

Information Literacy in Conversational Search: Developing a Categorization Scheme for Measuring Information-Literate Behavior in Human-AI Interaction

Joachim Griesbaum^a, Nicola Hoppe^a, Beatrix Kreß^a, Karen Malin Krüger^a, Bettina Lindner-Bornemann^a, Anna Mierzecka^b, Stephan Schlickau^a, and Karsten Senkbeil^a

^a University of Hildesheim, Hildesheim, Germany

^b University of Warsaw, Warsaw, Poland

1. Introduction

Information literacy can be defined as "the individual's ability to deal with information in an effective and ethical manner" (Griesbaum et al., 2024). The proliferation of generative AI systems is profoundly reshaping users' information behavior. Several AI literacy frameworks have been proposed (e.g., Miao & Shiohira, 2024; Kong et al., 2024; Long & Magerko, 2020), yet the relationship between AI literacy and information literacy remains underexplored. Griesbaum et al. (2026) report that AI will both facilitate and complicate information behavior. While AI makes information access easier and partially takes over functions of selecting, evaluating, and presenting information, it becomes much more demanding to check the quality and validity of AI-generated outputs, which are often not transparent regarding their underlying knowledge base. Information-literate behavior is most visible in real human-AI interaction. Instruments are needed that can a) assess the current state of information literacy in human-AI interaction and b) measure whether interventions lead to observable improvements. This paper reports the ongoing development of such an instrument: a categorization scheme designed to capture the extent to which users exhibit information-literate behavior in their interactions with AI. The work should be understood as an explorative concept paper. The presented results are tentative.

2. Related Work

Several strands of related work inform our approach. Tibau et al. (2024) adapted classical search tactics (Bates, 1979; Smith, 2012) for ChatGPT-like tools, developing a scheme of 45 tactics grouped into seven categories (tracking, information input/output, information retrieval, prompt formulation, prompting granularity, and evaluation). Their work suggests that many classic search strategies can be successfully translated for LLM-based tools and that their use contributes to more effective AI interaction. Lin et al. (2024) developed rubrics to automatically measure user satisfaction in LLM-based conversations, some of which relate to information-literate behavior (e.g., follow-up inquiries, error identification). Kim et al. (2024) developed a taxonomy of follow-up queries in conversational search with two axes (seven motivation categories, eleven action types) and found significant correlations between query types and user satisfaction. They argue that categories such as clarifying queries, broadening and deepening information, and evaluating results can be linked to information-literate behavior. Li et al. (2025) developed an observational framework for "Human-AI Interaction Capability" (HAI) with three dimensions (information retrieval & analysis, evaluation, integration) and eleven criteria. Their results indicate that HAI correlates with critical thinking and that the quality and logic of questions are central indicators of good interaction. Krakowska & Zych (2025) investigated LIS students' ChatGPT strategies and developed categories distinguishing surface evaluation from critical competence. Beyond the user perspective, frameworks such as LLM-Rubric (Hashemi et al., 2024) and FACTS (Cheng et al., 2025) evaluate AI output quality from a technical perspective. However, an

information literacy perspective on AI output (e.g., analyzing whether responses are transparent about uncertainty, invite verification, or encourage critical engagement) remains largely unexplored. In sum, current research indicates that the analysis of information literate behavior can be meaningfully transferred to LLM-based interactions. Still, methodological challenges remain. The evaluation of AI output through an information literacy lens is an understudied area.

3. Current state of the development of the categorization scheme

Building on these insights, we argue that measuring information literacy in human-AI interaction requires analyzing three levels simultaneously: (1) user behavior at the post level (e.g., how users control input and output and evaluate results), (2) AI behavior at the post level (e.g., the validity and comprehensibility of answers and elicitation of information-literate behavior in AI responses), and (3) the discourse level (e.g., the evolution of the information need, interaction course and output quality). Crucially, information-literate behavior cannot be attributed exclusively to the user. It emerges in relation to the AI's responses, making the AI a co-actor whose contributions must be analyzed both for quality and for their capacity to elicit information-literate engagement. To our knowledge, this dual-actor perspective has not been systematically addressed in existing research. The following figure illustrates this conceptual framework to analyze information literacy in human-AI interaction.

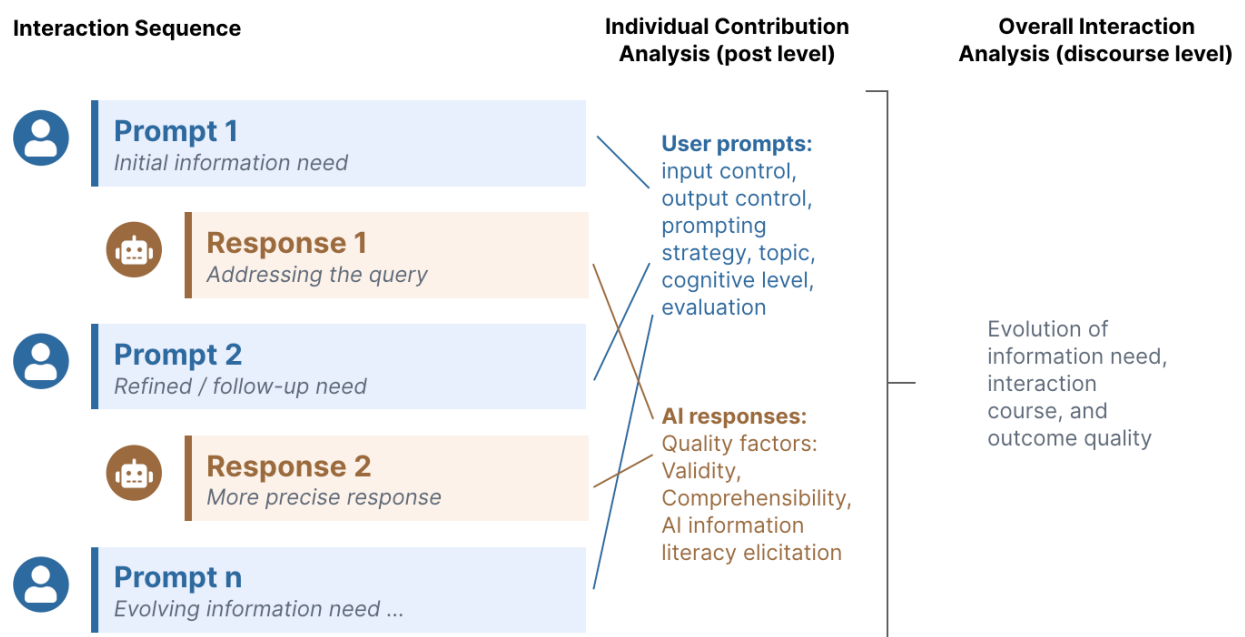


Figure 1: Conceptual framework to analyze information literacy in human-AI interaction

The next step is to develop the categories to determine how each of the three levels will be measured in concrete terms. We report on three iterative development phases.

In Phase 1, we adapted Tibau et al.'s (2024) 45-tactic scheme for a pilot study (Griesbaum et al., 2025) in which students completed nutrition-related counseling scenarios with an anthropomorphic AI in a virtual environment. The scheme proved impractical: only a few tactics (mainly prompt formulation and prompting granularity) were observable, categorization was difficult, and individual tactics were not clearly distinguishable.

In Phase 2, we simplified the scheme into five user prompts categories (Sequential Control, Input Control, Output Control, Prompting Strategy, Output Evaluation), three AI responses categories (Initiative, Quality/Validity, Explanation/Transparency), and a discourse-level "Mode" category

capturing topical breadth and depth. Testing this iteration on several hundred posts from student-AI interactions about AI-related topics yielded substantially more feasible and insightful results. For instance, AI frequently offered unsolicited quality explanations and output control options (e.g., summaries), while user output evaluation and process control remained notably rare. However, overlaps between prompting strategy subcategories and the absence of cognitive-level measures (cf. Li et al., 2025) indicate that further refinement is needed.

In the next phase, we aim to finalize the categorization scheme and build a manually coded dataset consisting of several hundred posts from a dozen human-AI interactions, accompanied by a comprehensive codebook. This dataset will serve as input for an LLM-based classifier to scale the analysis process. The resulting instrument is intended not only for analytical research but also as an evaluation tool that provides insights into the effectiveness of information literacy interventions by comparing pre- and post-intervention behavior.

4. Discussion

This paper argues for the need to develop an analytical instrument capable of capturing information-literate behavior in human-AI interaction: (1) User behavior at the post level, (2) AI behavior at the post level, and (3) the discourse level. To our knowledge, this dual-actor perspective, treating AI not merely as a tool but as a co-actor whose output shapes the quality of information-literate engagement, has not been systematically addressed in the current discussion. The paper thus contributes a conceptual framework that extends existing approaches, even though the concrete categorization scheme is still under development.

5. References

Bates, M. J. (1979). Information search tactics. *Journal of the Association for Information Science and Technology*, 30(4), 205–214. <https://doi.org/10.1002/asi.4630300406>

Cheng, A., Jacovi, A., Globerson, A., Golan, B., Kwong, C., Alberti, C., Tao, C., Ben-David, E., Tomar, G. S., Haas, L., Bitton, Y., Bloniarz, A., Bai, A., Wang, A., Siddiqui, A., Raghuveer, A., Bajuelos Castillo, A., Atias, A., Liu, C., . . . Das, D. (2025). The FACTS leaderboard: A comprehensive benchmark for large language model factuality. arXiv. <https://doi.org/10.48550/arXiv.2512.10791>

Griesbaum, J., Michel, A., Dreisiebner, S., Gäde, M., & Wittich, A. (2024). Information literacy: Concepts, research, and promotion. In P. Heisig (Ed.), *Handbook on information sciences* (pp. 312–324). Edward Elgar Publishing. <https://doi.org/10.4337/9781035343706>

Griesbaum, J., Hoppe, N., Kreß, B., Krüger, K. M., Lindner-Bornemann, B., Schlickau, S., & Senkbeil, K. (2025). Sprach- und informationswissenschaftliche Perspektiven auf eine KI-basierte Ernährungsberatung in Virtual Reality. *[Manuscript in preparation]*.

Griesbaum, J., Dreisiebner, S., Michel, A., Tappenbeck, I., & Wittich, A. (2026). Information literacy and artificial intelligence: A library and information science perspective on effects, research questions, challenges and opportunities. In S. Kurbanoğlu et al. (Eds.), *Information literacy in an AI-driven world. ECIL 2025. Communications in Computer and Information Science* (Vol. 2864, pp. 3–13). Springer. https://doi.org/10.1007/978-3-032-17272-3_1

Hashemi, H., Eisner, J., Rosset, C., Van Durme, B., & Kedzie, C. (2024). LLM-Rubric: A multidimensional, calibrated approach to automated evaluation of natural language texts. In L.-W. Ku, A. Martins, & V. Srikumar (Eds.), *Proceedings of the 62nd Annual Meeting of the Association for*

Computational Linguistics (Volume 1: Long Papers) (pp. 13806–13834). Association for Computational Linguistics. <https://doi.org/10.18653/v1/2024.acl-long.745>

Kim, H., Choi, Y., Yang, T., Lee, H., Park, C., Lee, Y., Kim, J. Y., & Kim, J. (2024). Using LLMs to investigate correlations of conversational follow-up queries with user satisfaction. In C. Siro et al. (Eds.), *Proceedings of the 1st Workshop on Large Language Models for Evaluation in Information Retrieval (LLM4Eval 2024)*, Washington DC, United States, 18 July, 2024. *CEUR Workshop Proceedings* (Vol. 3752, pp. 66–91).

Kong, S.-C., Cheung, M.-Y. W., & Tsang, O. (2024). Developing an artificial intelligence literacy framework: Evaluation of a literacy course for senior secondary students using a project-based learning approach. *Computers and Education: Artificial Intelligence*, 6, Article 100214. <https://doi.org/10.1016/j.caeai.2024.100214>

Krakowska, M. & Zych, M. (2025). (Un)conventional ways of dialogic information retrieval. *Information Research*, 30(CoLIS), 121–140.

Li, F., Yan, X., Su, H., Shen, R., & Mao, G. (2025). An assessment of human–AI interaction capability in the generative AI era: The influence of critical thinking. *Journal of Intelligence*, 13(6), Article 62. <https://doi.org/10.3390/jintelligence13060062>

Lin, Y.-C., Neville, J., Stokes, J. W., Yang, L., Safavi, T., Wan, M., Counts, S., Suri, S., Andersen, R., Xu, X., Gupta, D., Jauhar, S. K., Song, X., Buscher, G., Tiwary, S., Hecht, B., & Teevan, J. (2024). Interpretable user satisfaction estimation for conversational systems with large language models. In *Proceedings of the 62nd Annual Meeting of the Association for Computational Linguistics (Volume 1: Long Papers)* (pp. 11100–11115). Association for Computational Linguistics.

Long, D., & Magerko, B. (2020). What is AI literacy? Competencies and design considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (Paper 599). Association for Computing Machinery.

Miao, F. & Shiohira, K. (2024). AI competency framework for students. UNESCO.

Smith, A. G. (2012). Internet search tactics. *Online Information Review*, 36, 7–20.

Tibau, M., Siqueira, S. W. M., & Nunes, B. P. (2024). ChatGPT for chatting and searching: Repurposing search behavior. *Library & Information Science Research*, 46(4), Article 101331. <https://doi.org/10.1016/j.lisr.2024.101331>