

Correção em Escala — Um Ensaio Aleatorizado de IA vs. Corretores Humanos em Quatro Programas de Ciência Política

AI in Education · Boa prática · Estados Unidos da América

Pontuação de qualidade: 67/100

Força da evidência: Ensaio controlado aleatorizado

Sumário

Um ensaio clínico aleatorizado revisado por pares em quatro universidades atribuiu aleatoriamente 3.080 respostas curtas a correção por GPT-4 ou por humanos. As notas foram estatisticamente indistinguíveis, e a IA permitiu que docentes dessem a turmas grandes um feedback equivalente ao de turmas um terço do tamanho.

Dimensões-chave

Dimensão	Pontuação
Evidence of learning impact	1/3
Equity & inclusion	1/3
Transparency, ethics & data protection	2/3
Teacher capability & pedagogy	3/3
Scalability & sustainability	3/3

Avaliado por C-NAPSE

Fontes de dados

[C-NAPSE web research](#) — acedido a 2026-08-20

[University of Houston, with University of South Carolina, University of Michigan & Academia Sinica \(Taiwan\)](#) — acedido a 2026-08-20

[PMC full text \(Heinrich et al., PLoS One 2025\)](#) — acedido a 2026-08-20

[Semantic Scholar record](#) — acedido a 2026-08-20

Citar — Evidoria. *Correção em Escala — Um Ensaio Aleatorizado de IA vs. Corretores Humanos em Quatro Programas de Ciência Política*. ID persistente: 49bab669-e535-477b-b874-00b96311cc56.