
| **RESEARCH ARTICLE**

From Echo Chambers to Conversational Reinforcement: Rethinking Trust and Epistemic Closure in Conversational AI

Nishant Upadhyay¹ ✉ and **Mohd Aleem Khan²**

¹²*Department of Mass Communication and Media Studies, Central University of Punjab*

Corresponding Author: Nishant Upadhyay, **E-mail:** photobynishant@gmail.com

| **ABSTRACT**

Conversational artificial intelligence (AI) is changing information seeking by combining information retrieval, response generation, and follow-up questioning within a single interaction. Conversational systems enable users to refine questions and develop an inquiry through successive turns, in contrast to traditional search, where users navigate across multiple sources. This study looks at how this interactive format might influence engagement with conflicting interpretations, trust, and information seeking. The paper develops conversational reinforcement as a conceptual lens through a focused conceptual narrative review of literature on conversational AI, information seeking, trust, confirmation bias, selective exposure, sycophancy, stance adaptation, and epistemic environments. The framework suggests that while accommodating or adaptive AI responses may impact subsequent interaction and reinforce the current conversational frame, users' prior interpretive positions may shape their queries. According to current research, users' initial positions may be reinforced by sycophantic and stance-adaptive responses, while conversational search may increase confirmatory querying. However, since too much accommodation may diminish perceived authenticity, trust does not always rise with agreement. Conversational reinforcement is thus distinguished from the traditional notions of echo chambers, filter bubbles, and epistemic bubbles in the paper. It suggests that under certain circumstances, repeated accommodation, ongoing interaction, and minimal engagement with opposing interpretations may contribute to epistemic narrowing rather than viewing epistemic closure as an inevitable result of conversational AI. The paper makes the case that the cumulative trajectory of human-AI interaction should be examined in addition to conversational AI's generated responses.

| **KEYWORDS**

Conversational AI, conversational reinforcement, trust, epistemic closure, confirmation bias

| **ARTICLE INFORMATION**

ACCEPTED: August 11, 2026 **PUBLISHED:** September 07, 2026 **DOI:** <https://doi.org/10.61424/ijah.v4i3.997>

1. Introduction

Conversational AI is revolutionising the way people look for and access information. Systems like ChatGPT, Gemini, and Claude enable users to ask follow-up questions, refine requests, and continue the conversation within a single conversation, in contrast to traditional search engines that offer users a variety of sources to review. As a result, obtaining information can become a continuous conversation rather than a cycle of moving between various sources (Mo et al., 2025; Zhou & Li, 2026).

This change is crucial because users don't enter such conversations without preconceived notions. People's questions, preferred information, and responses to information that contradicts existing views can all be influenced

by confirmation bias (Lopez-Lopez et al., 2025). Additionally, according to recent experimental evidence and compared to traditional search, users may ask more confirmatory questions when using LLM-powered conversational search, particularly when the system supports their preexisting beliefs (Sharma et al., 2024).

This process is further complicated by trust. Reliance on automated systems can be influenced by trust, and users' opinions of the system can be shaped by recurring interactions and the features of interactive agents (Lee & See, 2004; Rheu et al., 2021). According to Carro (2024) and Sun & Wang (2026), excessive agreeableness may actually diminish perceived authenticity rather than increase trust. Conversational adjustment, trust, and users' preexisting beliefs are therefore not directly related.

In light of this, this paper develops conversational reinforcement as a conceptual lens to investigate the potential connections between user prior beliefs, conversational adjustments, trust, and information-seeking behaviour through repeated human- conversational AI interactions. This idea describes a possible process where the interaction is shaped by the user's preexisting interpretive position, the AI adapts or develops to that position, and further interactions proceed within a similar framework. In some cases, this could lead to epistemic narrowing by decreasing interaction with conflicting interpretations.

The paper addresses three research questions:

RQ1. How does conversational AI alter the conditions and dynamics of contemporary information seeking?

RQ2. How can trust, conversational accommodation, and users' pre-existing beliefs interact across repeated human-AI exchanges?

RQ3. How might conversational reinforcement contribute to epistemic narrowing, and how does this differ from conventional echo chambers, filter bubbles, and epistemic bubbles?

Through a targeted narrative synthesis of previous research on conversational AI, information search, trust, confirmation bias, sycophancy, and epistemic environments, this paper aims to develop a conceptual explanation of these relationships. This paper identifies situations in which frequent conversational accommodation may contribute to limiting information engagement rather than assuming conversational AI is intrinsically prone to epistemic closure.

2. Methodology

To examine the relationship between conversational AI, trust, conversational reinforcement, and epistemic closure, this paper employs a targeted conceptual narrative review. Rather than collecting primary data, the paper synthesises theoretical and empirical literature to develop a conceptual understanding of recurring interactions between humans and AI.

Using combinations of terms associated with conversational AI, conversational search, trust in AI, sycophancy, confirmation bias, selective exposure, echo chambers, filter bubbles, and epistemic closure, relevant studies were found through focused searches of academic databases and scholarly discovery platforms. The selection of foundational works and recent studies that are directly relevant to this argument was given priority.

The literature was organised thematically around information seeking, trust and reliance, sycophancy and stance adaptation, confirmation bias and selective exposure, and epistemic environments. Findings were compared to identify conceptual connections, tensions, and gaps. This synthesis informed the development of conversational reinforcement as a conceptual lens rather than a formally tested causal theory.

3. Findings

3.1 *Conversational AI and the Transformation of Information Seeking*

A simple example helps to clarify this distinction. For instance, a Google search for "how does climate change affect Indian agriculture?" may yield reports, news articles, research papers, and institutional webpages. The same user could use ChatGPT to ask this question, ask for a more straightforward explanation, focus on wheat production in India, and then ask for sources all within the same conversation. As a result, information seeking goes beyond

merely looking through specific documents to cultivating an interest or question through conversation. Instead of considering every search as a distinct request, conversational search encourages multi-turn interaction, asking pertinent questions, and producing feedback (Mo et al., 2025).

This format may also be preferred by users since it makes it easier to locate and comprehend multiple results. Intentions to use generative AI are linked to dissatisfaction with traditional search, specifically information overload and task-specific inaccuracy; the perceived value of the information obtained is influenced by its quality and interactivity (Zhou & Li, 2026). However, this does not imply that generative AI has entirely supplanted conventional search. While AI is increasingly working alongside search engines in users' larger information environments, search engines are still crucial information sources (Lund et al., 2026).

Users' search behaviours also clearly reflect this difference. Without having to go back to a new results page, conversational systems enable users to refine their questions, add context, and change the direction of their responses. According to research on ChatGPT, users employ these repetitive prompting techniques to make their information needs clear and work with responses, but they also raise issues with conflicting opinions and limited source visibility (Peng et al., 2026). Therefore, conventional search techniques do not vanish; rather, they are modified for the conversational setting, where refining questions, assessing answers, and determining credibility all become part of the continuous dialogue (Tibau et al., 2024).

There is a communication aspect to this change as well. Users are exposed to a variety of visible sources through traditional search, which can be contrasted according to authors, organisations, and opposing viewpoints. An interactive system, on the other hand, usually displays a synthesised response, which obscures the choice and combination of underlying information during the dialogue. Further evidence from people looking for health information indicates that, in comparison to other online sources, users who get information from LLM-based chatbots might not do as much cross-checking (Yun & Bickmore, 2025). Thus, the significance lies not in the fact that conversational AI has taken the place of search engines, but rather in the way that information access has evolved, with retrieval, synthesis, and subsequent questioning increasingly taking place within the same conversation. This begs the question of how the development of trust and dependence is affected by repeated encounters with the same conversational source.

3.2 Trust in Human–Conversational AI Interaction

Trust becomes important when users rely on a system whose internal processes they cannot fully observe. It is generally understood to be different from reliance: dependence refers to the behaviour that results, whereas trust refers to the expectation that an agent will assist in achieving a goal in the face of uncertainty. Previous studies have connected competence, kindness, integrity, predictability, and dependability with trust; these expectations grow with experience (Mayer et al., 1995; Rempel et al., 1985; Lee & See, 2004).

Due to the possibility of inaccurate calibration, this distinction is crucial. While excessive reliance can result in reliance that surpasses the system's actual capabilities, automation research separates appropriate reliance from misuse and disuse (Parasuraman & Riley, 1997; Lee & See, 2004). This problem is especially pertinent to conversational AI since users must deal with both a conversational interface and system performance. Although trust in automated systems should not be equated with interpersonal trust, research on human-computer interaction (HCI) demonstrates that people can react socially to computers (Reeves & Nass, 1996; Nass & Moon, 2000; Madhavan & Wiegmann, 2007). Conversational agents may still be viewed by users as instrumental, particularly in situations where task completion and performance are still crucial (Clark et al., 2019).

However, the manner in which a conversation is conducted can influence trust. Social intelligence, communication style, anthropomorphic presentation, nonverbal cues, and performance are identified as important factors in a review of experimental studies on conversational agents; in certain contexts, socio-emotional dialogue and culturally familiar communication can boost trust (Rheu et al., 2021). For instance, the review describes a study where participants were more receptive to an agent's recommendation when they interacted with a communication style that was culturally familiar. In addition to performance, accuracy and perceived competence are important, and transparency and consideration of the agent's error-handling techniques can impact trust (Rheu et al., 2021). This

implies that users may evaluate a conversational system based on both its features and the flow of the conversation rather than just its factual performance.

Trust has behavioural repercussions as well. While more general AI research links trust to system adoption and continued use, studies on conversational agents link trust to cooperation and dependability (Hayashi & Wakabayashi, 2017; Kunder et al., 2019; Yang & Wibowo, 2022; Sundar & Kim, 2019). Furthermore, trust is not static. Previous interactions shape subsequent decisions, and trust can be established, maintained, strengthened, weakened, or rebuilt over time (Lee & See, 2004; Langer et al., 2023; Skjuve et al., 2021; Ng & Zhang, 2025).

This temporal dimension draws attention to a weakness in the present work. While a number of predictors of trust in conversational AI have been found in previous research, the conversational processes that lead to the development of trust in repeated interactions have received relatively less attention. In their review, Ng & Zhang (2025), for instance, discovered that very few studies looked at the mediating mechanisms of trust. Therefore, the question is not whether or not users trust conversational AI, but rather how that trust can be shaped by repeated interactions and, consequently, how users continue to use the system as a source of information. This provides the conceptual entry point for 'conversational reinforcement.'

3.3 From Conversational Reinforcement to Epistemic Closure

Conversational AI raises more issues than just the fact that these systems occasionally concur with users. The more crucial question is what effects this repeated accommodation might have when a user's starting position interacts with the system's responses in turn. Conversational systems enable users to refine questions within the same conversation, formulating each new question based on prior responses, in contrast to traditional search, where users navigate between multiple sources. We refer to a possible process as "conversational reinforcement", in which the user's preexisting ideological stance influences the discussion while the AI adopts or develops that stance, making it easier to uphold and possibly lowering attention to opposing interpretations.

Instead of starting with the system, this process can start with the user. Confirmation bias can influence how we formulate questions, choose information that supports our beliefs, and react to information that contradicts our preconceived notions (Lopez-Lopez et al., 2025). According to their analysis, information seeking powered by GenAI can interact with these tendencies through three pressure points: question formulation, content selection, and resistance to information that challenges beliefs. The implication for conversational AI is that the user's initial framing can dictate the direction of the conversation, so the AI does not need explicit consent to start reinforcement.

This possibility is supported by evidence from conversational search. Compared to users of traditional web search, users of LLM-powered conversational search demonstrated more attitude-consistent information search, and an opinionated system that supported users' preexisting opinions further increased selective exposure (Sharma et al., 2024). Additionally, there was no discernible overall shift in the participants' self-reported attitudes following the search task. Conversational AI can therefore affect what users search for, pay attention to, and keep asking questions about, rather than automatically persuading them.

An explicit type of such accommodation is offered by sycophancy. According to users' stated beliefs, LLMs can alter their responses. For example, they can repeat factual errors that users have expressed and alter correct answers after users challenge them (Sharma et al., 2024). The results of Sun and Wang's research point to a more comprehensive conversational process: stance adaptation can strengthen initial positions and lessen psychological reactance, but its effects are contingent on conversational behaviour and do not correspond to a general shift in attitude (Sun & Wang, 2026). Thus, conversational reinforcement extends beyond overt praise. It can also happen when the user's questions already contain implicit beliefs that are repeatedly reinforced by the system.

This relationship is made more difficult by trust. From consensus to trust, there is no direct route. While attitude changes can result in varying trust outcomes depending on how the interaction is expressed, overt flattery can lower perceived and behavioural trust (Carro, 2024; Sun & Wang, 2026). Therefore, it is preferable to think of trust as a potential prerequisite for long-term interaction rather than as a necessary outcome of consensus. Although

apparent alignment can also raise suspicions, a user may keep using a system because its responses appear reliable, helpful, and consistent.

The epistemic concern is heightened when these results are considered in conjunction with past studies on ideological amplification and information filtering. Sunstein (2007) demonstrates how one-sided arguments and increased support can reinforce initial preferred positions in discussions among like-minded people. While Nguyen (2018) distinguishes between an echo chamber, where external sources are actively disqualified or questioned, and an epistemological bubble, where pertinent voices are excluded, Pariser (2011)'s filter-bubble description concentrates on technologically mediated personalisation. Interactions between humans and AI don't always closely resemble these structures. Neither a group of like-minded people nor a clear exclusion of outside sources is present. Instead, the direction of the discussion has the potential to restrict ideas.

Think about a user who feels that a specific treatment is dangerous. The user may pose a number of queries regarding the treatment's risks, unintended consequences, and alternatives rather than looking for evidence to support or refute it. The user may come across more in-depth information that seems to support the initial concern if the system consistently reinforces the same framing. Instead of being contained in a single response, this issue is cumulative. The "Chat-Chamber" effect, which links trust, alignment with pre-existing assumptions, and limited cross-validation to the acceptance of misinformation, is an AI-specific manifestation of this concern (Jacob et al., 2025).

The current framework does not attempt to rename the "chat-chamber" phenomenon. Instead, it focuses on the conversational process that can give rise to such an environment. The proposed sequence is as follows:

prior position → belief-consistent query → accommodating response → continued interaction → reinforcement of the existing frame → potentially weaker engagement with competing information.

Consequently, "conversational reinforcement" is suggested as a conceptual lens rather than an established cause-and-effect theory. Its constituent processes, confirmation bias, selective exposure, stance adaptation, sycophancy, and conversational trust, are supported by the literature, but it does not provide a general route to epistemic closure. However, conversational AI can aid in epistemic narrowing without mimicking the social structure of a conventional echo chamber when faced with frequent modifications and little interaction with alternatives.

4. Discussion

This paper's primary contribution is to shift the conversation about conversational AI from whether or not it produces an "echo chamber" to how conversational interaction can support preexisting ideologies. Existing research has examined confirmation bias, selective exposure, sycophancy, attitude adaptation, and trust as related but largely distinct processes. Overall, these results imply that information seeking via conversational AI is influenced by the path the system takes in subsequent turns as well as what it generates (Lopez-Lopez et al., 2025; Sharma et al., 2024; Sharma et al., 2024).

"Conversational reinforcement" offers a helpful conceptual distinction in this situation. While flattery illustrates how AI adjusts to a specific viewpoint, selective exposure characterises users' preference for information that aligns with prior beliefs. Another example of how the system's alignment with the user's position can affect the conversation's direction is stance adaptation (Sharma et al., 2024; Sun & Wang, 2026). These processes are linked by the current concept at the level of repeated interactions. Reinforcement can also happen when the user's prior beliefs are reinforced by subsequent responses; explicit consent is not necessary at every stage. Consequently, its conceptual contribution is to the interactional relationship between flattery and confirmation bias rather than to give them new names.

Similar conditions apply to the role of trust. Reliance is influenced by trust, which can be built through prior interactions. However, the degree to which users' perceptions align with the system's actual capabilities determines the appropriate level of trust (Lee & See, 2004). Additionally, while the effects of stance adaptation vary with conversational style and authenticity, overt flattery can diminish perceived trust (Carro, 2024; Sun & Wang, 2026). It follows that the framework should not be interpreted as agreement → trust → closure. Although trust is neither a

necessary prerequisite for epistemic narrowness nor an automatic result of rapport, it may aid in sustaining engagement with the conversational source.

Instead of replacing conventional explanations of information restriction, this framework expands upon them. While Pariser concentrates on individual filtering, Sunstein uses limited argument pools and corroboration to connect like-minded discussions to strong preconditions. Nguyen makes a distinction between echo chambers where external sources are actively disregarded and omission-based epistemological bubbles (Sunstein, 2007; Pariser, 2011; Nguyen, 2020). These structures are not always precisely replicated in a conversational exchange. Rather, narrowness may result from the dialogue's trajectory, where the same ideological framework continues to influence the questions and topics covered.

This distinction is especially pertinent to the growing body of "chat-chamber" literature. Jacob et al. (2025) associate the acceptance of false information with limited cross-verification, alignment with preconceptions, and trust. Conversational reinforcement provides a potential explanation for how such an environment might emerge during a conversation instead of renaming this phenomenon. This framework is purposefully provisional: excessive flattery may erode trust, and conversational exploration may increase confirmatory questions without immediately changing attitudes (Sharma et al., 2024; Carro, 2024).

As a result, there are conceptual and methodological ramifications for communication research. Instead of concentrating only on specific AI responses or overall levels of trust, future research should look at the conversation as the unit of analysis. Changes in verification behaviour, exposure to counterarguments, confidence in prior beliefs, and reliance on independent sources could all be observed in repeated exchanges. Instead of presuming that conversational reinforcement always results in closure, this would enable it to be investigated as a possible route to epistemic narrowing.

5. Conclusion

This paper highlights how conversational AI is transforming information seeking by integrating information retrieval, response generation, and follow-up queries into a single conversation. Conversational systems enable the pursuit of information through successive turns, in contrast to traditional search, where users navigate between competing and visible sources. As a result, a new environment is created where users' preconceived notions can influence the discussion, and answers can affect further questions.

This paper proposes "conversational reinforcement" as a conceptual lens to comprehend this process. This idea explains how repeated accommodations can reinforce a user's preexisting ideological frame by combining research on confirmation bias, selective exposure, flattery, attitude adaptation, trust, and conversational search. This argument does not imply that agreement inevitably breeds trust or that conversational AI inevitably produces echo chambers. Instead, conversational interaction can lead to epistemic narrowing when there is frequent accommodation and little interaction with opposing interpretations.

Additionally, conversational reinforcement is not the same as the well-known ideas of filter bubbles and echo chambers. It focuses on the recursive nature of human-AI interactions rather than just social homogeneity or personalisation. Therefore, rather than being an established cause-and-effect theory, this idea should be viewed as a tentative analytical lens. Finding a procedure that can be empirically tested is what makes it valuable.

Future studies should investigate how verification behaviour, engagement with counterarguments, confidence in prior beliefs, and reliance on independent sources are impacted by repeated human-AI interactions. Such studies could shed light on when conversational reinforcement is still a beneficial aspect of interpersonal communication and when it could lead to epistemic closure.

Disclosure statement

Funding: The author(s) received no financial support for the research, authorship, and/or publication of this article.

Conflict of interest: No potential conflict of interest was reported by the author(s).

ORCID

Nishant Upadhyay: <https://orcid.org/0009-0005-2464-7970>

Mohd Aleem Khan: <https://orcid.org/0000-0003-1694-7699>

Publisher's Note: All claims expressed in this article are solely those of the authors and do not necessarily represent those of their affiliated organizations, or those of the publisher, the editors and the reviewers

Artificial Intelligence (AI) Use Disclosure: The authors declare that no artificial intelligence tools were used in the preparation of this manuscript.

References

- [1] Carro, M. V. (2024). *Flattering to deceive: The impact of sycophantic behaviour on user trust in large language models*. arXiv. <https://arxiv.org/abs/2412.02802>
- [2] Clark, L., Pantidi, N., Cooney, O., Doyle, P., Garaialde, D., Edwards, J., Spillane, B., Gilmartin, E., Murad, C., Munteanu, C., Wade, V., & Cowan, B. R. (2019). What makes a good conversation? Challenges in designing truly conversational agents. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems* (pp. 1–12). Association for Computing Machinery. <https://doi.org/10.1145/3290605.3300705>
- [3] Hayashi, Y., & Wakabayashi, K. (2017). Can AI become a reliable source to support human decision making in a court scene? In *Companion of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing* (pp. 195–198). Association for Computing Machinery.
- [4] Jacob, C., Kerrigan, P., & Bastos, M. (2025). The chat-chamber effect: Trusting the AI hallucination. *Big Data & Society*, 12(1), Article 20539517241306345. <https://doi.org/10.1177/20539517241306345>
- [5] Kunderling, T., Wintersberger, P., & Riener, A. (2019). (Over) Trust in automated driving: The sleeping pill of tomorrow? In *Extended Abstracts of the 2019 CHI Conference on Human Factors in Computing Systems* (Article LBW2418). Association for Computing Machinery.
- [6] Langer, M., König, C. J., Back, C., & Hemsing, V. (2023). Trust in artificial intelligence: Comparing trust processes between human and automated trustees in light of unfair bias. *Journal of Business and Psychology*, 38(3), 493–508.
- [7] Lee, J. D., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors*, 46(1), 50–80. https://doi.org/10.1518/hfes.46.1.50_30392
- [8] Lopez-Lopez, E., Abels, C. M., Holford, D., Herzog, S. M., & Lewandowsky, S. (2025). Generative artificial intelligence–mediated confirmation bias in health information seeking. *Annals of the New York Academy of Sciences*, 1550(1), 23–36. <https://doi.org/10.1111/nyas.15413>
- [9] Lund, B. D., Teel, Z. A., Mohammed, Y., Jagathpally, A., & Wang, T. (2026). Artificial intelligence (AI) and information seeking: A comparative exploration of AI chatbots, search engines, and library resources as information sources among university students. *Journal of Librarianship and Information Science*. Advance online publication. <https://doi.org/10.1177/09610006261438484>
- [10] Madhavan, P., & Wiegmann, D. A. (2007). Similarities and differences between human–human and human–automation trust: An integrative review. *Theoretical Issues in Ergonomics Science*, 8(4), 277–301. <https://doi.org/10.1080/14639220500337708>
- [11] Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734. <https://doi.org/10.5465/amr.1995.9508080335>
- [12] Mo, F., Mao, K., Zhao, Z., Qian, H., Chen, H., Cheng, Y., ... Nie, J. Y. (2025). A survey of conversational search. *ACM Transactions on Information Systems*, 43(6), 1–50. <https://doi.org/10.1145/3759453>
- [13] Nass, C., & Moon, Y. (2000). Machines and mindlessness: Social responses to computers. *Journal of Social Issues*, 56(1), 81–103. <https://doi.org/10.1111/0022-4537.00153>
- [14] Ng, S. W. T., & Zhang, R. (2025). Trust in AI chatbots: A systematic review. *Telematics and Informatics*, 97, Article 102240. <https://doi.org/10.1016/j.tele.2025.102240>
- [15] Nguyen, C. T. (2020). Echo chambers and epistemic bubbles. *Episteme*, 17(2), 141–161.
- [16] Parasuraman, R., & Riley, V. (1997). Humans and automation: Use, misuse, disuse, abuse. *Human Factors*, 39(2), 230–253.
- [17] Pariser, E. (2011). *The filter bubble: What the internet is hiding from you*. Penguin UK.
- [18] Peng, W., Meng, J., Tang, L., Zou, W., & Issaka, B. (2026). How do lay users seek information with ChatGPT: An in-situ interview study. *International Journal of Human–Computer Interaction*, 42(13), 10520–10536.

- [19] Reeves, B., & Nass, C. I. (1996). *The media equation: How people treat computers, television, and new media like real people and places*. CSLI Publications.
- [20] Rempel, J. K., Holmes, J. G., & Zanna, M. P. (1985). Trust in close relationships. *Journal of Personality and Social Psychology*, 49(1), 95–112. <https://doi.org/10.1037/0022-3514.49.1.95>
- [21] Rheu, M., Shin, J. Y., Peng, W., & Huh-Yoo, J. (2021). Systematic review: Trust-building factors and implications for conversational agent design. *International Journal of Human–Computer Interaction*, 37(1), 81–96. <https://doi.org/10.1080/10447318.2020.1807710>
- [22] Sharma, M., Tong, M., Korbak, T., Duvenaud, D., Askeel, A., Bowman, S., ... Perez, E. (2024, May). Towards understanding sycophancy in language models [Conference presentation]. International Conference on Learning Representations 2024, Vienna, Austria.
- [23] Sharma, N., Liao, Q. V., & Xiao, Z. (2024, May). Generative echo chamber? Effect of LLM-powered search systems on diverse information seeking. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems* (pp. 1–17). Association for Computing Machinery. <https://doi.org/10.1145/3613904.3642459>
- [24] Skjuve, M., Følstad, A., Fostervold, K. I., & Brandtzaeg, P. B. (2021). My chatbot companion — A study of human-chatbot relationships. *International Journal of Human-Computer Studies*, 149, Article 102601. <https://doi.org/10.1016/j.ijhcs.2021.102601>
- [25] Sun, Y., & Wang, T. (2026, April). Be friendly, not friends: How LLM sycophancy shapes user trust. In *Proceedings of the 2026 CHI Conference on Human Factors in Computing Systems* (pp. 1–15). Association for Computing Machinery. <https://doi.org/10.1145/3772318.3791079>
- [26] Sundar, S. S., & Kim, J. (2019). Machine heuristic: When we trust computers more than humans with our personal information. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems* (Article 538). Association for Computing Machinery.
- [27] Sunstein, C. R. (2007). Ideological amplification. *Constellations*, 14(2), 273–279.
- [28] Tibau, M., Siqueira, S. W. M., & Nunes, B. P. (2024). ChatGPT for chatting and searching: Repurposing search behaviour. *Library & Information Science Research*, 46(4), Article 101331. <https://doi.org/10.1016/j.lisr.2024.101331>
- [29] Yang, R., & Wibowo, S. (2022). User trust in artificial intelligence: A comprehensive conceptual framework. *Electronic Markets*, 32(4), 2053–2077. <https://doi.org/10.1007/s12525-022-00592-6>
- [30] Yun, H. S., & Bickmore, T. (2025). Online health information-seeking in the era of large language models: Cross-sectional web-based survey study. *Journal of Medical Internet Research*, 27, Article e68560.
- [31] Zhou, T., & Li, S. (2026). Understanding user switch of information seeking: From search engines to generative AI. *Journal of Librarianship and Information Science*, 58(1), 696–708. <https://doi.org/10.1177/09610006241244800>