

ARTIGO

O CONFORTO QUE NOS DESAPRENDE

A armadilha da zona de conforto algorítmica e a influência
cognitiva estrutural

Jandislei Antonio Genova

*Advogado, autor e pesquisador independente em Direito, inteligência artificial e
governança institucional*

“A decisão não começa no clique. A dependência também não.”

São Paulo | Julho de 2026

Resumo

Este artigo, de natureza teórico-conceitual e caráter exploratório, examina a hipótese de que sistemas de inteligência artificial podem produzir uma armadilha de conforto cognitivo: ganhos imediatos de desempenho e redução de esforço que não se convertem necessariamente em aprendizagem, compreensão ou preservação da capacidade autônoma de decidir. A análise articula evidências sobre educação assistida por IA, descarregamento cognitivo, fluência metacognitiva, memória e viés de automação, sem generalizar causalmente para além dos contextos estudados. Em seguida, relaciona esses fenômenos à influência cognitiva estrutural, compreendida como a atuação das arquiteturas tecnológicas sobre condições anteriores à decisão — percepção, formulação do problema, seleção de fontes, verificação e confiança no próprio julgamento. Propõe-se uma matriz analítica preliminar de reversibilidade, com dimensões cognitiva e tecnológico-institucional, para distinguir assistência capacitadora de dependência. Sustenta-se que o conforto não constitui, por si só, um risco: a hipótese de captura exige a convergência entre redução persistente de capacidade própria, aumento relevante do custo de saída e assimetria de controle sobre conhecimento ou infraestrutura.

Palavras-chave: inteligência artificial; autonomia decisória; influência cognitiva estrutural; descarregamento cognitivo; viés de automação; reversibilidade cognitiva; dependência tecnológica.

Abstract

This theoretical-conceptual and exploratory article examines the hypothesis that artificial intelligence systems may create a cognitive comfort trap: immediate performance gains and reduced effort that do not necessarily translate into learning, understanding, or the preservation of autonomous decision-making capacity. The analysis connects evidence on AI-assisted education, cognitive offloading, metacognitive fluency, memory, and automation bias, without making causal generalizations beyond the settings studied. It then relates these phenomena to structural cognitive influence, understood as the operation of technological architectures upon conditions preceding a decision — perception, problem formulation, source selection, verification, and confidence in one's own judgment. A preliminary analytical reversibility matrix, comprising cognitive and technological-institutional dimensions, is proposed to distinguish capability-enhancing assistance from dependency. The article argues that comfort is not inherently harmful: the capture hypothesis requires a convergence of persistent loss of internal capability, materially higher exit costs, and asymmetric control over knowledge or infrastructure.

Keywords: artificial intelligence; decisional autonomy; structural cognitive influence; cognitive offloading; automation bias; cognitive reversibility; technological dependence.

Nota metodológica

Este texto é um ensaio teórico-conceitual, exploratório e interdisciplinar. Não constitui revisão sistemática, metanálise nem validação de instrumento psicométrico. A bibliografia foi selecionada por sua relevância para quatro eixos analíticos: desempenho e aprendizagem com IA generativa; descarregamento cognitivo e fluência metacognitiva; viés de automação e

supervisão humana; reversibilidade institucional e governança. Os estudos empíricos sustentam proposições circunscritas às populações, tarefas e períodos examinados.

A passagem da escala individual para as escalas organizacional e estatal é apresentada como hipótese derivada, a ser testada de forma independente, e não como generalização causal já demonstrada. As fontes normativas são empregadas comparativamente para identificar consequências de governança. O modelo e a matriz propostos devem, portanto, ser avaliados longitudinalmente e em contextos definidos.

A decisão não começa no clique. A dependência também não.

1. Melhores resultados, a mesma aprendizagem

Duzentos e setenta e cinco estudantes receberam a mesma tarefa: durante noventa minutos, deveriam resolver um exercício de programação sobre concorrência computacional. Foram distribuídos aleatoriamente em três grupos. O primeiro podia utilizar o ChatGPT sem restrições. O segundo recorria a um tutor de inteligência artificial chamado Iris, programado para oferecer pistas graduais sem entregar a solução completa. O terceiro utilizava recursos convencionais de pesquisa, sem assistência generativa.

Os dois grupos que utilizaram inteligência artificial obtiveram resultados melhores no exercício e relataram menos frustração. Os estudantes que usaram o ChatGPT concentraram-se entre as pontuações mais altas e consideraram a ferramenta mais fácil e útil. Aparentemente, a tecnologia havia cumprido a promessa habitual: mais produtividade, menos esforço e melhor desempenho.

Mas os testes realizados antes e depois da atividade mostraram algo diferente. Nenhum dos grupos assistidos por IA apresentou ganho superior de conhecimento ou vantagem na compreensão do código. Apenas o tutor que oferecia pistas, sem entregar respostas completas, aumentou a motivação intrínseca. Naquele contexto experimental, a IA funcionou principalmente como instrumento de desempenho, e não como instrumento de aprendizagem. Os pesquisadores chamaram de *comfort trap* a situação em que a preferência dos estudantes pela solução mais fácil deixou de corresponder à alternativa pedagogicamente mais eficaz (BASSNER et al., 2026).¹

O estudo possui limites claros. Foi realizado em um curso introdutório de programação, com uma tarefa específica e duração reduzida. Não permite concluir que toda utilização de IA diminua a aprendizagem, nem que os mesmos efeitos ocorram em todas as profissões. Seu resultado é importante por outra razão: oferece evidência, naquele contexto experimental, de que desempenho e compreensão podem se separar.

¹ Bassner et al. utilizam a expressão *comfort trap* para designar, naquele experimento, a preferência pela alternativa percebida como mais fácil, embora ela não tenha produzido aprendizagem superior.

Essa dissociação importa porque o indicador mais visível pode melhorar enquanto a capacidade que deveria sustentá-lo permanece estável. No experimento de Bassner et al. (2026), não se demonstrou perda generalizada de competência; a possibilidade de erosão em usos prolongados constitui hipótese adicional, que requer observação longitudinal. A entrega pronta, o código funcional e a meta cumprida são mensurados. A competência que deixou de ser exercida, entretanto, costuma permanecer fora do painel.

Essa dissociação oferece uma chave para compreender uma transformação mais ampla. A inteligência artificial não altera apenas a velocidade com que realizamos tarefas. Ela pode reorganizar a relação entre esforço, aprendizagem, confiança e autonomia.

2. O conforto não é o problema

A expressão “zona de conforto” costuma ser empregada como censura moral.² Quem nela permanece seria acomodado, resistente à mudança ou pouco disposto ao crescimento. Esse uso é insuficiente para compreender as relações entre seres humanos e inteligência artificial, porque desloca toda a responsabilidade para o indivíduo e ignora a arquitetura que organiza suas escolhas.

O conforto pode ser legítimo e produtivo. Reduzir trabalho repetitivo, organizar informações, superar barreiras de acessibilidade, traduzir conteúdos ou liberar memória de trabalho pode ampliar capacidades humanas. A história da cognição também é a história de instrumentos externos: linguagem escrita, mapas, calendários, bibliotecas, calculadoras e mecanismos de busca.

Os ganhos de produtividade, por sua vez, são reais e não devem ser reduzidos à aparência. Em experimento pré-registrado com 453 profissionais, Noy e Zhang (2023) observaram redução média de 40% no tempo de execução e aumento de 18% na qualidade de tarefas de escrita realizadas com ChatGPT. O estudo não mediu retenção de conhecimento ou capacidade posterior de atuação sem a ferramenta. Por isso, evidencia benefícios relevantes de desempenho, mas deixa aberta a questão da aprendizagem e da reversibilidade no tempo.

Risko e Gilbert (2016) denominam *cognitive offloading* — descarregamento cognitivo — o uso de ações ou recursos externos para reduzir as exigências internas de processamento de uma tarefa. Anotar um compromisso para não depender da memória é uma forma simples desse fenômeno. O descarregamento pode melhorar precisão, velocidade e desempenho, permitindo que recursos mentais limitados sejam empregados em atividades mais complexas.

² Neste artigo, “zona de conforto” é empregada como expressão descritiva de uso corrente, e não como diagnóstico psicológico ou categoria clínica padronizada.

Por isso, não se deve transformar esforço em virtude abstrata nem dificuldade em requisito obrigatório de autenticidade. Uma tecnologia não é emancipadora por ser difícil, assim como não é capturadora apenas por ser conveniente.

A armadilha surge sob condições mais específicas:

1. o desempenho imediato melhora ou o custo de execução diminui;
2. a aprendizagem, a transferência ou a competência própria não avançam em proporção equivalente — ou se deterioram;
3. a capacidade de monitorar, verificar e contestar enfraquece com o uso continuado;
4. a dependência e o custo de substituir ou abandonar o sistema tornam-se funcionalmente relevantes.

Proponho denominar esse processo **armadilha do conforto cognitivo algorítmico**: a dinâmica pela qual a redução continuada do esforço e da incerteza produz benefícios imediatos, mas pode substituir o exercício necessário à formação, à conservação e à recuperação autônoma de capacidades.

Para evitar que o conceito alcance qualquer uso recorrente de uma ferramenta, a assistência somente será classificada aqui como captura quando três condições convergirem: redução persistente e relevante da capacidade própria; aumento material do custo de reversão ou saída; e assimetria de controle sobre conhecimento, dados, critérios ou infraestrutura essenciais. Os limiares precisam ser definidos conforme a atividade e o risco. Preferência, conveniência ou dependência ordinária de instrumentos, isoladamente, não bastam.

O elemento decisivo não é a comodidade isolada. É a passagem silenciosa da assistência para a dependência.

3. Quando facilitar modifica aquilo que sabemos fazer

Experimentos anteriores à popularização das IAs generativas já haviam identificado uma tensão entre desempenho imediato e memória posterior. Grinschgl, Papenmeier e Meyerhoff (2021), em três experimentos com 172 participantes cada, observaram que o descarregamento cognitivo podia acelerar a execução de uma tarefa, mas reduzir a memória subsequente das informações externalizadas. O efeito não era inevitável: no terceiro experimento, participantes submetidos ao descarregamento máximo e previamente informados sobre o teste de memória conseguiram neutralizar quase completamente o impacto negativo. A finalidade consciente e o desenho da atividade importavam.

Esse resultado impede duas conclusões simplistas. A primeira seria afirmar que toda externalização produz empobrecimento cognitivo. A segunda seria presumir que recursos liberados pela ferramenta serão automaticamente direcionados a formas superiores de

pensamento. Eles podem ser utilizados para análise mais profunda, mas também podem simplesmente desaparecer do processo.

Com a IA generativa, a externalização alcança funções mais amplas. O sistema não apenas armazena uma informação que o usuário previamente escolheu registrar. Ele pode formular o problema, selecionar elementos relevantes, estabelecer relações, redigir a resposta, antecipar objeções e apresentar uma conclusão em linguagem fluente.

Essa fluência tem valor funcional, mas também atua como pista metacognitiva. A literatura mostra que a facilidade subjetiva de processamento pode influenciar julgamentos de verdade, familiaridade e confiança (ALTER; OPPENHEIMER, 2009). Isso não demonstra, por si só, que toda resposta fluente de IA produza aceitação acrítica. Sustenta, porém, uma hipótese verificável: respostas claras, rápidas e linguisticamente seguras podem reduzir a percepção de incerteza e antecipar a conclusão antes que o usuário tenha construído o percurso necessário para avaliá-la.

A literatura sobre automação já descrevia o risco da confiança inadequadamente calibrada. Parasuraman e Riley (1997) distinguiram uso, desuso, uso indevido e abuso da automação. A confiança excessiva pode produzir dependência inadequada, falhas de monitoramento e vieses decisórios. Mosier e Skitka (1999) associaram o viés de automação à tendência humana de seguir o caminho de menor esforço cognitivo, utilizando a recomendação automatizada como substituta da busca vigilante por informações.

Parasuraman, Sheridan e Wickens (2000) acrescentaram que a automação não apenas substitui trabalho: modifica a atividade humana e cria novas exigências de coordenação. Seu modelo distingue automação na aquisição de informações, na análise, na seleção de decisões e na implementação da ação. Essa decomposição é relevante porque a autonomia pode ser preservada em uma etapa e enfraquecida em outra; não existe um único nível abstrato de ‘uso de IA’ capaz de representar todas essas situações.

Em pesquisa com 319 trabalhadores do conhecimento e 936 exemplos reais de uso de IA generativa, Lee et al. (2025) encontraram associação entre maior confiança na capacidade da IA e menor exercício percebido de pensamento crítico. Maior confiança do trabalhador em sua própria capacidade esteve associada a maior reflexão. O estudo é baseado em autorrelatos e não estabelece causalidade, mas identifica um deslocamento relevante: parte do esforço deixa a execução e migra para a verificação, a integração da resposta e a supervisão da tarefa.

Trata-se de uma atualização das “ironias da automação” descritas por Bainbridge (1983). À medida que as operações rotineiras são automatizadas, o ser humano tende a ser convocado principalmente nas situações excepcionais — justamente aquelas em que precisará de competências que passou a exercer com menor frequência.

Essa migração só preserva a autonomia quando o usuário possui conhecimento suficiente, tempo disponível e incentivos reais para supervisionar. Caso contrário, surge um paradoxo: exige-se que a pessoa revise um resultado produzido por um sistema cuja utilização reiterada pode estar reduzindo justamente a competência necessária para revisá-lo.

4. A arquitetura do conforto

A armadilha não precisa ser planejada por um agente mal-intencionado. Pode resultar da convergência entre incentivos comerciais, métricas de produtividade e preferências humanas legítimas por reduzir tempo, dúvida e esforço.

Em modelos de negócio cuja adoção depende de utilidade percebida, recorrência ou permanência, existe incentivo para reduzir fricções e ampliar a satisfação imediata. Isso não autoriza presumir intenção de enfraquecer o usuário. Indica apenas que rapidez, fluência e conveniência podem ser recompensadas sem que aprendizagem, capacidade de contestação ou reversibilidade recebam mensuração equivalente.

Como modelo heurístico — e não como sequência causal universal — o mecanismo pode ser representado por sete momentos relacionados:

facilidade → fluência → confiança → delegação → menor exercício → perda de competência → dependência

No início, o usuário pode delegar tarefas periféricas. Depois, pode solicitar estruturas, argumentos, avaliações e decisões preliminares. Quando personalizada, a ferramenta preserva contexto, aprende preferências e pode se tornar o ponto inicial de acesso ao problema. Interações percebidas como bem-sucedidas tendem a tornar novas delegações mais prováveis, mas o efeito varia conforme a tarefa, o conhecimento prévio, o desenho da interface e os incentivos de verificação.

O processo é reforçado porque seus benefícios são imediatos, enquanto os custos cognitivos são diferidos e dificilmente mensurados. O tempo economizado aparece. A redução da frustração é sentida. A entrega concluída pode ser apresentada à organização. A perda de capacidade somente se revela quando o sistema falha, fica indisponível, encontra uma situação inédita ou precisa ser contestado.

Os momentos podem se sobrepor, interromper-se ou regredir; não se afirma que todo usuário percorra a cadeia completa. A proposta descreve um mecanismo candidato. Sua confirmação exige comparação ao longo do tempo entre capacidades anteriores e posteriores ao uso, tarefas assistidas e não assistidas, e diferentes formas de desenho da ferramenta.

Nesse momento, o usuário pode ocupar formalmente a posição de decisor, mas já não dominar materialmente as condições da decisão.

5. A influência cognitiva estrutural

O viés de automação costuma ser observado em um episódio: a máquina recomenda, o ser humano confia excessivamente e uma decisão inadequada é tomada. A influência cognitiva estrutural opera em uma camada anterior e temporalmente mais extensa.

Chamo de **influência cognitiva estrutural** a atuação das arquiteturas informacionais sobre as condições de percepção, compreensão, formulação do problema e escolha, antes que a decisão final seja exteriorizada. Ela não exige uma ordem, um engano deliberado ou a imposição direta de determinado resultado. Pode ocorrer pela seleção do que aparece, pela ordem das alternativas, pela linguagem de autoridade, pela supressão de incertezas e pela criação de hábitos de dependência.³

A armadilha do conforto cognitivo é uma de suas possíveis manifestações. A influência não incide apenas sobre o conteúdo que o usuário escolherá. Incide sobre quanto ele ainda consegue formular, investigar, comparar e decidir sem recorrer à arquitetura.

Isso explica por que até uma resposta correta pode participar de um processo de captura. O problema não está necessariamente no resultado isolado, mas na substituição reiterada do percurso cognitivo que permitiria ao usuário compreender suas razões, reconhecer seus limites e produzir alternativa própria.

Uma resposta correta, um uso frequente ou a simples preferência pelo sistema não comprovam captura. A inferência exige evidências reiteradas dos três critérios antes delimitados: perda de capacidade residual, custo material de saída e assimetria de controle. Sem essa convergência, o fenômeno deve ser descrito com categorias menos intensas — assistência, confiança, hábito, dependência funcional ou aprisionamento contratual — conforme o caso.

Em termos sintéticos:

A influência cognitiva estrutural orienta as condições da decisão; a armadilha do conforto pode enfraquecer as condições de decidir sem a própria infraestrutura que as orienta.

O usuário continua clicando voluntariamente. Continua solicitando respostas e aceitando sugestões. A aparência de liberdade permanece. O que se modifica é a estrutura das capacidades disponíveis e, conseqüentemente, a qualidade material dessa liberdade.

³ “Influência cognitiva estrutural” é empregada aqui como categoria analítica proposta pelo autor para examinar como arquiteturas informacionais condicionam etapas anteriores à decisão manifesta.

6. A matriz preliminar de reversibilidade

Para distinguir assistência capacitadora de dependência, proponho uma **matriz analítica preliminar de reversibilidade**. Ela não constitui escala psicométrica validada, certificação nem teste com ponto de corte universal. Sua função é organizar perguntas e evidências contextuais. A questão central é:

Se o sistema fosse retirado ou substituído, o usuário e a instituição ainda conseguiriam explicar, contestar, reconstruir e executar os elementos essenciais da tarefa?

A matriz possui duas dimensões complementares. Cada resposta deve ser acompanhada por evidência observável — demonstração prática, documentação, registro de auditoria, teste de transferência ou simulação de indisponibilidade — e não apenas por autopercepção.

Dimensão cognitiva

1. O usuário consegue explicar as razões do resultado sem repetir a linguagem produzida pela IA?
2. Consegue identificar fontes, incertezas, premissas, limitações e alternativas relevantes?
3. Possui conhecimento suficiente para rejeitar uma resposta plausível, porém incorreta?
4. Consegue executar os componentes críticos da tarefa e transferir o raciocínio para uma situação inédita quando a ferramenta está indisponível?

Dimensão tecnológico-institucional

5. É possível substituir o sistema ou o fornecedor sem perder acesso ao histórico, aos dados, aos critérios e às capacidades essenciais?
6. Dados, registros, instruções e resultados podem ser exportados em formato utilizável e interoperável?
7. A organização consegue auditar intervenções, reconstruir decisões e identificar responsabilidades humanas e técnicas?
8. Existem pessoas, documentação e procedimentos suficientes para manter a função crítica durante indisponibilidade ou migração?

Respostas positivas, apoiadas em evidências, sugerem preservação de capacidade e de saída. Respostas reiteradamente negativas indicam risco, mas não autorizam conclusão automática: a classificação depende da gravidade da função, da duração do uso, da existência de alternativas e da convergência dos critérios de captura.

As dimensões também podem divergir. Um profissional pode conservar conhecimento e, ainda assim, estar preso a uma infraestrutura sem portabilidade; uma organização pode trocar de

fornecedor e, ao mesmo tempo, ter perdido competência interna para revisar resultados. Essa separação evita confundir erosão cognitiva com aprisionamento tecnológico.

A matriz não exige que todas as tarefas sejam refeitas manualmente. Uma sociedade tecnologicamente complexa depende de especialização e ferramentas. Seu objetivo é verificar se permaneceu capacidade suficiente para compreender, supervisionar, contestar e substituir aquilo que foi delegado. Aplicações sérias devem registrar uma linha de base, repetir avaliações em intervalos definidos e documentar incidentes, erros detectados e testes de saída.

Essa capacidade residual é especialmente importante em contextos de alto impacto. Médicos, advogados, magistrados, gestores públicos, engenheiros e profissionais de segurança podem utilizar sistemas sofisticados sem dominar todos os detalhes computacionais. Mas precisam conservar competência substantiva para reconhecer quando a resposta tecnológica colide com fatos, direitos, evidências ou limites profissionais.

Sem isso, a supervisão humana transforma-se em uma formalidade: alguém permanece responsável por uma decisão que já não consegue reconstruir.

7. Do indivíduo ao Estado: escalas e consequências jurídico-institucionais

A armadilha do conforto pode ser investigada em três escalas. A estrutura comparativa abaixo é uma hipótese analítica: os estudos individuais citados não demonstram, por si mesmos, efeitos organizacionais ou estatais. Cada escala exige indicadores, séries temporais e evidências próprias.

Na escala **individual**, a pessoa produz mais, mas compreende, recorda ou contesta menos. A dependência torna-se visível quando precisa agir sem a ferramenta.

Na escala **organizacional**, equipes passam a entregar resultados assistidos por IA, enquanto conhecimento tácito, memória institucional e capacidade interna de auditoria deixam de ser exercidos. A empresa aparenta eficiência, mas sua continuidade passa a depender de sistemas que poucos conseguem avaliar.

Na escala estatal, a administração pública pode adotar soluções prontas para saúde, segurança, justiça, educação e gestão. O ganho imediato pode ser real. Entretanto, se não houver infraestrutura adequada, equipes qualificadas, documentação, auditabilidade, interoperabilidade e capacidade de substituição, a conveniência tecnológica pode converter-se em dependência com consequências para a autonomia institucional do Estado.

Nos três níveis, a estrutura é semelhante:

- o benefício imediato é atribuído ao usuário ou à instituição;
- o conhecimento necessário migra para uma infraestrutura externa;
- a responsabilidade permanece com quem utiliza;
- o poder efetivo desloca-se para quem controla o sistema.

Por essa razão, governança de IA não deve avaliar apenas precisão, segurança e produtividade. Precisa perguntar quais competências humanas e institucionais deverão permanecer disponíveis depois da automação.

No Brasil, o art. 20 da Lei Geral de Proteção de Dados Pessoais assegura ao titular o direito de solicitar revisão de decisões tomadas unicamente com base em tratamento automatizado que afetem seus interesses; o § 1º exige informações claras e adequadas sobre critérios e procedimentos, resguardados os segredos comercial e industrial (BRASIL, 2018). O direito formal à revisão, contudo, não garante revisão materialmente efetiva quando faltam competência, tempo, documentação ou acesso às evidências necessárias para contestar o resultado.

Como referência comparada, o art. 14 do Regulamento Europeu de Inteligência Artificial vincula supervisão humana à competência, ao treinamento e à autoridade do supervisor, à consciência sobre o viés de automação e à capacidade de interpretar, desconsiderar, reverter ou interromper a operação de sistemas de alto risco (UNIÃO EUROPEIA, 2024). A presença nominal de uma pessoa no fluxo decisório, portanto, não equivale a controle humano efetivo.

A reversibilidade institucional também possui correspondentes concretos de governança. O capítulo VI do Regulamento Europeu de Dados disciplina a remoção de obstáculos à troca entre serviços de processamento de dados e à portabilidade (UNIÃO EUROPEIA, 2023). O AI Risk Management Framework do NIST inclui governança, definição de funções humano-IA, monitoramento, documentação e desativação entre seus resultados de gestão (NIST, 2023). Embora esses instrumentos não sejam indistintamente aplicáveis ao contexto brasileiro, oferecem parâmetros comparativos para contratação, auditoria e planejamento de saída.

Em atividades de alto impacto, algumas salvaguardas tornam-se particularmente relevantes: exposição de incertezas; acesso às fontes; registro das intervenções; justificativa humana independente; testes periódicos sem assistência; treinamento para contestação; possibilidade de revisão efetiva; simulações de indisponibilidade; e cláusulas de contratação que assegurem documentação, interoperabilidade, portabilidade, continuidade do serviço e saída sem perda funcional desproporcional.

Não se trata de introduzir dificuldade artificial em todas as interações. Trata-se de preservar a fricção cognitiva necessária nos pontos em que compreender, justificar e contestar fazem parte da própria legitimidade da decisão.

8. Delimitação, indicadores e falsificabilidade

A hipótese da armadilha é condicional e deve poder ser contrariada em contextos definidos. Ela perderá força se, após uso continuado por período previamente estabelecido, os usuários conservarem ou ampliarem conhecimento, transferirem o aprendizado para tarefas inéditas, detectarem erros sem auxílio, mantiverem desempenho independente e substituírem a ferramenta sem perda funcional ou custo desproporcional.

Uma avaliação adequada deve combinar indicadores anteriores e posteriores ao uso: retenção, transferência, detecção de respostas incorretas, qualidade da justificativa própria, execução de componentes críticos sem assistência, tempo e custo de migração, portabilidade de dados e continuidade operacional. Autorrelatos isolados não bastam. Os limiares devem ser definidos antes da avaliação e calibrados conforme risco, complexidade e consequências do erro.

Também não haverá captura quando o descarregamento de tarefas inferiores liberar recursos efetivamente utilizados em reflexão superior. Um pesquisador pode empregar IA para organizar referências e dedicar mais tempo à análise. Um médico pode automatizar registros e ampliar a atenção ao paciente. Um servidor pode reduzir trabalho burocrático e qualificar a fundamentação da decisão. Ganhos documentados de produtividade, como os encontrados por Noy e Zhang (2023), estabelecem um ônus importante para a hipótese: eles não podem ser tratados como ilusórios. A questão adicional é saber se permanecem acompanhados de aprendizagem, supervisão e capacidade de saída.

O estudo de Bassner et al. (2026) fornece uma indicação importante nesse sentido. O tutor que oferecia pistas, mas não respostas completas, combinou assistência com participação ativa e preservou benefícios motivacionais. O desenho da ferramenta pode favorecer dependência ou aprendizagem.

Portanto, a pergunta correta não é se devemos aceitar ou rejeitar a inteligência artificial. É outra:

Nas condições concretas de uso, o desenho, a adoção e os incentivos da arquitetura ampliam progressivamente a capacidade humana ou produzem dependência difícil de reverter?

Essa pergunta desloca o debate da moralidade individual para a responsabilidade do desenho, da adoção e da governança.

Considerações finais

A zona de conforto deixa de ser apenas um estado psicológico quando passa a ser produzida, personalizada e administrada por uma infraestrutura. Em modelos cuja sustentabilidade depende de recorrência ou permanência, a conveniência pode convergir com incentivos econômicos sem que competências residuais e capacidade de saída sejam avaliadas com o mesmo cuidado.

O risco mais difícil de perceber não é a substituição imediata do ser humano pela máquina. É a preservação formal do humano dentro do processo, acompanhada pela perda gradual das capacidades necessárias para compreender e controlar aquilo que a máquina realiza.

A pessoa continua assinando. O profissional continua revisando. A autoridade continua decidindo. Mas a distância entre responsabilidade formal e domínio material cresce a cada nova delegação não examinada.

É por isso que a armadilha do conforto cognitivo se conecta diretamente à autonomia decisória. A decisão não começa quando o usuário aceita a resposta. Começa muito antes, na arquitetura que define quanto ele precisará pensar, quais elementos verá, que dúvidas encontrará e quais capacidades continuará exercendo.

Em *Autonomia Decisória na Era da IA*, sustento que decisões aparentemente individuais podem ser condicionadas por estruturas informacionais anteriores ao clique. A armadilha analisada neste artigo acrescenta uma dimensão temporal a essa tese: a mesma arquitetura que influencia uma decisão presente pode modificar a capacidade de produzir decisões futuras sem sua intermediação.

Este artigo integra o lançamento de *Autonomia Decisória na Era da IA*, obra já publicada pela Editora Dialética e disponível para aquisição na loja da editora. O livro desenvolve em profundidade a tese aqui condensada: a autonomia não se preserva apenas pela manutenção formal da escolha, mas pelo domínio material das condições que tornam a escolha possível.

O conforto não captura porque reduz o sofrimento de uma tarefa. Captura quando transforma alívio em desuso, desuso em dependência e dependência em incapacidade de saída.

A decisão não começa no clique. A dependência também não. Ela começa quando a facilidade deixa de apoiar o pensamento e passa silenciosamente a substituí-lo.

Referências

- ALTER, Adam L.; OPPENHEIMER, Daniel M. Uniting the tribes of fluency to form a metacognitive nation. *Personality and Social Psychology Review*, v. 13, n. 3, p. 219–235, 2009. DOI: [10.1177/1088868309341564](https://doi.org/10.1177/1088868309341564).
- BAINBRIDGE, Lisanne. Ironies of automation. *Automatica*, v. 19, n. 6, p. 775–779, 1983. DOI: [10.1016/0005-1098\(83\)90046-8](https://doi.org/10.1016/0005-1098(83)90046-8).
- BASSNER, Patrick; LENK-OSTENDORF, Ben; BEINSTINGEL, Ramona; WASNER, Tobias; KRUSCHE, Stephan. Less stress, better scores, same learning: the dissociation of performance and learning in AI-supported programming education. *Computers and Education: Artificial Intelligence*, v. 10, art. 100537, 2026. DOI: [10.1016/j.caeai.2025.100537](https://doi.org/10.1016/j.caeai.2025.100537).
- BRASIL. Lei nº 13.709, de 14 de agosto de 2018. Lei Geral de Proteção de Dados Pessoais (LGPD). Brasília, DF: Presidência da República, 2018. Disponível em: https://www.planalto.gov.br/ccivil_03/_ato2015-2018/2018/lei/13709compilado.htm. Acesso em: 22 jul. 2026.
- GENOVA, Jandislei Antonio. *Autonomia decisória na era da IA: formação da vontade, influência cognitiva estrutural e responsabilidade jurídica em ambientes algorítmicos*. São Paulo: Editora Dialética, 2026. ISBN 978-65-288-0467-2.
- GRINSCHGL, Sandra; PAPENMEIER, Frank; MEYERHOFF, Hauke S. Consequences of cognitive offloading: boosting performance but diminishing memory. *Quarterly Journal of Experimental Psychology*, v. 74, n. 9, p. 1477–1496, 2021. DOI: [10.1177/17470218211008060](https://doi.org/10.1177/17470218211008060).
- LEE, Hao-Ping et al. The impact of generative AI on critical thinking: self-reported reductions in cognitive effort and confidence effects from a survey of knowledge workers. In: *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems — CHI '25*. New York: Association for Computing Machinery, 2025. art. 1121, p. 1–22. DOI: [10.1145/3706598.3713778](https://doi.org/10.1145/3706598.3713778).
- MOSIER, Kathleen L.; SKITKA, Linda J. Automation use and automation bias. *Proceedings of the Human Factors and Ergonomics Society Annual Meeting*, v. 43, n. 3, p. 344–348, 1999. DOI: [10.1177/154193129904300346](https://doi.org/10.1177/154193129904300346).
- NATIONAL INSTITUTE OF STANDARDS AND TECHNOLOGY (NIST). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*. Gaithersburg, MD: NIST, 2023. DOI: [10.6028/NIST.AI.100-1](https://doi.org/10.6028/NIST.AI.100-1).
- NOY, Shakked; ZHANG, Whitney. Experimental evidence on the productivity effects of generative artificial intelligence. *Science*, v. 381, n. 6654, p. 187–192, 2023. DOI: [10.1126/science.adh2586](https://doi.org/10.1126/science.adh2586).
- PARASURAMAN, Raja; RILEY, Victor. Humans and automation: use, misuse, disuse, abuse. *Human Factors*, v. 39, n. 2, p. 230–253, 1997. DOI: [10.1518/001872097778543886](https://doi.org/10.1518/001872097778543886).
- PARASURAMAN, Raja; SHERIDAN, Thomas B.; WICKENS, Christopher D. A model for types and levels of human interaction with automation. *IEEE Transactions on Systems, Man, and Cybernetics — Part A: Systems and Humans*, v. 30, n. 3, p. 286–297, 2000. DOI: [10.1109/3468.844354](https://doi.org/10.1109/3468.844354).
- RISKO, Evan F.; GILBERT, Sam J. Cognitive offloading. *Trends in Cognitive Sciences*, v. 20, n. 9, p. 676–688, 2016. DOI: [10.1016/j.tics.2016.07.002](https://doi.org/10.1016/j.tics.2016.07.002).
- UNIÃO EUROPEIA. Regulamento (UE) 2023/2854 do Parlamento Europeu e do Conselho, de 13 de dezembro de 2023, relativo a regras harmonizadas sobre o acesso equitativo aos dados e a sua utilização (Regulamento dos Dados). *Jornal Oficial da União Europeia*, 22 dez. 2023. Disponível em: <https://eur-lex.europa.eu/eli/reg/2023/2854/oj>. Acesso em: 22 jul. 2026.
- UNIÃO EUROPEIA. Regulamento (UE) 2024/1689 do Parlamento Europeu e do Conselho, de 13 de junho de 2024, que cria regras harmonizadas em matéria de inteligência artificial. *Jornal Oficial da União Europeia*, 12 jul. 2024. Disponível em: <https://eur-lex.europa.eu/eli/reg/2024/1689/oj>. Acesso em: 22 jul. 2026.