



Anais do Simpósio Acadêmico de Engenharia de Produção (SAEPRO) da EEL-USP

IX SAEPRO – 25 e 26 de novembro de 2025

Ergonomia Cognitiva na Interação Humano IA: Princípios para Sustentabilidade Cognitiva

Cognitive Ergonomics in Human-AI Interaction: Principles for Cognitive Sustainability

Ergonomía Cognitiva en la Interacción Humano-IA: Principios para la Sostenibilidad Cognitiva

Carlos Gabriel Pires Cavaliere

Universidade de São Paulo / (carlos.cgabriel@usp.br)

Humberto Felipe da Silva

Universidade de São Paulo / (Humberto.felipe@usp.br)

Resumo

Este artigo analisa os impactos da interação humano IA no trabalho do conhecimento 4.0, com foco nos riscos cognitivos emergentes e nas implicações para o design de sistemas inteligentes. A partir de uma revisão de escopo estruturada segundo as diretrizes *PRISMA-ScR* e normas recentes (*ISO* e *NIST*), foram identificadas seis categorias de riscos, incluindo sobrecarga informacional, fadiga decisória e oscilação de confiança. Com base nessa síntese, propõem-se princípios de design centrados no humano — como redução da carga extrínseca, explicações orientadas à ação com exposição de incerteza e automação graduada — e um *framework* de avaliação multimodal que combina métricas subjetivas, objetivas, comportamentais e fisiológicas. Os resultados indicam que integrar ergonomia cognitiva à governança de IA é essencial para equilibrar eficiência tecnológica e sustentabilidade humana, garantindo decisões mais rápidas, consistentes e seguras.



Anais do Simpósio Acadêmico de Engenharia de Produção (SAEPRO) da EEL-USP

IX SAEPRO – 25 e 26 de novembro de 2025

Palavras-chave: ergonomia cognitiva; inteligência artificial; carga mental; interação humano-*IA*; confiança em automação.

Abstract

This paper examines the cognitive challenges arising from human AI interaction in knowledge work 4.0 and proposes actionable design principles for intelligent systems. Based on a structured scoping review following PRISMA-ScR guidelines and recent standards (ISO and NIST), six categories of cognitive risks were identified, including information overload, decision fatigue, and trust oscillation. From this synthesis, we propose human-centered design principles—such as reducing extraneous load, providing action-oriented explanations with uncertainty disclosure, and implementing graduated automation—and introduce a multimodal evaluation *framework* combining subjective, objective, behavioral, and physiological metrics. Findings indicate that integrating cognitive ergonomics into *AI* governance is essential to balance technological efficiency and human sustainability, ensuring faster, more consistent, and safer decisions.

Keywords: cognitive ergonomics, artificial intelligence, mental workload, human-AI interaction, trust in automation.

Resumen

Este artículo analiza los desafíos cognitivos derivados de la interacción humano-*IA* en el trabajo del conocimiento 4.0 y propone principios de diseño aplicables a sistemas inteligentes. A partir de una revisión de alcance estructurada según las directrices *PRISMA-ScR* y normas recientes (*ISO* y *NIST*), se identificaron seis categorías de riesgos cognitivos, incluyendo sobrecarga informativa, fatiga decisoria y oscilación de confianza. Con base en esta síntesis, se proponen principios de diseño centrados en el humano — como reducción de carga extrínseca, explicaciones orientadas a la acción con exposición de incertidumbre y automatización gradual— y se presenta un *framework* de evaluación multimodal que combina métricas subjetivas, objetivas, conductuales y fisiológicas. Los resultados indican que integrar la ergonomía cognitiva en la gobernanza de la *IA* es



esencial para equilibrar eficiencia tecnológica y sostenibilidad humana, garantizando decisiones más rápidas, consistentes y seguras.

Palabras clave: ergonomía cognitiva, inteligencia artificial, carga mental, interacción humano-IA, confianza en automatización.

1. INTRODUÇÃO

A digitalização acelerada do trabalho do conhecimento consolidou ecossistemas intensivos em dados, colaboração distribuída e decisões em tempo quase real — agora mediados por sistemas de inteligência artificial (IA) incluindo IA generativa. Esse arranjo ampliou ganhos de produtividade, sobretudo para profissionais menos experientes, mas deslocou o principal gargalo do desempenho para o domínio cognitivo: gerir sobrecarga informacional, manter atenção sustentada e decidir sob incerteza enquanto se supervisiona e audita recomendações algorítmicas (Noy; Zhang, 2023; Brynjolfsson; Li; Raymond, 2025; Marsh et al., 2024).

Nesse contexto, a ergonomia cognitiva reassume papel central. A revisão e atualização das normas ISO 10075-2:2024 (princípios de design para carga mental) e ISO 9241-112:2025 (apresentação da informação) reforçam que decisões seguras e sustentáveis dependem do desenho de tarefas, interfaces e condições organizacionais que minimizem carga extrínseca e previnam tanto sobrecarga quanto a subutilização (ISO, 2024; ISO, 2025). Em paralelo, a neuro ergonomia vem demonstrando o valor de avaliações multimodais (autorrelato, desempenho e medidas fisiológicas como EEG/fNIRS/ocular) para monitorar carga mental em tarefas complexas, informando intervenções de projeto (Mark et al., 2024). Fenômenos correlatos — como *alarm fatigue* em saúde, cibersegurança e processos industriais — ilustram o risco de dessensibilização, atrasos de resposta e quedas de desempenho quando o volume de sinais e notificações excede a capacidade humana de triagem (AACN, 2023; Tariq et al., 2025; Mustafa et al., 2023).

A interação humana IA adiciona outra camada: decisões assistidas por IA dependem de confiança apropriada (nem subutilização, nem delegação excessiva). Pesquisas recentes identificam dois riscos complementares: *overreliance*, caracterizado pela aceitação

acrítica das recomendações algorítmicas, especialmente quando o custo de verificação é elevado ou as explicações são complexas (Vasconcelos et al., 2023), e *algorithm aversion*, que ocorre após erros salientes e é modulada por benefícios percebidos e pelo grau de controle humano (Schaap et al., 2024; Jussupow et al., 2024). Diretrizes contemporâneas de transparência e *XAI (Explainable Artificial Intelligence)* centrada no humano apontam que explicações orientadas à ação, com incerteza calibrada e baixo esforço de auditoria, melhoram a calibração da confiança e reduzem fadiga decisória (Liao; Vaughan, 2024).

Em paralelo, o ecossistema normativo avança no eixo de risco e governança, o *NIST AI Risk Management Framework 1.0* (2023) fornece resultados-alvo para confiabilidade e gestão de riscos sociotécnicos; a ISO/IEC 23894:2023 adapta princípios de gestão de riscos à *AI* ao longo do ciclo de vida; e a ISO/IEC 42001:2023 define requisitos para sistemas de gestão de *AI (AIMS)*, conectando políticas, papéis e monitoramento contínuo. Tais instrumentos, porém, ainda carecem de operacionalização ergonômico-cognitiva para o trabalho do conhecimento cotidiano (NIST, 2023; ISO/IEC, 2023).

Apesar dos avanços, persiste uma lacuna: faltam sínteses que integrem ergonomia cognitiva, interação humano IA e *AI governance* em princípios de design acionáveis e avaliáveis para o trabalho do conhecimento 4.0. Este artigo busca preencher essa lacuna por meio de uma revisão de escopo estruturada, seguindo diretrizes atualizadas para revisões sistemáticas e de escopo (*PRISMA-ScR*) (Page et al., 2021; Peters et al., 2020), complementada por normas recentes (ISO, 2024; ISO/IEC, 2023; NIST, 2023). Os objetivos específicos são: (i) mapear uma taxonomia de riscos cognitivos típicos na interação *human-AI*; (ii) sintetizar princípios de design centrados no humano, alinhados a padrões internacionais; e (iii) propor um *framework* de avaliação que combine métricas subjetivas, objetivas, comportamentais e normativas, visando equilibrar eficiência tecnológica e sustentabilidade cognitiva.

2. FUNDAMENTAÇÃO TEÓRICA

A ergonomia cognitiva surge como resposta ao deslocamento das exigências do trabalho moderno, que deixou de ser predominantemente físico para se tornar intensamente mental. Em ambientes digitais e automatizados, o desempenho humano depende da capacidade de gerenciar informações, tomar decisões sob incerteza e interagir com sistemas inteligentes. Esse cenário exige compreender conceitos fundamentais como carga mental, confiança em automação e consciência situacional.

2.1. Ergonomia Cognitiva e Carga Mental

A carga mental é um conceito central na ergonomia cognitiva e refere-se ao esforço necessário para processar informações e executar tarefas, considerando tanto a complexidade da atividade quanto o design do sistema. Em ambientes digitais complexos, característicos do trabalho do conhecimento 4.0, esse esforço é amplificado por fatores como múltiplos fluxos de informação, alertas simultâneos e necessidade de supervisão de sistemas automatizados (Marsh et al., 2024).

Ferramentas clássicas como o NASA-TLX continuam sendo utilizadas para mensuração subjetiva, mas pesquisas recentes recomendam abordagens multimodais que combinem indicadores de desempenho (tempo, erros), autorrelato e medidas fisiológicas. Por exemplo, Mark et al. (2024) demonstraram que a integração de eletroencefalografia (EEG), espectroscopia funcional no infravermelho próximo (fNIRS) e rastreamento ocular permite capturar variações dinâmicas da carga mental em tarefas complexas, especialmente quando há interação com sistemas de inteligência artificial.

Normas atualizadas, como a ISO 10075-2:2024, estabelecem princípios para prevenir fadiga cognitiva e orientar o design de sistemas que reduzam carga extrínseca — isto é, aquela gerada por elementos desnecessários ou mal estruturados na interface (ISO, 2024). Essa redução é essencial para liberar recursos cognitivos e manter a qualidade da tomada de decisão (Page et al., 2021).

Pesquisas recentes, como Tariq et al. (2025), AACN (2023) e Mustafa et al. (2023), evidenciam os riscos da sobrecarga informacional em setores críticos. Essas pesquisas mostram que excesso de alertas não acionáveis compromete a consciência situacional e



aumenta o risco de erros, tanto em ambientes hospitalares quanto em operações industriais e cibersegurança.

Por fim, pesquisas em neuro ergonomia, como Mark et al. (2024), indicam que métricas fisiológicas podem antecipar estados de sobrecarga antes que se manifestem em erros ou fadiga decisória. Esses trabalhos demonstram que padrões de atividade cerebral e sinais fisiológicos permitem ajustes adaptativos no nível de automação ou na apresentação da informação.

2.2. Trabalho do Conhecimento 4.0

O trabalho do conhecimento evoluiu significativamente com a transformação digital e a adoção de tecnologias da Indústria 4.0. Essa mudança intensificou a conectividade, a automação e a integração de sistemas inteligentes, criando ambientes caracterizados por fluxos contínuos de dados, múltiplos canais de comunicação e decisões em tempo quase real (Marsh et al., 2024). Embora esses avanços ampliem a produtividade e a capacidade analítica, eles também introduzem riscos cognitivos relevantes, como sobrecarga informacional, interrupções frequentes e fadiga decisória.

Pesquisas recentes, como Marsh et al. (2024), demonstram que a hiperconectividade e a multiplicidade de tarefas simultâneas aumentam a carga mental e reduzem a capacidade de manter atenção sustentada. Essas pesquisas indicam que trabalhadores expostos a ambientes digitais complexos apresentam maior vulnerabilidade a erros e estresse, especialmente quando precisam alternar entre tarefas críticas e monitoramento de sistemas automatizados.

Outro fator emergente é a adoção de inteligência artificial generativa (*generative AI*) no suporte a atividades cognitivas. Pesquisas como Noy e Zhang (2023) e Brynjolfsson et al. (2025) mostram que ferramentas baseadas em IA generativa aumentam a produtividade em tarefas de escrita e atendimento, com ganhos mais expressivos para profissionais menos experientes. No entanto, essas mesmas pesquisas alertam para a redução do esforço crítico e para o risco de dependência excessiva, o que pode comprometer a qualidade da tomada de decisão e a capacidade de auditoria.



Além disso, pesquisas sobre *tecnostress*, como as conduzidas por Marsh et al. (2024), evidenciam que a exposição prolongada a sistemas digitais complexos está associada a sintomas de exaustão mental e queda de desempenho. Esses achados reforçam a necessidade de estratégias de design que reduzam a carga extrínseca, priorizem informações relevantes e promovam mecanismos de controle sobre a automação.

Em síntese, o trabalho do conhecimento 4.0 não se limita à ampliação do acesso a dados e ferramentas digitais; ele implica uma reconfiguração das demandas cognitivas. Para garantir sustentabilidade e desempenho, é essencial adotar práticas que integrem ergonomia cognitiva, gestão da carga mental e princípios de design centrados no humano, alinhados às normas internacionais mais recentes.

2.3. Interação Humano IA e Confiança

A interação entre humanos e sistemas de inteligência artificial (IA) é um dos elementos mais críticos para garantir decisões seguras e eficazes. Essa relação depende da calibragem adequada da confiança: quando insuficiente, ocorre subutilização da tecnologia; quando excessiva, há risco de delegação indevida e perda de controle.

Pesquisas recentes, como Vasconcelos et al. (2023), identificam o fenômeno de *overreliance*, caracterizado pela aceitação acrítica das recomendações algorítmicas, especialmente quando o custo de verificação é elevado ou as explicações fornecidas pelo sistema são complexas. Em contraste, pesquisas como Schaap et al. (2024) e Jussupow et al. (2024) analisam a *algorithm aversion*, que ocorre após erros salientes e leva usuários a rejeitar sistematicamente sistemas automatizados, mesmo quando eles apresentam desempenho superior à média humana.

Outro aspecto relevante é a explicabilidade. O conceito de *Explainable Artificial Intelligence (XAI)* refere-se a métodos e práticas que tornam as decisões da IA compreensíveis para o usuário. Diretrizes contemporâneas, como as propostas por Liao e Vaughan (2024), indicam que explicações orientadas à ação, com granularidade adequada e exposição de incerteza, reduzem fadiga decisória e melhoram a calibragem da confiança. Esses achados reforçam que a transparência não deve ser apenas técnica, mas

cognitiva, permitindo que o usuário entenda não apenas “como” a IA chegou à decisão, mas “por que” essa decisão é relevante para a tarefa.

Além disso, normas recentes, como a ISO/IEC 23894:2023 e o *NIST AI Risk Management Framework* (2023), enfatizam a necessidade de mecanismos que promovam confiança apropriada, incluindo feedback recuperável, níveis graduais de automação e políticas claras de delegação. Esses elementos são fundamentais para evitar tanto a complacência quanto a rejeição injustificada, garantindo que a colaboração humano-IA seja eficiente e segura.

2.4. Psicologia da Decisão e Fadiga Cognitiva

A tomada de decisão em ambientes complexos é influenciada por limitações cognitivas, vieses e pela necessidade de lidar com incertezas. Em cenários mediados por inteligência artificial, esses fatores se intensificam devido à frequência de micro decisões e à necessidade de revisar recomendações algorítmicas.

Pesquisas recentes, como Vasconcelos et al. (2023), indicam que a exposição contínua a sistemas de IA pode induzir *overreliance*, levando usuários a aceitar recomendações sem verificação crítica, especialmente quando o custo de auditoria é elevado. Em contraste, pesquisas como Schaap et al. (2024) e Jussupow et al. (2024) mostram que erros salientes geram *algorithm aversion*, fenômeno em que usuários rejeitam sistematicamente sistemas automatizados, mesmo quando estes apresentam desempenho superior à média humana.

Outro aspecto relevante é a fadiga decisória. pesquisas como Marsh et al. (2024) evidenciam que ambientes digitais com alta densidade de alertas e múltiplas tarefas simultâneas aumentam a carga mental e reduzem a qualidade das decisões. Essa fadiga é agravada quando os sistemas exigem constante supervisão e validação de outputs, criando um ciclo de esforço cognitivo elevado.

Diretrizes recentes sobre explicabilidade, como as propostas por Liao e Vaughan (2024), apontam que explicações orientadas à ação, com granularidade adequada e exposição de incerteza, reduzem fadiga decisória e melhoram a calibragem da confiança. Esses achados



reforçam que a transparência deve ser cognitiva, permitindo que o usuário compreenda não apenas como a IA chegou à decisão, mas porque essa decisão é relevante para a tarefa.

Em síntese, a psicologia da decisão aplicada ao contexto da IA revela que a interação humano-algoritmo não elimina o esforço mental; ela apenas o transforma. Para mitigar fadiga e vieses, é essencial adotar estratégias que combinem design centrado no humano, explicabilidade efetiva e políticas claras de delegação.

2.5. Diretrizes e Normas

A evolução das normas internacionais reflete a necessidade de integrar princípios ergonômicos e requisitos de governança para sistemas inteligentes. No contexto da interação humano IA, essas diretrizes não apenas orientam o design centrado no humano, mas também estabelecem parâmetros para reduzir riscos cognitivos e garantir confiabilidade.

Normas recentes, como a ISO 9241-112:2025, atualizam princípios para apresentação da informação, enfatizando clareza, relevância e redução da carga extrínseca. Esses elementos são fundamentais para evitar sobrecarga informacional e preservar a consciência situacional em ambientes digitais complexos (ISO, 2025). De forma complementar, a ISO 10075-2:2024 define princípios para prevenção da fadiga cognitiva, orientando práticas de design que considerem limitações atencionais e memória de trabalho (ISO, 2024).

No eixo da governança, a ISO/IEC 23894:2023 introduz diretrizes para gestão de riscos em sistemas de IA cobrindo aspectos como transparência, explicabilidade e mitigação de vieses. Já a ISO/IEC 42001:2023 estabelece requisitos para sistemas de gestão de IA (*Artificial Intelligence Management Systems – AIMS*), conectando políticas organizacionais, papéis e monitoramento contínuo (ISO/IEC, 2023). Em paralelo, o *NIST AI Risk Management Framework* (2023) propõe resultados alvo para confiabilidade, segurança e *accountability*, oferecendo um referencial prático para implementação de controles sociotécnicos.

Pesquisas recentes, como Liao e Vaughan (2024), reforçam que a aplicação dessas normas deve ser acompanhada por práticas de explicabilidade centradas no humano, incluindo explicações orientadas à ação, exposição de incerteza e mecanismos de correção. Esses elementos, quando integrados às diretrizes normativas, contribuem para reduzir riscos cognitivos, calibrar confiança e promover decisões mais consistentes.

Em síntese, as normas e diretrizes atuais não se limitam à conformidade regulatória; elas constituem um arcabouço estratégico para alinhar eficiência tecnológica e sustentabilidade cognitiva, garantindo que a automação inteligente amplie (e não comprometa) as capacidades humanas.

3. MÉTODO

Esta pesquisa foi conduzida por meio de uma revisão de escopo, seguindo as diretrizes do *PRISMA-ScR*, recomendadas para mapear conceitos e identificar lacunas em áreas emergentes (Page et al., 2021; Peters et al., 2020; Munn et al., 2022). Segundo Munn et al. (2022), revisões de escopo são apropriadas quando o objetivo é explorar a extensão, a natureza e as características de um corpo de conhecimento, sem restringir-se a perguntas de eficácia ou análises quantitativas. Essa abordagem é indicada para temas complexos e multidimensionais, pois permite integrar evidências de diferentes tipos de pesquisas, normas e diretrizes, oferecendo uma visão abrangente e fundamentada do estado da arte.

A busca focou em pesquisas publicados entre 2020 e 2025, período definido para garantir atualidade das evidências, conforme recomendações para revisões em engenharia de produção (Page et al., 2021). Foram consultadas bases de dados multidisciplinares e especializadas, incluindo *Scopus*, *Web of Science*, e *Google Scholar*, além de documentos normativos disponíveis nos portais da ISO e do NIST. A estratégia de busca combinou descritores em português e inglês, como “ergonomia cognitiva”, “carga mental”, “*human-AI interaction*”, “*trust*”, “*Explainable AI*” e “*AI risk management*”.

Os critérios de inclusão contemplaram artigos publicados entre 2020 e 2025, revisões sistemáticas ou de escopo, artigos revisados por pares, diretrizes normativas e relatórios técnicos que abordassem carga mental, confiança em automação, explicabilidade, fadiga decisória ou princípios de design aplicáveis à IA. Foram excluídos pesquisas puramente técnicos sobre algoritmos sem relação com fatores humanos e publicações anteriores a 2020, salvo normas ou conceitos clássicos citados em revisões recentes.

O processo de seleção seguiu as etapas recomendadas pelo *PRISMA-ScR*: identificação, triagem, elegibilidade e inclusão. Essas etapas foram realizadas conforme as diretrizes do *PRISMA-ScR* (Page et al., 2021; Peters et al., 2020; Munn et al., 2022).

A extração de dados foi realizada em planilha estruturada, incluindo informações sobre autores, ano, contexto, métricas utilizadas (p.ex., NASA-TLX, indicadores fisiológicos), principais achados e implicações para design. A síntese seguiu abordagem temática, agrupando evidências em três dimensões: cognição individual (memória de trabalho, atenção, fadiga decisória), interação (elementos de interface, feedback, explicações) e organização (políticas de delegação, governança e métricas ergonômico-cognitivas). Essa estrutura permitiu integrar resultados de forma coerente com os objetivos da pesquisa e com as recomendações metodológicas para revisões de escopo (Munn et al., 2022).

4. RESULTADOS E DISCUSSÃO

A revisão realizada permitiu identificar padrões consistentes na literatura recente sobre ergonomia cognitiva aplicada à interação com sistemas de inteligência artificial. Os achados foram organizados em três partes: (i) taxonomia de riscos cognitivos, (ii) cenários ilustrativos fundamentados e (iii) princípios de design e proposta de avaliação.

4.1 Taxonomia de Riscos Cognitivos

A análise das pesquisas incluídas revelou seis categorias principais de riscos cognitivos associados ao trabalho do conhecimento em ambientes mediados por IA:

- **Sobrecarga informacional**

Pesquisas como Marsh et al. (2024) demonstram que a exposição contínua a múltiplos fluxos de dados e alertas simultâneos aumenta a carga mental e compromete a atenção sustentada. Interfaces com excesso de indicadores e notificações não acionáveis reduzem a velocidade e a precisão das decisões.

- **Vigilância passiva**

Pesquisas como Tariq et al. (2025) e AACN (2023) evidenciam que a automação excessiva pode levar à perda de consciência situacional, dificultando a retomada do controle em situações críticas. Esse fenômeno, associado à *alarm fatigue*, é observado tanto em centros de operações de segurança quanto em ambientes hospitalares.

- **Oscilação de confiança**

Pesquisas recentes, como Schaap et al. (2024) e Jussupow et al. (2024), mostram que erros salientes induzem rejeição sistemática a sistemas automatizados (*algorithm aversion*), enquanto benefícios percebidos e explicações complexas favorecem *overreliance*. Essa oscilação compromete a consistência das decisões e aumenta a carga metacognitiva.

- **Dissonância explicativa**

Pesquisas como Liao e Vaughan (2024) apontam que explicações excessivamente técnicas ou pouco relevantes aumentam a carga extrínseca e prolongam o tempo de decisão. Diretrizes recentes recomendam explicações orientadas à ação e ajustadas ao contexto da tarefa para mitigar esse risco.

- **Fadiga decisória**

Pesquisas como Marsh et al. (2024) indicam que ambientes digitais com alta densidade de micro decisões — como validação contínua de outputs

algorítmicos — reduzem a qualidade das escolhas e aumentam a exaustão mental.

- **Aversão ou complacência algorítmica**

Pesquisas como Vasconcelos et al. (2023) reforçam que a falta de mecanismos para calibrar confiança leva a extremos: rejeição injustificada ou delegação acrítica, ambos prejudiciais à segurança e à eficiência.

Esses riscos raramente ocorrem isolados; sobrecarga informacional e dissonância explicativa, por exemplo, tendem a coexistir, amplificando fadiga decisória e erros.

4.2 Cenários Ilustrativos Fundamentados

Para contextualizar os achados e demonstrar sua aplicabilidade, foram elaborados cenários com base em duas fontes: evidências empíricas extraídas diretamente das pesquisas revisadas e exemplos hipotéticos fundamentados, construídos a partir da síntese das evidências e das normas apresentadas na Seção 2.5.

- **Atendimento assistido por IA generativa (empírico)**

Pesquisas como Noy e Zhang (2023) e Brynjolfsson et al. (2025) mostram ganhos expressivos de produtividade em tarefas de escrita e atendimento, especialmente para profissionais menos experientes. Em paralelo, Vasconcelos et al. (2023) indicam risco de *overreliance* quando o custo de verificação é elevado. Implicações: automação graduada, revisão obrigatória em interações sensíveis e explicações orientadas à ação com exposição de incerteza (Liao; Vaughan, 2024), alinhadas às diretrizes ISO/IEC 23894:2023 e NIST AI RMF (2023).

- **Centros de Operações de Segurança – SOC (empírico)**

Pesquisas como Tariq et al. (2025) relatam *alert fatigue* por alto volume de alertas e falsos positivos, prejudicando consciência situacional e tempo de resposta; evidências em saúde apontam padrão análogo com alarmes clínicos (AACN, 2023). Implicações: agregação e priorização de sinais, eliminação de

notificações não acionáveis, rótulos de incerteza e rotas de escalonamento (ISO, 2025; ISO, 2024; NIST, 2023).

- **Controle de processos industriais (empírico)**

Análises como Mustafa et al. (2023) documentam *alarm flooding* e desenho deficiente de alarmes elevando carga e risco operacional. Implicações: racionalização do ciclo de alarmes e conformidade com princípios de apresentação da informação (ISO, 2025) e de projeto para carga mental (ISO, 2024).

- **Painel executivo com IA preditiva (hipotético fundamentado)**

Com base em padrões de sobrecarga (Marsh et al., 2024), *alert fatigue* (Tariq et al., 2025; AACN, 2023) e confiança/explicabilidade (Vasconcelos et al., 2023; Liao; Vaughan, 2024), um painel que consolida KPIs, alertas e previsões deve aplicar: (i) divulgação progressiva de explicações orientadas à ação; (ii) exposição de incerteza e limites de validade; (iii) agregação de alertas; (iv) níveis graduais de automação com reversibilidade. Esse exemplo combina evidências empíricas e normas (ISO, 2025; ISO, 2024; ISO/IEC, 2023; NIST, 2023).

4.3 Princípios de Design e Avaliação

Da convergência entre evidências (Seções 2.1–2.4) e diretrizes (Seção 2.5), emergem quatro princípios centrais:

1. Reduzir carga extrínseca na interface e no fluxo de trabalho por meio de organização por objetivos, priorização e limitação de notificações (Marsh et al., 2024; ISO, 2025; ISO, 2024).
2. Fornecer explicações orientadas à ação com incerteza exposta, evitando aceitação acrítica ou rejeição injustificada (Liao; Vaughan, 2024; ISO/IEC, 2023; NIST, 2023).

3. Implementar automação graduada com reversibilidade e políticas claras de delegação, preservando controle humano significativo (Noy; Zhang, 2023; Brynjolfsson et al., 2025; ISO/IEC, 2023; NIST, 2023).
4. Medir e calibrar continuamente carga e confiança, combinando indicadores subjetivos, objetivos e fisiológicos (Mark et al., 2024; ISO, 2024).

Para operacionalizar esses princípios, propõe-se um *framework* de avaliação com quatro camadas: subjetiva (NASA-TLX), objetiva (tempo, erros, retrabalho), comportamental (taxa de aceitação vs. *override*, uso de explicações) e fisiológica (EEG/fNIRS/ocular, quando aplicável), complementadas por um checklist normativo (ISO 9241-112:2025; ISO 10075-2:2024; ISO/IEC 23894:2023; ISO/IEC 42001:2023; NIST AI RMF, 2023). Pesquisas como Mark et al. (2024) mostram que a triangulação com sinais fisiológicos antecipa sobrecarga, enquanto indicadores comportamentais ajudam a monitorar calibragem de confiança (Vasconcelos et al., 2023).

5. CONSIDERAÇÕES FINAIS

Esta pesquisa reforça que a ergonomia cognitiva é estratégica para o trabalho do conhecimento na era da inteligência artificial, não apenas como fator de conforto, mas como elemento crítico para decisões seguras, consistentes e sustentáveis. A revisão realizada permitiu mapear uma taxonomia de riscos cognitivos — incluindo sobrecarga informacional, fadiga decisória, dissonância explicativa e oscilação de confiança — que emergem em ambientes mediados por IA. Além disso, foram sintetizados princípios de design centrados no humano, como redução da carga extrínseca, explicações orientadas à ação com exposição de incerteza e automação graduada com reversibilidade, todos alinhados às normas ISO e ao NIST AI RMF.

A principal contribuição prática reside na proposta de um *framework* de avaliação multimodal, que combina métricas subjetivas (*NASA-TLX*), objetivas (tempo, erros),



Anais do Simpósio Acadêmico de Engenharia de Produção (SAEPRO) da EEL-USP

IX SAEPRO – 25 e 26 de novembro de 2025

comportamentais (taxa de aceitação vs. *override*) e fisiológicas (EEG/fNIRS/ocular), complementadas por um checklist normativo. Essa abordagem oferece um caminho para operacionalizar a ergonomia cognitiva em sistemas inteligentes, equilibrando eficiência tecnológica e sustentabilidade humana.

Por tratar-se de uma revisão de escopo, esta pesquisa não valida empiricamente os princípios e o *framework* propostos. Recomenda-se avançar em três frentes: validação experimental do *framework* em setores críticos, como saúde, segurança cibernética e operações industriais; integração com IA adaptativa e mecanismos de neuro *feedback* para ajuste dinâmico da carga mental; e desenvolvimento de métricas normativas que conectem ergonomia cognitiva à governança de IA ampliando a aplicabilidade das normas ISO/IEC e do NIST AI RMF.

Em síntese, projetar sistemas que respeitem limites cognitivos, promovam confiança calibrada e preservem a consciência situacional é condição para um futuro do trabalho que seja, ao mesmo tempo, eficiente e humano.

6. REFERÊNCIAS

AACN – American Association of Critical-Care Nurses. (2023). Ten Years Later, Alarm Fatigue Is Still a Safety Concern. *AACN Advanced Critical Care*, 34(3), 189–197. Disponível em: <https://aacnjournals.org/aacnacconline/article/34/3/189/32176/Ten-Years-Later-Alarm-Fatigue-Is-Still-a-Safety>

Brynjolfsson, E., Li, D., & Raymond, L. (2025). Generative AI at Work. *The Quarterly Journal of Economics*, 140(2), 889–942. Disponível em: <https://academic.oup.com/qje/article/140/2/889/7990658>

ISO. (2019). ISO 9241-210:2019 — Ergonomics of human-system interaction — Human-centred design for interactive systems.



ISO. (2024). ISO 10075-2:2024 — Ergonomic principles related to mental workload — Part 2: Design principles. Disponível em: <https://www.iso.org/standard/76686.html>

ISO. (2025). ISO 9241-112:2025 — Ergonomics of human-system interaction — Principles for the presentation of information.

ISO/IEC JTC 1/SC 42. (2023). ISO/IEC 23894:2023 — Information technology — Artificial intelligence — Guidance on risk management.

ISO/IEC JTC 1/SC 42. (2023). ISO/IEC 42001:2023 — Information technology — Artificial intelligence — Management system.

Jussupow, E., Benbasat, I., & Heinzl, A. (2024). An Integrative Perspective on Algorithm Aversion and Appreciation in Decision-Making. *MIS Quarterly*, 48(4), 1575–1590. Disponível em: <https://misq.umn.edu/misq/article/48/4/1575/2300/An-Integrative-Perspective-on-Algorithm-Aversion>

Lee, H. P., et al. (2025). The Impact of Generative AI on Critical Thinking: Self-Reported Reductions in Cognitive Effort and Confidence Effects From a Survey of Knowledge Workers. In *CHI Conference on Human Factors in Computing Systems (CHI '25)*. Disponível em: https://www.microsoft.com/en-us/research/wp-content/uploads/2025/01/lee_2025_ai_critical_thinking_survey.pdf

Liao, Q. V., & Vaughan, J. W. (2024). AI Transparency in the Age of LLMs: A Human-Centered Research Roadmap. *Harvard Data Science Review*. Disponível em: <https://hdsr.mitpress.mit.edu/pub/aelql9qy>

Mark, J. A., Curtin, A., Kraft, A. E., Ziegler, M. D., & Ayaz, H. (2024). Mental workload assessment by monitoring brain, heart, and eye with six biomedical modalities during six cognitive tasks. *Frontiers in Neuroergonomics*, 5, 1345507. Disponível em: <https://www.frontiersin.org/articles/10.3389/fnrgo.2024.1345507/full>



Marsh, E., Perez Vallejos, E., & Spence, A. (2024). Mindfully and confidently digital: A mixed-methods study on personal resources to mitigate the dark side of digital working. *PLOS ONE*, 19(2), e0295631. Disponível em: <https://journals.plos.org/plosone/article?id=10.1371/journal.pone.0295631>

Mustafa, F. E., Ahmed, I., Basit, A., Alvi, U. H., Mahmood, A., & Ali, P. R. (2023). A review on effective alarm management systems for industrial process control: Barriers and opportunities. *International Journal of Critical Infrastructure Protection*, 41, 100599. Disponível em: <https://doi.org/10.1016/j.ijcip.2023.100599>

NIST. (2023). Artificial Intelligence Risk Management Framework (AI RMF 1.0). NIST AI 100-1. Disponível em: <https://nvlpubs.nist.gov/nistpubs/ai/NIST.AI.100-1.pdf>

NIST. (2024). *Generative Artificial Intelligence Profile*. Disponível em: https://tsapps.nist.gov/publication/get_pdf.cfm?pub_id=958388

Noy, S., & Zhang, W. (2023). Experimental evidence on the productivity effects of generative artificial intelligence. *Science*, 381(6654), 187–192. Disponível em: <https://www.science.org/doi/10.1126/science.adh2586>

Schaap, G., Bosse, T., & Vettehen, P. H. (2024). The ABC of algorithmic aversion: Not agent, but benefits and control determine the acceptance of automated decision-making. *AI & Society*, 39, 1947–1960. Disponível em: <https://link.springer.com/article/10.1007/s00146-023-01649-6>

Shneiderman, B. (2020). Human-Centered Artificial Intelligence: Reliable, Safe & Trustworthy. *International Journal of Human-Computer Interaction*, 36(6), 495–504.

Tariq, S., Baruwat Chhetri, M., Nepal, S., & Paris, C. (2025). Alert Fatigue in Security Operations Centres: Research Challenges and Opportunities. *ACM Computing Surveys*, 57(9), Article 224. Disponível em: <https://dl.acm.org/doi/10.1145/3723158>



Vasconcelos, H., Jörke, M., Grunde-McLaughlin, M., Gerstenberg, T., Bernstein, M. S., & Krishna, R. (2023). Explanations Can Reduce Overreliance on AI Systems During Decision-Making. *Proceedings of the ACM on Human–Computer Interaction (CSCW)*, 7(CSCW1), Article 129. Disponível em: <https://arxiv.org/abs/2212.06823>

Page, M. J., et al. (2021). *PRISMA 2020 explanation and elaboration: updated guidance and exemplars for reporting systematic reviews*. *BMJ*, 372:n160. Disponível em: <https://www.bmj.com/content/372/bmj.n160>

Peters, M. D. J., et al. (2020). *Updated methodological guidance for the conduct of scoping reviews*. *JBIM Evidence Synthesis*, 18(10), 2119–2126. Disponível em: <https://rgu-repository.worktribe.com/preview/998095/PETERS%202020%20Updated%20methodological.pdf>

Munn, Z., et al. (2022). *Systematic review or scoping review? Guidance for authors when choosing between a systematic or scoping review approach*. *BMC Medical Research Methodology*, 22(1), 123. Disponível em: <https://rgu-repository.worktribe.com/preview/1699608/MUNN%20What%20are%20scoping%20reviews%20%28AAM%29.pdf>