

Correction à Grande Échelle — Un Essai Randomisé Comparant l'IA aux Correcteurs Humains dans Quatre Cours de Science Politique

AI in Education · Bonne pratique · États-Unis d'Amérique

Score de qualité : 67/100

Niveau de preuve : Essai contrôlé randomisé

Résumé

Un essai randomisé évalué par des pairs dans quatre universités a assigné aléatoirement 3 080 réponses courtes à une correction par GPT-4 ou humaine. Les notes étaient statistiquement indiscernables, et l'IA a permis aux enseignants d'offrir à de grandes classes un retour équivalent à celui de classes trois fois plus petites.

Dimensions clés

Dimension	Score
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

Évalué par C-NAPSE

Sources de données

[C-NAPSE web research](#) — consulté le 2026-08-20

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

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

[Semantic Scholar record](#) — consulté le 2026-08-20

Citer — Evidoria. *Correction à Grande Échelle — Un Essai Randomisé Comparant l'IA aux Correcteurs Humains dans Quatre Cours de Science Politique*. ID persistant : 49bab669-e535-477b-b874-00b96311cc56.