

LETTER TO THE EDITOR

Correspondence “Exploring the potential of artificial intelligence in traumatology: Conversational answers to specific questions”



Correspondencia “Explorando el potencial de la inteligencia artificial en traumatología: respuestas conversacionales a preguntas específicas”

Dear Editor,

We would like to discuss on the publication “Exploring the potential of artificial intelligence in traumatology: Conversational answers to specific questions”.¹ When it came to answering medical queries, ChatGPT had the best accuracy (72.81%), followed by Perplexity (67.54%) and BARD (60.53%) in a study that compared the three chatbot models. Although BARD offered the most accessible and thorough answers, in 14% of the questions all three models failed at the same time. Conversational bots’ inability to effectively handle medical queries was demonstrated by the identification of errors in information and logical reasoning in the responses.

The study found that one of the chatbot models’ weaknesses was its dependence on accuracy as the main performance indicator. In assessing the efficacy of the bots, readability, logical reasoning, and the utilization of outside data should all be taken into account in addition to accuracy. Furthermore, the evaluation’s breadth may have been restricted by the technique employed to evaluate the chatbots, which was limited to responding to particular medical queries rather than having a more comprehensive dialogue or offering context-based answers.

Further study in this field may focus on creating better chatbot models that give precedence to external information retrieval and logical reasoning in their responses. Furthermore, investigating methods to incorporate human supervision and input into the chatbot exchanges may assist reduce errors and guarantee the accuracy of the data returned. In order to evaluate the continued advancement and efficacy of conversational bots in the healthcare industry, longitudinal studies may also be carried out, incorporating user feedback and fine-tuning the models according to actual usage scenarios.

Level of evidence

Level of evidence V.

Ethics of approval statement

Not applicable.

Funding statement

There is no funding.

Authors’ contributions

HP 50% ideas, writing, analyzing, approval.
VW 50% ideas, supervision, approval.

Patient consent statement

Not applicable.

Permission to reproduce material from other sources

Not applicable.

Clinical trial registration

Not applicable.

Conflict of interest

The authors declare no conflict of interest.

Data availability statement

There is no new data generated.

Reference

1. Del Rey FC, Arias MC. Exploring the potential of artificial intelligence in traumatology: conversational answers to

specific questions. *Rev Esp Cir Ortop Traumatol.* 2024. S1888-4415(24)00086-9.

H. Daungsupawong^{a,*}, V. Wiwanitkit^b

^a *Private Academic Consultant, Phonhong, Lao Democratic People's Republic*

^b *Department of Research Analytics, Saveetha Dental College and Hospitals, Saveetha Institute of Medical and Technical Sciences Saveetha University, Chennai, India*

* Corresponding author.

E-mail address: hinpetchdaung@gmail.com (H. Daungsupawong).