

ASSESSING THE REASONING CAPABILITIES OF STATE-OF-THE-ART LARGE LANGUAGE MODELS IN COMPLEX SURGICAL DECISION-MAKING FOLLOWING ADMISSION FOR INGUINAL HERNIA REPAIR

Lee Ying, MD, PhD¹; Elizabeth Schaefer²; Benlu Wang²; Xinyue Li²; Samuel Butensky, MD¹; Emily Flom, MD¹; Randal Zhou, MD, FACS, FASMB¹; Andrew Duffy, MD, FACS, FASMB¹; John Morton, MD, FACS, FASMB¹; Karen Gibbs, MD, FACS, FASMB¹; Eric Schneider, PhD¹; Arman Cohan, PhD²; Chin Siang Ong, MBBS, PhD, MPH¹; ¹Yale School of Medicine; ²Yale University

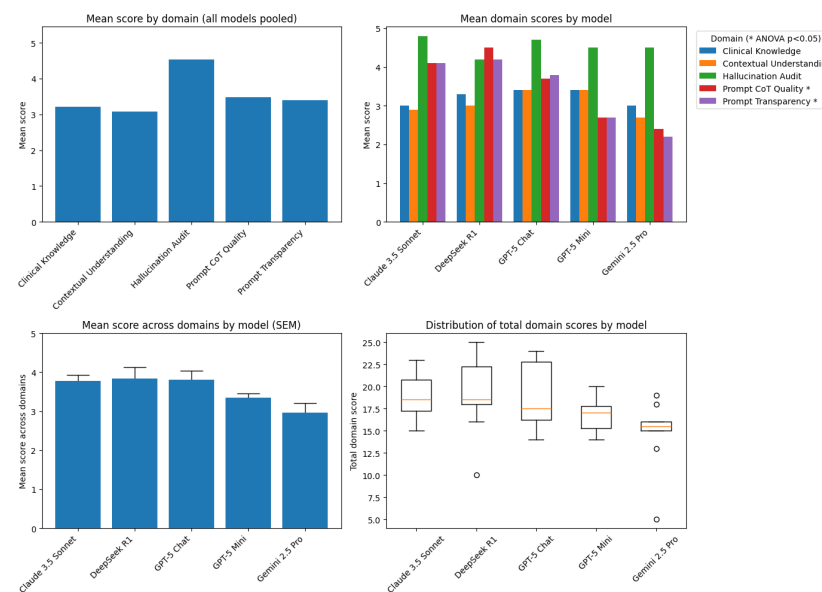
Background

- Large language models (LLMs) may augment surgical decision-making by reducing cognitive burden and supporting adherence to current guidelines
- Inguinal hernia repair is a commonly performed surgical procedure in both emergent and elective settings and requires careful consideration to determine the optimal surgical approach (i.e., open versus laparoscopic, with or without mesh).
- We evaluated the clinical reasoning capabilities of state-of-the-art models in patients admitted following inguinal hernia repair.

Methods

- Five advanced large language models were assessed: Anthropic Claude 3.5 Sonnet, DeepSeek R1, Google Gemini 2.5 Pro, OpenAI GPT-5 Chat, and OpenAI GPT-5 Mini.
- Models were prompted using a standardized case-specific clinical reasoning schema
- Outputs were evaluated across five domains using a 5-point scale with pre-defined anchor points: Clinical Knowledge, Contextual Understanding, Hallucination Audit, Prompt Chain-of-Thought (CoT) Quality, and Prompt Transparency.

Figures



Discussion

- The models accurately summarized a wide range of patient contexts, from healthy patients undergoing elective laparoscopic bilateral inguinal hernia repair with mesh to patients with cirrhosis presenting emergently with incarcerated bowel undergoing primary repair without mesh.
- Between the models, there were no significant differences in Clinical Knowledge, Contextual Understanding, or Hallucination Audit.
- Significant differences were observed for Prompt CoT Quality ($p<0.01$) and Prompt Transparency ($p<0.01$). Claude 3.5 Sonnet demonstrated the highest overall reasoning performance.
- Pairwise comparisons showed Claude 3.5 Sonnet outperformed Gemini 2.5 Pro ($p=0.01$) and GPT-5 Mini ($p=0.03$). GPT-5 Chat and DeepSeek R1 also outperformed Gemini 2.5 Pro ($p=0.02$ and $p=0.03$, respectively).
- Within each model, Hallucination Audit scores were higher than Clinical Knowledge and Contextual Understanding ($p<0.01$ for all models).

Conclusion

- With their low hallucination rates, all models were able to evaluate patient notes and provide succinct, accurate summaries, which may have utility in creating new hand-off tools to ensure safe transitions in patient care.