

Corrección a Gran Escala — Un Ensayo Aleatorizado de IA frente a Correctores Humanos en Cuatro Programas de Ciencia Política

AI in Education · Buena práctica · Estados Unidos

Puntuación de calidad: 67/100

Solidez de la evidencia: Ensayo controlado aleatorizado

Resumen

Un ensayo aleatorizado revisado por pares en cuatro universidades asignó al azar 3.080 respuestas cortas a corrección por GPT-4 o por humanos. Las notas fueron estadísticamente indistinguibles, y la IA permitió a los docentes dar a clases grandes una retroalimentación equivalente a la de clases un tercio de su tamaño.

Dimensiones clave

Dimensión	Puntuación
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

Evaluable por C-NAPSE

Fuentes de datos

[C-NAPSE web research](#) — consultado el 2026-08-20

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

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

[Semantic Scholar record](#) — consultado el 2026-08-20

Citar — Evidoria. *Corrección a Gran Escala — Un Ensayo Aleatorizado de IA frente a Correctores Humanos en Cuatro Programas de Ciencia Política*. ID persistente: 49bab669-e535-477b-b874-00b96311cc56.