

On the Quality and Impact of Google's AI Overviews

Lewandowski, Dirk^a

^a *Hamburg University of Applied Sciences, Finkenau 35, 22081 Hamburg, Germany*

1. Theoretical considerations

The integration of AI-generated answers marks a fundamental shift in information retrieval in general (Breuer et al., 2025), and search engines in particular. Apart from the obvious shift from document lists to answers, this also affects the search engine's role. Traditionally, search engines acted as information intermediaries, bridging the gap between users and external content providers. With AI Overviews, they transition into information producers, generating new information objects on the fly (Lewandowski, 2025b). This creates a "closed-loop" environment in which the search system no longer merely points to information but synthesises it, potentially making visits to the original source unnecessary. The interaction paradigm has evolved from a simple query-response model to natural-language, multi-turn conversations. Finally, information seeking and use are converging. Search systems now summarise and process documents, making them directly applicable to the user's tasks. This expands the "task frontier" (White, 2024), enabling systems to handle complex activities that were previously the user's sole responsibility.

An important consideration is the accuracy of AI Overviews. The fallibility of these answers is not a temporary bug but is system-immanent. Since Large Language Models (LLMs) operate on statistical probabilities rather than truth values, they are prone to specific types of errors:

1. "Hallucinations": The fabrication of facts or the invention of non-existent sources
2. Inconsistency: Providing different or even contradictory answers to the same query over time
3. Source Misalignment: The system may attribute statements to sources that do not actually contain that information.

All these errors occur even when a system uses Retrieval-Augmented Generation (RAG), i.e., the grounding of AI Overviews in retrieved documents.

2. Empirical results

Empirical data indicate that Google often uses a different set of sources for AI Overviews than for its organic top results. In the context of health queries, while organic results prioritise high-quality, authoritative medical domains, AI Overviews frequently draw from a broader, less curated set of sources (Sünkler et al., 2026). This can have a huge impact on users' health information-seeking and on the information they base medical decisions on.

Eye-tracking shows that the prominent placement of AI Overviews pushes traditional organic results down, drastically reducing their visibility (Schultheiß & Lewandowski, submitted). Furthermore, most task completions begin with an initial view of the AI-generated overview. Another eye-tracking study (Brüning Genannt Wolter et al., 2026) evaluated how citation visuals, source credibility, and factual accuracy influence the adoption of AI-generated health answers. Findings indicate that answer accuracy is the primary driver of adoption, while source credibility and citation format (logos vs. footnotes) show no significant impact. Although logos attract more visual attention, they do not increase trust or adoption of content from the AI answers. Notably, participants adopted false but confidently phrased answers more often than vague ones, emphasising that perceived quality outweighs superficial design signals.

Another study compared the influence of Google's AI Overviews versus Featured Snippets on user opinion formation regarding controversial topics (Haase et al., submitted). Experimental results show that while Featured Snippets significantly influence user attitudes, AI Overviews do not have a similarly persuasive impact.

3. Policy

Regarding diversity of opinion, AI Overviews are often criticised for their potential to narrow diversity of opinion by showing only a reduced set of opinions, compared to traditional search results. However, one could also hypothesise that AI Overviews strengthen the diversity of opinion by presenting more aspects of a topic than a user would consider in traditional search results. While the full result set would obviously contain more aspects than the much shorter AI answer, compared to the few results users usually consider, an AI answer could offer an improvement (Lewandowski, 2025a).

The impact of AI Overviews on content producers is more pronounced, as the "zero-click" trend is significantly amplified by them. Findings suggest a 20%-50 % loss of traffic for content providers (Lewandowski, 2025a). This poses a structural threat to publishers' business models and may reduce opinion plurality, as publishers will be unable to finance new content if they are deprived of the ability to monetise it. On the other hand, content producers not dependent on directly monetising their content will profit, bringing content, for instance, from PR agencies, large corporations, NGOs, and political parties to the forefront.

As AI Overviews constitute new information objects (rather than merely reproducing external information objects or parts thereof), there is ongoing discussion about who should be responsible for the accuracy and non-discriminatory nature of these answers. As the process of assembling the AI Overviews is similar to (automated) journalism, it is argued that the producers of the AI Overviews, namely the search engines, should be held responsible.

Last but not least, AI Overviews affect the established model of the web, in which content producers give their content to search engines for free. In exchange, they receive traffic (visitors) that they can monetise through advertisements or by selling products and services. We need to ensure that AI systems' extraction of value from these content producers does not destroy the very ecosystem they rely on for training data.

4. Suggestions for further research

Major areas for further research are

1. Quality evaluations: It is of tremendous importance to systematically assess the quality of AI Overviews, providing a solid basis for discussions on how AI Overviews promote inaccurate information, misinformation, disinformation, and the like.
2. The role of sources: On what sources are AI Overviews built? Are these sources accurate and appropriate to answer users' questions? What role do users assign to these sources?
3. Impact on diversity of opinion: Studies on whether the synthesis of information into a single "answer" leads to a reduction in the diversity of perspectives encountered by the average user.

As with other topics in search engine research, the question is what kind of studies are feasible. Lewandowski et al. (2020) argued for mid-size studies, which can be supported by tools for automatic collection and analysis of search results. In that vein, the Result Assessment Tool (RAT) Sünkler et al., (2025) has been amended to also allow studies of AI answers.

5. References

Breuer, T., Frihat, S., Fuhr, N., Lewandowski, D., Schaer, P., & Schenkel, R. (2025). Large Language Models for Information Retrieval: Challenges and Chances. *Datenbank-Spektrum*, 25(2), 71–

81. <https://doi.org/10.1007/s13222-025-00503-x>

SEASON 2026, September 15–17, 2026, Hamburg, Germany

EMAIL: dirk.lewandowski@haw-hamburg.de

ORCID: 0000-0002-2674-9509

DOI: <https://doi.org/10.48441/4427.3708>



© 2026 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).



- Brüning Genannt Wolter, M., Haase, L. E., Lewandowski, D., Michaelson, B., & Schott, F. (2026). Visual Cues and Trust in AI-Generated Search Results: An Eye-Tracking Study on Google AI Overviews. *SEASON 2026*. SEASON 2026, September 15–17, 2026, Hamburg, Germany.
- Haase, H., Görlitz, L., Naruz, B., & Lewandowski, D. (submitted). *An Investigation Into How Google's AI Overviews and Featured Snippets Influence Opinion Formation*.
- Lewandowski, D. (2025a). *Integration von KI-Anwendungen in Suchmaschinen und ihre Auswirkungen auf die Meinungsvielfalt. die Medienanstalten*. https://www.die-medienanstalten.de/fileadmin/user_upload/die_medienanstalten/Service/Studien_und_Gutachten/KI-Gutachten_2025_final.pdf
- Lewandowski, D. (2025b). KI-Chatbots als Suchsysteme, Suchmaschinen als KI-Bots? In M. Eibl (Ed.), *Datenströme und Kulturoasen—Die Informationswissenschaft als Bindeglied zwischen den Informationswelten. Proceedings des 18. Internationalen Symposiums für Informationswissenschaft (ISI 2025), Chemnitz, Deutschland, 18.—20. März 2025*. (pp. 150–167). Verlag Werner Hülsbusch. <https://doi.org/10.5281/zenodo.14925636>
- Lewandowski, D., Sünkler, S., & Schultheiß, S. (2020). Studies on Search: Designing Meaningful IIR Studies on Commercial Search Engines. *Datenbank-Spektrum*, 20(1), 5–15. <https://doi.org/10.1007/s13222-020-00331-1>
- Schultheiß, S., & Lewandowski, D. (submitted). *How Google's AI Overviews Capture User Attention: An Eye-Tracking Study*.
- Sünkler, S., Lewandowski, D., Schultheiß, S., & Koop, O. (2026). Information Diversity and Authority Shifts in Google's AI Overviews: An Audit of Health Queries. *WebSci Companion '26*. Web Science Conference.
- Sünkler, S., Lewandowski, D., Schultheiß, S., & Yagci, N. (2025). Result Assessment Tool (RAT): Empowering search engine data analysis. *PeerJ Computer Science*, 11, e2962. <https://doi.org/10.7717/peerj-cs.2962>



White, R. W. (2024). Advancing the Search Frontier with AI Agents. *Communications of the ACM*, 67(9), 54–65. <https://doi.org/10.1145/3655615>

SEASON 2026, September 15–17, 2026, Hamburg, Germany

EMAIL: dirk.lewandowski@haw-hamburg.de

ORCID: 0000-0002-2674-9509

DOI: <https://doi.org/10.48441/4427.3708>



© 2026 Copyright for this paper by its authors.

Use permitted under Creative Commons License Attribution 4.0 International (CC BY 4.0).