

Avaliação Crítica de Revisões Sistemáticas: Investigadores versus Inteligência Artificial

Carla Duarte de Mendonça^{1,2,3}, Inês Cardoso¹, Bruno Rosa^{1,2,3}, Susana Dias^{1,2}, Ruben Pereira^{1,2}, António Mata^{1,2,3}

¹ Faculdade de Medicina Dentária (FMDUL), Universidade de Lisboa, Lisboa, Portugal;
² Oral Biology and Biochemistry Research Group at FMDPhys-FCT UIDB/04559/2020 da FMDUL, Lisboa, Portugal;
³ Center for Evidence-Based Dental Medicine da FMDUL, Lisboa, Portugal.



INTRODUÇÃO E OBJETIVO

A avaliação crítica de revisões sistemáticas é essencial para a Medicina Dentária Baseada na Evidência. Contudo, a aplicação de instrumentos como o *Risk Of Bias in Systematic Reviews (ROBIS)*¹ exige tempo, formação metodológica e julgamento especializado, tornando o processo de avaliação crítica moroso e dependente da interpretação humana²⁻⁹. Os modelos de linguagem de grande escala (LLMs), devido ao seu fácil acesso e à capacidade de processar grandes volumes e diferentes formatos de informação científica, poderão apoiar esta fase da elaboração de sínteses de evidência^{10,11}. No entanto, a sua fiabilidade na aplicação do ROBIS⁵ e a concordância com investigadores calibrados permanece insufficientemente esclarecidas.

O objetivo deste estudo foi avaliar a concordância entre investigadores calibrados e um modelo de linguagem de grande escala (LLM) na avaliação crítica de revisões sistemáticas em Medicina Dentária, utilizando o instrumento ROBIS.

MÉTODOS



RESULTADOS

PERFIL DE RESPOSTAS "YES"	TESTE DE HOMOGENEIDADE MARGINAL	REPRODUTIBILIDADE ChatGPT Pro 5.4 Kappa ponderado (wk) Teste-reteste	CONCORDÂNCIA ChatGPT Pro 5.4 vs. Gold Standard Kappa ponderado (wk)	CLASSIFICAÇÃO FINAL DO RISCO DE VIÉS Kappa ponderado (wk)
60,9% Gold standard	12 em 29 % concordância (Pa) 45,4% p<0,001	0,81 IC 95%: 0,78-0,83 p<0,001	0,43 IC 95%: 0,39-0,47	0,08 IC 95%: -0,08-0,25
40,9% LLM	p<0,05, diferenças sistemáticas entre as classificações do ChatGPT e do gold standard	Concordância Excelente	Concordância Moderada	Concordância Ligeira

DISCUSSÃO

- O modelo LLM não reproduz de forma consistente o raciocínio dos investigadores ao longo dos diferentes domínios do ROBIS.
- A reprodutibilidade interna do LLM, contrasta com níveis moderados de concordância com o *gold standard*. Estas divergências resultam de diferenças na interpretação dos critérios metodológicos, e não de inconsistência interna dos modelos.
- No entanto, estas divergências na forma como o modelo interpreta e aplica critérios metodológicos sugerem que, na sua configuração atual, não deverão substituir a avaliação humana, constituindo antes uma ferramenta de apoio complementar em tarefas de triagem inicial sob supervisão humana.
- O tempo reduzido de aplicação da ferramenta ROBIS pelo LLM reforça o seu potencial contributo para a eficiência do processo, ainda que condicionado pelas limitações de fiabilidade identificadas.

CONCLUSÃO

Na avaliação metodológica de revisões sistemáticas através do instrumento ROBIS, o LLM demonstrou elevada reprodutibilidade nas classificações atribuídas, porém apenas uma concordância moderada com investigadores calibrados.

1. Whiting P, Savovic J, Higgins JP, Caldwell DM, Roever C, Shea B, et al. ROBIS: A new tool to assess risk of bias in systematic reviews was developed. J Clin Epidemiol. 2016;69:225-34. 2. Sackel D, Rothenberg WM, Gray JA, Haynes RB, Richardson WS. Evidence-based medicine: what it is and what it isn't. BMJ. 1996;312(7023):71-2. 3. Samal A, Butler D, A Raju A. ADA Code, Division of S, Journal of the American Dental Association. Evidence-based dentistry in clinical practice. J Am Dent Assoc. 2004;135(1):78-82. 4. Straus SE, Glasziou P, Richardson WS, Haynes RB. Evidence-based Medicine: How to Practice and Teach EBM. 5th ed. Elsevier; 2015. 5. Greenhalgh T, Peacock R, Macleod M, Garsden R, Vandenbroucke JP, Altman DG, et al. Evidence-based medicine: how to practice and teach it. 2nd ed. London: Churchill Livingstone; 2005. 6. Guyatt GH, Oxman AD, Vist GE, Rivas R, Falck-Ytter Y, Alonso-Coello P, et al. GRADE: an emerging consensus on rating quality of evidence and strength of recommendations. BMJ. 2008;336(7690):924-6. 7. Shoaib M, Greenhalgh T, Peacock R, Macleod M, Garsden R, Vandenbroucke JP, Altman DG, et al. Evidence-based medicine: how to practice and teach it. 2nd ed. London: Churchill Livingstone; 2005. 8. Guyatt GH, Oxman AD, Vist GE, Rivas R, Falck-Ytter Y, Alonso-Coello P, et al. GRADE: an emerging consensus on rating quality of evidence and strength of recommendations. BMJ. 2008;336(7690):924-6. 9. Shoaib M, Greenhalgh T, Peacock R, Macleod M, Garsden R, Vandenbroucke JP, Altman DG, et al. Evidence-based medicine: how to practice and teach it. 2nd ed. London: Churchill Livingstone; 2005. 10. Sackel D, Rothenberg WM, Gray JA, Haynes RB, Richardson WS. Evidence-based medicine: what it is and what it isn't. BMJ. 1996;312(7023):71-2. 11. Samal A, Butler D, A Raju A. ADA Code, Division of S, Journal of the American Dental Association. Evidence-based dentistry in clinical practice. J Am Dent Assoc. 2004;135(1):78-82.