos ainda instáveis, para os quais as categorias explicativas são elaboradas indutivamente a partir do próprio material empírico, e não fixadas de antemão (Strauss; Corbin, 2008; Creswell; Lopes; Silva, 2021).

Como recursos operacionais para o tratamento do material, o percurso analítico foi organizado em três procedimentos de codificação inspirados na Grounded Theory, conforme Strauss e Corbin (2008) — codificação aberta, axial e seletiva — sem que sua utilização isolada caracterize a adoção integral desse método. No primeiro,

correspondente à codificação aberta, o exame minucioso das transcrições permitiu isolar noções recorrentes ligadas à confiança na IA, à leitura de riscos, aos modos de interação e às práticas de co-criação, com preferência por rótulos extraídos da própria fala dos participantes quando estes carregavam valor analítico. O segundo movimento, de codificação axial, reuniu tais noções em categorias de maior alcance, tomando como eixo a relação entre condicionantes institucionais, formação da confiança e conduta gerencial; em lugar da aplicação convencional do paradigma condicional, essa fase seguiu a lógica relacional definida pelo estudo, centrada na articulação entre contexto, confiança e co-criação. O terceiro movimento, de codificação seletiva, voltou-se à integração das categorias em arranjo explicativo unificado, capaz de tornar visíveis os mecanismos pelos quais as variações de contexto afetam a confiança e, por decorrência, os padrões de co-criação humano-IA.

A análise partiu do entendimento de que a confiança é construída de forma dinâmica ao longo das interações com a tecnologia, podendo se alterar em função da experiência, do desempenho percebido da IA e das condições institucionais. Essa perspectiva permite compreender a co-criação não como estado fixo, mas como processo contínuo de ajuste entre humano e sistema tecnológico. O princípio da comparação constante foi aplicado de forma sistemática, tanto internamente a cada contexto quanto entre os contextos brasileiro e chinês, assegurando a consistência dos códigos e a robustez das categorias emergentes (Strauss; Corbin, 2008).

As categorias foram desenvolvidas inicialmente de forma independente nos contextos brasileiro e chinês, preservando as especificidades de cada conjunto empírico. As entrevistas dos dois contextos foram submetidas ao mesmo modelo relacional de codificação, organizado em quatro dimensões: contexto institucional e cultural (CTX), confiança na IA (TRUST), cocriação humano-IA (COCRI) e elo relacional (LINK). Em seguida, procedeu-se à análise comparativa entre os conjuntos de categorias, com vistas a identificar recorrências, variações, casos divergentes e mecanismos explicativos relacionados à influência do contexto institucional e cultural sobre a confiança e a cocriação.

Mangoni (2025) foi utilizado como referência de continuidade analítica, e não como substituto do corpus brasileiro desta pesquisa. Seus resultados foram mobilizados para discutir em que medida os achados das entrevistas brasileiras convergem, ampliam ou tensionam padrões anteriormente identificados, especialmente aqueles relacionados ao refinamento iterativo, à validação crítica e ao repertório profissional dos gestores.

Cabe registrar uma limitação assumida no diálogo com o estudo de referência. Os corpora brasileiro e chinês desta pesquisa foram submetidos ao mesmo modelo relacional de codificação (CTX, TRUST, COCRI e LINK), condição que assegura equivalência procedimental à comparação entre os dois contextos. Os achados de Mangoni (2025), entretanto, foram originalmente codificados segundo o paradigma condicional de Strauss e Corbin (causal, contexto, condição interveniente, estratégia e consequência), razão pela qual sua mobilização neste trabalho consiste em reinterpretação analítica à luz das categorias do modelo relacional, e não em recodificação a partir dos dados brutos originais da autora. Essa opção é compatível com o caráter analítico, e não replicativo, do diálogo proposto, mas impõe cautela na leitura das convergências identificadas, que devem ser compreendidas como aproximação interpretativa entre esquemas categóricos distintos.

O percurso metodológico adotado estabelece relação entre os objetivos específicos da pesquisa, a estrutura do roteiro de entrevistas e os procedimentos de codificação utilizados. Embora a análise tenha sido orientada pelas dimensões do modelo analítico, os códigos e as categorias interpretativas foram desenvolvidos a partir do material empírico, preservando abertura para relações e elementos não previstos inicialmente.

3.6 USO DE INTELIGÊNCIA ARTIFICIAL COMO SUPORTE METODOLÓGICO

Em articulação com o estudo de Mangoni (2025), esta pesquisa incorpora o uso de ferramentas de inteligência artificial generativa como suporte ao processo analítico, com explicitação dos critérios de uso e dos mecanismos de validação adotados. As ferramentas foram utilizadas como apoio à organização dos dados, à sugestão de agrupamentos temáticos e à elaboração de sínteses preliminares, sem substituição do julgamento analítico da pesquisadora.

Para operacionalizar esse suporte, foi desenvolvido um assistente de IA próprio, com instruções fixas (*system prompt*) para conduzir a codificação aberta, estruturado a partir do modelo relacional que orienta a presente pesquisa, e não a partir do paradigma condicional de Strauss e Corbin, empregado por Mangoni (2025) para o mesmo fim. Em lugar das cinco dimensões desse paradigma, o assistente foi instruído a identificar, em cada trecho de entrevista, evidências associadas às três dimensões substantivas do modelo, contexto institucional e cultural (CTX), confiança na IA (TRUST) e co-criação

humano-IA (*COCRI*), além de um quarto eixo, denominado elo relacional (*LINK*), destinado a capturar os trechos em que o próprio gestor articula explicitamente a relação entre duas ou mais dessas dimensões. Como no procedimento de referência, cada código gerado foi acompanhado do trecho literal que o originou, submetido a comparação com códigos já utilizados em entrevistas anteriores para assegurar consistência semântica, e validado criticamente pela pesquisadora antes de sua incorporação à planilha de análise. Todas as interações foram registradas em log de prompts, garantindo rastreabilidade equivalente à empregada por Mangoni (2025). Dado o desenho comparativo entre contextos distintos, o assistente foi instruído a preservar o texto original em inglês dos excertos provenientes do contexto chinês, mantendo os códigos analíticos em português para assegurar comparabilidade entre os dois conjuntos empíricos.

Prompt 1 – *System prompt* para assistente de IA em codificação aberta (modelo relacional *CTX/TRUST/COCRI/LINK*)

You are a qualitative coding assistant supporting open coding procedures inspired by Grounded Theory (Strauss & Corbin strand). You support a study that does NOT use the conditional paradigm (causal/context/intervening/strategy/consequence). Instead, the analysis follows a relational model with three substantive dimensions plus one connective axis: *CTX* - Institutional and Cultural Context (structuring condition): stable, external conditions that shape how AI is perceived and regulated - national governance model, organizational culture, sectoral norms, regulatory environment, professional expectations. *TRUST* - Trust in AI (mediating variable): the manager's cognitive, affective, or behavioral stance toward the reliability of AI - criteria for believing or doubting an output, validation habits, risk perception, expressions of confidence, distrust, or calibration of trust over time. *COCRI* - Human-AI Co-creation (observable outcome): concrete joint-work practices between manager and AI - formulating, delegating, iterating, refining, validating, or applying AI-generated outputs in real managerial tasks. *LINK* - Relational mechanism: excerpts where the participant explicitly connects two or more of the above dimensions in a single causal or conditional statement (e.g., "because my company requires X, I only trust the AI when Y" or "that trust is what lets me hand the first draft to it"). Tag as *LINK* only when the

connection is explicit in the participant's own reasoning, not when you infer it yourself. **DECISION HEURISTIC** 1. Is it a stable, external condition (rule, culture, governance, sector norm)? -> **CTX** 2. Is it a belief, expectation, doubt, or evaluation about AI's reliability? -> **TRUST** 3. Is it a concrete joint action or practice with the AI? -> **COCRI** 4. Does the excerpt itself state how one of these dimensions produces or conditions another? -> **LINK** (this takes priority over 1-3 when present) 5. If none of the above applies, respond with "No relevant analytical content."

CODING RULES - Each code must be a short, concrete, descriptive phrase, close to the participant's own words. - Use in vivo codes (the participant's own wording) whenever possible. - Avoid abstract single words such as "technology," "trust," or "culture" alone; codes must describe a specific stance, action, or condition. - Check every candidate code against the set of codes already assigned in earlier interviews of this project. Only when an earlier code carries the identical underlying meaning should you keep its original wording; never force a match when the substance differs, even slightly. - Attach to every code the original passage it was drawn from, framed in single quotes, right before the code label and its tag (**CTX** / **TRUST** / **COCRI** / **LINK**).

MULTILINGUAL HANDLING (Chinese-context interviews) Interviews from the Chinese sample were conducted in English. Preserve the exact English wording of the excerpt whenever an in vivo quotation is used. Codes should be written in Portuguese so that Brazilian and Chinese-context codes remain comparable during axial coding. When translation into Portuguese is needed for analytical integration or presentation, preserve the substantive meaning of the original excerpt and flag any expression whose interpretation is ambiguous rather than inferring a meaning not explicit in the participant's statement.

OUTPUT FORMAT Present only markdown code blocks, one per code, in the form: 'exact excerpt' (then) **CODE** (**TAG**). Do not include summaries, explanations, or classifications beyond the tag. If there is no analytical content in the excerpt, respond only with: "No relevant analytical content."

RESEARCH OBJECTIVES (for reference, do not restate them in your answers) General objective: to analyze how trust in AI influences co-creation between managers and generative AI in daily work routines, comparing the Brazilian and Chinese contexts. Specific objectives: a) map how generative AI is incorporated into managers' routines in the Brazilian and Chinese contexts; b) examine how trust in AI is built, perceived, and mobilized in these practices; c)

characterize the interaction and co-creation processes between managers and generative AI systems; d) compare similarities and differences between Brazil and China regarding trust and its effects on co-creation. Hold until an excerpt is submitted. Once you receive it, work through it attentively, following every rule above without exception.

Prompt elaborado pela autora (2026).

Como procedimento de avaliação da consistência da classificação assistida por IA, uma parcela do corpus, composta por excertos dos contextos brasileiro e chinês, foi classificada de forma independente pela pesquisadora e pelo assistente nas dimensões CTX, TRUST, COCRI e LINK. A classificação da pesquisadora foi realizada antes da consulta às respostas produzidas pela ferramenta, e os mesmos excertos foram posteriormente submetidos ao prompt padronizado apresentado nesta seção. A concordância foi examinada por meio do percentual de concordância e do coeficiente Kappa de Cohen, calculado separadamente para cada dimensão. As divergências foram revistas com base nos excertos originais, cabendo à pesquisadora a decisão final sobre sua classificação e interpretação.

Os 24 excertos utilizados nessa avaliação foram selecionados de modo a contemplar os dois contextos nacionais e diferentes tipos de conteúdo associados às categorias CTX, TRUST, COCRI e LINK, incluindo trechos de classificação mais direta e trechos com maior proximidade semântica entre categorias.

O ChatGPT também foi utilizado como apoio à elaboração da representação visual integrada dos resultados. O conteúdo analítico e as relações representadas foram definidos pela pesquisadora com base nos dados da pesquisa, e o prompt utilizado foi registrado no log apresentado no Apêndice A.

3.7 CONSIDERAÇÕES SOBRE RIGOR METODOLÓGICO

O rigor metodológico foi buscado por meio da aplicação sistemática dos procedimentos de codificação adotados, da comparação constante entre os dados, do registro detalhado das decisões analíticas e da explicitação do uso de IA como ferramenta auxiliar. Para tornar esses procedimentos verificáveis, os critérios de rigor

qualitativo foram associados às estratégias metodológicas empregadas ao longo da pesquisa, conforme apresentado no Quadro 5.

Quadro 5 - Critérios de rigor qualitativo e procedimentos adotados

Critério de rigor	Procedimentos adotados na pesquisa	Evidência o estudo
Credibilidade	Comparação constante entre entrevistas, casos e contextos; retorno aos excertos originais durante a análise; preservação de convergências e divergências; análise de casos que não reproduziram integralmente os padrões predominantes	Categorias e interpretações confrontadas continuamente com o material empírico; manutenção de variações e casos divergentes
Dependabilidade	Aplicação sistemática do mesmo procedimento de codificação; estrutura comum de análise CTX, TRUST, COCRI e LINK; registro das decisões analíticas e dos prompts utilizados no apoio da IA	Procedimentos de codificação e classificação explicitados na metodologia e registros preservados para rastreabilidade
Confirmabilidade	Codificação e interpretação sob responsabilidade da pesquisadora; classificação independente pesquisadora × assistente de IA; análise das divergências e prevalência da decisão humana; manutenção dos excertos associados aos códigos	Comparação independente apresentou concordância elevada; divergências foram examinadas individualmente e decididas pela pesquisadora
Transferibilidade	Descrição dos critérios de seleção, perfil dos participantes, contextos investigados e condições de utilização da IAG; explicitação dos limites da amostra e do desenho comparativo	Caracterização dos participantes e dos contextos permite ao leitor avaliar a aplicabilidade dos achados a outras situações
Comparabilidade intercultural	Núcleo conceitual comum do roteiro; pré-teste; adaptação linguística para o inglês; manutenção das mesmas dimensões analíticas nos dois contextos	Redução do risco de atribuir a diferenças contextuais efeitos decorrentes apenas de formulações distintas do instrumento
Suficiência/saturação teórica	Comparação sequencial das entrevistas; acompanhamento da emergência, recorrência e densificação das categorias; verificação da ausência de	CN04 e CN05 não produziram novas categorias estruturantes e aprofundaram/reiteraram mecanismos já identificados

	novas categorias estruturantes nas entrevistas finais	
--	---	--

Fonte: Elaborado pela Autora com base em Lincoln e Guba (1985) e Corbin e Strauss (2008), a partir dos procedimentos adotados na pesquisa.

A credibilidade e a confirmabilidade foram reforçadas pela comparação constante entre os dados e pelo retorno aos excertos originais durante a construção e revisão das categorias. A classificação independente realizada pela pesquisadora e pelo assistente de IA constituiu procedimento adicional de contraste analítico, sem transferência da autoridade interpretativa à ferramenta. As divergências identificadas foram examinadas individualmente, permanecendo sob responsabilidade da pesquisadora a decisão sobre a classificação final.

A dependabilidade foi buscada pela aplicação sistemática dos procedimentos de codificação e pelo registro das decisões analíticas, enquanto a transferibilidade foi favorecida pela descrição dos participantes, dos contextos e das condições em que a IAG foi incorporada às práticas gerenciais. A padronização conceitual do instrumento entre Brasil e China, apoiada pelo pré-teste e pela adaptação linguística do roteiro para o inglês, contribuiu adicionalmente para a comparabilidade intercultural, reduzindo a possibilidade de que diferenças de formulação fossem interpretadas como diferenças substantivas entre os casos.

A suficiência do material empírico foi examinada por meio do processo de saturação descrito na Seção 3.3.4, considerando a emergência, recorrência e estabilização das categorias ao longo das entrevistas. Dessa forma, o rigor da pesquisa não se apoiou em um único procedimento de validação, mas na convergência entre rastreabilidade analítica, contraste entre casos, classificação independente, preservação de divergências, comparabilidade entre os contextos e demonstração da estabilização das categorias.

A padronização conceitual do instrumento entre os contextos, apoiada pelo pré-teste e pela adaptação linguística do roteiro para o inglês, contribuiu para a comparabilidade da pesquisa, ao reduzir a possibilidade de que diferenças de formulação fossem confundidas com diferenças substantivas entre os casos.

A abordagem comparativa entre Brasil e China contribuiu adicionalmente para a robustez da análise, ao permitir a identificação de padrões recorrentes e de variações contextuais. A manutenção de um núcleo analítico comum entre os dois contextos,

associada à preservação das particularidades emergentes em cada corpus, favoreceu a interpretação comparativa dos achados sem pressupor previamente que diferenças nacionais correspondessem a diferenças culturais.

4 RESULTADOS E DISCUSSÃO

Os resultados foram apresentados a partir do modelo analítico proposto, articulando as três dimensões centrais do estudo, contexto institucional e cultural, confiança na IA e co-criação humano-IA, e organizados de forma a responder aos objetivos específicos. A exposição segue estrutura comparativa, contrastando sistematicamente os contextos brasileiro e chinês, e dialoga com os resultados de Mangoni (2025), utilizados como referência para o contexto brasileiro. Cada subseção associa os achados empíricos às categorias emergentes da codificação e aos construtos teóricos discutidos na fundamentação.

4.1 AVALIAÇÃO DA CONCORDÂNCIA DO PROCEDIMENTO DE CODIFICAÇÃO ASSISTIDA POR IA

A avaliação da concordância do procedimento de classificação assistida por IA foi realizada com 24 excertos, dos quais 12 foram provenientes das entrevistas realizadas no contexto brasileiro e 12 no contexto chinês. Os excertos foram classificados de forma independente pela pesquisadora e pelo assistente de IA nas dimensões CTX, TRUST, COCRI e LINK, conforme o procedimento descrito na seção 3.6. O Quadro 3 apresenta os percentuais de concordância e os coeficientes Kappa de Cohen obtidos para as quatro categorias de classificação utilizadas no procedimento (CTX, TRUST, COCRI e LINK), bem como a concordância global.

Quadro 6 - Concordância entre as classificações da pesquisadora e do assistente de IA

Dimensão	N de excertos avaliados	Concordância (%)	Kappa de Cohen
CTX	24	91,7%	0,800
TRUST	24	95,8%	0,864
COCRI	24	87,5%	0,684

LINK	24	83,3%	0,560
------	----	-------	-------

Fonte: Elaborado pela Autora (2026)

Os resultados evidenciaram níveis satisfatórios de concordância entre a classificação realizada pela pesquisadora e aquela produzida pelo assistente de IA. A maior concordância foi observada em TRUST, com 95,8% e Kappa de Cohen de 0,864, seguida por CTX, com 91,7% e Kappa de 0,800, COCRI, com 87,5% e Kappa de 0,684, e LINK, com 83,3% e Kappa de 0,560.

Considerando conjuntamente as quatro categorias utilizadas no procedimento de classificação (CTX, TRUST, COCRI e LINK), foram observadas 19 concordâncias entre os 24 excertos avaliados, correspondentes a uma concordância global de 79,2% e Kappa global de 0,722. As cinco divergências identificadas concentraram-se principalmente em trechos que simultaneamente descreviam práticas de interação e explicitavam relações entre condições e comportamentos, evidenciando proximidade analítica entre algumas categorias.

As divergências foram revistas com base nos excertos originais e nas definições adotadas para cada dimensão. Nos casos de discordância, prevaleceu a classificação da pesquisadora, mantendo-se a IA como ferramenta de suporte ao processo analítico.

4.2 PERFIL DOS PARTICIPANTES E CARACTERIZAÇÃO DOS CONTEXTOS

Esta subseção apresentará a caracterização dos participantes de cada contexto nacional, incluindo área de atuação, tempo de experiência gerencial, tempo de uso de ferramentas de IAG e tipos de sistemas utilizados. Foram também sintetizados os elementos institucionais e regulatórios de cada contexto, retomando a caracterização da governança *state-led* no caso chinês (Zou; Zhang, 2025; Wang et al., 2025) e da configuração institucional em consolidação no caso brasileiro (Serra; Scafuto; Serra Vieira, 2025; Ribeiro; Segatto, 2025), de modo a situar os achados subsequentes.

Quadro 7 - Visão geral da amostra

Brasil	China
5 gestores	5 gestores
Setores: Participantes dos setores público e privado, atuantes em diferentes áreas	Setores: Participantes dos setores de tecnologia, finanças e marketing/comunicação, em diferentes

organizacionais	funções gerenciais.
Distribuição da experiência gerencial: 1 participante com menos de 1 ano; 2 entre 5 e 10 anos; 2 com mais de 10 anos	Distribuição da experiência gerencial: 1 com menos de 1 ano; 2 entre 1 e 5 anos; 1 entre 5 e 10 anos; 1 não informou duração.
Tempo de uso de IAG: Predominantemente entre 1 e 2 anos	Tempo de uso de IAG: Predominantemente entre aproximadamente 2 e 4 anos; um participante informou apenas que utiliza há alguns anos.
Ferramentas predominantes: Claude, Gemini, ChatGPT, Perplexity e NotebookLM	Ferramentas predominantes: DeepSeek é a ferramenta mais recorrente; também aparecem ChatGPT, Claude, Kimi, Doubao e Qwen.

Fonte: Elaborado pela Autora (2026)

4.2.1 CARACTERIZAÇÃO DOS PARTICIPANTES

Quadro 8 - Caracterização dos participantes

Código	País	Experiência gerencial	Tempo aproximado de uso da IAG
BR01	Brasil	Entre 5 e 10 anos	Entre 1 e 2 anos
BR02	Brasil	Mais de 10 anos	Entre 1 e 2 anos
BR03	Brasil	Mais de 10 anos	Aproximadamente 2 anos
BR04	Brasil	Entre 5 e 10 anos	Aproximadamente 2 anos
BR05	Brasil	Menos de 1 ano	Aproximadamente 2 anos
CN01	China	Entre 5 e 10 anos	Há alguns anos; duração exata não informada
CN02	China	Entre 1 e 5 anos	Aproximadamente 2 anos
CN03	China	Menos de 1 ano	Aproximadamente 2 anos
CN04	China	Não informado	Aproximadamente 3 a 4 anos
CN05	China	Entre 1 e 5 anos	Aproximadamente 4 anos

Fonte: Elaborado pela Autora (2026)

Participaram da pesquisa 10 gestores, sendo 5 brasileiros e 5 chineses, pertencentes a diferentes setores econômicos e níveis hierárquicos. Buscou-se contemplar diversidade de experiências gerenciais e distintos níveis de familiaridade com ferramentas de IAG, favorecendo a identificação de diferentes formas de incorporação da tecnologia às rotinas organizacionais.

4.2.2 PERFIL DE UTILIZAÇÃO DA IA GENERATIVA

A caracterização do perfil de utilização da IAG entre os participantes buscou contextualizar o nível de familiaridade dos gestores com essas tecnologias antes da análise das categorias emergentes. Foram considerados aspectos como tempo de utilização de IAG, frequência de uso, tipos de ferramentas empregadas e principais finalidades atribuídas à tecnologia no contexto das atividades gerenciais.

De modo geral, observou-se que os participantes apresentaram diferentes níveis de experiência com sistemas de IAG, variando desde usuários com utilização recente até gestores que incorporaram essas ferramentas de forma contínua às suas rotinas de trabalho. Essa heterogeneidade constitui característica relevante da amostra, pois permite examinar diferentes estágios de incorporação da tecnologia e suas possíveis relações com a construção da confiança e com os processos de co-criação analisados posteriormente.

No que se refere às ferramentas utilizadas, verificou-se predominância de plataformas de IAG voltadas à produção e organização de informações, elaboração de textos, apoio à tomada de decisão e desenvolvimento de soluções para problemas cotidianos de gestão. Embora existam diferenças entre os contextos nacionais analisados quanto à disponibilidade e ao ecossistema tecnológico de cada país, o levantamento teve caráter exclusivamente descritivo nesta etapa, buscando apenas caracterizar os recursos efetivamente utilizados pelos participantes.

No conjunto brasileiro, foram mencionadas diferentes plataformas de IAG, entre elas Claude, Gemini, ChatGPT, Perplexity e NotebookLM, além de ferramentas e ambientes de uso mais especializados. A apresentação agregada dessas informações busca caracterizar o ecossistema de uso sem associar combinações específicas de ferramentas a participantes individuais.

A frequência de utilização também apresentou variações entre os gestores, abrangendo desde usos ocasionais até utilização intensiva integrada às atividades diárias. Da mesma forma, as finalidades atribuídas à IAG envolvem diferentes tipos de tarefas, incluindo apoio à elaboração de documentos, síntese de informações, geração de ideias, planejamento, comunicação organizacional, análise de dados e suporte a decisões gerenciais. Essas informações constituem um panorama inicial da presença da IAG nas rotinas profissionais dos participantes, servindo como base contextual para interpretação das categorias analíticas desenvolvidas nas seções subsequentes.

Os dados detalhados referentes ao perfil de utilização da IAG pelos participantes encontram-se sintetizados no Quadro 9, que apresenta, para cada contexto nacional, as principais características relacionadas ao tempo de utilização, frequência de uso, ferramentas empregadas e aplicações predominantes.

Quadro 9 - Perfil de utilização da IA

Aspecto	Brasil	China
Tempo médio de uso	Concentração em cerca de 1,5 a 2 anos de uso.	Em torno de 2 a 4 anos entre os participantes que informaram duração precisa; CN01 relatou apenas uso há alguns anos.
Ferramentas predominantes	Claude, Gemini, ChatGPT, Perplexity e NotebookLM; uso avançado de modelos locais/APIs em um caso.	DeepSeek, ChatGPT, Claude e Kimi; também foram citadas Doubao e Qwen.
Frequência de utilização	Predominantemente diária ou intensiva, com variação entre ferramentas e tarefas.	Predominantemente diária entre os casos que informaram frequência; alguns relatam uso várias vezes ao dia.
Principais aplicações	Pesquisa, síntese, conteúdo, organização, análise de dados, programação, planejamento e apoio à decisão.	Pesquisa e síntese, conteúdo, gestão de projetos, relatórios, automação/agentes, análise de produto/mercado e apoio à decisão.

Fonte: Elaborado pela Autora (2026)

4.2.3 CARACTERIZAÇÃO DE CONTEXTO

Os gestores participantes desta pesquisa atuam em contextos institucionais distintos, cujas características contribuem para compreender diferentes formas de incorporação da IAG às práticas organizacionais. Embora o foco da investigação recaia sobre as experiências dos gestores, a caracterização desses contextos fornece elementos importantes para a interpretação dos resultados apresentados nas seções seguintes.

O contexto brasileiro caracteriza-se por um ambiente regulatório em consolidação, no qual coexistem iniciativas governamentais, setoriais e organizacionais voltadas ao uso da Inteligência Artificial. Observa-se a construção gradual de diretrizes para governança, ética e regulação da IA, coexistindo diferentes níveis de maturidade tecnológica e autonomia organizacional na adoção dessas ferramentas (Serra; Scafuto; Serra Vieira, 2025; Ribeiro; Segatto, 2025).

Em contraste, o contexto chinês apresenta um modelo de governança *state-led*, no qual políticas públicas, investimentos estratégicos e regulamentações específicas orientam o desenvolvimento e a disseminação da Inteligência Artificial. Esse ambiente favorece uma incorporação mais institucionalizada da tecnologia, articulando governo, universidades e setor produtivo em torno de objetivos nacionais de inovação e transformação digital (Zou; Zhang, 2025; Wang et al.,2025).

A caracterização dos contextos considera, ainda, as condições estruturais que diferenciam os dois países no ecossistema global de IA. Jiang et al. (2026) documentam assimetrias de infraestrutura digital, formação de talentos e capacidade industrial entre os membros do BRICS, bem como posições distintas quanto à dependência de tecnologias ocidentais, com o Brasil orientado à diversificação de parcerias com provedores estrangeiros e a China orientada ao desenvolvimento de ecossistema tecnológico próprio. Esses elementos auxiliam a interpretar as diferenças observadas nas ferramentas de IAG utilizadas pelos participantes e nas condições materiais de uso da tecnologia em cada contexto.

As diferenças entre esses contextos não são tratadas como fatores determinantes do comportamento dos gestores, mas como condições institucionais capazes de influenciar a construção da confiança e os processos de co-criação com a Inteligência Artificial Generativa. A partir dessa contextualização, a seção seguinte apresenta os resultados da análise qualitativa, iniciando pela incorporação da IAG às rotinas gerenciais, conforme emergiu das entrevistas realizadas com gestores brasileiros e chineses.

4.2.4 SÍNTESE COMPARATIVA

Quadro 10 - Síntese dos contextos analisados

Aspecto	Brasil	China
Ambiente regulatório	Percebido de forma heterogênea e, por alguns participantes, como ainda difuso; regras organizacionais aparecem mais diretamente na rotina.	Segurança de dados, compliance e limites de uso de ferramentas externas aparecem de forma mais explícita nos relatos.
Estratégia nacional	Pouco mencionada de forma concreta pelos participantes; a governança é percebida sobretudo por meio das organizações.	Maior visibilidade de ferramentas locais e de restrições/condições nacionais que afetam disponibilidade e escolha de sistemas.

Papel das organizações	Incentivo à produtividade combinado a restrições sobre dados confidenciais e, em alguns casos, ferramentas autorizadas.	Forte incentivo ao uso, sistemas internos/aprovados, treinamento e regras claras para dados sensíveis e outputs de maior risco.
Adoção da IA	Heterogênea, de usos básicos a arquiteturas avançadas com agentes e modelos locais.	Cotidiana e orientada à eficiência, com adoção facilitada por ferramentas locais e sistemas corporativos.
Governança	Descentralizada e organizacionalmente variável; autocontrole e julgamento do usuário permanecem relevantes.	Mais explicitamente formalizada na escolha de ferramentas, segurança de dados, compliance e revisão humana.

Fonte: Elaborado pela Autora (2026)

4.3 INCORPORAÇÃO DA IAG NAS ROTINAS GERENCIAIS (OBJETIVO ESPECÍFICO A)

Esta subseção mapeará as formas de incorporação da IAG nas rotinas de gestores em cada contexto, identificando as principais finalidades de uso, tais como a agilização de tarefas operacionais, o apoio ao aprendizado, a qualificação da comunicação e o subsídio à decisão estratégica. Os achados foram contrastados com os padrões identificados por Mangoni (2025) no contexto brasileiro, verificando-se convergências e divergências no contexto chinês, e foram interpretados à luz da concepção de IAG como tecnologia sociotécnica que reconfigura práticas de trabalho (Brynjolfsson; Li; Raymond, 2025; Johri; Schleiss; Ranade, 2025).

As categorias apresentadas a seguir no Quadro 11, partem da análise das entrevistas e foram interpretadas à luz da literatura sobre tecnologias sociotécnicas (Brynjolfsson; Li; Raymond, 2025; Johri; Schleiss; Ranade, 2025). Sempre que pertinente, os resultados foram comparados aos achados de Mangoni (2025), buscando identificar convergências, divergências e novas evidências da perspectiva intercultural adotada nesta pesquisa.

Quadro 11 - Categorias emergentes da incorporação da IAG nas rotinas gerenciais

Categoria	Brasil	China
Pesquisa, síntese e contextualização	Categoria transversal aos cinco casos.	Categoria transversal aos cinco casos.

Comunicação e produção de conteúdo	Recorrente em todos os casos, com intensidades distintas.	Muito recorrente, especialmente em atualizações, reports e marketing; menos central em CN04.
Organização cognitiva e operacional	Inclui organização de ideias, documentos, tarefas e fluxos.	Inclui roadmaps, projetos, reportings, social listening e organização de análises.
Apoio à análise e decisão	IA fornece dados, alternativas e estrutura, sem substituir a decisão final.	IA estrutura comparações, cenários e alternativas; decisão permanece humana.
Agentes, automação e integração de fluxos	Presente em casos de maior maturidade técnica, incluindo agentes, APIs e ambientes locais.	Presente de forma mais forte em CN02 e em sistemas internos; não é transversal a todos os casos.
Aprendizado e sofisticação progressiva do uso	Evolução de consultas simples para melhor contextualização, agentes e fluxos especializados.	Evolução de prompts básicos para contexto mais rico, automações e uso incorporado ao workflow.
IAG como ponto de partida, não substituto do julgamento	Panoramas, rascunhos e alternativas são posteriormente avaliados e aprofundados.	Rascunhos e análises funcionam como ponto inicial, com revisão humana especialmente em tarefas críticas.
Incorporação seletiva conforme criticidade	Maior uso direto em tarefas simples e verificáveis; maior supervisão em tarefas críticas.	Mesmo padrão, com ênfase adicional em dados sensíveis, compliance e contexto local.

Fonte: Elaborado pela Autora (2026)

Os relatos mostram que a incorporação da IAG não se restringe ao aumento da frequência de uso, mas envolve progressiva reorganização da forma como os gestores estruturam o trabalho. No conjunto brasileiro, observaram-se trajetórias que partiram de consultas e pesquisas pontuais e passaram a incluir agentes, ambientes especializados e rotinas automatizadas. Em um dos casos, por exemplo, o participante descreveu a transição de atividades antes realizadas manualmente para agentes capazes de executar rotinas recorrentes de pesquisa e monitoramento (BR01). Em outros casos, a IAG permaneceu predominantemente como instrumento de pesquisa, síntese e ponto de partida para aprofundamento humano.

No conjunto chinês, também se observou progressão de usos simples para integração mais sistemática ao fluxo de trabalho. CN01 relatou uma curva de aprendizagem marcada pelo aumento do contexto fornecido à ferramenta e pela incorporação da IAG à organização de projetos; CN02 descreveu a passagem de usos experimentais para automações e agentes; e CN05 indicou que a utilização deixou de

ser percebida como atividade separada e passou a constituir parte ordinária do workflow. A heterogeneidade interna dos dois grupos indica, portanto, que a maturidade de uso varia substancialmente entre participantes e não pode ser atribuída exclusivamente ao contexto nacional.

Os dados também evidenciam que o aumento da integração não corresponde à substituição do julgamento gerencial. A IAG é utilizada predominantemente para ampliar velocidade, organizar informação, gerar alternativas e estruturar pontos de partida. Quanto maior a criticidade, a ambiguidade ou a consequência potencial da tarefa, maior tende a ser a intervenção humana posterior. Esse padrão prepara a relação entre incorporação e confiança examinada na seção seguinte.

4.4 CONSTRUÇÃO, PERCEPÇÃO E MOBILIZAÇÃO DA CONFIANÇA NA IA (OBJETIVO ESPECÍFICO B)

Considerando os padrões de incorporação da IAG identificados anteriormente, esta subseção examinará como a confiança é construída, percebida e mobilizada pelos gestores nos contextos brasileiro e chinês. A análise considerará suas dimensões cognitiva, afetiva e comportamental (GLIKSON; WOOLLEY, 2020; YANG; WIBOWO, 2022), bem como a influência da experiência de uso e de fatores institucionais em sua formação (BUNDUCHI; SITAR-TĂUT; MICAN, 2025; RIBEIRO; SEGATTO, 2025). A comparação buscou identificar convergências e diferenças entre os contextos e verificar em que medida os padrões observados se aproximam ou se afastam das relações discutidas na fundamentação teórica. Os resultados foram sintetizados no Quadro 12.

Os resultados foram comparados entre os contextos brasileiro e chinês, como pode ser observado no Quadro 12, buscando compreender em que medida a mediação institucional da confiança se manifesta na forma antecipada pela fundamentação teórica e se, no contexto brasileiro, a confiança emerge de modo mais situado e dependente da experiência de uso (RIBEIRO; SEGATTO, 2025).

Quadro 12 - Comparação da construção, percepção e mobilização da confiança na IAG entre gestores brasileiros e chineses

Aspecto analisado	Brasil	China	Síntese comparativa
-------------------	--------	-------	---------------------

Construção da confiança	Experiencial, dinâmica e diferenciada por ferramenta e tarefa; acertos e erros recalibram expectativas.	Experiencial e progressiva, também influenciada por testes, adequação local e legitimidade de ferramentas/ambientes aprovados.	Em ambos, confiança é construída e revisada ao longo do uso; no corpus chinês, limites organizacionais/institucionais aparecem com maior saliência.
Critérios para confiar	Aderência à tarefa, fontes, dados objetivos, coerência, previsibilidade e possibilidade de verificação.	Referências, precisão factual, consistência lógica, adequação ao contexto de negócio/local e clareza sobre tratamento de dados.	Verificabilidade e fit são comuns; China adiciona ênfase recorrente em contexto local, segurança e compliance.
Critérios para validar respostas	Double-check, texto integral, comparação entre ferramentas, fontes públicas, especialistas e revisão manual.	Documentos originais e internos, papers/relatórios, links, primeira saída de automações e revisão humana.	Validação independente é transversal aos dois contextos.
Situações de maior confiança	Tarefas simples, reversíveis, bem delimitadas e outputs ancorados em fontes.	Tarefas internas/estruturais de baixo risco, automações já testadas e tarefas com contexto bem delimitado.	Baixa criticidade e alta verificabilidade ampliam a confiança comportamental.
Situações de menor confiança	Tarefas complexas, dados inexistentes/desatualizados, outputs pouco verificáveis e alta consequência.	Números/benchmarks, documentos importantes, legal/compliance, dados confidenciais e uso de sistemas inadequados ao contexto.	Maior risco, sensibilidade ou baixa verificabilidade reduzem delegação.
Influência da organização	Incentivo à eficiência com restrições de confidencialidade e, em alguns casos, ambientes autorizados.	Forte incentivo, sistemas internos/aprovados, workshops e revisão humana obrigatória em outputs de maior risco.	Organizações simultaneamente promovem e limitam o uso; a formalização aparece mais uniforme nos casos chineses.
Influência do contexto institucional	Percepções heterogêneas; a regulação nacional aparece pouco na prática cotidiana de parte da amostra.	Disponibilidade de ferramentas, segurança de dados, compliance e preferência/uso de sistemas locais afetam diretamente escolhas.	O contexto institucional aparece mais explicitamente operacionalizado no corpus chinês.
Impacto da confiança na tomada de decisão	A IA pode fornecer fontes, alternativas ou challenge, mas a decisão final permanece humana.	A IA organiza referências, compara opções e aponta riscos, mas estratégia, recursos e compliance permanecem sob	Confiança aumenta a influência informacional da IA, não transfere a autoridade decisória final.

		juízo humano.	
Impacto da confiança na co-criação	Regula quanto delegar, quantas rodadas realizar e quando retomar manualmente.	Regula automação, profundidade de revisão e necessidade de supervisão; falhas podem levar a decomposição ou retirada.	Confiança opera como regulador da participação da IAG nos dois contextos.

Fonte: Elaborado pela Autora (2026)

A análise evidencia que a confiança não se apresenta nos relatos como avaliação binária ou permanente da IAG. Os participantes diferenciam a confiabilidade da ferramenta conforme a tarefa, o tipo de informação, o grau de risco e a possibilidade de verificar o resultado. Assim, a experiência acumulada não produz necessariamente aumento linear da confiança; produz maior discriminação sobre onde, quando e em que condições delegar.

A verificabilidade emergiu como mecanismo central nos dois conjuntos. Quando o output envolve números, benchmarks, fatos, regulamentação ou informações passíveis de rastreamento, os gestores ampliam a checagem e recorrem a documentos originais, fontes externas, outras ferramentas ou conhecimento profissional. CN01, por exemplo, relatou verificar sistematicamente números e *benchmarks*; CN03 afirmou não aceitar números citados pela IA sem confrontá-los com documentos originais e relatórios confiáveis; e CN05 descreveu a necessidade de retornar à fonte original quando tendências ou percentuais apresentados pela ferramenta não podiam ser confirmados.

Os erros não produziram, contudo, abandono automático da tecnologia. Em vários casos, a reação consistiu em modificar o prompt, aumentar restrições, decompor a tarefa ou retomar a execução manual. O efeito predominante é, portanto, de recalibração. Essa dinâmica é particularmente clara quando a confiança diminuída em um episódio pode ser restabelecida em interações subsequentes satisfatórias. O objeto da confiança também se mostrou localizado: um participante pode desconfiar da capacidade da IAG para dados factuais e, simultaneamente, confiar amplamente em sua utilização para geração de alternativas, estruturação ou ideação.

Essa distinção leva a um resultado relevante: confiança na IAG não equivale à transferência de autoridade decisória. Mesmo quando a ferramenta exerce influência informacional, os participantes preservam o juízo final. CN01, por exemplo, condiciona a aceitação à compatibilidade entre a sugestão e sua expertise; CN03 utiliza

a IAG para organizar alternativas, mas mantém sob julgamento próprio adequação ao roadmap, recursos e compliance; e CN05 utiliza a ferramenta para desafiar diferentes direções criativas sem transferir a ela a decisão sobre a direção final.

A confiança pode, portanto, ser interpretada menos como autorização geral para o uso da tecnologia e mais como mecanismo de calibração de sua participação. Ela regula quanto pode ser delegado, quanto precisa ser verificado, quantas interações adicionais são justificáveis e em que momento o gestor interrompe a colaboração e retoma a tarefa manualmente.

4.5 PROCESSOS DE INTERAÇÃO E CO-CRIAÇÃO HUMANO-IA (OBJETIVO ESPECÍFICO C)

Esta subseção caracterizará os processos de interação e co-criação entre gestores e sistemas de IAG, descrevendo os ciclos iterativos de formulação, geração e refinamento de *outputs* e o papel do julgamento humano na direção, interpretação e validação das respostas (Prahalad; Ramaswamy, 2004; Mangoni, 2025). Foram identificadas as estratégias de co-criação mobilizadas em cada contexto e analisada a relação entre o nível de confiança e o grau de delegação cognitiva, incluindo a eventual manifestação de padrões de *overreliance* (Keding; Meissner, 2021).

Com objetivo de compreender como se desenvolvem os processos de interação e co-criação entre gestores e sistemas de IAG, a análise foi organizada em duas perspectivas complementares. A primeira, apresentada no Quadro 13, sintetiza, de forma comparativa, as principais características observadas nos contextos brasileiro e chinês, permitindo identificar convergências e particularidades entre os dois grupos de participantes. A segunda, apresentada no Quadro 14, representa o processo de cocriação propriamente dito, evidenciando as etapas de interação entre gestor e IAG, o papel desempenhado por cada agente ao longo do processo e as categorias analíticas emergentes que sustentam a interpretação dos resultados. Os Quadros 13 e 14, apresentam essa estrutura de análise que fundamenta a discussão desenvolvida nesta seção.

Quadro 13 - Comparação dos processos de interação e co-criação entre gestores e IAG

Aspecto analisado	Brasil	China	Síntese comparativa
Finalidade predominante da interação	Pesquisa, comunicação, organização, análise, criação e solução de problemas.	Eficiência operacional, conteúdo, projetos, análise de produto/mercado, pesquisa e automação.	Finalidades amplas e convergentes, com variação por função e maturidade.
Forma de interação com IAG	Conversação iterativa; em casos avançados, agentes, APIs e ambientes integrados.	Conversação iterativa, <i>trial-and-error</i> , sistemas internos e automações em alguns casos.	Iteratividade é comum; arquitetura técnica varia entre participantes.
Estratégias da formulação de prompts	Contexto prévio, perguntas genéricas seguidas de especificação, restrições de fonte e testes.	Contexto detalhado, escopo, audiência/objetivo, restrições locais/regulatórias e exemplos positivos/negativos.	Maior qualidade de contexto é tratada como condição para outputs utilizáveis.
Processo de refinamento das respostas	Reescrita de prompts, correções incrementais, comparação de modelos e explicação da linha de raciocínio.	Feedback, adição de restrições, decomposição em subtarefas, novas rodadas e checagem de fontes.	Refinamento é componente estrutural da co-criação.
Participação do gestor na construção da solução	Define intenção, critérios, validação e decisão de aplicação.	Define escopo, contexto local, restrições, seleção de alternativas, validação e decisão final.	O gestor dirige o processo em ambos os contextos.
Papel do julgamento humano	Autoridade final; verifica, transforma, rejeita ou interrompe.	Autoridade final; avalia precisão, contexto, compliance, estratégia e viabilidade.	Julgamento humano é mecanismo final de autorização.
Grau de delegação cognitiva	Seletivo e dependente de criticidade, verificabilidade, domínio e tempo.	Seletivo; maior em baixo risco e fluxos testados, menor em fatos críticos, dados sensíveis e decisões estratégicas.	Delegação é calibrada, não irrestrita.
Evidências de <i>overreliance</i>	Risco episódico, com pelo menos um caso concreto de quase aceitação de recomendação inadequada e outras vulnerabilidades percebidas.	Não se configurou padrão empírico; participantes reconhecem explicitamente o risco e relatam mecanismos de checagem.	<i>Overreliance</i> não é predominante; a evidência concreta foi mais clara no corpus brasileiro.

Forma predominante de co-criação	Produção assistida, refinamento iterativo, orquestração de agentes e retomada manual.	Geração de alternativas, <i>trial-and-error</i> , refinamento contextual, automação testada e retomada manual em casos de falha.	Co-criação alterna geração algorítmica e julgamento humano.
Impacto na tomada de decisão	Influência indireta por fontes e alternativas validadas; decisão permanece humana.	Influência por comparação de opções, challenge e informação estruturada; decisão permanece humana.	IA amplia o espaço informacional, sem assumir a decisão final.

Fonte: Elaborado pela Autora (2026)

Quadro 14 - Etapas do processo de co-criação humano-IA identificadas na pesquisa

Etapa do processo	Papel do gestor	Papel da IAG	Evidências Empíricas
Formulação do problema	Define necessidade, objetivo, criticidade e o que não pode ser delegado.	Pode sugerir questões ou alternativas iniciais quando acionada.	Recorrente em BR01–BR05 e CN01–CN05; tarefas mais críticas são delimitadas antes da interação.
Construção do prompt	Seleciona contexto, fontes, restrições, audiência, parâmetros e dados permitidos.	Recebe e processa instruções e contexto.	BR01, BR03–BR05; CN01–CN03 e CN05 descrevem explicitamente aumento de contexto e restrições.
Geração de resposta	Solicita panorama, rascunho, alternativas, síntese ou análise.	Gera texto, estrutura, dados, ideias, comparação, código ou proposta inicial.	Presente nos dez casos, com diferentes níveis de integração.
Avaliação crítica	Verifica aderência, factualidade, fontes, risco, contexto e viabilidade.	Pode fornecer referências, explicar ou revisar a própria resposta.	Fontes/verificabilidade aparecem em BR01–BR05 e CN01–CN05.
Refinamento iterativo	Reformula, corrige, adiciona critérios, seleciona partes úteis e fornece feedback.	Gera novas versões e incorpora feedback.	Recorrente nos dois corpora; CN03 e CN05 também relatam decomposição ou abandono após várias falhas.
Validação da solução	Confirma em fontes, documentos, outra ferramenta, especialista ou conhecimento próprio.	Fornece versão refinada e, quando possível, lastro documental.	A intensidade da validação cresce com criticidade, dados factuais e sensibilidade.

Aplicação da decisão	Aceita, edita, aplica, automatiza ou retoma a execução manualmente.	Participa do resultado final ou é retirada do processo.	Nos dois contextos, a ação final permanece sob responsabilidade humana.
----------------------	---	---	---

Fonte: Elaborado pela Autora (2026)

Embora o Quadro 14 organize o processo em etapas analíticas, os dados não indicam uma sequência rígida ou universal. As etapas podem ser abreviadas, repetidas ou interrompidas conforme o tipo de tarefa. Em atividades simples e facilmente verificáveis, geração e aplicação podem ocorrer com poucas rodadas; em atividades complexas, o processo envolve ciclos sucessivos de contextualização, geração, avaliação e refinamento. Em automações, surge ainda uma etapa de teste do fluxo antes de sua utilização recorrente.

Também emergiu uma trajetória de retirada da IAG. Quando sucessivos refinamentos não produzem resultado satisfatório, alguns participantes reduzem a delegação, decompõem o problema ou abandonam a assistência da ferramenta e retomam a execução manual. Essa possibilidade de reversão é analiticamente importante porque demonstra que a co-criação não corresponde a uma progressão contínua em direção a maior automação. A participação da IAG é continuamente renegociada a partir de seu desempenho percebido.

Assim, mais do que uma cadeia linear de etapas, a co-criação observada pode ser compreendida como processo iterativo e reversível, no qual geração algorítmica e julgamento humano se alternam até que o gestor considere o resultado suficientemente adequado, ou determine que a continuidade da interação deixou de ser produtiva.

4.6 ANÁLISE COMPARATIVA E DISCUSSÃO INTEGRADA BRASIL-CHINA (OBJETIVO ESPECÍFICO D)

Esta subseção integra os resultados relativos à incorporação da IAG, à construção da confiança e aos processos de cocriação, examinando as relações identificadas entre essas dimensões nos contextos brasileiro e chinês. A análise buscou compreender como as formas de incorporação da tecnologia se relacionam com as experiências e os critérios de confiança dos gestores e como, por sua vez, esses processos se refletem nas práticas de validação, delegação, refinamento e construção conjunta de resultados com a IAG. A comparação entre os contextos permitiu identificar

em que medida fatores culturais, institucionais e organizacionais participam dessas relações, sem pressupor previamente a direção ou a intensidade de seus efeitos.

A análise dos casos divergentes também limita interpretações excessivamente generalizantes. Nem todos os participantes relataram processos sofisticados de co-criação ou influência direta da IAG sobre decisões. Da mesma forma, quando questionados explicitamente sobre características culturais nacionais, alguns participantes do contexto chinês não identificaram fatores culturais específicos e deslocaram sua explicação para características do ambiente de trabalho, disponibilidade de ferramentas e regras organizacionais. Esses casos negativos foram preservados na interpretação e constituem evidência contra uma leitura cultural nacional determinista.

Quadro 15 - Síntese comparativa dos resultados entre Brasil e China

Dimensão analítica	Brasil	China	Síntese comparativa
Incorporação da IAG	Heterogênea, de consultas conversacionais a agentes e ambientes locais; orientada por produtividade e ampliação cognitiva.	Cotidiana e orientada à eficiência, de prompts e drafts a automações e sistemas internos; forte recorrência de ferramentas locais.	Nos dois contextos há passagem de experimentação para integração ao fluxo; a forma técnica varia mais por maturidade e ambiente do que por país isoladamente.
Finalidades predominantes de utilização	Pesquisa, síntese, comunicação, organização, análise, programação e desenvolvimento de soluções.	Pesquisa, síntese, conteúdo, gestão de projetos, automação, análise de produto/mercado e apoio à decisão.	Predomina uso de <i>augmentation</i> : a IAG amplia capacidade, velocidade e espaço de alternativas.
Construção da confiança	Experiencial, dinâmica e diferenciada por ferramenta e tarefa.	Experiencial e dinâmica, combinando desempenho percebido, adequação local e legitimidade do ambiente/ferramenta.	A confiança é recalibrada em ambos; no corpus chinês, condições institucionais e organizacionais aparecem mais explicitamente no processo.
Critérios de validação das respostas	Fontes, dados, texto integral, comparação de modelos, especialistas e revisão manual.	Documentos originais/internos, referências, links, precisão factual, contexto local, compliance e revisão humana.	Verificabilidade é transversal; a China adiciona maior recorrência de adequação local e compliance como critérios explícitos.
Influência do contexto institucional	Regulação percebida de modo heterogêneo;	Disponibilidade de ferramentas, segurança	O contexto influencia os dois grupos, mas sua

	na rotina, organizações e autocontrole aparecem como principais delimitadores de uso e dados.	de dados, sistemas aprovados, compliance e preferência por soluções locais delimitam escolhas de forma mais visível.	operacionalização é mais explícita e uniforme entre os casos chineses.
Papel do julgamento humano	Autoridade final e componente indispensável da confiabilidade.	Autoridade final para validar fatos, adequação local, estratégia, compliance e viabilidade.	Convergência forte: a IA influencia, mas não recebe autoridade final.
Grau de delegação cognitiva	Seletivo, maior em tarefas simples e menor em tarefas críticas ou pouco verificáveis.	Seletivo, maior em baixo risco e fluxos testados; menor em dados sensíveis, fatos críticos e decisões estratégicas.	A criticidade e a verificabilidade regulam a delegação nos dois contextos.
Padrões de co-criação	Produção assistida, refinamento iterativo, orquestração de agentes e retomada manual.	Geração de alternativas, <i>trial-and-error</i> , refinamento contextual, automação testada e retomada manual.	Co-criação é predominantemente iterativa e reversível, com redistribuição dinâmica de tarefas.
Evidências de over reliance	Risco episódico, sobretudo em temas desconhecidos ou sob conveniência/pressão de tempo; há caso concreto de quase aceitação inadequada.	Não emergiu como padrão concreto; participantes demonstram consciência do risco e práticas sistemáticas de checagem.	Não há suporte para tratar <i>overreliance</i> como padrão dominante; o risco existe, mas é geralmente mitigado por validação.
Síntese interpretativa	A confiança regula a participação da IAG por meio de delegação, supervisão, refinamento e possibilidade de retirada.	A confiança exerce a mesma função reguladora, com maior saliência de fronteiras institucionais, segurança e ferramentas aprovadas.	A confiança funciona como mecanismo de calibração da co-criação; diferenças contextuais alteram as condições e fronteiras dessa calibração, não a centralidade do julgamento humano.

Fonte: Elaborado pela Autora (2026)

A comparação evidencia que as diferenças entre os dois contextos são menores no mecanismo básico da co-criação do que nas condições que delimitam sua realização. Nos dois grupos, a participação da IAG tende a aumentar em tarefas de menor criticidade e maior verificabilidade e a diminuir quando a atividade envolve informações sensíveis, consequências relevantes ou outputs de difícil confirmação. Também em ambos os conjuntos, a utilização mais intensa da tecnologia não implica

transferência da autoridade final, permanecendo com o gestor a decisão de validar, modificar, rejeitar ou aplicar o resultado.

Nesse sentido, os dados sustentam apenas parcialmente a expectativa inicial de que diferenças nacionais produziriam formas substantivamente distintas de confiança. O mecanismo central identificado é amplamente convergente: a confiança é produzida pela experiência e recalibrada pelo desempenho percebido; a confiança recalibrada, por sua vez, redefine a quantidade de delegação, supervisão e refinamento empregada na interação seguinte. O processo assume, portanto, caráter recursivo, e não estritamente linear.

A principal diferença entre os contextos aparece nas fronteiras dentro das quais essa calibração ocorre. Entre os participantes chineses, ferramentas aprovadas, sistemas internos, segurança de dados, adequação ao contexto local e compliance são mencionados de forma particularmente recorrente. CN01 relata que informações sensíveis devem ser tratadas em ferramentas aprovadas; CN02 descreve sistemas internos e canais protegidos; CN03 associa explicitamente a utilização de material confidencial a plataformas internas autorizadas; CN04 enfatiza as limitações ligadas a dados no setor em que atua; e CN05 diferencia ferramentas permitidas conforme a sensibilidade da informação.

Essa recorrência não autoriza, contudo, afirmar que gestores chineses possuem maior ou menor confiança na IAG. O que os dados demonstram é que o contexto institucional e organizacional aparece mais explicitamente incorporado aos critérios de uso. Em outras palavras, as fronteiras de utilização são mais frequentemente verbalizadas por meio da adequação da ferramenta, da segurança do ambiente e das exigências de compliance.

No conjunto brasileiro, também aparecem restrições relativas a confidencialidade, proteção de dados e ambientes autorizados, o que impede tratar esses fatores como exclusivos do contexto chinês. BR01 descreve simultaneamente forte incentivo organizacional e proibição de inserção de informações confidenciais; outro participante relata que ferramentas não licenciadas não podem receber informações corporativas; e BR05 menciona a exigência de ambiente controlado para conteúdos confidenciais. A diferença observada é de saliência e uniformidade, e não de presença versus ausência dessas preocupações.

A dimensão cultural apresentou evidência mais limitada do que a dimensão organizacional e institucional. Alguns participantes brasileiros recorreram

espontaneamente a interpretações sobre características nacionais, enquanto outros não identificaram influência clara. Entre os chineses ocorreu padrão semelhante: alguns associaram o contexto a pragmatismo, velocidade de trabalho e atenção a risco, enquanto outros recusaram atribuir diretamente seus comportamentos à “cultura chinesa”. Assim, os resultados não permitem inferir que valores nacionais sejam responsáveis pelas diferenças observadas. A comparação indica, de forma mais robusta, que condições organizacionais, regulatórias, profissionais e tecnológicas participam conjuntamente da construção das fronteiras de confiança.

Esse resultado modifica o modelo analítico inicial. O contexto não atua apenas como antecedente que molda a confiança, e a confiança não produz de maneira unidirecional a co-criação. As próprias interações retroalimentam o processo: a IAG é utilizada em determinada tarefa; seu desempenho é avaliado; erros, acertos e grau de verificabilidade recalibram a confiança; essa confiança redefine a delegação e a supervisão na interação subsequente; e os novos resultados voltam a informar avaliações futuras. A co-criação configura-se, portanto, como processo recursivo de calibração.

A contribuição analítica central do estudo reside nessa passagem de uma concepção estática de divisão de tarefas para uma concepção dinâmica de participação tecnológica. A pergunta gerencial não é apenas se a IAG será utilizada, mas quanto espaço ela receberá em determinado episódio de trabalho e sob quais mecanismos de controle. Criticidade, verificabilidade, sensibilidade da informação, desempenho acumulado e condições organizacionais funcionam como parâmetros dessa calibração. Desse modo, a confiança representa o mecanismo pelo qual os gestores ajustam continuamente as fronteiras entre contribuição algorítmica e julgamento humano.

Figura 2 - Trajetórias e mecanismos de articulação entre incorporação da IAG, confiança e cocriação humano-IA nos contextos brasileiro e chinês



Fonte: Elaborado pela Autora (2026) com base nos resultados da pesquisa, com apoio do ChatGPT na geração da representação visual. O prompt integral encontra-se no Apêndice A.

A Figura 2 sintetiza o modelo analítico revisado a partir dos resultados empíricos. A representação evidencia o caráter recursivo das relações entre contexto, experiência de uso, confiança e co-criação, mostrando que erros, acertos e condições de verificabilidade recalibram a confiança e, consequentemente, redefinem o grau de delegação, supervisão e refinamento nas interações subsequentes. As diferenças observadas entre Brasil e China incidem principalmente sobre as condições e fronteiras dessa calibração, sem alterar a centralidade do julgamento humano identificada nos dois conjuntos.

5 CONSIDERAÇÕES FINAIS

A crescente incorporação da Inteligência Artificial Generativa (IAG) às práticas organizacionais desloca o debate da simples adoção tecnológica para a compreensão de como humanos e sistemas inteligentes passam a compartilhar atividades cognitivas. Nesse contexto, esta pesquisa investigou como a confiança na IAG influencia os processos de co-criação entre gestores e sistemas generativos, em perspectiva comparativa entre Brasil e China.

A análise temática, apoiada por procedimentos de codificação inspirados na Grounded Theory, mostrou que a confiança não constitui um atributo fixo da tecnologia nem uma disposição estável do gestor. Ela é construída e recalibrada ao longo das interações, a partir da experiência acumulada, do desempenho percebido, da verificabilidade dos outputs, da criticidade das tarefas e das condições organizacionais e institucionais em que a tecnologia é utilizada.

O principal achado da pesquisa é que a confiança atua como mecanismo de calibração da participação da IAG no trabalho gerencial. Ao invés de simplesmente confiar ou não confiar, os gestores ajustam quanto delegam à tecnologia, quanto verificam seus resultados, quantas rodadas de refinamento realizam e quando aceitam, modificam, rejeitam ou retomam manualmente uma tarefa. Tarefas de menor criticidade e facilmente verificáveis favorecem maior delegação, enquanto atividades de maior consequência, informações sensíveis ou *outputs* pouco verificáveis ampliam a supervisão humana. Erros e alucinações, por sua vez, tendem a recalibrar a confiança, sem necessariamente eliminar o uso da tecnologia.

A comparação revelou maior convergência do que oposição entre os dois contextos. Tanto entre brasileiros quanto entre chineses, a co-criação ocorre predominantemente de forma iterativa, combinando geração algorítmica, avaliação, feedback e refinamento, enquanto o julgamento humano permanece como instância final de validação e decisão. A principal diferença observada está nas condições que delimitam essa relação: entre os participantes chineses, ferramentas aprovadas, segurança de dados, compliance e restrições institucionais aparecem de forma mais recorrente e explícita; entre os brasileiros, esses limites surgem de maneira mais heterogênea e associados principalmente às políticas organizacionais, à experiência individual e ao julgamento crítico. Os dados, portanto, não sustentam uma explicação

cultural nacional determinista, mas indicam que fatores institucionais, organizacionais, profissionais e tecnológicos participam conjuntamente da formação da confiança.

Como principal contribuição analítica, o estudo propõe compreender a confiança não como um estado antecedente ao uso da IAG, mas como mecanismo recursivo de calibração de sua participação no trabalho gerencial. A cada interação, desempenho percebido, verificabilidade, criticidade da tarefa e condições contextuais atualizam a confiança atribuída ao sistema; essa confiança recalibrada redefine, por sua vez, o grau de delegação, supervisão, refinamento e eventual retirada da IAG nas interações subsequentes.

Os achados também oferecem implicações para a gestão da IAG nas organizações. Ao indicar que a confiança é continuamente calibrada em função da criticidade, da verificabilidade, do desempenho percebido e das condições institucionais de uso, os resultados permitem traduzir os mecanismos identificados empiricamente em práticas de gestão voltadas à definição dos limites de delegação, dos níveis de supervisão e dos critérios de validação dos outputs. O Quadro 16 sintetiza essas implicações, relacionando os principais achados da pesquisa a recomendações gerenciais e aos objetivos associados à sua aplicação.

Quadro 16 - Implicações gerenciais decorrentes dos resultados da pesquisa

Achado da pesquisa	Recomendação gerencial	Objetivo
A confiança varia conforme criticidade e verificabilidade	Classificar as tarefas segundo criticidade, risco e possibilidade de verificação antes de ampliar a participação da IAG	Adequar o nível de delegação e supervisão
Outputs factuais demandam validação	Estabelecer procedimentos de conferência em fontes independentes e originais	Reduzir a incorporação de informações incorretas
A delegação ocorre de forma seletiva	Definir limites de participação da IAG conforme o tipo e a consequência da atividade	Evitar delegação excessiva em tarefas críticas
O refinamento iterativo integra a co-criação	Incorporar ciclos de geração, avaliação, feedback e refinamento	Aumentar a adequação e a qualidade dos outputs
O julgamento humano permanece como instância final	Definir responsabilidade humana pela aprovação dos resultados e das decisões	Preservar julgamento, responsabilidade e accountability
Segurança, compliance e condições institucionais	Estabelecer ferramentas autorizadas e regras para	Reduzir riscos informacionais e institucionais

delimitam o uso	tratamento de informações sensíveis	
Falhas provocam recalibração e eventual retomada manual	Estabelecer critérios para interromper a interação e retomar a execução humana	Evitar persistência improdutiva e dependência inadequada da ferramenta
Condições organizacionais e institucionais delimitam a confiança e a participação da IAG	Desenvolver políticas organizacionais de adoção da IAG que estabeleçam ferramentas autorizadas, tipos de dados permitidos, níveis de autonomia, critérios de validação e responsabilidades humanas conforme a criticidade das atividades	Institucionalizar o uso responsável da IAG e reduzir decisões individuais inconsistentes sobre delegação, segurança e supervisão

Fonte: Elaborado pela Autora com base nos resultados da pesquisa (2026).

Essas recomendações não pressupõem a adoção de um nível uniforme de confiança ou autonomia da IAG. Ao contrário, os resultados sugerem que práticas gerenciais de uso responsável devem permitir a calibração da participação da tecnologia conforme as características da tarefa e do contexto, preservando mecanismos de validação, possibilidade de retomada manual e responsabilidade humana sobre as decisões.

Os resultados devem ser interpretados à luz das limitações desta pesquisa. A amostra qualitativa e intencional, composta por cinco gestores em cada contexto nacional, não permite generalizações estatísticas para as populações de gestores brasileiros ou chineses. Além do número restrito de participantes, os dois conjuntos apresentam heterogeneidade quanto à distribuição setorial, à trajetória profissional e à maturidade no uso da IAG. Assim, as diferenças observadas não podem ser atribuídas exclusivamente ao contexto nacional, podendo também refletir características do setor, da organização, da função exercida e do grau de familiaridade dos participantes com a tecnologia. A comparação deve, portanto, ser compreendida como analítica e exploratória, orientada à identificação de mecanismos, convergências e variações contextuais situadas, e não como representação das populações nacionais investigadas.

Uma limitação adicional decorre da dimensão linguística da pesquisa. As entrevistas com os participantes chineses foram conduzidas em inglês, idioma que pode não corresponder à língua de maior domínio ou espontaneidade de todos os entrevistados. Essa condição pode ter afetado nuances de expressão, precisão semântica e elaboração das respostas. A posterior integração dos resultados em português

acrescentou uma segunda camada interpretativa associada à tradução. Como estratégia de mitigação, a codificação e a análise consideraram os registros originais em inglês, enquanto as traduções foram utilizadas prioritariamente para apresentação e integração comparativa dos resultados. Ainda assim, não se pode excluir integralmente a possibilidade de perdas ou deslocamentos de sentido decorrentes da mediação linguística.

Ademais, o uso de IA como apoio à codificação constitui outra limitação metodológica, mitigada pela validação independente das classificações e pela manutenção das decisões interpretativas sob responsabilidade da pesquisadora. Por fim, o modelo analítico privilegia as relações entre contexto, confiança e co-criação e, embora tenha preservado abertura para categorias emergentes, pode ter reduzido a visibilidade de mecanismos situados fora desse enquadramento.

A partir dessas limitações, estudos futuros podem avançar em diferentes direções. Primeiramente, recomenda-se ampliar a amostra e construir conjuntos comparativos mais equivalentes entre os contextos nacionais quanto a setor de atuação, função gerencial, experiência profissional e maturidade no uso da IAG, permitindo examinar com maior precisão quais variações estão associadas ao contexto institucional e quais decorrem de características profissionais ou organizacionais. Em pesquisas envolvendo o contexto chinês, a realização das entrevistas em mandarim, associada a procedimentos sistemáticos de tradução e retrotradução, poderá reduzir os efeitos da mediação linguística. Desenhos longitudinais também se mostram particularmente relevantes para acompanhar como episódios de acerto, erro, correção e mudanças nas próprias tecnologias modificam, ao longo do tempo, a confiança, os limites de delegação e a intensidade da supervisão humana. Estudos observacionais, experimentais ou multimétodos poderão ainda confrontar as práticas relatadas pelos gestores com interações efetivas com sistemas de IAG e testar as relações identificadas qualitativamente entre verificabilidade, criticidade, confiança e delegação. Por fim, pesquisas futuras podem examinar como diferentes políticas organizacionais de adoção da IAG — incluindo regras de segurança de dados, ferramentas autorizadas, mecanismos de validação e definição de responsabilidades — influenciam a formação e a calibração da confiança em diferentes contextos institucionais.

Quadro 17 - Limitações da pesquisa e agenda de estudos futuros

Limitação identificada	Alcance/ efeito sobre o estudo	Direção para estudos futuros
Amostra qualitativa e intencional de 10 gestores	Não permite generalização estatística	Ampliar a amostra e testar quantitativamente as relações identificadas
Heterogeneidade setorial e profissional	Não permite isolar efeitos de contexto nacional, setor, função ou maturidade em IAG	Realizar estudos setoriais e comparações entre níveis gerenciais e graus de maturidade
Comparação restrita a Brasil e China	Limita a extensão das conclusões para outros contextos institucionais	Incorporar outros países do BRICS e do Sul Global
Entrevistas chinesas realizadas em inglês	Possibilidade de perda de nuances e mediação linguística	Realizar entrevistas em língua nativa, com tradução e retrotradução sistemáticas
Dados baseados em relatos dos participantes	Não permite observar diretamente a prática de cocriação	Desenvolver estudos observacionais e experimentais de interações gestor-IAG
Desenho transversal	Não acompanha a recalibração da confiança ao longo do tempo	Realizar estudos longitudinais sobre confiança, delegação e supervisão
Uso de IA como apoio à codificação	Introduz uma camada adicional de mediação analítica	Comparar processos de codificação humana, assistida por IA e entre múltiplos codificadores humanos
Modelo centrado em contexto–confiança–cocriação	Pode reduzir a visibilidade de mecanismos externos ao enquadramento analítico	Investigar outros mecanismos emergentes e testar modelos alternativos

Fonte: Elaborado pela Autora com base nas limitações identificadas na pesquisa (2026).

Nesse sentido, a agenda de estudos futuros concentra-se em três direções principais: ampliar a diversidade e a comparabilidade das amostras, inclusive por meio da incorporação de outros contextos institucionais e de entrevistas em língua nativa; adotar desenhos longitudinais que permitam acompanhar a recalibração da confiança e da delegação ao longo do tempo; e observar diretamente os processos de co-criação humano-IAG, por meio de abordagens observacionais, experimentais ou multimétodos. Estudos futuros também podem aprofundar como diferentes políticas organizacionais de adoção da IAG influenciam a construção e a calibração da confiança.

Por fim, os achados também contribuem para discussões mais amplas sobre governança e cooperação em IA. Ao evidenciar como a confiança é construída, testada e recalibrada nas práticas cotidianas de gestores, a pesquisa mostra que agendas

institucionais de governança tecnológica também dependem de processos microorganizacionais de validação, supervisão e definição das fronteiras de atuação da tecnologia. Essa perspectiva complementa discussões sobre cooperação tecnológica entre economias emergentes, como aquelas apresentadas por Jiang et al. (2026), ao demonstrar que a confiança necessária à adoção e à cooperação em IA também é produzida no nível das práticas organizacionais. Nesse contexto, o estudo oferece uma perspectiva situada sobre como os gestores negociam, na prática, os limites da participação da IAG no trabalho, mantendo o julgamento humano como elemento central desse processo.

A pesquisa mostra, portanto, que os limites da participação da IAG no trabalho são definidos na própria interação entre gestores e tecnologia, sob diferentes condições institucionais e organizacionais, enquanto o julgamento humano permanece como elemento central desse processo.

REFERÊNCIAS

BENK, Michaela; KERSTAN, Sophie; VON WANGENHEIM, Florian; FERRARIO, Andrea. Twenty-four years of empirical research on trust in AI: a bibliometric review of trends, overlooked issues, and future directions. *AI & Society*, v. 40, n. 4, p. 2083–2106, 2025. DOI: 10.1007/s00146-024-02059-y.

BRYNJOLFSSON, Erik; LI, Danielle; RAYMOND, Lindsey. Generative AI at work. *The Quarterly Journal of Economics*, v. 140, n. 2, p. 889–942, 2025. DOI: 10.1093/qje/qjae044.

BRYNJOLFSSON, Erik; MCAFEE, Andrew. The second machine age: work, progress, and prosperity in a time of brilliant technologies. New York: W. W. Norton & Company, 2014.

BUNDUCHI, Raluca; SITAR-TĂUT, Dan-Andrei; MICAN, Daniel. A legitimacy-based explanation for user acceptance of controversial technologies: the case of Generative AI. *Technological Forecasting and Social Change*, v. 215, p. 124095, 2025. DOI: 10.1016/j.techfore.2025.124095.

CARDOSO, A. S.; GAMMARANO, I. D. L. Guia para o uso ético e responsável da inteligência artificial generativa no âmbito acadêmico. *Revista de Ciências da Administração*, v. 27, n. 67, 2025.

CHARMAZ, Kathy. *Constructing grounded theory*. 2. ed. Los Angeles: Sage, 2014.

CONSELHO NACIONAL DE DESENVOLVIMENTO CIENTÍFICO E TECNOLÓGICO – CNPq. Portaria nº 2.664/2026: estabelece a Política de Integridade na Atividade Científica do Conselho Nacional de Desenvolvimento Científico e Tecnológico. Brasília, 2026.

CRESWELL, John W.; LOPES, Magda França; SILVA, Dirceu da. *Projeto de pesquisa: métodos qualitativo, quantitativo e misto*. [S. l.]: Penso, 2021.

CVETKOVIC, Anica et al. User trust in AI and major tech companies in twelve countries. *Behaviour & Information Technology*, 2026. DOI: 10.1080/0144929X.2026.2619648.

DALY, Sarah J.; WIEWIORA, Anna; HEARN, Greg. Shifting attitudes and trust in AI: influences on organizational AI adoption. *Technological Forecasting and Social Change*, v. 215, art. 124108, 2025.

DANG, Qinqu; LI, Guiquan. Unveiling trust in AI: the interplay of antecedents, consequences, and cultural dynamics. *AI & Society*, v. 41, p. 669–692, 2026. DOI: 10.1007/s00146-025-02477-6.

DAVENPORT, Thomas H.; RONANKI, Rajeev. Artificial intelligence for the real world. *Harvard Business Review*, v. 96, n. 1, p. 108–116, jan./fev. 2018.

DELLERMANN, Dominik et al. The future of human-AI collaboration: a taxonomy of design knowledge for hybrid intelligence systems. *Business & Information Systems Engineering*, v. 61, n. 5, p. 637–643, 2019.

DONEY, Patricia M.; CANNON, Joseph P.; MULLEN, Michael R. Understanding the influence of national culture on the development of trust. *Academy of Management Review*, v. 23, n. 3, p. 601–620, 1998. DOI: 10.5465/amr.1998.926629.

GIL, Antonio Carlos. *Como elaborar projetos de pesquisa*. 4. ed. São Paulo: Atlas, 2002.

GLIKSON, Ella; WOOLLEY, Anita Williams. Human trust in artificial intelligence: review of empirical research. *Academy of Management Annals*, v. 14, n. 2, p. 627–660, 2020.

GMYREK, Paweł et al. *Generative AI and Jobs: A Refined Global Index of Occupational Exposure*. Geneva: International Labour Organization, 2025. ILO Working Paper 140. Disponível em:

<https://www.ilo.org/publications/generative-ai-and-jobs-refined-global-index-occupational-exposure>. Acesso em: 11 ago. 2026.

HALLIKAINEN, Heli; LAUKKANEN, Tommi. National culture and consumer trust in e-commerce. *International Journal of Information Management*, v. 38, n. 1, p. 97–106, 2018.

HOFSTEDE, Geert. *Culture's consequences: comparing values, behaviors, institutions and organizations across nations*. 2. ed. Thousand Oaks: Sage Publications, 2001.

HOUSE, Robert J.; HANGES, Paul J.; JAVIDAN, Mansour; DORFMAN, Peter W.; GUPTA, Vipin (ed.). *Culture, leadership, and organizations: the GLOBE study of 62 societies*. Thousand Oaks: Sage Publications, 2004.

IKKATAI, Yuko et al. The relationship between the attitudes of the use of AI and diversity awareness: comparisons between Japan, the US, Germany, and South Korea. *AI & Society*, 2024. DOI: 10.1007/s00146-024-01982-4.

JARRAHI, Mohammad Hossein. Artificial intelligence and the future of work: human-AI symbiosis in organizational decision making. *Business Horizons*, v. 61, n. 4, p. 577–586, 2018.

JIANG, Tianjiao et al. Greater BRICS AI cooperation from a Global South perspective: opportunities and challenges. Xangai: Center for Global AI Innovative Governance, Fudan Development Institute, Fudan University, 2026. (Fudan Report Series, WAIC 2026 Forum).

JOHRI, Aditya; SCHLEISS, Johannes; RANADE, Nupoor. Lessons for GenAI literacy from a field study of human-GenAI augmentation in the workplace. *arXiv*, arXiv:2502.00567, 2025.

KAPLAN, Alexandra D.; KESSLER, Theresa T.; BRILL, J. Christopher; HANCOCK, P. A. Trust in artificial intelligence: meta-analytic findings. *Human Factors*, v. 65, n. 2, p. 337–359, 2023. DOI: 10.1177/00187208211013988.

KEDING, Christoph; MEISSNER, Philip. Managerial overreliance on AI-augmented decision-making processes: how the use of AI-based advisory systems shapes choice behavior in R&D investment decisions. *Technological Forecasting and Social Change*, v. 171, art. 120970, 2021.

LEIDNER, Dorothy E.; KAYWORTH, Timothy. Review: a review of culture in information systems research: toward a theory of information technology culture conflict. *MIS Quarterly*, v. 30, n. 2, p. 357–399, 2006. DOI: 10.2307/25148735.

LEONARDI, Paul M. When flexible routines meet flexible technologies: affordance, constraint, and the imbrication of human and material agencies. *MIS Quarterly*, v. 35, n. 1, p. 147–167, 2011.

LEWIS, Marianne W. Exploring paradox: toward a more comprehensive guide. *Academy of Management Review*, v. 25, n. 4, p. 760–776, 2000. DOI: 10.5465/amr.2000.3707712.

LUSCH, Robert F.; NAMBISAN, Satish. Service innovation: a service-dominant logic perspective. *MIS Quarterly*, v. 39, n. 1, p. 155–176, 2015. DOI: 10.25300/MISQ/2015/39.1.07.

MICROSOFT; LINKEDIN. *AI at Work Is Here. Now Comes the Hard Part: 2024 Work Trend Index Annual Report*. 2024. Disponível em: <https://www.microsoft.com/en-us/worklab/work-trend-index/ai-at-work-is-here-now-comes-the-hard-part>. Acesso em: 11 ago. 2026.

MANGONI, Giovanna Karla. *Cocriação humano-IA: o uso prático da IA generativa na rotina de gestores*. Florianópolis: UFSC, 2025.

NG, Davy Tsz Kit et al. AI literacy: definition, teaching, evaluation and ethical issues. *Computers and Education: Artificial Intelligence*, v. 2, 2021.

NORDHEIM, Christina B.; FØLSTAD, Asbjørn; BJERKAN, Anne Mette. An initial model of trust in chatbots for customer service. *Interacting with Computers*, v. 31, n. 3, p. 317–335, 2019.

ORLIKOWSKI, Wanda J. Sociomaterial practices: exploring technology at work. *Organization Studies*, v. 28, n. 9, p. 1435–1448, 2007. DOI: 10.1177/0170840607081138.

PRAHALAD, C. K.; RAMASWAMY, Venkat. Co-creation experiences: the next practice in value creation. *Journal of Interactive Marketing*, v. 18, n. 3, p. 5–14, 2004.

RHEU, Minji et al. Systematic review of trust in automation: distinguishing trust and trustworthiness. *ACM Computing Surveys*, v. 54, n. 3, p. 1–36, 2021.

RIBEIRO, Manuella Maia; SEGATTO, Catarina Ianni. Inteligência artificial nas organizações públicas brasileiras: heterogeneidades e capacidades em tecnologia da informação. *Revista de Administração Pública*, v. 59, n. 1, e2024-0066, 2025.

ROBINSON, Linda. Trust in artificial intelligence: cultural influences. *AI & Society*, 2020.

SCHAD, Jonathan; LEWIS, Marianne W.; RAISCH, Sebastian; SMITH, Wendy K. Paradox research in management science: looking back to move forward. *Academy of Management Annals*, v. 10, n. 1, p. 5–64, 2016. DOI: 10.5465/19416520.2016.1162422.

SERRA, Fernando Antonio Ribeiro; SCAFUTO, Isabel Cristina; SERRA VIEIRA, Patricia. Regulation of artificial intelligence: between the switch and innovation. Is Brazil prepared? *BAR – Brazilian Administration Review*, v. 22, n. 4, e250015, 2025. DOI: 10.1590/1807-7692bar2025250015.

SHERIDAN, Thomas B. Individual differences in attributes of trust in automation: measurement and application. *Human Factors*, v. 61, n. 6, p. 935–949, 2019.

SMITH, Wendy K.; LEWIS, Marianne W. Toward a theory of paradox: a dynamic equilibrium model of organizing. *Academy of Management Review*, v. 36, n. 2, p. 381–403, 2011. DOI: 10.5465/amr.2009.0223.

SÖLLNER, Matthias et al. Understanding the formation of trust in IT artifacts. *Journal of the Association for Information Systems*, v. 17, n. 10, p. 1–40, 2016.

STRAUSS, Anselm L.; CORBIN, Juliet. *Pesquisa qualitativa: técnicas e procedimentos para o desenvolvimento de teoria fundamentada*. Tradução de Luciane de Oliveira da Rocha. 2. ed. Porto Alegre: Artmed, 2008.

TEECE, David J. Explicating dynamic capabilities: the nature and microfoundations of (sustainable) enterprise performance. *Strategic Management Journal*, v. 28, n. 13, p. 1319–1350, 2007.

TEECE, David J.; PISANO, Gary; SHUEN, Amy. Dynamic capabilities and strategic management. *Strategic Management Journal*, v. 18, n. 7, p. 509–533, 1997.

THIEBES, Scott et al. Trustworthy artificial intelligence. *Electronic Markets*, v. 31, p. 447–464, 2021.

VARGO, Stephen L.; LUSCH, Robert F. Evolving to a new dominant logic for marketing. *Journal of Marketing*, v. 68, n. 1, p. 1–17, 2004. DOI: 10.1509/jmkg.68.1.1.24036.

VARGO, Stephen L.; LUSCH, Robert F. Institutions and axioms: an extension and update of service-dominant logic. *Journal of the Academy of Marketing Science*, v. 44, n. 1, p. 5–23, 2016. DOI: 10.1007/s11747-015-0456-3.

WANG, Shangrui; ZHANG, Yuanmeng; XIAO, Yiming; LIANG, Zheng. Artificial intelligence policy frameworks in China, the European Union and the United States: an analysis based on structure topic model. *Technological Forecasting and Social Change*, v. 212, 2025. DOI: 10.1016/j.techfore.2025.123971.

WOOLLEY, Anita Williams. Generative AI and collaboration: opportunities for cultivating collective intelligence. *Journal of Organization Design*, 2025. DOI: 10.1007/s41469-025-00199-z.

YANG, R.; WIBOWO, S. User trust in artificial intelligence: a comprehensive conceptual framework. *Electronic Markets*, v. 32, n. 4, p. 2053–2077, 2022. DOI: 10.1007/s12525-022-00592-6.

ZOU, Mimi; ZHANG, Lu. Navigating China's regulatory approach to generative artificial intelligence and large language models. *Cambridge Forum on AI: Law and Governance*, v. 1, e8, 2025.

APÊNDICES

APÊNDICE A - PLANILHA LOG DE PROMPTS

Em razão do volume de informações, a planilha integral utilizada nesta pesquisa foi disponibilizada em formato digital no endereço indicado abaixo. O arquivo reúne o registro dos prompts empregados no apoio às diferentes etapas da análise, incluindo codificação, comparação e integração das categorias, procedimentos de validação e elaboração da representação visual dos resultados, além das abas referentes aos códigos e às entrevistas analisadas. Nos casos em que o registro textual integral de determinada interação não estava disponível, o prompt foi reconstruído a partir do procedimento realizado e identificado como tal na planilha. Em caso de dificuldade de acesso, o arquivo poderá ser solicitado diretamente à autora.

LOG_Prompts_TCC

ID	Etapla analítica	Ferramenta	Finalidade	Status do registro	Prompt / instrução principal
P01	Configuração do apoio de IA	ChatGPT	Definir o papel, os limites analíticos e as dimensões relacionais utilizadas no apoio à análise.	Reconstrução fiel do procedimento	Você atuará como assistente de análise qualitativa para apoiar a codificação de entrevistas sobre confiança e cocriação entre gestores e Inteligência Artificial Generativa. A análise utiliza procedimentos de codificação inspirados na Grounded Theory, sem pressupor categorias não sustentadas pelos dados. Considere quatro dimensões relacionais: CTX – condições contextuais explicitadas pelo participante, incluindo aspectos organizacionais, institucionais, regulatórios, culturais, profissionais ou relacionados à tarefa; TRUST – avaliações, expectativas, dúvidas ou percepções sobre a confiabilidade da IAG, incluindo confiança, desconfiança, verificação e recalibração; COCHI – práticas de interação e construção conjunta entre gestor e IAG, incluindo formulação de prompts, refinamento, divisão de tarefas, supervisão, delegação e intervenção humana; LINK – utilizar somente quando o excerto explicitar uma relação entre dimensões, indicando como uma condição influencia a confiança, a interação ou o comportamento do gestor. Preserve a proximidade com o significado expresso pelo participante. Não atribua causalidade, interação ou significado que não estejam sustentados pelo excerto. A decisão analítica final pertence à pesquisadora.
P02	Codificação aberta	ChatGPT	Gerar códigos candidatos próximos ao conteúdo empírico dos excertos.	Reconstrução fiel do procedimento	Analise o excerto abaixo considerando também a pergunta da entrevista à qual o participante está respondendo. Identifique unidades de significado e sugira códigos abertos curtos, específicos e próximos ao conteúdo empírico. Evite códigos excessivamente genéricos e não transforme interpretações teóricas em dados empíricos. Quando o participante estiver relatando um exemplo, diferencie o acontecimento concreto do significado que ele atribui ao episódio. Para cada código sugerido, indique brevemente qual elemento do excerto o sustenta. Os códigos são sugeridos para posterior avaliação da pesquisadora.
P03	Refinamento e comparação constante	ChatGPT	Comparar códigos gerados entre excertos, reduzir redundâncias e identificar variações ou casos negativos.	Reconstrução fiel do procedimento	Compare os códigos sugeridos para este excerto com os códigos já identificados nas entrevistas anteriores. Avalie se o fechamento já está representado por um código existente, se exige refinamento ou se constitui um significado distinto. Não agrupe códigos apenas porque utilizam palavras semelhantes. Considere a função do trecho, o conteúdo da resposta e a diferença entre condição, avaliação de confiança e prática de interação. Identifique também possíveis casos que contradigam ou tensionem os padrões já encontrados.
P04	Organização relacional dos dados	ChatGPT	Classificar excertos nas dimensões CTX, TRUST, COCHI e LINK de forma consistente.	Reconstrução fiel do procedimento	Classifique cada excerto em uma dimensão predominante: CTX, TRUST, COCHI ou LINK. Utilize CTX quando o foco estiver em uma condição organizacional, institucional, regulatória, cultural, profissional ou relacionada à tarefa; TRUST quando o foco estiver na avaliação da confiabilidade da IA; COCHI quando o trecho descrever como humano e IA interagem ou dividem o trabalho; e LINK apenas quando o participante explicitar uma relação entre essas dimensões. Não utilize LINK somente porque duas dimensões aparecem no mesmo trecho. Deve existir uma conexão explícita ou claramente expressa pelo participante. Atribua uma única categoria predominante e apresente uma justificativa breve baseada no excerto.
P05	Validação da codificação	ChatGPT	Produzir classificação independente dos 24 excertos selecionados para comparação com a pesquisadora.	Reconstrução fiel do procedimento	Classifique independentemente cada um dos 24 excertos abaixo em CTX, TRUST, COCHI ou LINK, utilizando exclusivamente as definições analíticas fornecidas. Não tente inferir nem reproduzir a classificação realizada previamente pela pesquisadora. Considere cada excerto de maneira independente e atribua somente uma categoria predominante. Apresente os resultados na ordem B01-C04, sem alterar os excertos e sem realizar conciliação entre possíveis interpretações neste momento.
P06	Análise comparativa Brasil-China	ChatGPT	Identificar convergências e diferenças entre os conjuntos sem atribuir causalidade automática à cultura nacional.	Reconstrução fiel do procedimento	Compare os padrões identificados nas entrevistas B001-B005 e C001-C005, considerando incorporação da IAG, construção da confiança, critérios de validação, formas de delegação, refinamento, supervisão e julgamento humano. Organize a comparação em: (1) padrões recorrentes nos dois contextos; (2) padrões mais recorrentes entre os casos brasileiros; (3) padrões mais recorrentes entre os casos chineses; e (4) variações ou casos específicos. Não atribua automaticamente diferenças à cultura nacional. Diferencie, sempre que os dados permitirem, efeitos relacionados a organização, setor, regulação, segurança de dados, acesso às ferramentas, experiência profissional e características da tarefa. Quando os dados não permitirem identificar a origem de uma diferença, apresente-a como diferença observada, e não como explicação causal.
P07	Integração interpretativa	ChatGPT	Examinar o mecanismo que melhor explica como a confiança influencia a participação da IAG no trabalho gerencial.	Reconstrução fiel do procedimento	Considerando conjuntamente os padrões de CTX, TRUST, COCHI e LINK identificados nas dez entrevistas, examine qual mecanismo explica de forma mais consistente como a confiança influencia a participação da IAG no trabalho gerencial. Considere especialmente as variações de delegação, verificação, supervisão, refinamento, aceitação ou rejeição de outputs e retomada manual da tarefa. Verifique como criticidade da tarefa, verificabilidade, experiências anteriores, erros da IA, expertise profissional e condições organizacionais modificam essa relação. Teste a interpretação nos casos brasileiros e chineses e procure também evidências que a contradigam ou limitem. Não trate confiança como uma variável binária de confiar ou não confiar e não formule relações determinísticas que não estejam sustentadas pelos dados.
P08	Representação visual dos resultados	ChatGPT / ferramenta de geração visual	Construir uma arquitetura conceitual integrada para representar incorporação da IAG, confiança e cocriação.	Registro literal	Create a high-rigor academic conceptual figure for a university thesis on human-Generative AI co-creation, trust in AI, and intercultural comparison between Brazil and China. The figure should function as an analytical model derived from empirical findings, not as a decorative infographic. Its purpose is to visually integrate three previously separated result dimensions: (1) incorporation of Generative AI into managerial routines, (2) construction and recalibration of trust in AI, and (3) human-AI co-creation practices. Use a clean, sophisticated, publication-quality academic visual style suitable for a thesis or peer-reviewed journal. Avoid playful icons, cartoon aesthetics, gradients, excessive decoration, and corporate infographic clichés. Use a restrained neutral palette, precise typography, generous white space, and strong visual hierarchy. Structure the figure horizontally in three analytical layers: contextual conditions, empirical trajectories and mechanisms, and trust inflection points and feedback. Include four central analytical zones: Incorporation of Generative AI; construction and recalibration of trust; regulation of human-AI interaction; and configurations of human-AI co-creation. Do not imply deterministic causality. Allow multiple empirical trajectories and optional BR cases, CH cases and Both contexts markers. Include empirical traceability through case identifiers and a methodological note stating that connections represent interpretive relationships identified in the data and do not imply deterministic causality.
P09	Refinamento da representação visual	ChatGPT / ferramenta de geração visual	Ajustar a figura conceitual aos achados empíricos consolidados e garantir rastreabilidade.	Reconstrução fiel do procedimento	Revise a arquitetura conceitual da figura a partir dos resultados empíricos já consolidados. Mantenha apenas relações sustentadas pelas entrevistas. Represente a incorporação da IAG, a construção e recalibração da confiança, a regulação da interação e as configurações de cocriação como dimensões relacionadas, incluindo feedback entre experiência e confiança. Diferencie padrões compartilhados por BR e CN de diferenças de recorência entre os dois conjuntos. Não represente diferenças nacionais como causalidade cultural. Utilize somente os identificadores B001-B005 e C001-C005 para indicar rastreabilidade empírica.

APÊNDICE B - CONSENTIMENTO DE ENTREVISTAS

BR01

Pesquisadora: “Como eu expliquei, a pesquisa é anônima e suas informações aparecerão somente por código. Para que eu possa transcrever nossa conversa depois, preciso da sua autorização para gravar esta entrevista. Você autoriza?”

Participante: “Confirmado, autorizado.”

BR02

Pesquisadora: “Isso aqui é para o meu TCC e será tratado de forma anônima, então o que você falar aparecerá apenas por código. Para não perder nada e poder transcrever depois, eu preciso da sua autorização para gravar o áudio. Você autoriza?”

Participante: “Sim.”

BR03

Pesquisadora: “Primeiramente, muito obrigada por aceitar participar. Como eu já te contei um pouco sobre o TCC, para não perder nada e poder transcrever depois eu preciso da sua autorização para gravar o áudio. Você autoriza?”

Participante: “Autorizo. Autorizo esta entrevista.”

BR04

Pesquisadora: “Antes de começarmos, você autoriza a gravação desta entrevista para que ela possa ser posteriormente transcrita e utilizada na pesquisa acadêmica?”

Participante: “Autorizo a gravação.”

BR05

Pesquisadora: “Basicamente, já te contei sobre qual é o contexto do meu TCC e, para poder transcrever depois, eu preciso da sua autorização para gravar o áudio. Você autoriza?”

Participante: “Autorizo.”

CN01

Pesquisadora: “Como eu já expliquei, estou pesquisando como gestores trabalham com Inteligência Artificial Generativa em suas rotinas e como a confiança nessas

ferramentas influencia essa relação, comparando Brasil e China. A pesquisa é 100% anônima e confidencial e seu nome será representado apenas por um código. Você concorda com a gravação desta conversa?”

Participante: “Sim.”

CN02

Pesquisadora: “Antes de começarmos, preciso da sua confirmação de que você concorda com a gravação desta entrevista. Seu nome e seus dados pessoais não serão utilizados e qualquer informação mencionada será codificada na pesquisa. Você autoriza?”

Participante: “Sim, autorizo.”

CN03

Pesquisadora: “Antes de começarmos, gostaria de confirmar seu consentimento para participar desta entrevista. Você concorda voluntariamente em participar desta pesquisa e autoriza a coleta e o uso de suas respostas para fins acadêmicos, mantendo sua identidade confidencial e utilizando as informações exclusivamente neste estudo?”

Participante: “Sim, concordo e autorizo a coleta e o uso das minhas respostas para esta pesquisa nessas condições.”

CN04

Pesquisadora: “Antes de começarmos, você autoriza a gravação desta entrevista e o uso das suas respostas para fins desta pesquisa acadêmica, preservando sua identidade?”

Participante: “Autorizo esta gravação.”

CN05

Pesquisadora: “Antes de começarmos, você autoriza a gravação desta entrevista e a utilização das suas respostas para fins da pesquisa, mantendo sua identidade preservada?”

Participante: “Sim, autorizo a gravação.”

APÊNDICE C - ROTEIRO DE ENTREVISTAS SEMI-ESTRUTURADAS

Roteiro de Entrevista Semiestruturada

Informações Gerais

Título da Pesquisa: Co-criação entre gestores e Inteligência Artificial Generativa: a influência da confiança na IA em perspectiva intercultural.

Objetivo Geral: Analisar como a confiança na IA influencia a co-criação entre gestores e Inteligência Artificial Generativa em rotinas de trabalho, considerando a comparação entre os contextos brasileiro e chinês.

Público-alvo: Gestores com experiência mínima de seis meses no uso de ferramentas conversacionais de IAG.

Duração estimada da entrevista: 30 minutos

Instruções ao entrevistador(a)

- Agradecer ao participante pela disponibilidade;
- Apresentar brevemente o tema e os objetivos da pesquisa;
- Solicitar consentimento para a gravação da entrevista.

Quadro 18 - Estrutura do Roteiro de entrevistas

Bloco	Objetivo Específico	Finalidade do bloco
Bloco 1 Incorporação da IAG	Objetivo específico 1 Mapear a incorporação da IAG nas rotinas de gestores nos contextos brasileiro e chinês.	Compreender como a IA passou a integrar as rotinas gerenciais, em quais atividades é utilizada e como seu uso evoluiu ao longo do tempo.
Bloco 2 Confiança na IAG	Objetivo específico 2 Examinar como a confiança na IA é construída, percebida e mobilizada nessas práticas.	Investigar os fatores que favorecem ou reduzem a confiança na IA, bem como os mecanismos de validação utilizados pelos gestores.
Bloco 3 Co-criação Humano-IA	Objetivo específico 3 Caracterizar os processos de interação e co-criação entre gestores e sistemas de IAG.	Identificar como ocorre a colaboração entre gestores e IA, os processos de refinamento, delegação cognitiva e construção conjunta de soluções.
Bloco 4 Perspectiva intercultural	Objetivo específico 4 Comparar similaridades e diferenças entre Brasil e China quanto à confiança e a seus efeitos na co-criação.	Compreender de que forma o contexto institucional e cultural influencia a construção da confiança e os processos de co-criação.

Fonte: Elaborado pela Autora (2026)

Quadro 19 - Roteiro de entrevista semiestruturada

Bloco Temático	Objetivo específico relacionado	Pergunta norteadora	Possíveis aprofundamentos
Bloco 1 Incorporação da IAG	Objetivo 1	Conte como a IA passou a fazer parte da sua rotina de trabalho.	Quando isso começou? Como você começou a utilizá-la? Houve algum acontecimento ou necessidade que motivou esse uso?
		Em quais atividades você utiliza a IA com maior frequência?	Se preferir, me diga os três principais casos de uso. O que faz você recorrer à IAG nessas atividades e não em outras?
		Como esse uso evoluiu desde que você começou a utilizá-la?	Houve mudanças ao longo do tempo? Há alguma funcionalidade que você gostaria de ter e que hoje não encontra?
Bloco 2 Confiança	Objetivo 2	Conte uma situação recente em que você decidiu utilizar uma resposta da IAG sem fazer muitas alterações. O que levou você a considerar aquela resposta como confiável?	
		Agora o contrário: pense em uma situação recente em que você decidiu conferir uma resposta antes de utilizá-la. O que fez você sentir necessidade de verificar aquela resposta, e como você fez essa verificação?	O que costuma acionar essa necessidade de validação?
		Conte uma situação em que uma resposta da IAG mudou o seu nível de confiança na ferramenta, para mais ou para menos. O que aconteceu e como você reagiu?	Como você lida quando a ferramenta apresenta informações incorretas ou inventadas? Isso afeta sua confiança de forma temporária ou definitiva? O fato de colegas ou de outras empresas usarem a ferramenta influencia a sua confiança nela?
		Para você, o que uma resposta da IAG precisa ter para ser considerada confiável?	Você consegue dar um exemplo? A empresa que desenvolve a ferramenta pesa nessa avaliação?
		Como você decide quando aceitar ou rejeitar uma sugestão da IAG, e o que	Você lembra de uma situação concreta em que

		você faz em cada um desses casos?	tomou essa decisão? O que considerou naquele momento?
Bloco 3 Co-criação	Objetivo 3	Pense em uma tarefa complexa que você realizou recentemente com apoio da IAG. Como foi o passo a passo da sua interação com a ferramenta?	Quem iniciou o processo? Como a interação foi se desenvolvendo? Em que momento você considerou que tinha chegado a um resultado?
		Conte uma situação em que você utilizou a IAG para desenvolver uma solução ou um resultado ao longo de várias interações. Como esse processo aconteceu?	
		Se o resultado obtido ao final de uma interação não atende às suas necessidades, o que você faz?	Como você conduz o refinamento das respostas?
		Em que tipos de tarefa a interação com a IAG costuma ser mais intensa ou envolver mais trocas?	
		Conte uma situação em que uma resposta da IAG influenciou uma decisão sua. O que a ferramenta trouxe para o processo e o que permaneceu sob o seu julgamento?	
Bloco 4 Contexto organizacional e cultural	Objetivo 4	Como as políticas, as regras ou a cultura da sua organização influenciam a forma como você e sua equipe utilizam a IAG no dia a dia?	Existem normas ou orientações sobre o uso de IAG? Há incentivos ou restrições?
		Na sua opinião, existem características da cultura do seu país que facilitam ou dificultam a adoção e a confiança na IAG no ambiente de trabalho? Se puder, me dê exemplos concretos.	Você percebe influência de valores, costumes ou formas de trabalhar? E de instituições, normas ou políticas?
		Em relação à privacidade de dados, à segurança das informações ou às regulamentações do seu país, como esses fatores afetam a sua decisão de usar e de confiar na IAG no trabalho?	Você lembra de uma situação concreta? Existem referências externas que influenciam essa percepção?
<i>LINK</i>		Antes de encerrarmos, há algo a mais sobre a forma como você usa a IAG no seu trabalho, ou sobre o ambiente em que você trabalha, que você considera importante mencionar e que a gente não tocou até agora?	

Fonte: Elaborado pela Autora (2026)

APÊNDICE D – DECLARAÇÃO DE USO DE INTELIGÊNCIA ARTIFICIAL

Em conformidade com as Diretrizes para o Uso de Inteligência Artificial da Revista de Ciências da Administração da Universidade Federal de Santa Catarina (RCA/UFSC) e com a Política de Integridade na Atividade Científica do CNPq, declara-se o uso de ferramentas de Inteligência Artificial Generativa no desenvolvimento desta pesquisa.

A utilização dessas ferramentas ocorreu de forma assistiva, sob supervisão intelectual da pesquisadora, mantendo-se sob sua responsabilidade todas as decisões teóricas, metodológicas, analíticas e interpretativas. A Inteligência Artificial Generativa não foi considerada autora ou coautora deste trabalho.

Quadro 20 - Declaração de uso de Inteligência Artificial Generativa

Ferramenta	Etapa da pesquisa	Finalidade do uso	Procedimento de validação humana
ChatGPT, OpenAI	Estruturação e revisão textual; apoio à análise qualitativa; representação visual dos resultados	Apoio à organização e revisão de textos; comparação e síntese de conteúdos; classificação independente de parcela delimitada do corpus para avaliação de concordância; apoio à estruturação da representação visual integrada dos resultados	Revisão integral pela pesquisadora; conferência das informações e referências em fontes originais; classificação humana realizada independentemente e anteriormente à comparação com a classificação da IA; resolução das divergências e decisão interpretativa final exclusivamente pela pesquisadora
Gemini, Google	Pesquisa exploratória e revisão de literatura	Exploração de conceitos, identificação de possíveis relações entre temas e comparação exploratória de interpretações	Verificação dos conteúdos nas fontes originais e avaliação crítica pela pesquisadora; respostas da ferramenta não foram utilizadas como evidência científica
Claude; Anthropic	Leitura crítica e revisão textual	Apoio à avaliação de coerência argumentativa, clareza, organização textual e alternativas de formulação	Sugestões avaliadas, selecionadas, modificadas ou rejeitadas pela pesquisadora; responsabilidade integral pela redação

			final
Perplexity	Pesquisa exploratória e rastreamento bibliográfico	Localização inicial de publicações, documentos e possíveis referências relacionadas ao tema da pesquisa	Consulta e verificação das publicações e documentos originais antes de sua utilização; respostas e sínteses da plataforma não foram utilizadas como fontes bibliográficas
SciSpace	Busca e exploração da literatura científica	Apoio à localização, leitura exploratória e identificação de literatura potencialmente relevante ao referencial teórico e à discussão	Consulta dos artigos e documentos originais; verificação de autoria, conteúdo, contexto e informações bibliográficas antes da incorporação ao trabalho
NotebookLM, Google	Organização e consulta de materiais documentais	Apoio à navegação, organização, recuperação e síntese exploratória de documentos selecionados e inseridos pela pesquisadora	Conferência das respostas nos documentos de origem e interpretação final realizada exclusivamente pela pesquisadora

Fonte: Elaborado pela Autora (2026)

Durante o desenvolvimento da pesquisa, ChatGPT, Gemini, Claude, Perplexity, SciSpace e NotebookLM foram utilizados em etapas distintas e com diferentes níveis de participação. Seu emprego ocorreu como apoio à pesquisa, à organização de materiais, à revisão textual e à exploração de alternativas analíticas, sem substituição do julgamento da pesquisadora.

As respostas, sínteses, sugestões textuais, indicações bibliográficas e interpretações produzidas ou mediadas por essas ferramentas foram submetidas à avaliação crítica da autora. Quando envolviam literatura científica ou informações externas, as fontes originais foram consultadas antes da incorporação do conteúdo ao trabalho. Resultados produzidos pelas ferramentas não foram tratados como evidência científica independente.

Entre as ferramentas declaradas, o ChatGPT teve também uma aplicação metodológica específica e delimitada no procedimento de avaliação da classificação do corpus. Uma parcela composta por 24 excertos, sendo 12 referentes ao contexto brasileiro e 12 ao contexto chinês, foi classificada de forma independente pela

pesquisadora e pelo assistente de IA segundo as dimensões CTX, TRUST, COCRI e LINK.

A classificação realizada pela pesquisadora ocorreu anteriormente à consulta da classificação produzida pela ferramenta. Posteriormente, as duas classificações foram comparadas para avaliação de concordância, sem utilização da resposta da IA como critério para modificação automática da codificação humana. As divergências foram examinadas individualmente a partir dos excertos originais e das definições operacionais das categorias, permanecendo com a pesquisadora a decisão final de classificação e interpretação.

O procedimento resultou em concordância global de 79,2% e coeficiente Kappa de 0,722. Nos casos de discordância entre a classificação humana e aquela produzida pela IA, a decisão analítica final permaneceu sob responsabilidade da pesquisadora. Dessa forma, o uso da ferramenta nesse procedimento teve função de contraste metodológico e avaliação de consistência, e não de substituição da análise qualitativa humana.

O ChatGPT também foi empregado como apoio à elaboração da representação visual integrada dos resultados. Nesse procedimento, os elementos analíticos, as categorias e as relações representadas foram definidos previamente pela pesquisadora com base nos resultados da análise. A ferramenta foi utilizada para apoiar a organização e a representação visual desse conteúdo, não para determinar autonomamente as relações teóricas apresentadas. O respectivo registro de interação foi mantido no log metodológico da pesquisa.

Responsabilidade autoral

A pesquisadora declara que:

- a) todas as decisões teóricas, metodológicas, analíticas e interpretativas permaneceram sob sua responsabilidade;
- b) conteúdos produzidos, processados ou reorganizados com auxílio de ferramentas de Inteligência Artificial foram submetidos à avaliação e validação humana;
- c) nenhuma ferramenta de Inteligência Artificial foi utilizada como autora, coautora, fonte autônoma de evidência científica ou instância decisória final;
- d) referências bibliográficas localizadas ou sugeridas com auxílio dessas ferramentas somente foram incorporadas ao trabalho após consulta e verificação das fontes originais;

e) a classificação realizada por Inteligência Artificial na parcela delimitada do corpus foi utilizada exclusivamente para fins de comparação e avaliação metodológica, não substituindo a codificação e a interpretação da pesquisadora;

f) as aplicações substantivas de Inteligência Artificial utilizadas no procedimento metodológico foram registradas de modo a permitir sua rastreabilidade; e

g) a pesquisadora assume integral responsabilidade pela precisão, originalidade, integridade, interpretação e conteúdo final deste trabalho, inclusive por eventuais imprecisões decorrentes do uso das ferramentas declaradas.