

Elsevier Clinical Decision Support Symposium

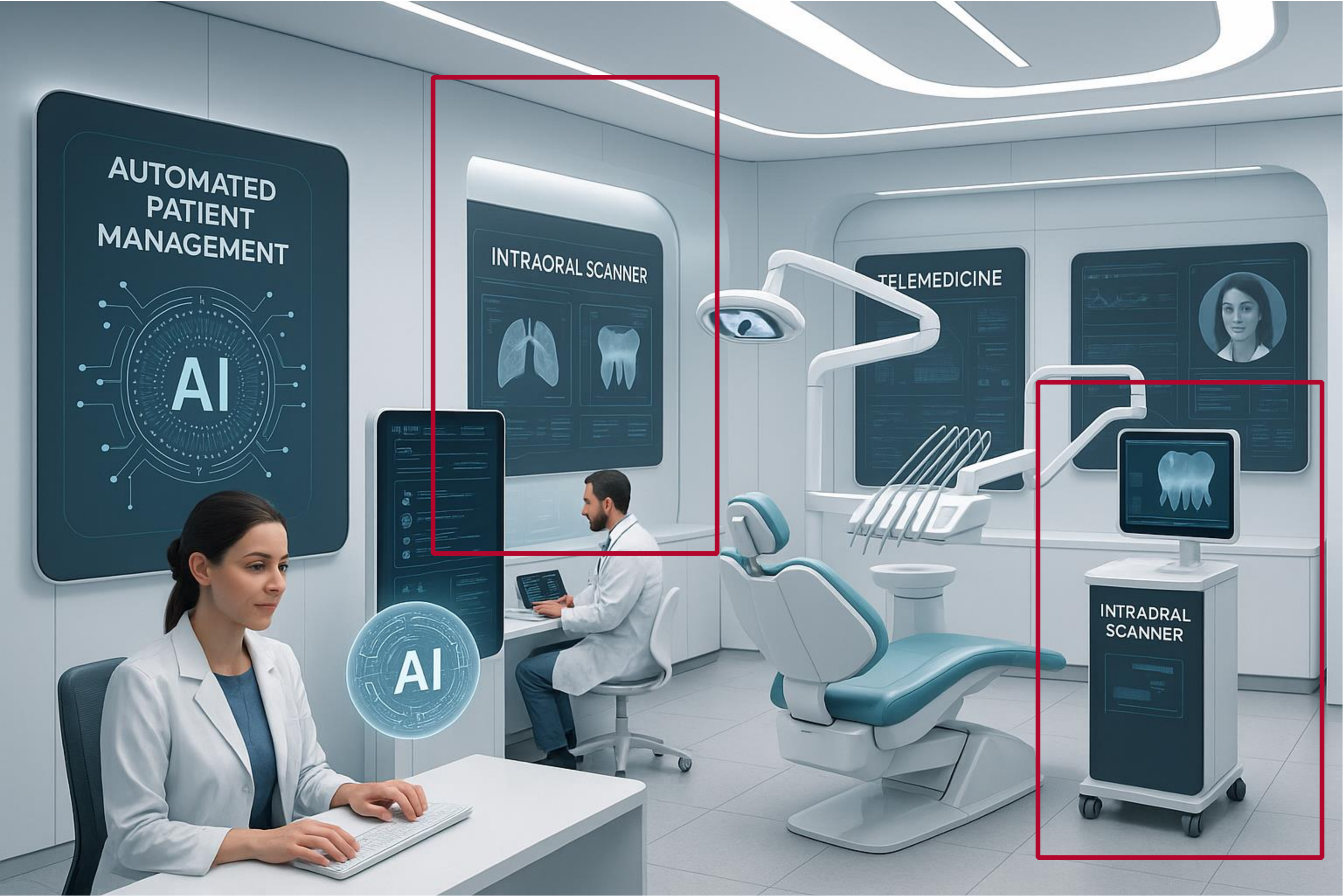
MedBot vs. RealDoc: Effizienz großer Sprachmodelle in der Arzt-Patient-Kommunikation bei seltenen Krankheiten

Prof. Dr. Jannik Schaaf, Institut für Medizininformatik, Goethe-Universität Frankfurt



INSTITUT FÜR
MEDIZIN-
INFORMATIK
UNIVERSITÄTSMEDIZIN FRANKFURT

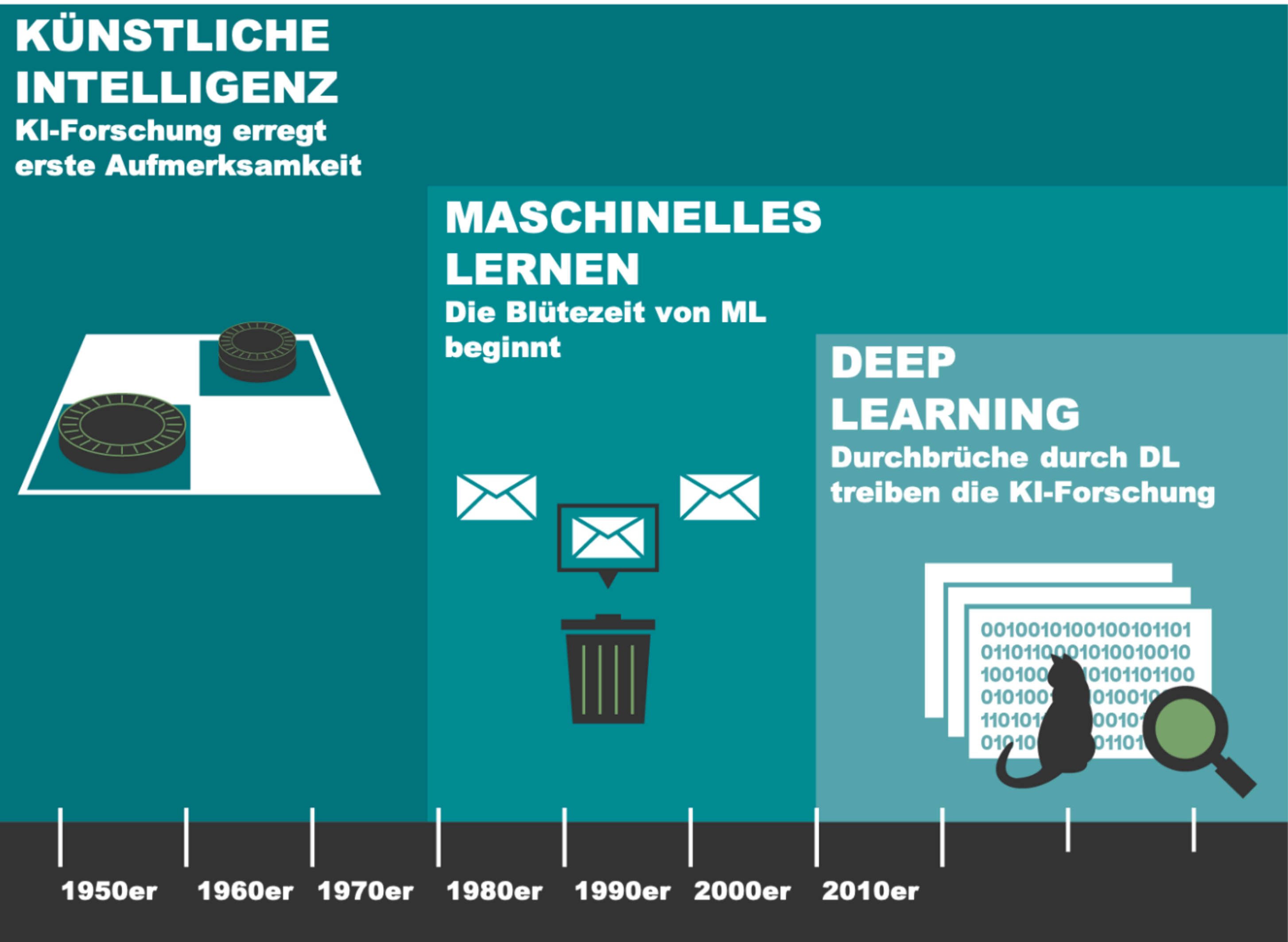
Relevanz des Themas

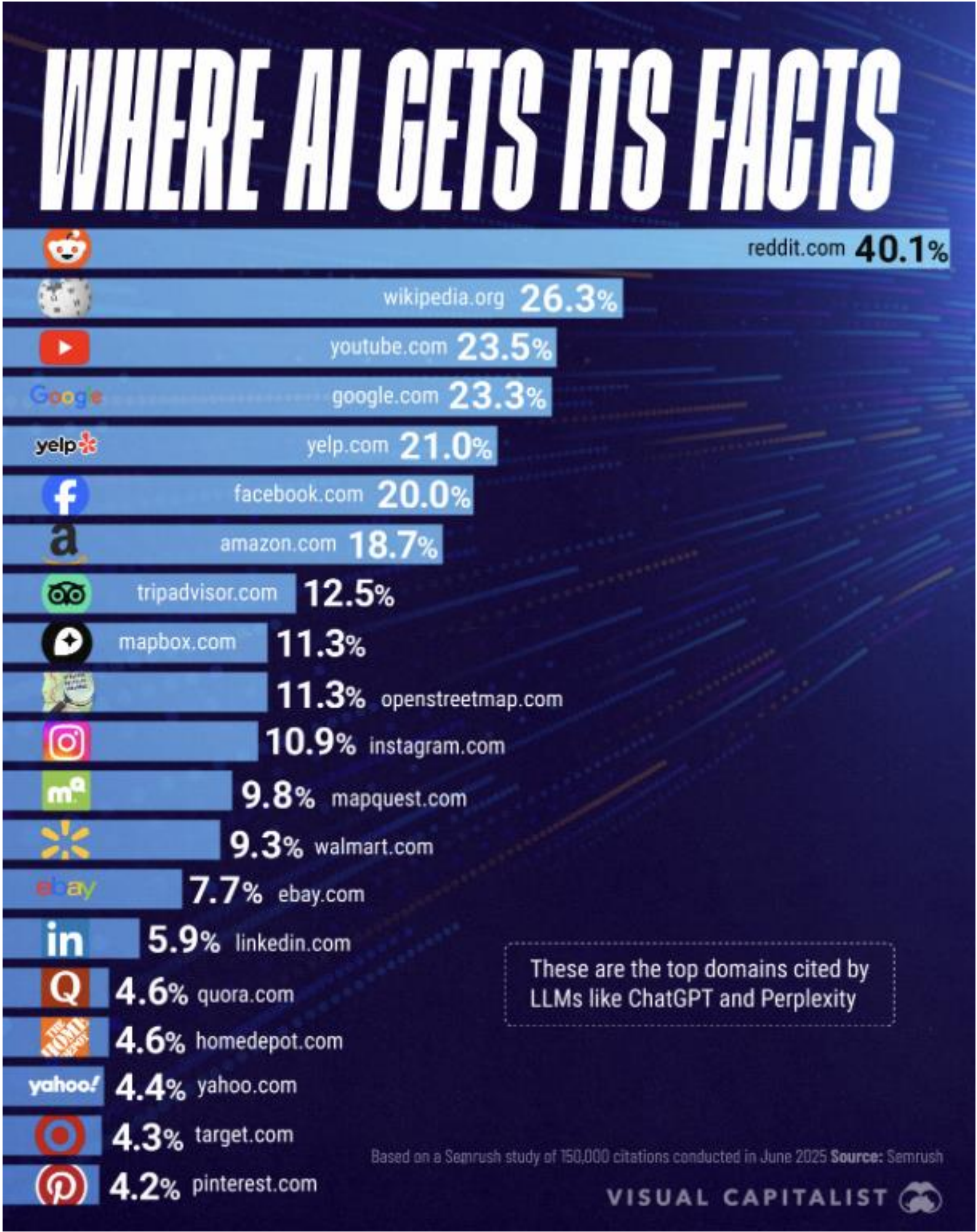


[Hörzu, 14. Ausgabe, 28.03.24]



Relevanz des Themas





Künstliche Intelligenz

Time to Reach 100M Users

Months to get to 100 million global Monthly Active Users

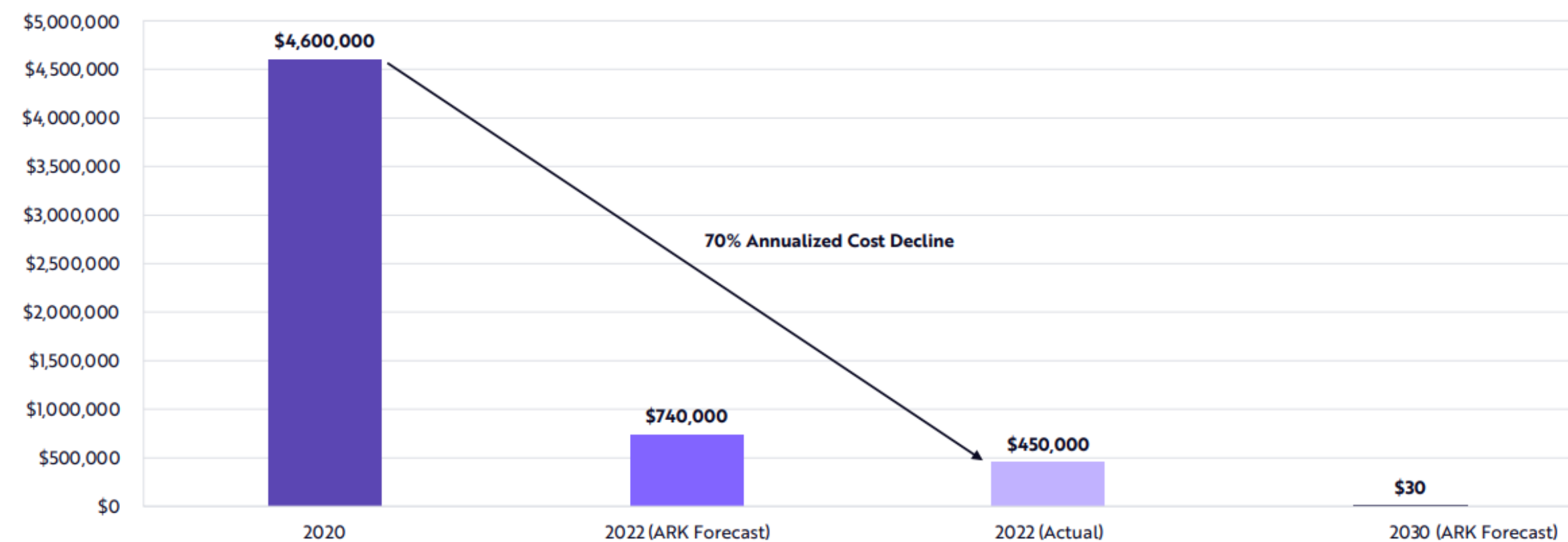


Source: UBS / Yahoo Finance

@EconomyApp

APP ECONOMY INSIGHTS

Cost To Train GPT-3 Level Performance



Künstliche Intelligenz

Meta will 600 Milliarden in KI-Infrastruktur in USA investieren

Um im Rennen um Künstliche Intelligenz mit Amazon, Microsoft und Co. mitzuhalten, will der Facebook-Mutterkonzern 600 Milliarden Dollar in US-Infrastruktur und Arbeitsplätze stecken.

07.11.2025 - 21:29 Uhr

Artikel anhören 02:04

🔗 ✉️ ✕ in f 📄 📄



Meta-Chef Zuckerberg: „Es ist die richtige Strategie, die Kapazitäten aggressiv vorab aufzubauen.“ Foto: AP



INSTITUT FÜR
MEDIZIN-
INFORMATIK
UNIVERSITÄTSMEDIZIN FRANKFURT

Research and Applications

MedBot vs RealDoc: efficacy of large language modeling in physician-patient communication for rare diseases

Magdalena T. Weber, MSc*,¹ Richard Noll , MSc*,¹ Alexandra Marchl, MSc¹, Carlo Facchinello, PhD², Achim Grünewaldt, PhD³, Christian Hügel, PhD⁴, Khader Musleh, PhD³, Thomas O.F. Wagner, PhD⁵, Holger Storf, PhD¹, Jannik Schaaf, PhD¹

¹Institute of Medical Informatics, University Medicine Frankfurt, Goethe University Frankfurt, Frankfurt 60590, Germany, ²Santagostino Medical Center, Bologna 40138, Italy, ³Department of Respiratory Medicine and Allergology, University Medicine Frankfurt, Goethe University Frankfurt, Frankfurt 60590, Germany, ⁴HELIOS Dr Horst Schmidt Kliniken Wiesbaden, Klinik für Pneumologie, Wiesbaden 65199, Germany, ⁵European Reference Network for Rare Respiratory Diseases (ERN-LUNG), University Medicine Frankfurt, Frankfurt 60590, Germany

*Corresponding authors: Magdalena T. Weber, MSc, Institute of Medical Informatics, University Medicine Frankfurt, Goethe University Frankfurt, Theodor-Stern-Kai 7, 60590 Frankfurt am Main, Germany (ma.weber@med.uni-frankfurt.de) and Richard Noll, MSc, Institute of Medical Informatics, University Medicine Frankfurt, Goethe University Frankfurt, Theodor-Stern-Kai 7, 60590 Frankfurt am Main, Germany (noll@med.uni-frankfurt.de)

Abstract

Objectives: This study assesses the abilities of 2 large language models (LLMs), GPT-4 and BioMistral 7B, in responding to patient queries, particularly concerning rare diseases, and compares their performance with that of physicians.

Materials and Methods: A total of 103 patient queries and corresponding physician answers were extracted from EXABO, a question-answering forum dedicated to rare respiratory diseases. The responses provided by physicians and generated by LLMs were ranked on a Likert scale by a panel of 4 experts based on 4 key quality criteria for health communication: correctness, comprehensibility, relevance, and empathy.

Results: The performance of generative pretrained transformer 4 (GPT-4) was significantly better than the performance of the physicians and BioMistral 7B. While the overall ranking considers GPT-4's responses to be mostly correct, comprehensive, relevant, and emphatic, the responses provided by BioMistral 7B were only partially correct and empathetic. The responses given by physicians rank in between. The experts concur that an LLM could lighten the load for physicians, rigorous validation is considered essential to guarantee dependability and efficacy.

Discussion: Open-source models such as BioMistral 7B offer the advantage of privacy by running locally in health-care settings. GPT-4, on the other hand, demonstrates proficiency in communication and knowledge depth. However, challenges persist, including the management of response variability, the balancing of comprehensibility with medical accuracy, and the assurance of consistent performance across different languages.

Conclusion: The performance of GPT-4 underscores the potential of LLMs in facilitating physician-patient communication. However, it is imperative that these systems are handled with care, as erroneous responses have the potential to cause harm without the requisite validation procedures.

Key words: artificial intelligence; medical informatics; natural language processing; health communication; rare diseases.

Background and significance

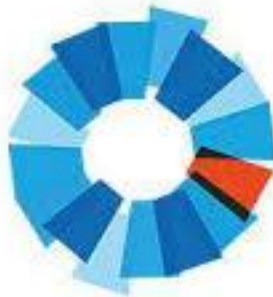
The advent of large language models (LLMs) led to a public comeback of artificial intelligence (AI), with profound implications for a variety of sectors, including health care.¹ Models such as OpenAI's generative pretrained transformer 4 (GPT-4) are capable of understanding and generating human-like text and performing various tasks, such as summarization, translation, and even coding.² Given their broad applicability, it seems feasible to advance automation in a variety of fields by employing LLMs, a prospect that is cur-

memorize prior conversations.⁴ One of the principal benefits of employing LLMs in the health-care sector is their capacity to facilitate improved access to medical information. These models can provide round-the-clock support, tending to patients with immediate responses to common medical questions and helping to alleviate disparities in health-care access, particularly in underserved regions.⁵ By augmenting traditional health-care services, LLMs can reduce the burden on health-care professionals.⁶ In addition to convenience, LLMs can deliver personalized advice, drawing on vast databases of

Downloaded from https://academic.oup.com/jamia/advance-article/doi/10.1093/jamia/ocaf034/8042190 by guest on 04 April 2025

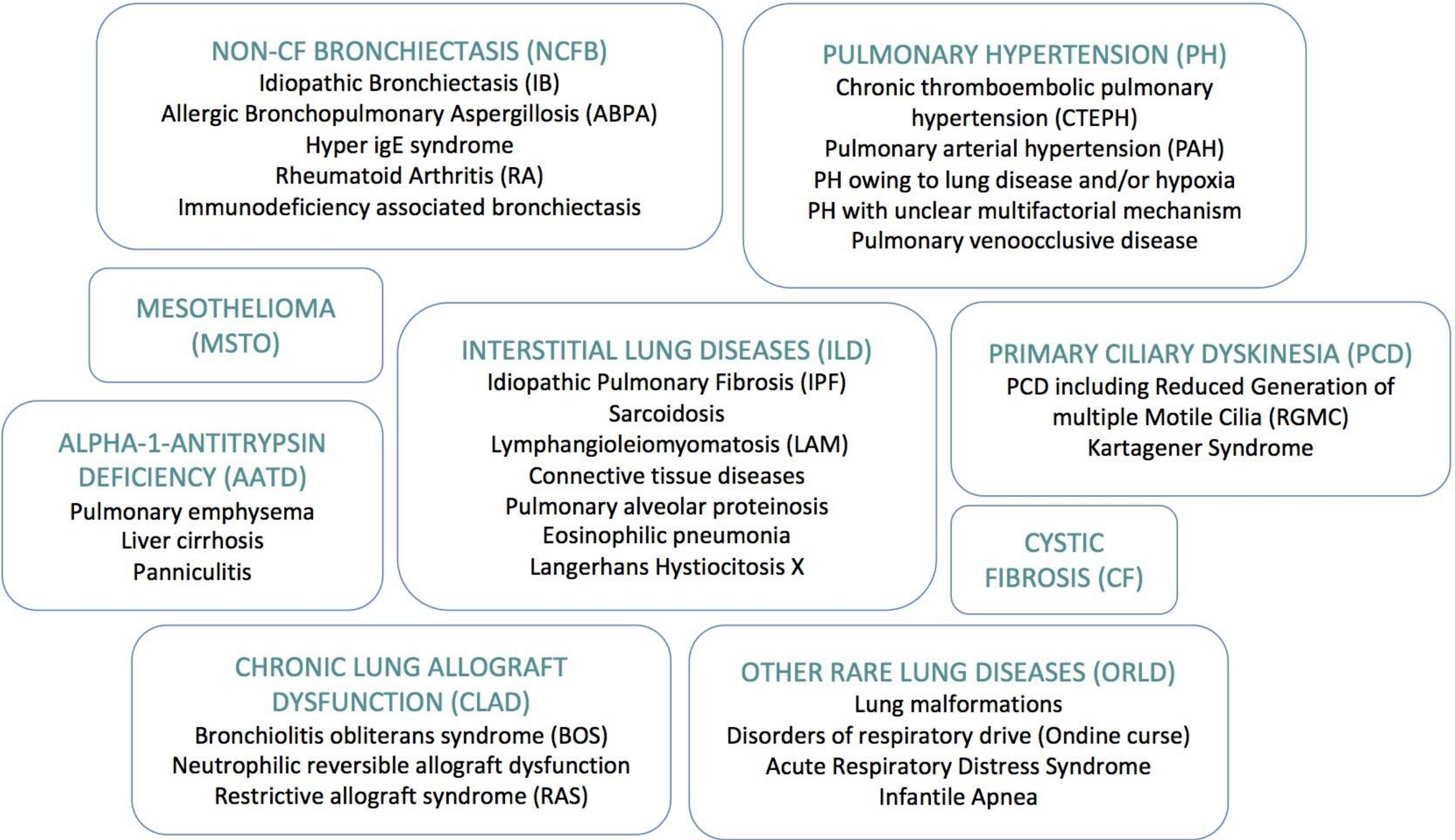


Respiratory Diseases
(ERN-LUNG)



Frankfurter Referenzzentrum
für Seltene Erkrankungen

ERN-LUNG addresses the following rare and complex pulmonary diseases and rare respiratory tumours



INSTITUT FÜR
MEDIZIN-
INFORMATIK
UNIVERSITÄTSMEDIZIN FRANKFURT

EXABO

3 Interstitial Lung Disease

language ▾

✉ Subject: Pulmonary and mediastinal sarcoidosis with extra-pulmonary symptoms

I was diagnosed with sarcoidosis in 2020 after suffering for about 10 years with symptoms. I was treated with steroids and methotrexate for a year and a half, but every time I reduced my steroid dose my symptoms returned. I was subsequently treated with Infliximab and that completely got rid of my symptoms and improved my x-ray results. I moved to France in August 2022 and the doctors here say that while I still have sarcoidosis, it's not severe enough to treat. I still have skin lesions, extreme fatigue, brain fog, tinnitus, and some respiratory symptoms without treatment. I went from a 90% average as a student to failing my most recent semester, and I can barely make a living because of my current state. The doctors here told me to do some yoga. I would like to know if there are any treatment options that I can explore that are a middle ground between no treatment at all and Infliximab.

Hello,

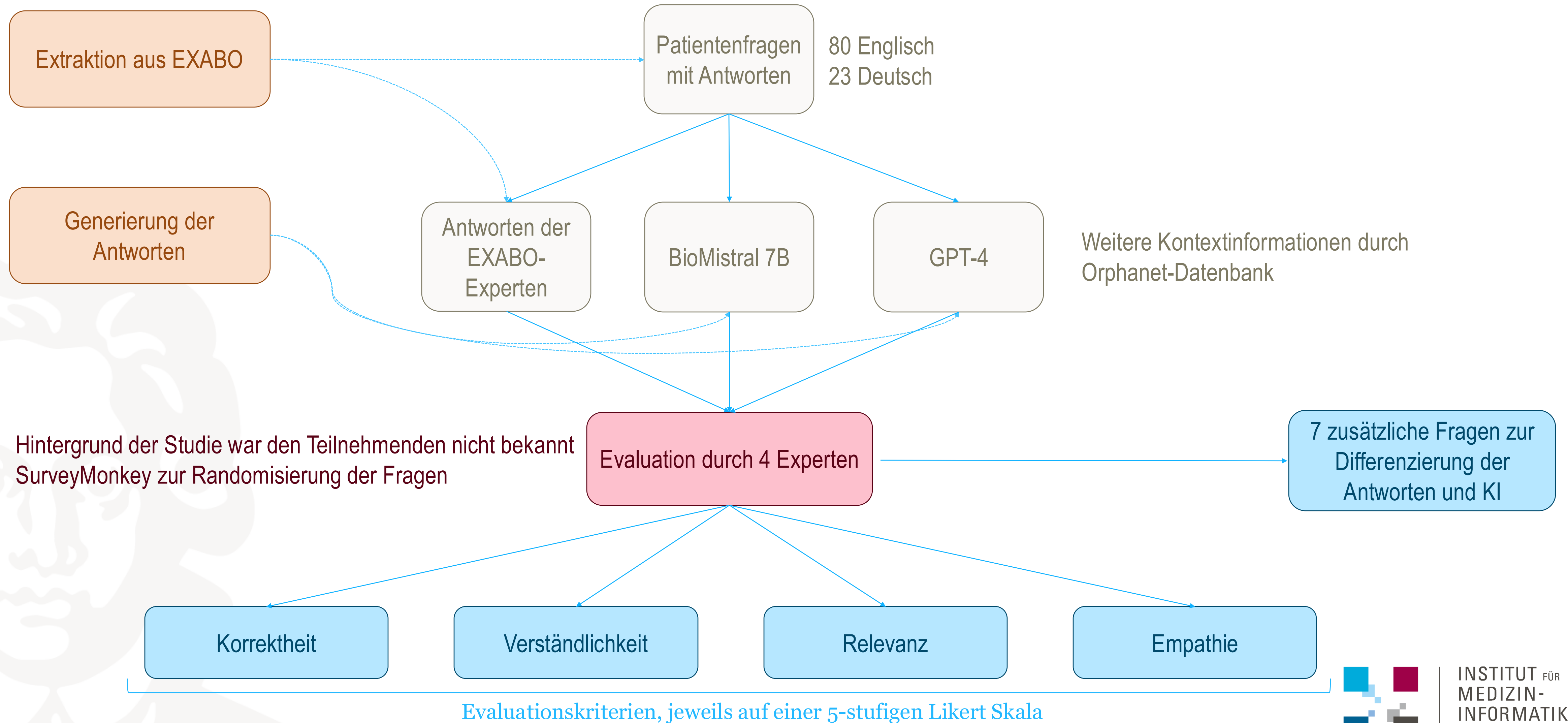
Prof. F. Bonella, Essen, Germany, wrote.

"According to the recent guideline on sarcoidosis management and treatment, published in 2022, indications for treatment are symptoms affecting quality of life, functional impairment and risk of permanent disability or death. Rehabilitation for fatigue is a very helpful treatment option. Hydroxychloroquine, an anti-malarial drug, can improve skin lesions. I suggest to be followed at a reference center for sarcoidosis, where all these aspects and alternative treatment options should be considered."

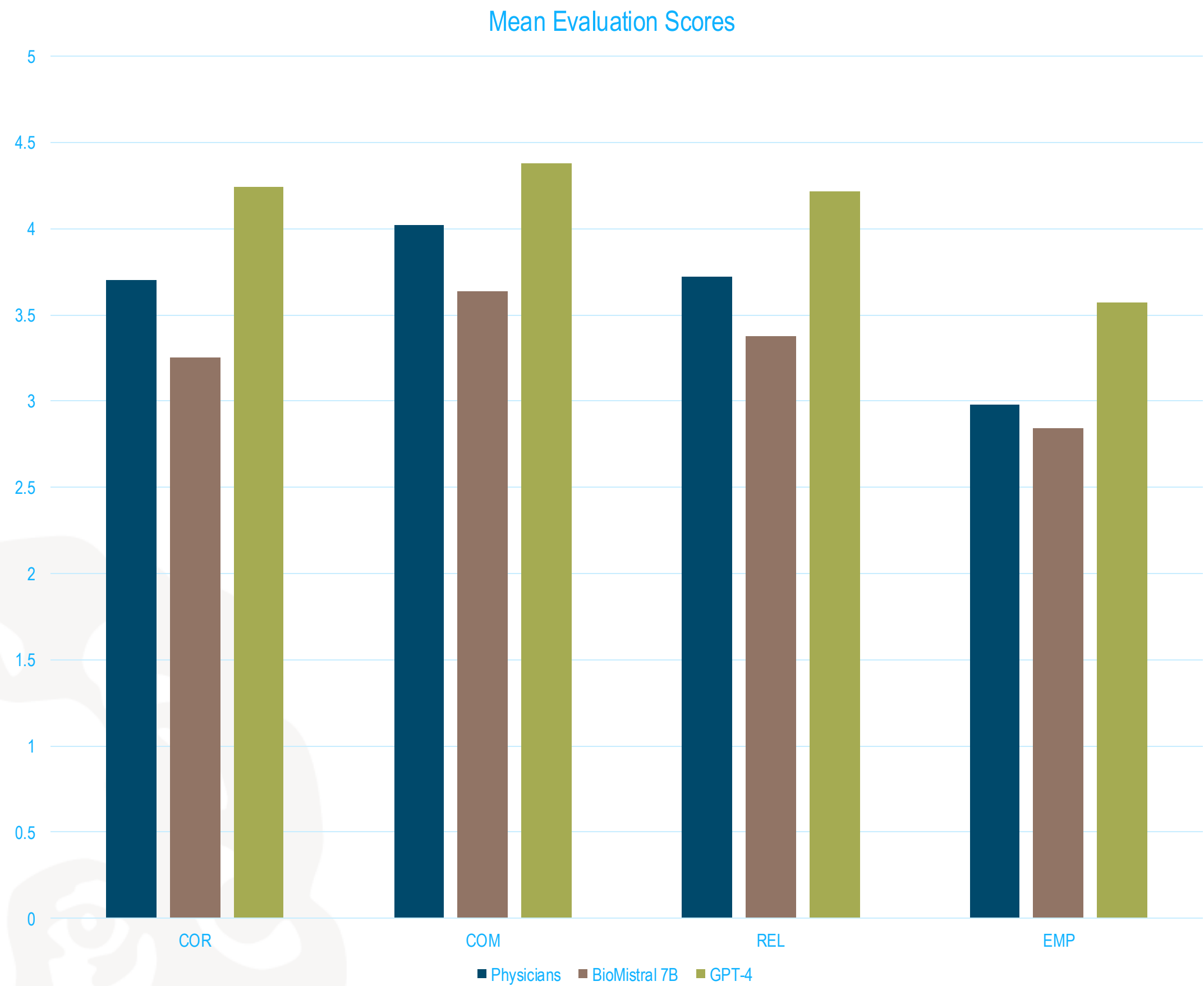
Best regards

Dr TOF Wagner, Coordinator ERN-LUNG

MedBot vs. RealDoc



Studie zur Large Language Modells



COR: Correctness
REL: Relevance

COM: Comprehensibility
EMP: Empathy

	GPT-4		BioMistral 7B		Physicians	
	Median	Mean (SD)	Median	Mean (SD)	Median	Mean (SD)
Correctness	4	4.24 (0.77)	3	3.25 (1.10)	4	3.70 (0.96)
Comprehensibility	4	4.38 (0.64)	4	3.64 (1.00)	4	4.02 (0.87)
Relevance	4	4.22 (0.75)	4	3.38 (1.06)	4	3.72 (0.93)
Empathy	4	3.57 (0.89)	3	2.84 (0.93)	3	2.98 (0.92)
Total		4.1 (0.76)		3.3 (1.02)		3.6 (0.92)

Intraklassenkorrelationskoeffizient
(Reliabilität)

Correctness (0.728)
Comprehensibility (0.629)
Relevance (0.701)
Empathy (0.663)



Antworten im Fragebogen – Zusammenfassung

Frage	Antwort
Inwieweit könnten Chatbot-Antworten dazu beitragen, die Arbeitsbelastung von Ärzten zu verringern? Würden Sie eine solche Antwort als ersten Entwurf betrachten?	Alle Teilnehmer empfinden die vorgefertigten Antworten eines Chatbots für hilfreich und würden sie nutzen. Sie erwarten, dass er grundlegende Informationen aus der Literatur bereitstellt.
Welche potenziellen Risiken sehen Sie bei der Verwendung von Chatbots, insbesondere im Hinblick auf Fehlinformationen oder das Übersehen wichtiger Informationen?	Die Teilnehmer sind sich weitgehend einig, dass diese Risiken vermeidbar sind, wenn das Tool den Fachpublikationen und Richtlinien entspricht und einer gründlichen Validierung unterzogen wird.
Welche Überlegungen würden Sie anstellen, bevor Sie den Einsatz eines Chatbots in der Praxis in Betracht ziehen?	Die Teilnehmer würden das Risiko von Halluzinationen überprüfen, Tests durchführen und Aspekte wie Datenschutz, Benutzerfreundlichkeit und Kosten berücksichtigen.
Wie beurteilen Sie die ethischen Aspekte der Verwendung eines Chatbots, insbesondere im Hinblick auf die Vertraulichkeit von Patientendaten?	Die Teilnehmer sind sich einig, dass der Schutz der Privatsphäre der Patienten von größter Bedeutung ist. Sie gehen davon aus, dass man diesen Schutz gewährleisten kann.
Gibt es in den Antworten irgendwelche rückblickenden Merkmale oder Hinweise, die Sie vermuten lassen, dass sie von einem Arzt oder einem Chatbot stammen? Wenn ja, welche?	Die Teilnehmer geben formale Korrektheit, aber Unlogik, Unverständlichkeit, unvollständige Sätze, Allgemeinheit und Studienzitate als Hinweise für Chatbot-Antworten an.
Welche rechtlichen Herausforderungen sehen Sie bei der Implementierung eines KI-basierten Chatbots in der Arzt-Patient-Kommunikation?	Die Meinungen der Teilnehmer zu diesem Thema gehen auseinander. Einige verweisen auf Datenschutzprobleme, während andere optimistisch sind, dass der Chatbot als Unterstützungssystem umgesetzt werden kann.

Diskussionen

- GPT4 erzielte signifikant bessere Ergebnisse
- Hohe Empathie; jedoch kulturell eingeschränkt
- Biomistral 7B signifikant schlechter; mögliche Ursache geringe Parameteranzahl und Defizite im Deutschen
- Biomistral 7B kann jedoch lokal betrieben werden
- Gefahr von Halluzinationen sind groß
- LLMs in der Patientenkommunikation könnten Zeitdruck minimieren und fehlendes Wissen kompensieren
- Begrenzte Repräsentativität: Nur 4 Ärzte als Bewertende, jedoch moderate bis hohe Inter-Rater-Reliabilitäten



Vielen Dank für Ihre Aufmerksamkeit!

Prof. Dr. Jannik Schaaf

Johann Wolfgang Goethe-Universität Frankfurt
Universitätsmedizin Frankfurt
Institut für Medizininformatik – IMI

Haus 4
Theodor-Stern-Kai 7
60590 Frankfurt

schaaf@med.uni-frankfurt.de
www.imi-frankfurt.de

