

Artificial Intelligence in Periprosthetic Joint Infection: Are All Large Language Models Equally Reliable?

R. Rojas-Sayol¹, O. Pujol², R. Ferrer³, M. Vicente², M. Monfort-Mira⁴, M. Sabater-Martos⁵, D. Pérez-Prieto⁶, A. Coelho¹

¹ Hospital Universitari Parc Taulí; ² Hospital Universitari Vall d'Hebron; ³ Hospital Sant Rafael; ⁴ Clínic Barcelona Hospital;

⁵ Hospital Universitari Joan XXIII; ⁶ Hospital del Mar. Orthopaedic Surgery Departments, Spain.

Methods

Introduction & Aim

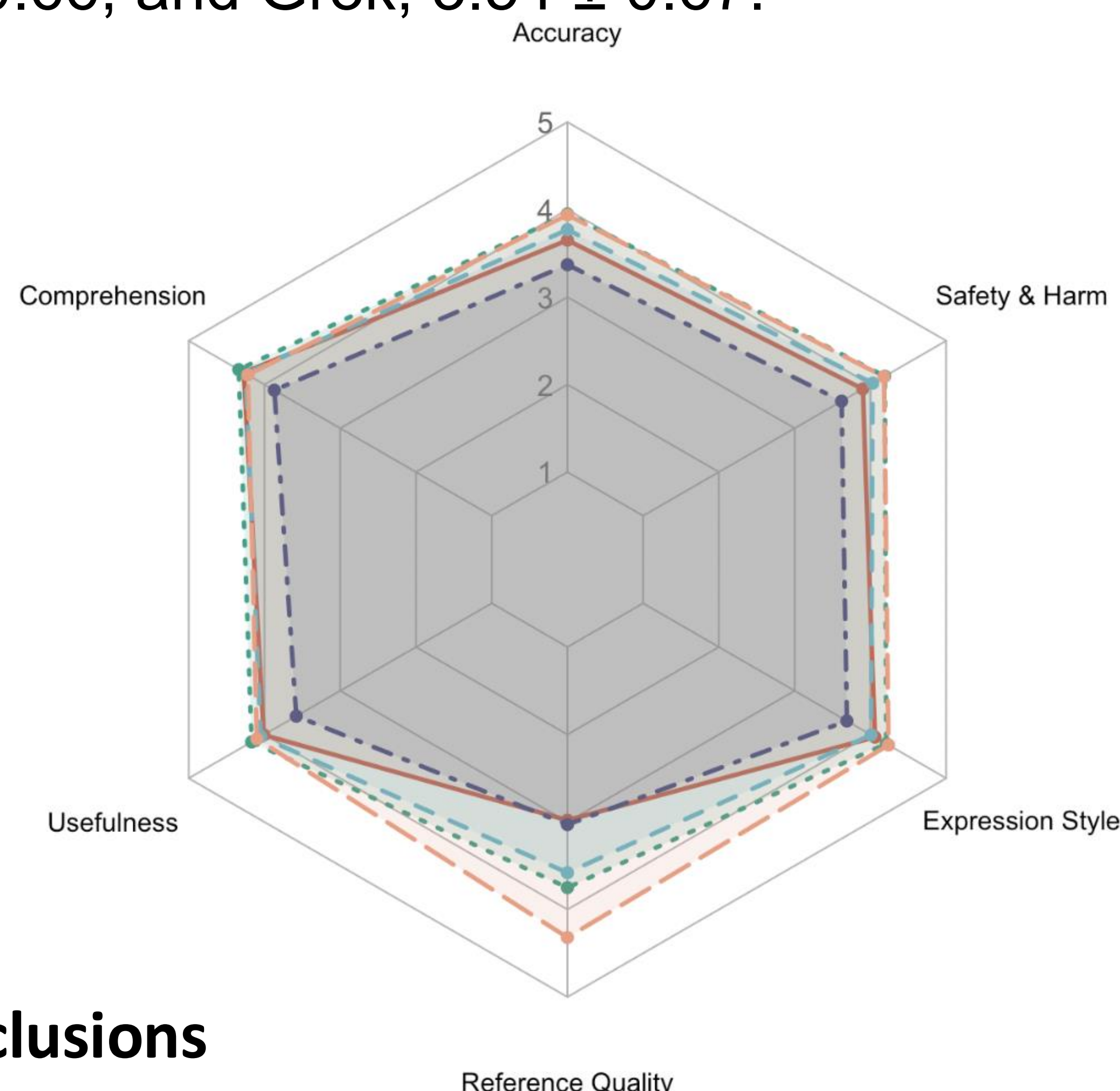
- Periprosthetic joint infection (PJI) remains a complex, high-impact complication requiring accurate and evidence-based clinical decisions.
- Large language models (LLMs) are increasingly used in medicine, but their reliability for expert, consensus-based PJI questions remains uncertain.
- Aim: to evaluate and compare five LLMs using clinical questions derived from the 2025 International Consensus Meeting (ICM).

- Cross-sectional comparative study using 20 clinically relevant ICM 2025 questions covering PJI prevention, diagnosis and management.
- Five LLMs were assessed: ChatGPT, Claude, Grok, Gemini and OpenEvidence.
- Five orthopaedic infection specialists independently performed blinded evaluations. Responses were scored from 1 to 5 across six adapted QUEST domains: accuracy, comprehensiveness, clinical utility, reference quality, writing style and safety.
- Statistical analysis: Friedman tests with post-hoc paired comparisons and mixed-effects models

Results

20 questions • 5 LLMs • 5 reviewers

- Significant differences in overall performance were observed ($p < 0.001$). Mean $\pm$ SD QUEST scores: OpenEvidence; 4.17 ± 0.56 , ChatGPT; 4.10 ± 0.58 , Claude; 3.95 ± 0.57 , Gemini; 3.87 ± 0.66 , and Grok; 3.54 ± 0.67 .



- ChatGPT achieved the highest Accuracy (3.95) and Comprehensiveness (4.34).
- OpenEvidence achieved the highest Reference Quality (4.32) and received the most best-answer nominations (31%).
- Overall inter-rater reliability was moderate (ICC = 0.635). Differences were especially relevant in Accuracy and Reference Quality

	LLM-A	LLM-B	LLM-C	LLM-D	LLM-E
Safety & Harm	3.9	4.03	4.19	3.62	4.18
Expression Style	4.06	4.01	4.21	3.69	4.24
Reference Quality	2.98	3.58	3.75	3.03	4.32
Usefulness	4.01	4.05	4.17	3.58	4.1
Comprehension	4.27	4.25	4.34	3.87	4.22
Accuracy	3.65	3.77	3.95	3.37	3.94

Conclusions

- LLMs demonstrated strong performance in answering expert-level PJI questions, although meaningful differences were observed between models, particularly in accuracy and reference quality. This variability highlights limitations in high-stakes clinical settings.
- LLMs may serve as useful adjuncts in orthopaedic infection practice, but their outputs require careful interpretation and should not replace expert clinical judgement.