

DOCUMENTO TÉCNICO-PEDAGÓGICO INTERNO

Fricção Cognitiva *Ótima* e aprendizagem

Referencial teórico para o design de ferramentas de inteligência artificial para a educação básica.

CONTEÚDO

Sumário

01	Apresentação	03
02	Fundamentação teórica: dificuldade desejável, carga cognitiva e o ponto ótimo	04
03	Os conceitos nas diferentes línguas	08
04	Os quatro efeitos de Bjork	12
05	Fracasso produtivo: o esforço de tentar antes de aprender	15
06	Evidência empírica: IA generativa, descarga cognitiva e dívida cognitiva	17
07	Riscos da IA generativa na educação básica	19
08	O princípio: retenção do produto cognitivo final por padrão	20
09	Implicações: equidade, reversão de expertise e ética	24
10	Da fricção ao sistema: a ontologia operacional educativa	26
11	Correção assistida por IA: configurações do humano no circuito	28
12	Seis hábitos mentais para educadores na era da IA	31
—	Referências	33

01

SEÇÃO UM

Apresentação

Este documento estabelece o referencial teórico que fundamenta uma decisão de design transversal às ferramentas do **IA para escolas**: por padrão, os recursos baseados em inteligência artificial não entregam ao estudante o produto cognitivo final de uma tarefa de aprendizagem. Em vez de devolver a resposta pronta (a redação corrigida e reescrita, a solução completa do problema, o resumo acabado), as ferramentas privilegiam a devolução de perguntas diagnósticas, critérios, pistas graduadas e estruturas de apoio que mantêm sob responsabilidade do estudante o esforço cognitivo essencial à aprendizagem.

O fundamento dessa decisão se apoia na expressão **Fricção Cognitiva Ótima**: a ideia de que existe um ponto ideal de dificuldade em uma tarefa de aprendizagem. Esse ponto ideal exige nem fricção de menos, que produz aprendizagem rasa e dependência, nem fricção de mais, que gera sobrecarga e frustração.

ESCLARECIMENTO TERMINOLÓGICO

Fricção Cognitiva Ótima é uma expressão-síntese criada por Guilherme Cintra e relacionada a dois conceitos aplicados atualmente no debate sobre IA na educação: *dificuldades desejáveis* (desirable difficulties), da psicologia cognitiva da aprendizagem, e *carga cognitiva*, da psicologia instrucional.

Por essa razão, ao longo do documento, essa expressão é ancorada explicitamente nesses corpos teóricos, recuperando suas formulações originais e mostrando como convergem para uma mesma diretriz de projeto.

02

SEÇÃO DOIS

Dificuldade desejável, carga cognitiva e o ponto ótimo

O conceito mais próximo de uma base empírica para a fricção ótima é o de **dificuldades desejáveis**, formulado por Robert A. Bjork e desenvolvido com Elizabeth L. Bjork. Refere-se a desafios intencionais que tornam a aprendizagem mais exigente no curto prazo, mas significativamente mais robusta e duradoura no longo prazo. As quatro dificuldades desejáveis mais bem documentadas empiricamente são a prática espaçada, a prática intercalada, o efeito de teste (recuperação ativa) e o efeito de geração. O princípio explica por que atividades que produzem sensação de domínio, como a releitura passiva, frequentemente pouco contribuem para a aprendizagem, ao passo que o esforço de recuperar da memória, mais custoso, consolida a aprendizagem. Vale notar que esses não são achados de contextos artificiais: são apresentados por seus próprios autores como resultados validados no mundo real da educação (Bjork e Bjork, em Gernsbacher e col., 2011).

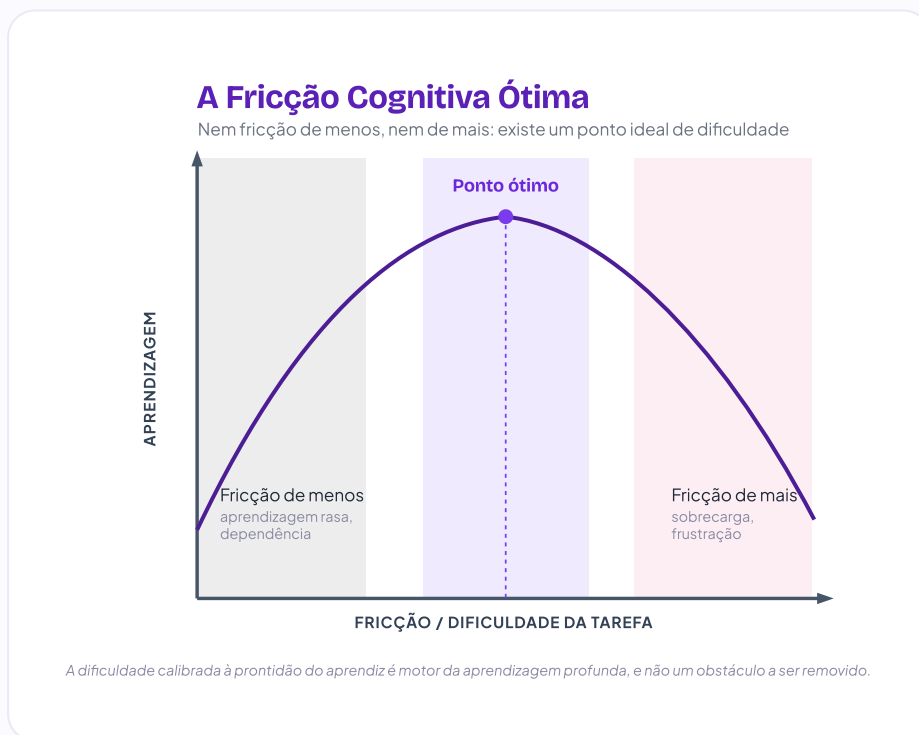


FIGURA 1 A Fricção Cognitiva Ótima: o ponto ideal entre fricção de menos e de mais.

A **Teoria da Carga Cognitiva**, de John Sweller, fornece o segundo pilar. Partindo do fato de que a memória de trabalho é limitada, distingue duas formas de carga. A *carga extrínseca* é o esforço improdutivo, imposto por má apresentação da informação, que deve ser eliminado; a *carga germane* é o esforço produtivo, dedicado à construção de esquemas mentais. Esta distinção é decisiva para o design: a mesma ferramenta pode remover fricção útil (o que prejudica a aprendizagem) ou remover fricção inútil (o que a beneficia). O objetivo não é minimizar o esforço, mas alocá-lo onde ele constrói conhecimento. Essa distinção tem raiz no estudo seminal de Sweller (1988): ao resolver problemas convencionais por análise meios-fins, o aprendiz consome quase toda a memória de trabalho perseguindo a meta, restando pouca capacidade para a construção de esquemas, razão pela qual resolver muitos problemas nem sempre ensina, achado que está na origem do efeito do exemplo trabalhado.

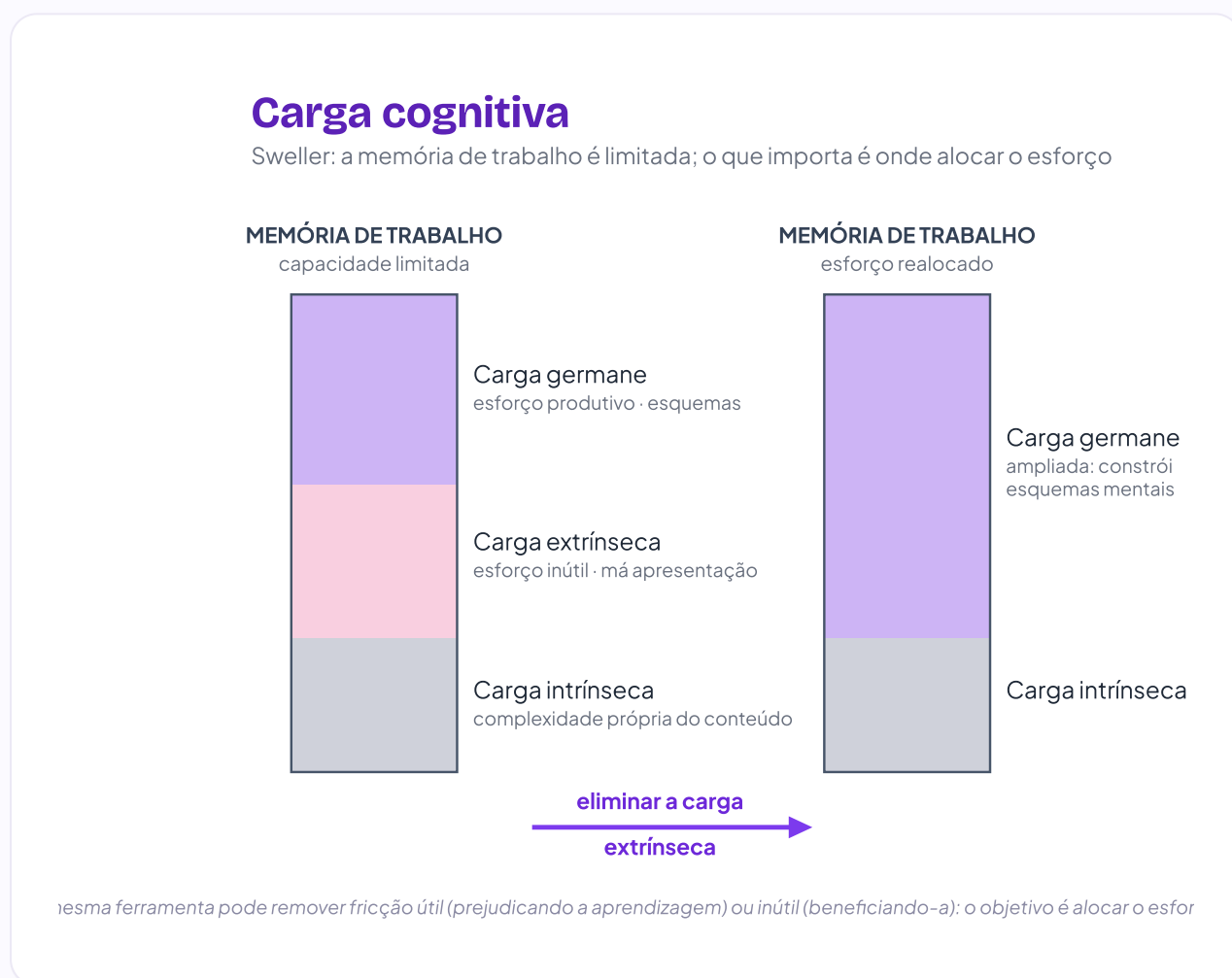


FIGURA 2 Carga cognitiva: eliminar a carga extrínseca e preservar a germane (Sweller).

A noção de **ponto ótimo** é convergente em diversas tradições: a Zona de Desenvolvimento Proximal de Vygotsky, o estado de fluxo de Csikszentmihalyi (equilíbrio entre desafio e habilidade) e, em formulação quantitativa, a “Regra dos 85%” de Wilson e colaboradores, segundo a qual a taxa de erro ótima para a aprendizagem situa-se em torno de 15,87%, treino nem fácil nem difícil demais. Convém, porém, tratar esse número como ilustrativo, e não prescritivo: ele foi derivado matematicamente para tarefas de classificação binária e algoritmos de aprendizagem do tipo gradiente-descendente, incluindo redes neurais artificiais e biológicas, e não medido em salas de aula reais. Em conjunto, esses trabalhos sustentam que a dificuldade, calibrada à prontidão do aprendiz, é condição da aprendizagem profunda, e não um obstáculo a ser removido.

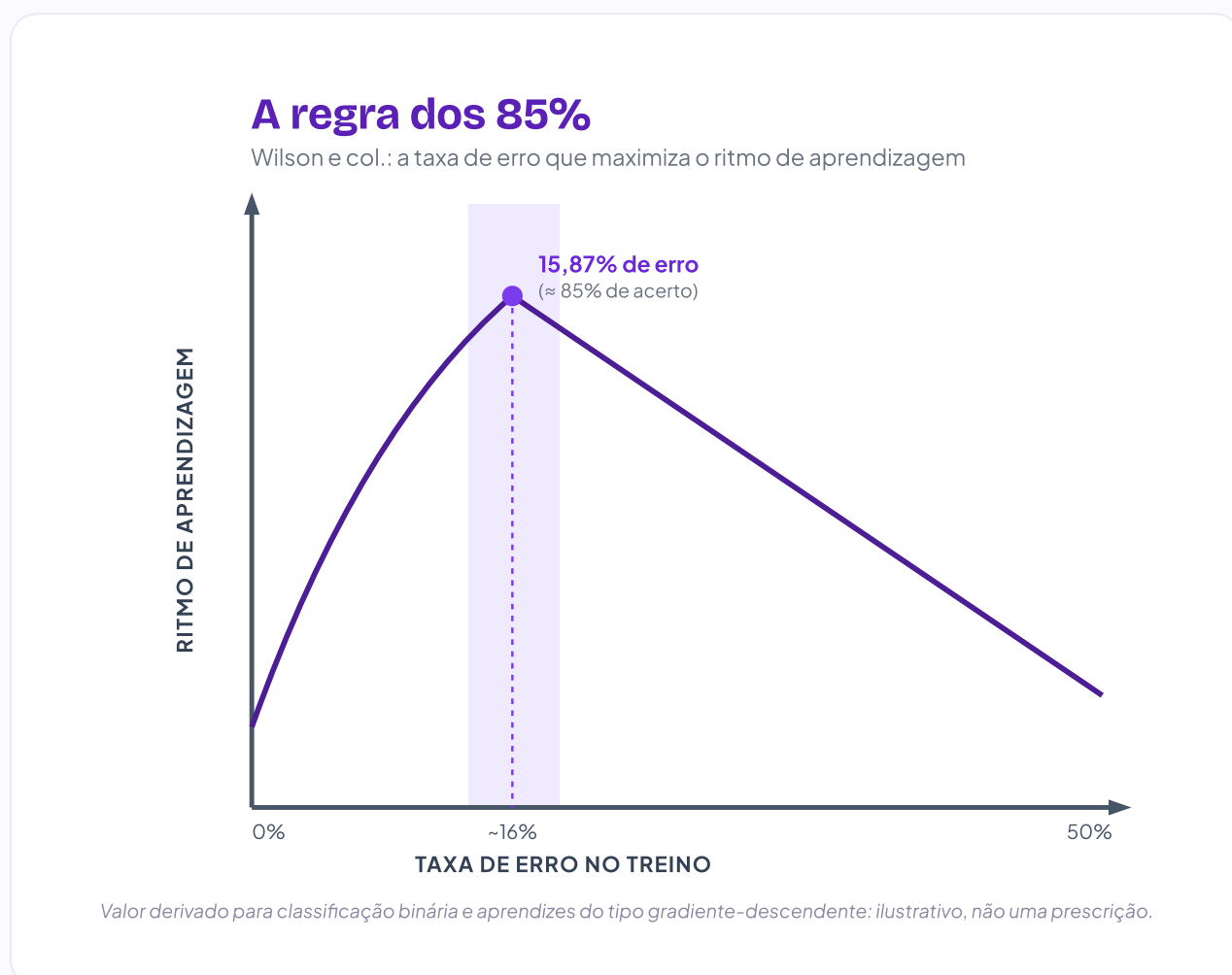


FIGURA 3 A regra dos 85%: a taxa de erro que maximiza o ritmo de aprendizagem (Wilson).

O **estado de fluxo** detalha esse ponto ótimo. Csikszentmihalyi descreve um canal de equilíbrio entre desafio e habilidade: quando o desafio supera a habilidade instala-se a ansiedade; quando a habilidade supera o desafio, o tédio. Duas qualidades desse modelo importam ao design. Primeiro, o equilíbrio é **dinâmico**: à medida que a habilidade cresce, o desafio precisa subir para manter o aprendiz no canal, o que fundamenta o desvanecimento gradual do apoio. Segundo, o que determina a experiência é a relação *percebida* entre desafio e habilidade, e não a objetiva (não são as habilidades que de fato temos, mas as que julgamos ter), o que liga o ponto ótimo à dimensão metacognitiva e à necessidade de calibragem individual.

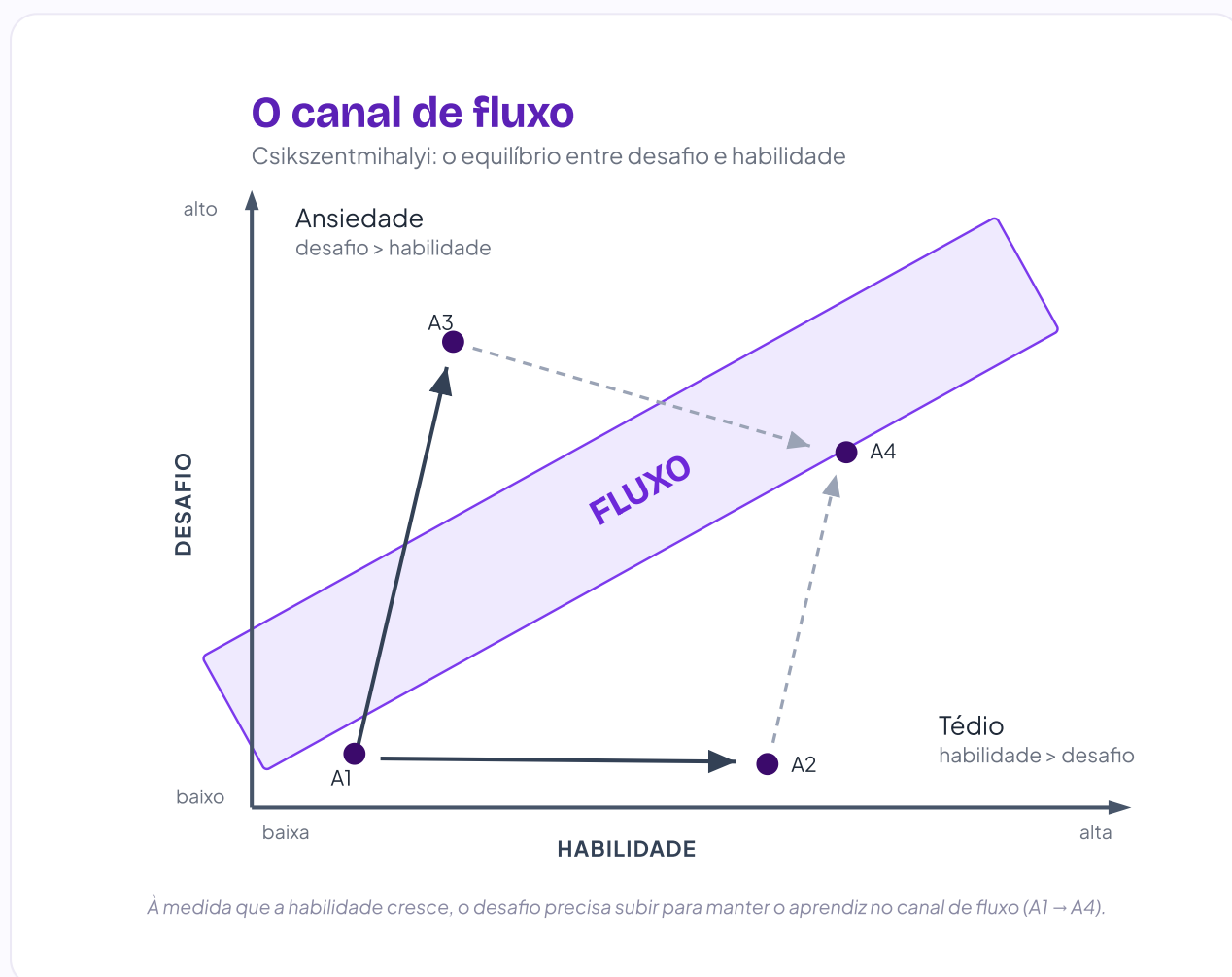


FIGURA 4 O canal de fluxo: o equilíbrio entre desafio e habilidade (Csikszentmihalyi).

03 SEÇÃO TRÊS

Os conceitos nas diferentes línguas

A expressão **Fricção Cognitiva Ótima** não possui tradução nem equivalente único entre as línguas em que o tema é discutido. Cada tradição linguística desenvolveu vocabulário próprio, com ênfases distintas, e a convergência só se dá no núcleo comum: a ideia de que uma dose calibrada de dificuldade é motor, e não obstáculo, da aprendizagem. O diagrama da Figura 5 sintetiza essas relações.



FIGURA 5 Os conceitos por tradição linguística e seus pontos de convergência.

Em **inglês**, coexistem duas linhagens. A primeira, de origem em design e interação humano-computador, é *cognitive friction*, termo cunhado por Alan Cooper (1999) para designar a resistência que o intelecto enfrenta diante de sistemas complexos, originalmente com conotação negativa, algo a eliminar. A segunda, da psicologia cognitiva, é *desirable difficulties* (Bjork), de conotação positiva: dificuldades intencionais que fortalecem a aprendizagem. A noção de “fricção ótima” nasce justamente do encontro dessas duas linhagens anglófonas.

Em **francês**, emprega-se *friction cognitive* sobretudo na bibliografia dos ambientes informatizados para a aprendizagem (EIAH: *Environnements Informatiques pour l'Apprentissage Humain*), em geral importando a acepção de Cooper. O conceito francófono mais distintivo, porém, é o de *conflit cognitif* (Piaget) e, sobretudo, *conflit sociocognitif* (Doise e Mugny), que enfatiza o desequilíbrio gerado na interação social como motor da reestruturação cognitiva, ênfase ausente nas demais tradições. Nessa mesma tradição francófona dos EIAH discute-se ainda a relação entre fricção cognitiva, emoções e a escolha entre trabalho individual e colaborativo (EduTech Wiki), antecipando a dimensão afetiva e social da dificuldade.

Em **alemão**, o termo consagrado é *wünschenswerte Erschwernisse*, que significa “dificuldades desejáveis”. É a tradição mais institucionalizada empiricamente, com linha de pesquisa dedicada (o projeto LOEWE) e autores de referência como Lipowsky e Richter, organizada em torno dos quatro efeitos propostos por Bjork: aprendizagem espaçada, intercalada, efeito de teste e efeito de geração.

Em síntese, as diferenças podem ser assim resumidas: o inglês reúne a origem do termo (UX) e a base empírica da memória e da aprendizagem; o francês acrescenta a dimensão social do conflito cognitivo; o alemão consolida e institucionaliza a vertente empírica das dificuldades desejáveis. O ponto de interseção das três tradições, a Fricção Cognitiva Ótima, não é, portanto, um conceito de um único autor, mas o nome que se dá a esse núcleo compartilhado.

NOTA TERMINOLÓGICA

Ao falar em **Fricção Cognitiva Ótima**, este documento ressignifica deliberadamente o termo de Cooper, originalmente algo a eliminar, em sentido positivo, como uma dose de dificuldade a *calibrar*, e não a suprimir.

Vale desenvolver a contribuição francófona, porque ela aponta uma dimensão ausente das demais tradições e fértil para o design. Para Doise e Mugny, o progresso cognitivo nasce com frequência do confronto entre pontos de vista divergentes: quando dois aprendizes sustentam respostas diferentes para o mesmo problema, o desequilíbrio social que daí resulta os obriga a coordenar perspectivas e a reconstruir, num patamar superior, a própria estrutura de raciocínio.

A fricção, nesse caso, não é apenas interna ao indivíduo, mas se origina na interação. A implicação de design é que as ferramentas não precisam tratar o estudante como um aprendiz isolado: podem provocar conflito sociocognitivo produtivo, em vez de encerrar a tarefa numa troca individual com a máquina, ao confrontar o aluno com respostas alternativas de colegas, ao pedir que defenda ou critique uma solução, ou ao promover a comparação entre raciocínios.

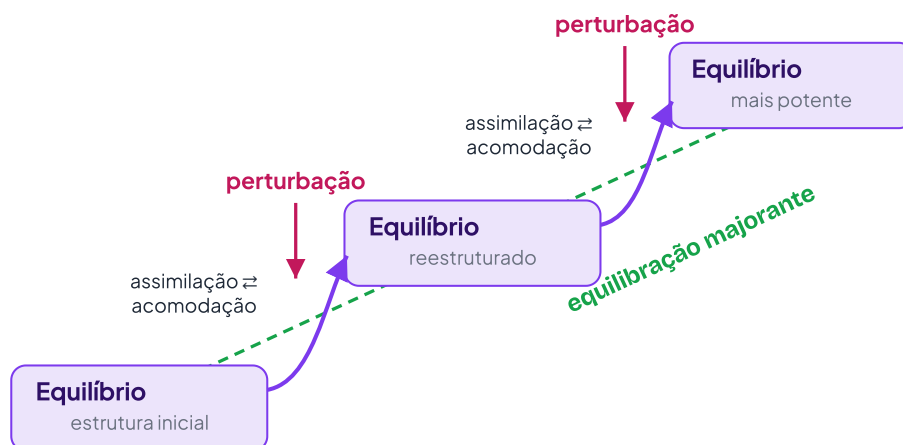


FIGURA 6 Conflito sociocognitivo: a fricção que nasce na interação entre pares (Doise e Mugny).

Convém recuar à raiz teórica mais profunda dessa ideia de fricção como motor: a **equilibração das estruturas cognitivas**, de Piaget. Para ele, a aprendizagem avança pela tensão entre assimilação (incorporar o novo aos esquemas existentes) e acomodação (modificar os esquemas diante do que resiste); uma perturbação, isto é, um desequilíbrio, força regulações e compensações que reconstróem o esquema. O ponto decisivo é a *equilibração majorante*: a boa perturbação não apenas restaura o equilíbrio anterior, mas eleva o sistema a um patamar mais potente. Tanto as dificuldades desejáveis quanto o conflito sociocognitivo de Doise descendem desse desequilíbrio produtivo piagetiano, o que dá unidade ao referencial.

Equilibração das estruturas cognitivas

Piaget: assimilação, acomodação e o equilíbrio majorante



Uma perturbação (desequilíbrio) força regulações e compensações que reconstróem o esquema num patamar superior. Raiz comum das dificuldades desejáveis e do conflito sociocognitivo: a boa perturbação eleva o sistema, não apenas o abala.

FIGURA 7 Equilibração das estruturas cognitivas e o equilíbrio majorante (Piaget).

04 SEÇÃO QUATRO

Os quatro efeitos de Bjork

Os quatro tipos de dificuldade desejável mais bem documentados na pesquisa de Bjork e aprofundados por Lipowsky e Richter são listados a seguir. O fio condutor dos quatro é o mesmo: **fluência e desempenho durante o estudo não são bons indicadores de aprendizagem**. As condições que produzem sensação de facilidade (reler, estudar em bloco, ver a resposta pronta) tendem a não implicar aprendizagem; as que geram um pouco de atrito e esforço (espaçar, intercalar, testar-se, gerar) constroem aprendizagem. É exatamente esse o princípio que sustenta a decisão de design das ferramentas: devolver ao estudante o esforço de gerar e recuperar, em vez de entregar o produto pronto. A razão de fundo é metacognitiva: como mostram os estudos sobre metacognição e metamemória (Metcalf e Shimamura), os aprendizes tendem a julgar mal o próprio aprendizado: a fluência de uma releitura produz uma ilusão de domínio que o esforço de recuperar e gerar desfaz.

4.1 — Aprendizagem espaçada (spacing / Verteiltes Lernen)

Em vez de estudar um tópico de uma vez só (*massed practice*, a “virada de noite”), você distribui o mesmo tempo total de estudo em sessões separadas por intervalos. Estudar 3 vezes por 30 minutos ao longo de uma semana produz retenção muito superior a estudar 90 minutos seguidos, mesmo sendo o mesmo tempo total. O mecanismo: quando há um intervalo, parte do conteúdo começa a ser esquecida, e o esforço de reativá-la na sessão seguinte fortalece a aprendizagem significativa.

Os quatro efeitos de Bjork

Dificuldades desejáveis: custam esforço no curto prazo, fortalecem a aprendizagem no longo prazo

1 · Aprendizagem espaçada

Distribuir o estudo em sessões separadas por intervalos.

Estudar tudo de uma vez dá uma sensação de domínio que evapora.

2 · Aprendizagem intercalada

Misturar os tipos de conteúdo (ABCA CBAC).

Com os conteúdos em bloco, o aluno entra no piloto automático.

3 · Efeito de teste

Autoteste, tentar lembrar sem olhar.

Reler é fluente, mas consolida muito menos.

4 · Efeito de geração

Gerar a resposta por conta própria antes de vê-la pronta.

Consultar a resposta priva de uma oportunidade de aprender.

**Fluência durante o estudo ≠ aprendizagem:
o que gera um pouco de atrito constrói aprendizagem significativa.**

FIGURA 8 Os quatro efeitos de Bjork.

4.2 — Aprendizagem intercalada (interleaving / Verschachteltes Lernen)

Em vez de praticar um tipo de problema em bloco até dominá-lo e só então passar ao próximo (AAAA BBBB CCCC), você mistura os tipos (ABCA CBAC...). Por exemplo: numa lista de matemática, alternar problemas de áreas, volumes e perímetros em vez de fazer dez de cada por vez. Funciona por dois motivos: força o cérebro a recuperar a estratégia certa do zero a cada problema e obriga a comparar as diferenças entre os tipos, que é justamente a habilidade exigida na vida real, quando os problemas vêm fora de ordem. Um exemplo empírico em contexto escolar é o estudo de Borromeo Ferri, Pede e Lipowsky (2021): alunos do 7º ano que estudaram proporções (*Zuordnungen*) de forma intercalada superaram, sobretudo em conhecimento conceitual, os que o estudaram de forma blocada, evidência de que a intercalação favorece não apenas a fluência procedural, mas a compreensão, justamente na educação básica.

Espaçar e intercalar

Duas dificuldades desejáveis que reorganizam a prática

Aprendizagem espaçada

Concentrado  *sensação de domínio que evapora*

Espaçado 
intervalos exigem reativar o que começou a ser esquecido: retenção superior

Aprendizagem intercalada

Blocado 

Intercalado 
força a recuperar a estratégia certa e a comparar os tipos: a habilidade exigida na vida real

A · B · C = tipos distintos de conteúdo

FIGURA 9 Aprendizagem espaçada e intercalada.

4.3 — Efeito de teste (testing / Testungseffekt)

Recuperar uma informação da memória pode significar responder a uma pergunta, fazer um autoteste ou tentar lembrar sem olhar. Esse ato fortalece a memória mais do que reestudar o mesmo material. A implicação é contraintuitiva: o teste não serve só para medir o que você sabe, mas para **consolidar** o que você sabe. Rer o capítulo passivamente dá sensação de fluência, mas fechar o livro e tentar recuperar o conteúdo (mesmo errando) produz aprendizagem bem mais significativa. É por isso que *flashcards* e perguntas de autoavaliação superam o grifar e rler.

4.4 — Efeito de geração (generation effect / Generierungseffekt)

Uma informação que você gerou por conta própria faz mais sentido do que uma que apenas leu ou recebeu pronta. Bjork formula isso quase como um alerta: toda vez que você consulta a resposta ou pede que alguém a mostre, em vez de tentar gerá-la a partir do que já sabe, você se priva de uma oportunidade poderosa de aprender.

05

SEÇÃO CINCO

Fracasso produtivo: o esforço de tentar antes de aprender

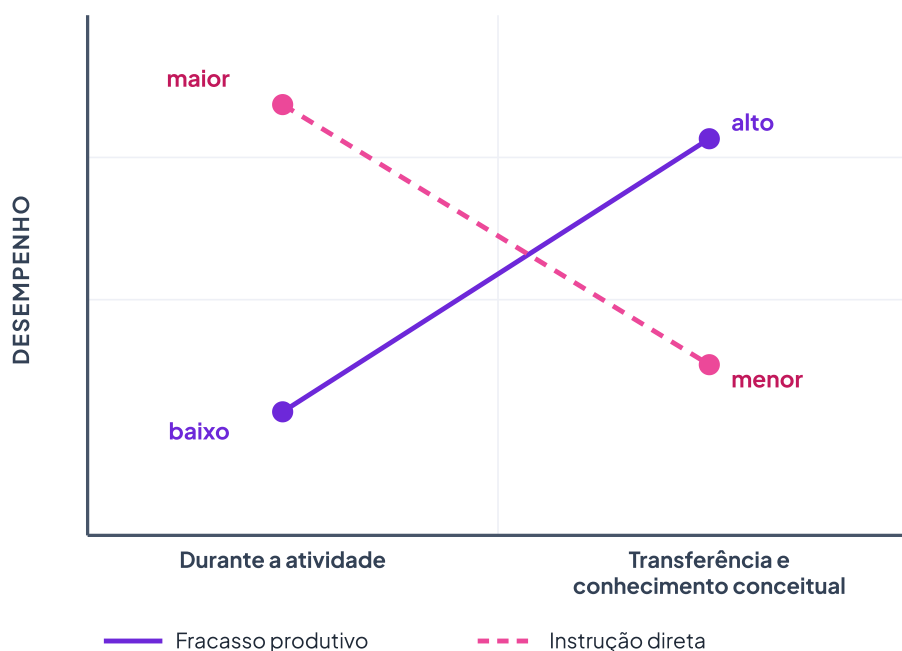
Às quatro dificuldades desejáveis soma-se um princípio convergente e particularmente fértil para o design das ferramentas: o **fracasso produtivo** (*productive failure*), formulado por Manu Kapur. A ideia, contraintuitiva, é que deixar o estudante lutar com um problema complexo antes de receber a instrução formal, ainda que não chegue à solução correta, produz aprendizagem mais profunda e mais transferível do que apresentar de imediato o método pronto. O fracasso inicial não é desperdício: ativa o conhecimento prévio, expõe as lacunas do aprendiz e prepara o terreno para que a instrução posterior faça sentido.

No estudo seminal de Kapur (2008), estudantes que enfrentaram problemas mal estruturados sem apoio inicial tiveram desempenho pior durante a atividade, mas superaram, em conhecimento conceitual e em transferência, os colegas que receberam instrução direta desde o começo. A síntese posterior de Kapur (2024) organiza o mecanismo em torno de quatro processos: ativação do conhecimento prévio, consciência das próprias lacunas, engajamento afetivo e montagem das peças conceituais. A partir deles, oferece princípios para desenhar deliberadamente bons impasses.

A implicação para as ferramentas do **IA para escolas** é direta: o valor pedagógico está em preservar a fase de tentativa autônoma. Entregar a solução pronta no primeiro pedido suprime justamente o fracasso produtivo que prepara a aprendizagem. Por isso, exigir uma tentativa registrada antes de liberar qualquer apoio, diretriz detalhada na seção sobre o princípio de design, não é um obstáculo burocrático, mas a operacionalização de um mecanismo de aprendizagem bem documentado.

Fracasso produtivo

Kapur: falhar antes de aprender constrói compreensão mais profunda



Tentar e falhar antes da instrução ativa o conhecimento prévio e expõe lacunas; base da diretriz de exigir tentativa.

FIGURA 10 Fracasso produtivo: desempenho durante a tarefa e em transferência (Kapur).

06

SEÇÃO SEIS

IA generativa, descarga cognitiva e dívida cognitiva

A IA generativa torna esse debate urgente porque produz uma **dissociação entre desempenho e aprendizagem**: durante a tarefa, com a IA ao lado, o estudante parece dominar o conteúdo; quando a IA é retirada, verifica-se que a aprendizagem não se consolidou. O mecanismo central é a *descarga cognitiva*: a delegação a uma ferramenta externa do processamento que deveria ocorrer na mente do aprendiz.

A evidência mais forte é o ensaio controlado randomizado de Bastani e colaboradores (*Generative AI Without Guardrails Can Harm Learning*), com cerca de mil estudantes do ensino médio. Comparando um “GPT Base” (tipo ChatGPT comum) a um “GPT Tutor” com salvaguardas pedagógicas, observou-se ganho expressivo de desempenho durante a prática (48% com o GPT Base; 127% com o GPT Tutor). Contudo, ao remover-se o acesso à IA em prova subsequente, os estudantes que haviam usado o GPT Base saíram-se **17% pior** que o grupo de controle que nunca teve acesso. As salvaguardas do GPT Tutor mitigaram em grande parte esse efeito negativo. É a demonstração direta do princípio: o desenho da ferramenta determina se a IA destrói ou preserva a aprendizagem.

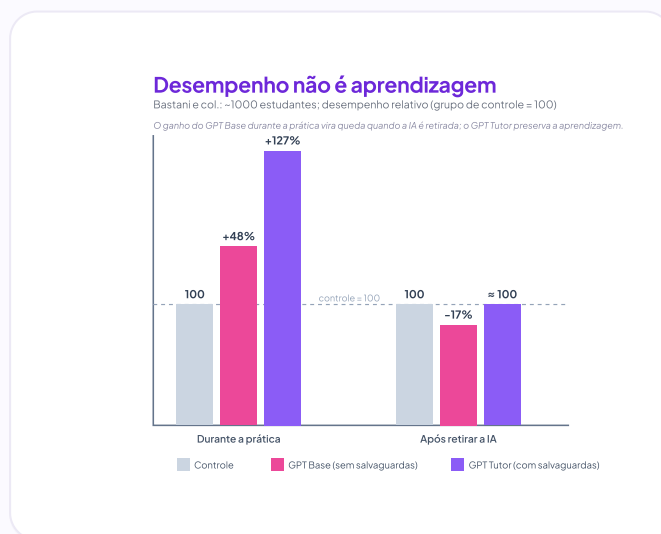


FIGURA 11 Desempenho não é aprendizagem: o experimento de Bastani e colaboradores.

Um segundo ensaio controlado randomizado reforça essa convergência. Barcaui (2025), com 120 universitários, comparou o uso livre do ChatGPT como apoio ao estudo a métodos tradicionais e mediu a retenção em um teste-surpresa aplicado 45 dias depois: o grupo que usara a IA reteve significativamente menos (57,5% de acertos) do que o grupo tradicional (68,5%), com tamanho de efeito médio ($d = 0,68$). A metáfora que dá título ao estudo, a IA como *muleta cognitiva* (cognitive crutch), converge com o achado de Bastani: o uso irrestrito alivia o esforço durante a tarefa, mas corrói os processos que sustentam a memória durável.

Em complemento, o estudo de Kosmyna e colaboradores no MIT Media Lab (*Your Brain on ChatGPT*) mediu, por eletroencefalografia, a atividade cerebral durante a escrita de redações: o grupo sem ferramentas exibiu as redes neurais mais fortes e distribuídas, e o grupo com LLM, a conectividade mais fraca, introduzindo a noção de *dívida cognitiva* (cognitive debt). Registre-se a cautela metodológica devida: trata-se de *preprint* não revisado por pares, com amostra pequena. Um comentário formal de Stanković e colaboradores (2026) pede interpretação mais conservadora, apontando o tamanho reduzido da amostra, problemas na análise de EEG, limites de reprodutibilidade, inconsistências no relato dos resultados e baixa transparência metodológica. Recomenda-se, portanto, utilizá-lo como ilustração sugestiva, e não como prova fechada. Um dado adicional do mesmo estudo aponta uma diretriz de sequenciamento: participantes que primeiro escreveram sem IA e só depois recorreram ao LLM exibiram conectividade neural mais forte do que os que fizeram o caminho inverso, indício de que a IA tende a ajudar mais quando entra após o esforço independente, e não antes dele.

A pesquisa qualitativa converge com esse quadro. Em um estudo de teoria fundamentada sobre programação assistida por IA, Liu e colaboradores (2026) distinguem o domínio do conteúdo do mero domínio da ferramenta e descrevem mecanismos recorrentes: a *ilusão de diálogo*, em que o estudante confunde a fluência da troca com compreensão; o padrão *confiar sem poder verificar*; e a *calibragem metacognitiva atenuada*: um descompasso entre a prontidão percebida e a capacidade real de resolver sozinho. Esse descompasso é o nome preciso da dissociação entre desempenho e aprendizagem que percorre todo este documento.

07

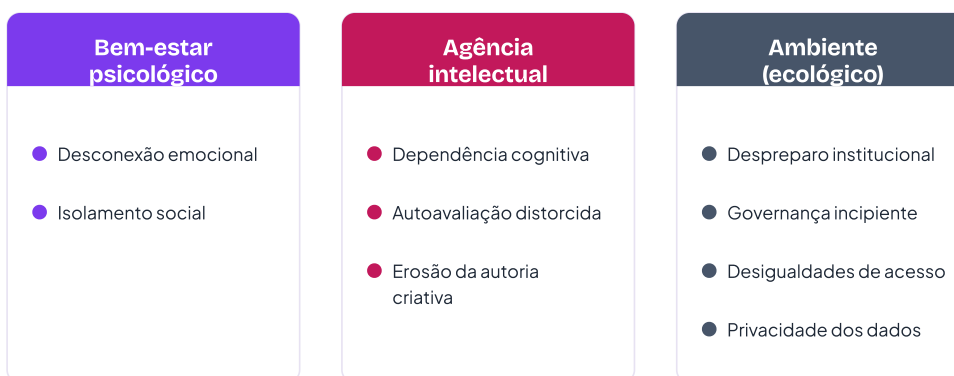
SEÇÃO SETE

Riscos da IA generativa na educação básica

Por se dirigir à educação básica, este referencial precisa considerar riscos que vão além do desempenho cognitivo imediato. A revisão de escopo de Tao e colaboradores (2026), que sintetiza estudos empíricos especificamente em contextos de educação básica (K-12), organiza esses riscos em três domínios. O primeiro é o do **bem-estar psicológico**, com indícios de desconexão emocional e isolamento social quando a interação com a máquina passa a substituir a interação humana. O segundo é o da **agência intelectual**, que reúne a dependência cognitiva, a erosão da autoria criativa e a autoavaliação distorcida, em que o estudante superestima o que de fato aprendeu. O terceiro é o **ecológico**, e abrange o despreparo institucional, a governança incipiente, as desigualdades de acesso e as questões de privacidade dos dados de crianças e adolescentes.

Riscos da IA generativa na escola básica

Tao e col.: revisão de escopo em educação básica: três domínios de risco



*A autoavaliação distorcida é a face metacognitiva do problema;
a erosão da autoria é o que está em jogo quando a ferramenta entrega o produto pronto.*

FIGURA 12 Riscos da IA generativa na escola básica (Tao e colaboradores).

Dois desses riscos dialogam diretamente com a fundamentação deste documento. A **autoavaliação distorcida** é a face metacognitiva do problema: a mesma fluência enganosa que faz a releitura parecer aprendizagem leva o estudante a confiar em uma competência que não construiu. E a **erosão da autoria criativa** é exatamente o que está em jogo quando a ferramenta entrega o produto final em vez do andaime. O princípio de retenção do produto cognitivo, exposto a seguir, é também uma resposta de design a esses riscos. A calibragem por perfil, tratada nas implicações, atende em particular à dimensão de equidade apontada por Tao.

Revisões recentes convergem: a delegação à IA tende a funcionar como *andaime produtivo* quando posicionada para amplificar competências já existentes, e como *descarga cognitiva nociva* quando substitui a aquisição de habilidades fundamentais. Em síntese, a IA generativa, sem mediação de design, remove justamente a fricção que constitui a dificuldade desejável.

08

SEÇÃO OITO

Retenção do produto cognitivo final por padrão

Da fundamentação acima decorre o princípio de design adotado: as ferramentas do **IA para escolas** retêm, por padrão, o produto cognitivo final, devolvendo ao estudante o *andaime* e não a solução. Operacionalmente, isso se traduz nas seguintes diretrizes:

8.1

Reter a resposta e usar questionamento socrático. Diante de um pedido de solução, a ferramenta responde com perguntas que ativam o conhecimento prévio do estudante, em vez de fornecer o resultado pronto.

8.2

Dicas graduadas com desvanecimento (fading). O apoio começa por perguntas, evolui para pistas e só então, quando necessário, para fragmentos parciais, e diminui à medida que a competência cresce, em coerência com a Zona de Desenvolvimento Proximal.

- 8.3 Exigir input e reflexão antes de liberar ajuda.** A ferramenta requer uma tentativa registrada e/ou uma pergunta bem formulada do estudante como condição para avançar, evitando o pedido de ajuda vago.
- 8.4 Oferecer estrutura, não solução.** Em redação, devolver critérios e diagnósticos (especificações do tema, recorte obrigatório, tangenciamento, foco distorcido), e não o texto reescrito; em problemas, devolver o caminho de raciocínio, não o gabarito.
- 8.5 Antecipar os contornos (workarounds).** Como proteções baseadas apenas em prompt são facilmente burladas (insistir “não sei” até obter a resposta), a retenção precisa ser implementada na lógica do produto: limites de uso, detecção de padrões de manipulação e exigência de etapas intermediárias.

Esse desenho mantém a IA como instrumento que amplifica o esforço produtivo (carga gerane) e elimina apenas o esforço improdutivo (carga extrínseca), preservando a dificuldade desejável que sustenta a aprendizagem. Cada diretriz é desenvolvida, com padrões de prompt e salvaguardas técnicas, no material complementar *Diretrizes de design de ferramentas com IA*.

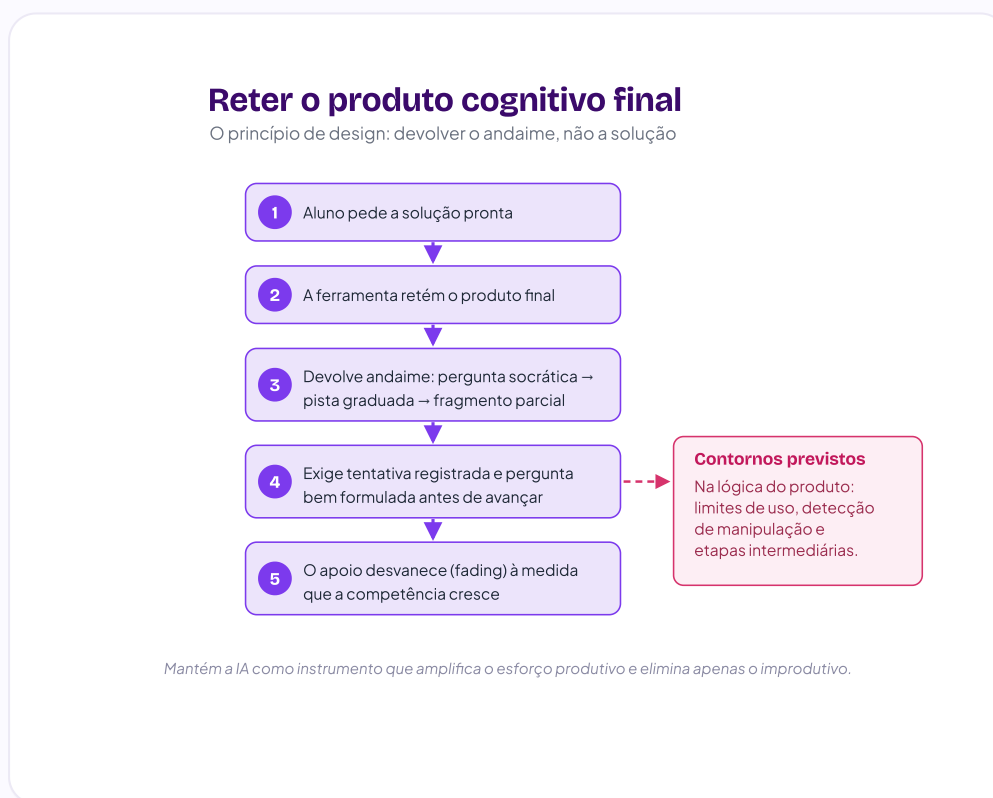
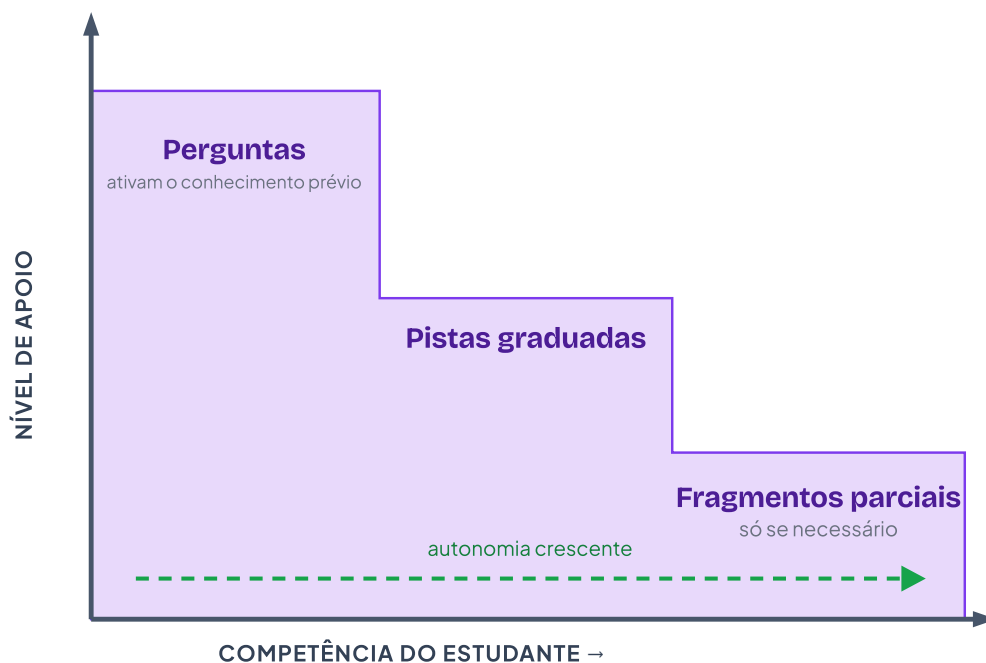


FIGURA 13 O princípio de retenção do produto cognitivo final.

Andaime com desvanecimento

A redução do apoio, coerente com a Zona de Desenvolvimento Proximal



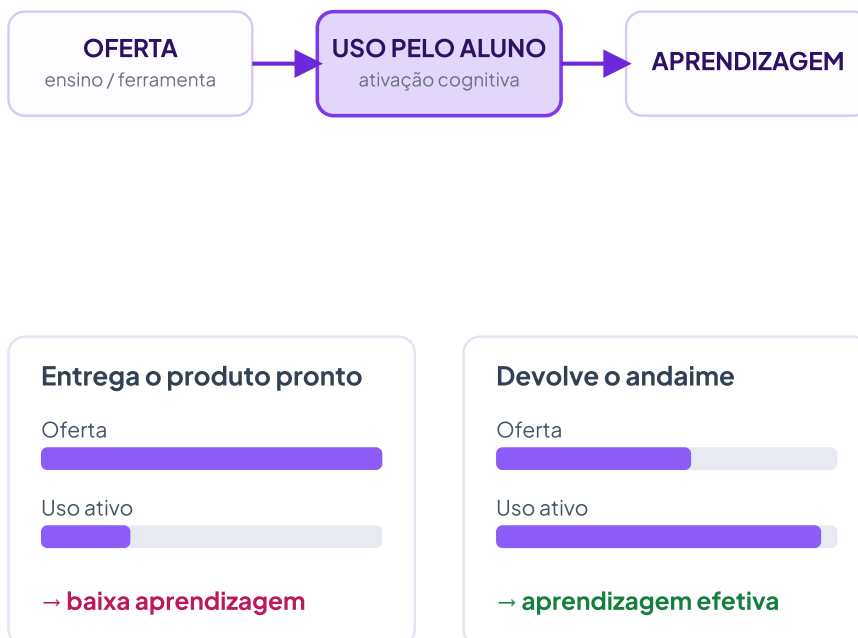
*Começa por perguntas, evolui para pistas e só então para fragmentos;
diminui conforme a competência do estudante cresce.*

FIGURA 14 Andaime com desvanecimento e a Zona de Desenvolvimento Proximal.

Esse princípio encontra respaldo na pesquisa sobre qualidade do ensino. O **modelo oferta-uso** (*Angebots-Nutzungs-Modell*), discutido por Hess e Lipowsky (2023), sustenta que nenhum ensino produz aprendizagem de forma direta: o ensino, assim como qualquer ferramenta, apenas oferece oportunidades, que precisam ser ativamente aproveitadas pelo estudante. O que converte a oferta em aprendizagem é o uso que o aluno faz dela, tendo a ativação cognitiva como um de seus motores centrais. Uma ferramenta que entrega o produto pronto oferece muito e exige pouco uso; uma que devolve perguntas, critérios e pistas graduadas mantém alta a exigência de uso ativo.

O modelo oferta-uso

Hess e Lipowsky: nenhum ensino produz aprendizagem de forma direta



O que converte a oferta em aprendizagem é o uso que o aluno faz dela, e é a esse uso que o design precisa obrigar.

FIGURA 15 O modelo oferta-uso (Hess e Lipowsky).

Esse desenho já é viável na prática. O **SocraticAI** (Sunil e Thakkar, 2025) é um tutor de programação que, em vez de proibir a IA, a submete a salvaguardas técnicas: exige perguntas bem formuladas, requer engajamento reflexivo, impõe limites diários de uso e ancora as respostas no material do curso. É a operacionalização concreta das diretrizes acima: questionamento socrático (8.1), exigência de pergunta e reflexão (8.3) e salvaguardas na lógica do produto, não apenas no prompt (8.5). Os resultados iniciais relatados são animadores: em duas a três semanas, os estudantes passaram de pedidos de ajuda vagos a uma decomposição estruturada dos problemas.

09

SEÇÃO NOVE

Equidade, reversão de expertise e ética

As implicações de adotar a fricção ótima como princípio extrapolam a cognição. O *framework* de Favero e colaboradores (2026) organiza os efeitos da IA na educação em quatro dimensões inter-relacionadas: cognição, agência, bem-estar emocional e ética. Além de preservar o esforço cognitivo, o design precisa sustentar a **agência** do estudante (sua capacidade de conduzir o próprio aprendizado, em vez de delegá-lo), cuidar do **bem-estar emocional** (evitando tanto a frustração da fricção excessiva quanto o desengajamento da facilidade) e responder a exigências éticas e de governança.

A calibragem da fricção é **individual**. A dificuldade desejável para um estudante pode ser uma barreira incapacitante para outro, fenômeno análogo ao *efeito de reversão de expertise* da teoria da carga cognitiva, segundo o qual apoios benéficos a iniciantes podem tornar-se redundantes ou prejudiciais a aprendizes avançados. Para estudantes com dificuldades específicas de aprendizagem, a literatura sugere que o objetivo não é maximizar a fricção, mas apoiar uma *descarga cognitiva* estratégica, ajustada às necessidades de cada perfil. A retenção do produto cognitivo deve, portanto, ser **adaptativa, e não uniforme**. A calibragem é também desenvolvimental: como evidencia o projeto LOEWE, uma dificuldade desejável para uma faixa etária pode não ser apropriada para outra, pois depende de pré-requisitos cognitivos que se constroem ao longo da escolarização.

O próprio efeito de reversão de expertise, desenvolvido por Sweller, Ayres e Kalyuga (2011), confirma que apoios benéficos a iniciantes podem tornar-se redundantes a avançados; e os demais efeitos da carga cognitiva oferecem critérios concretos para decidir o que a interface deve mostrar ou reter: exemplo trabalhado, atenção dividida e redundância.

Viés humano e viés algorítmico

A equidade depende também de reconhecer o viés. Toda criação humana tem viés e, por isso, também o tem toda IA, tanto quanto qualquer professor ou material didático. Vieses não se eliminam: reconhecem-se e administram-se. A correção assistida por IA (Seção onze) tem, como efeito colateral valioso, tornar visíveis alguns vieses humanos da correção, que então podem ser melhor administrados, desde que se administre igualmente o viés algorítmico, de natureza distinta, com origem, por exemplo, nos dados de treinamento. O ganho não é substituir o juízo humano por um árbitro neutro, mas calibrar um viés contra o outro, sob auditoria (Seção dez), preservando a primazia da relação humana apontada por Horvath. Cada viés a seguir é desenvolvido, com exemplos e o contraponto da IA, no material complementar *Viés humano e viés algorítmico*.

- . **Simpatia.** O professor tende a ser mais brando com o estudante de quem gosta e mais severo com os demais, o que é humano e não errado em si. A correção pela IA reduz a influência da simpatia pessoal.
- . **Rotulação.** Trocados os nomes, é provável que o estudante tido como “bom aluno” receba nota alta mesmo com desempenho fraco, e o tido como “mau aluno”, nota baixa mesmo com bom desempenho: é o efeito halo, ou viés de expectativa do professor. A IA ajuda a expor esse descompasso.
- . **Cansaço.** O rigor varia no início e no fim de uma maratona de correção, ou conforme o horário. A IA ajuda a evidenciar essas oscilações.
- . **Letra.** Letras mais legíveis tendem a ser mais bem avaliadas, porque o professor compreende melhor o que o estudante quis dizer. A IA decifra com mais facilidade as letras difíceis e penaliza menos por elas.
- . **Reclamação.** O estudante visto como “reclamão” tende a receber correção mais branda, para evitar o conflito. A IA ajuda a perceber esse desvio.

Há ainda um limite que nenhum design de ferramenta supera. Horvath (2024) argumenta que a relação empática entre professor e aluno, ancorada em processos socioafetivos que algoritmos não reproduzem, pesa mais sobre a aprendizagem do que a personalização, e que poupar o aluno do esforço de memorizar e elaborar pode inviabilizar o pensamento de ordem superior. A IA, portanto, deve ser pensada como apoio ao trabalho do professor, e não como seu substituto.

10

SEÇÃO DEZ

Da fricção ao sistema: a ontologia operacional educativa

As Seções oito e nove definem o *que* a ferramenta deve fazer (reter o produto cognitivo final) e *para quem* calibrar (equidade e reversão de expertise). Falta o *como institucional*: para que a fricção ótima não dependa do prompt nem da boa vontade do usuário, ela precisa estar inscrita na arquitetura do sistema. É o que a diretriz 8.5 antecipa, ao notar que proteções baseadas apenas em prompt são facilmente burladas, e o que o conceito de **ontologia operacional educativa** torna explícito: a especificação dos objetos de um domínio, das regras que os tornam confiáveis e das ações que neles se podem executar, em formato executável e auditável. Cada camada a seguir é desenvolvida, com exemplos e critérios, no material complementar *A ontologia operacional educativa*.

10.1

Objetos do domínio. Tarefa, rascunho, devolutiva, redação, critério de correção, material aprovado, cada um com as propriedades que o tornam confiável: ano/série, habilidade da BNCC e o projeto pedagógico da escola ao qual se vincula.

10.2

Regras de verdade operacional. É aqui que a Fricção Cognitiva Ótima deixa de ser recomendação e vira regra executável: o princípio de retenção do produto cognitivo (Seção oito), os quatro efeitos de Bjork (Seção quatro), o fracasso produtivo (Seção cinco) e a calibragem por prontidão (Seção nove) são codificados como o *que a IA nunca entrega*, não pedidos em linguagem natural, mas lógica do produto, imune ao *workaround* descrito em 8.5.

10.3

Ações governadas. Com permissão, validação e auditoria: o que exige aprovação humana, o que fica registrado e o que a IA pode sugerir, mas não executar sozinha.

O elemento humano é decisivo. Quem detém autoridade para definir essa verdade operacional, seja o coordenador ou o professor, é o *analista no circuito*, e não um revisor que apenas homologa. Isso reencontra Horvath (Seção nove), para quem o vínculo humano pesa mais que a personalização, e antecipa os seis hábitos da Seção doze: o **hábito 1** (*para terceirizar, é necessário ser capaz de avaliar*) é a condição de quem governa a IA; o **hábito 4** (*quanto maior o custo do erro, maior a fricção*) é, em rigor, um princípio das camadas de regras e ações: ação mais governada onde o risco é maior; e o **hábito 6** (*o que era ótimo ontem pode não ser hoje*) implica que essa ontologia operacional educativa é versionada, ajustando-se à medida que a expertise do estudante cresce.

Codificar o projeto pedagógico desta forma é o que transforma a fricção ótima de conselho em atributo do sistema, pré-condição concreta para o compromisso de design enunciado adiante.

11

SEÇÃO ONZE

Correção assistida por IA: configurações do humano no circuito

As decisões de design discutidas até aqui preservam o esforço do *estudante*. Resta o outro lado da relação pedagógica: o trabalho do professor. Para o docente, a correção é, em larga medida, carga extrínseca: esforço rotineiro que pouco constrói e que, delegado à IA, libera tempo para a carga gerane do ensino, isto é, desenhar tarefas com fricção ótima e cultivar o vínculo humano que Horvath mostra ser decisivo (Seção nove). Em todos os tipos a seguir, a IA corrige com base em critérios definidos pelos professores, ou propostos por ela e por eles revisados. Explicitar esses critérios é, em si, formativo: obriga a tornar pública a verdade operacional da correção (Seção dez), antes tácita no bom senso de cada docente. Os tipos diferenciam-se sobretudo pelo grau de supervisão humana, isto é, por onde se situa o *analista no circuito*.

11.1

Correção revisada. A IA corrige; os professores revisam todas as notas e comentários e editam o que julgarem necessário antes de liberá-los ao estudante. O docente emprega o mesmo tempo de uma correção manual, mas o estudante recebe um retorno substantivo sobre o seu desempenho, e não apenas uma nota e um comentário curto.

11.2

Correção compartilhada. A IA corrige e os professores revisam apenas uma parte das correções, segundo critérios definidos: por exemplo, todos os instrumentos a que a IA atribuiu nota baixa, ou apenas aqueles em que o estudante abre recurso. Reduz o tempo de correção e concentra a supervisão onde o custo do erro é maior: ação mais governada onde o risco é maior (Seção dez).

11.3

Correção autônoma. A IA corrige e os professores não revisam. Reserva-se a materiais que já não eram corrigidos, como a lição de casa, em que uma correção da IA é melhor que nenhuma; por não passar por revisão, costuma atribuir comentários, não notas. Não consome tempo docente e ainda alimenta o banco de dados, que o professor pode consultar para apoiar o estudante. Convém limitá-la a tarefas de baixo risco.

- 11.4 Metacorreção, parcial ou total.** O professor corrige manualmente e a IA transcreve as notas e amplia os comentários. Na versão total, em todas as turmas, não economiza tempo, mas eleva a qualidade do retorno; na parcial, a IA infere os critérios do professor em uma turma e os aplica às demais de forma autônoma, revisada ou compartilhada, economizando tempo sem perder qualidade. É a operacionalização da lógica de domínio (Seção dez): o conhecimento tácito do professor torna-se regra executável.
- 11.5 Correção gêmea (assessment twins).** Para conter o uso indevido de IA pelo estudante, Roe, Perkins e Giray (2026) propõem dividir a avaliação em duas fases: uma que impede o uso de IA e outra em que ele é plenamente liberado. Mede-se assim o que o estudante produz sozinho e o quanto aprimora com a IA, ou o quanto planeja com IA e executa sem ela. É, em si, um dispositivo de retenção do produto cognitivo (Seção oito): a fase sem IA preserva e avalia o esforço que ensina.
- 11.6 Correção cruzada.** Como o tempo de correção diminui, torna-se viável avaliar um mesmo conteúdo, competência ou habilidade com mais de um instrumento: o professor corrige autonomamente um instrumento, a IA corrige autonomamente outro equivalente e os resultados são comparados. A triangulação aumenta a validade da avaliação.
- 11.7 Metacorreção por pares.** A IA distribui, anônima e aleatoriamente, a cada estudante a correção do trabalho de um colega, e cabe à IA corrigir a correção feita pelos estudantes. Corrigir o trabalho do outro exige uma reflexão sofisticada sobre o que se aprende, um eco do conflito sociocognitivo de Doise e Mugny: a correção converte-se em aprendizagem.
- 11.8 Correção com reações.** O estudante produz uma primeira versão do instrumento de avaliação e recebe nota e comentários da IA; com base neles, elabora uma segunda versão, e tantas outras quantas o professor solicitar ou o estudante desejar. É especialmente útil para ensinar a elaboração de textos no Ensino Fundamental ou para introduzir a lógica da redação de vestibular no primeiro ano do Ensino Médio. A estratégia é anterior à IA, mas, além de poupar tempo na execução, a IA agiliza e facilita a análise dos dados da evolução individual e do grupo ao longo das reações.

Lidos em conjunto, esses tipos formam uma escala de supervisão, da autônoma à revisada, com a metacorreção como ponte entre o juízo humano e a execução pela máquina, em coerência com a Escala de Avaliação com IA, apresentada a seguir. Qualquer que seja a configuração, ela só é confiável sob a governança da Seção dez: permissão, validação e auditoria. O ativo central é o tempo que a correção devolve ao professor, que pode ser reinvestido no acompanhamento dos dados (em um painel e na consulta em linguagem natural ao banco de dados das correções) e, sobretudo, no que a IA não faz: desenhar a fricção e sustentar a relação. Cada tipo é desenvolvido em detalhe, com fluxo de trabalho, critérios de decisão e exemplos por etapa de escolaridade, no material complementar *Tipos de correção com IA*.

A Escala de Avaliação com IA

Os dois eixos deste documento, a autonomia do professor na correção (esta seção) e a autonomia do estudante diante do produto cognitivo (Seção oito), convergem em um único instrumento: a Escala de Avaliação com IA, adaptada do *AI Assessment Scale* de Perkins, Roe e Furze. Ela cruza cinco níveis de uso de IA pelo professor (A–E) com cinco níveis de uso pelo estudante (1–5), gerando vinte e cinco combinações, cada uma apontando o tipo de avaliação mais adequado (Figura 16).

O anel do professor (A–E) recupera a escala de supervisão da Seção onze: A, sem IA; B, IA no planejamento, não na correção; C, correção com revisão humana obrigatória; D, correção autônoma; E, co-desenho da avaliação com os estudantes. O anel do estudante (1–5) gradua a retenção do produto cognitivo (Seção oito): 1, sem IA; 2, IA apenas para ideias; 3, IA só na edição, com entrega do rascunho original; 4, IA em partes citadas; 5, uso livre, em que a avaliação recai sobre voz autoral, originalidade e metacognição. Quanto mais alto o nível, mais fricção é preciso preservar deliberadamente para que a autonomia não se converta em dependência.

PROFESSOR ↓	ALUNO →				
	1 Sem IA	2 Ideias	3 Edição	4 Execução	5 IA total
A Sem IA	A1	A2	A3	A4	A5
B Planejamento	B1	B2	B3	B4	B5
C Colaboração	C1	C2	C3	C4	C5
D IA total	D1	D2	D3	D4	D5
E Exploração	E1	E2	E3	E4	E5

FIGURA 16 A Escala de Avaliação com IA: cinco níveis de uso pelo professor (A–E) cruzam cinco níveis de uso pelo estudante (1–5). Adaptada do AIAS (Perkins, Roe e Furze), licença CC BY-NC-SA 4.0.

12 SEÇÃO DOZE

Seis hábitos mentais para educadores na era da IA

As onze seções anteriores estabelecem o referencial; resta condensá-lo em heurísticas que orientem decisões concretas de adoção, configuração e governança da IA. Os seis enunciados a seguir são adaptados do quadro *Seis hábitos mentais para educadores na era da IA*, de Guilherme Cintra, aqui reancorados no aparato teórico deste documento.

- 1 **Para terceirizar, é necessário ser capaz de avaliar.** Delegar uma tarefa cognitiva à IA só é seguro quando o sujeito preserva competência para julgar o resultado. Liga-se ao risco de *confiar sem poder verificar* (Seção sete) e à dimensão da agência no *framework* de Favero e colaboradores (Seção nove): sem capacidade avaliativa, a terceirização vira dependência, não autonomia.
- 2 **Para aprender, é necessário fazer.** Reafirma o modelo oferta-uso de Hess e Lipowsky (Seção oito): nenhuma oferta, de ensino ou de ferramenta, produz aprendizagem sem uso ativo. É a tradução, em hábito mental, da distinção entre carga *germane* e *extrínseca* (Seção dois) e do efeito de geração (Seção quatro).
- 3 **Para formar um ser humano, é necessário preservar o que o constitui.** Eleva a decisão de design ao plano formativo: automatizar esforço, reflexão e autonomia corrói as próprias capacidades que a educação busca cultivar. É o que o compromisso de design desenvolve a seguir, e ecoa Horvath (Seção nove) sobre a primazia do vínculo humano sobre a personalização.
- 4 **Quanto maior o custo do erro, maior a fricção necessária.** Critério para decidir quando reter o produto cognitivo (Seção oito): a dose de dificuldade desejável deve crescer com o que está em jogo na aprendizagem. Articula a calibragem individual (Seção nove) com a noção de ponto ótimo (Seção dois) e com a camada de ação da ontologia operacional educativa (Seção dez).
- 5 **Fluência é diferente de acerto.** Recupera o cerne metacognitivo das dificuldades desejáveis (Seções dois e quatro): a sensação de domínio produzida pela releitura passiva, ou pela conversa fluente com a IA, não é evidência de aprendizagem. Distingue desempenho de aprendizagem e fundamenta a desconfiança da *ilusão de diálogo* (Seção sete).
- 6 **O que era ótimo ontem pode não ser hoje.** Consagra o caráter dinâmico do ponto ótimo: o equilíbrio de Csikszentmihalyi desloca-se à medida que a habilidade cresce (Seção dois), o que fundamenta o desvanecimento do apoio (Seção oito) e o efeito de reversão de expertise (Seção nove). A fricção ótima é um alvo móvel, não um parâmetro fixo.

COMPROMISSO DE DESIGN

Quando sistemas de IA priorizam eficiência, fluência e automação em detrimento de esforço, reflexão e autonomia, arriscam minar as próprias capacidades que a educação busca cultivar. O princípio de **retenção do produto cognitivo** é, nesse sentido, também um compromisso com a formação da autonomia intelectual do estudante.

BIBLIOGRAFIA

Referências

- BARCAUI, A.** ChatGPT as a cognitive crutch: evidence from a randomized controlled trial on knowledge retention. *Computers and Education: Artificial Intelligence*, 2025.
- BASTANI, H. et al.** Generative AI without guardrails can harm learning. *SSRN*, 2024.
- BJORK, E. L.; BJORK, R. A.** Making things hard on yourself, but in a good way: creating desirable difficulties to enhance learning. In: GERNSBACHER, M. A. et al. (Ed.). *Psychology and the real world*. New York: Worth, 2011. p. 56–64.
- BJORK, R. A.** Memory and metamemory considerations in the training of human beings. In: METCALFE, J.; SHIMAMURA, A. (Ed.). *Metacognition: knowing about knowing*. Cambridge: MIT Press, 1994.
- BORROMEO FERRI, R.; PEDE, S.; LIPOWSKY, F.** Auswirkungen verschachtelten Lernens auf das prozedurale und konzeptuelle Wissen von Lernenden über Zuordnungen. *J Math Didakt*, v. 42, p. 1–23, 2021.
- CINTRA, G.** Seis hábitos mentais para educadores na era da IA: Fricção Cognitiva Ótima. LinkedIn, 2026. Disponível em: [ver publicação no LinkedIn](#).
- COOPER, A.** The inmates are running the asylum: why high tech products drive us crazy and how to restore the sanity. Indianapolis: Sams, 1999.
- CSIKSZENTMIHALYI, M.** Flow: the psychology of optimal experience. New York: Harper & Row, 1990.
- DOISE, W.; MUGNY, G.** Le développement social de l'intelligence. Paris: InterÉditions, 1981.

- EDUTECH WIKI.** Friction cognitive, émotions et choix de la modalité de travail: individuelle ou collaborative. Université de Genève.
- FAVERO et al.** AI in education beyond learning outcomes: cognition, agency, emotion, and ethics. *arXiv:2602.04598*, 2026.
- FURZE, L.; PERKINS, M.; ROE, J.** The AI Assessment Scale (AIAS) in Australian K–12 Education. *Teachers' Frontiers*, v. 1, n. 1, 2024.
- HESS, M.; LIPOWSKY, F.** Unterrichtsqualität und das Lernen der Schülerinnen und Schüler. In: ROTHLAND, M. (Hrsg.). *Beruf Lehrer:in — ein Studienbuch*. 2. Aufl. Münster: Waxmann, 2023. S. 171–194.
- HORVATH, J. C.** 3 critical problems Gen AI poses for learning. *Harvard Business Publishing*, 2024.
- KAPUR, M.** Productive failure. *Cognition and Instruction*, v. 26, n. 3, p. 379–424, 2008.
- KAPUR, M.** Productive failure: unlocking deeper learning through the science of failing. Wiley and Sons, 2024.
- KOSMYNA, N. et al.** Your brain on ChatGPT: accumulation of cognitive debt when using an AI assistant for essay writing task. *arXiv:2506.08872*, 2025.
- LIPOWSKY, F. et al.** Wünschenswerte Erschwernisse beim Lernen. *Schulpädagogik Heute*, v. 6, n. 11, 2015.
- LIU, D.; FAN, G.; PAN, L.** Tool, tutor, or crutch? A grounded theory of cognitive scaffolding and offloading in AI-assisted programming education. *International Journal of STEM Education*, v. 12, 2026.
- LOEWE-Forschungsprojekt.** “Wünschenswerte Erschwernisse beim Lernen” (2015–2018). Universität Kassel.
- PERKINS, M.; ROE, J.; FURZE, L.** Reimagining the Artificial Intelligence Assessment Scale (AIAS): a refined framework for educational assessment. *Journal of University Teaching and Learning Practice*, v. 22, n. 7, 2025.
- PIAGET, J.** L'équilibration des structures cognitives: problème central du développement. Paris: PUF, 1975.
- ROE, J.; PERKINS, M.; GIRAY, L.** Assessment twins: an approach for strengthening assessment validity in the age of generative AI. *Journal of Applied Learning & Teaching*, v. 9, n. 2, 2026.

- STANKOVIĆ, M. et al.** Comment on: Your brain on ChatGPT. *arXiv:2601.00856*, 2026.
- SUNIL, K.; THAKKAR, A.** SocraticAI: transforming LLMs into guided CS tutors through scaffolded interaction. *arXiv:2512.03501*, 2025.
- SWELLER, J.** Cognitive load during problem solving: effects on learning. *Cognitive Science*, v. 12, n. 2, p. 257–285, 1988.
- SWELLER, J.; AYRES, P.; KALYUGA, S.** Cognitive load theory. New York: Springer, 2011.
- TAO, S.; LAN, M.; WANG, M.; LI, H.** Potential risks of generative artificial intelligence integration into K-12 education: a scoping review. *Computers and Education: Artificial Intelligence*, 2026.
- WILSON, R. C. et al.** The eighty five percent rule for optimal learning. *Nature Communications*, v. 10, art. 4646, 2019.