

Ensaio AHEAD — Perguntas de Exame Médico Geradas por IA Igualam as Escritas por Humanos na Maioria das Medidas, com uma Lacuna de Equidade Assinalada

AI in Education · Boa prática · Canadá

Pontuação de qualidade: 47/100

Força da evidência: Moderada

Sumário

Um ECR cego com 258 estudantes de medicina do primeiro ano da UBC constatou que perguntas de exame geradas por ChatGPT-4/Gemini foram produzidas 5,6 vezes mais depressa do que as escritas por estudantes, com notas e aceitação semelhantes, embora os itens humanos discriminassem ligeiramente melhor e tenha surgido uma diferença por género apenas no exame de IA.

Dimensões-chave

Dimensão	Pontuação
Evidence of learning impact	1/3
Equity & inclusion	2/3
Transparency, ethics & data protection	3/3
Teacher capability & pedagogy	0/3
Scalability & sustainability	1/3

Avaliado por C-NAPSE

Fontes de dados

[C-NAPSE web research](#) — acedido a 2026-09-06

[University of British Columbia Faculty of Medicine](#) — acedido a 2026-09-06

[Zenodo - sample question dataset](#) — acedido a 2026-09-06

[BMC Medical Education - full text](#) — acedido a 2026-09-06

[ClinicalTrials.gov registration NCT07481162](#) — acedido a 2026-09-06

Citar — Evidoria. *Ensaio AHEAD — Perguntas de Exame Médico Geradas por IA Igualam as Escritas por Humanos na Maioria das Medidas, com uma Lacuna de Equidade Assinalada*. ID persistente: 6e5d3873-5ba1-4b93-b6a1-329127631218.