

Health Inquiry with AI: How Empathetic Expression and Conversational Contexts Shape Users' Communicative Acts

Xi Zheng

Department of Computer Science
City University of Hong Kong
Hong Kong SAR, China
zheng.xi@my.cityu.edu.hk

Can Liu

School of Creative Media
City University of Hong Kong
Hong Kong SAR, China
canliu@cityu.edu.hk

Xuyu Yang

School of Creative Media
City University of Hong Kong
Hong Kong SAR, China
xuyuyang2-c@my.cityu.edu.hk

Yuhan Luo*

Department of Computer Science
City University of Hong Kong
Hong Kong SAR, China
yuhanluo@cityu.edu.hk

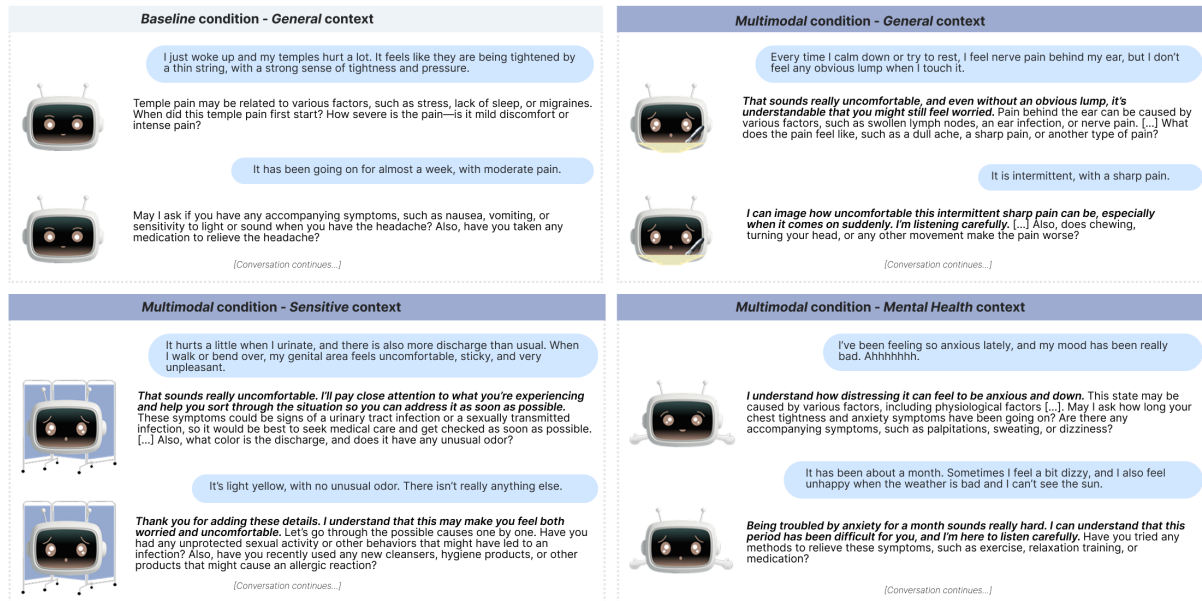


Figure 1: Example conversations across modality and conversational contexts. Empathetic verbal responses are bold-italicized.

Abstract

As online health information-seeking shifts to conversational AI, high-quality information retrieval increasingly relies on users' "communicative acts" (proactively sharing and seeking information)—similar to how effective diagnosis and personalized guidance are elicited in patient-clinician communication. Drawing on health communication research, this study examines how a chatbot's modality

of empathetic expression (Verbal, Visual, Multimodal) and the conversational context (General, Sensitive, Mental Health) influence these acts through a $2 \times 2 \times 3$ within-subjects experiment ($N = 48$). The results revealed that while verbal and multimodal empathy significantly increased reply length, communicative acts were largely shaped by conversational context, with Sensitive context triggering more question-asking and Mental Health context leading to heightened concerns, assertive responses, and unprompted information disclosure. Combined with qualitative findings, we discuss design implications for building context-sensitive AI health inquiry systems that can encourage active user participation.

*Corresponding Author.



This work is licensed under a Creative Commons Attribution 4.0 International License.
UbiComp Companion '26, Shanghai, China
© 2026 Copyright held by the owner/author(s).
ACM ISBN 979-8-4007-2533-3/2026/10
<https://doi.org/10.1145/3798063.3838813>

CCS Concepts

• Human-centered computing → HCI design and evaluation methods.

Keywords

Multimodal interaction, Health communication, Empathy, Chatbot

ACM Reference Format:

Xi Zheng, Xuyu Yang, Can Liu, and Yuhua Luo. 2026. Health Inquiry with AI: How Empathetic Expression and Conversational Contexts Shape Users' Communicative Acts. In *Companion of the 2026 ACM International Joint Conference on Pervasive and Ubiquitous Computing (UbiComp Companion '26), October 11–15, 2026, Shanghai, China*. ACM, New York, NY, USA, 5 pages. <https://doi.org/10.1145/3798063.3838813>

1 Introduction and Backgrounds

Driven by the rapid evolution of large language models (LLMs), online information-seeking, particularly health inquiries, is shifting away from traditional search engines to conversational AI [1, 35, 39]. Yet, the quality of information retrieved from AI depends not only on how well the model is designed and trained [10, 28], but also individual's ability to communicate with (or prompt) the AI, such as asking follow-up questions or adding contextual details about their health conditions [13, 15, 18, 22]. This is similar to effective communication with human clinicians, where the patient's active inquiry and disclosure are essential for the clinician to fully understand their needs and concerns, and deliver more appropriate, personalized guidance [4, 6, 31].

Health communication researchers have characterized this active participation as “**communicative acts**,” which include asking questions, expressing concerns, and making assertive responses [31]. These acts allow patients to articulate their needs, worries, preferences, and interpretations during clinical encounters [31]. For example, asking questions signals patient's knowledge gap; expressing concerns reveals their emotional distress and worries; and assertive responses clarifies their beliefs and treatment expectations [6, 31]. Prior work has shown that patients who engage in these communicative acts are more likely to elicit informative and accommodating responses, as these acts provide clinicians with cues needed to tailor explanations and decision-making support [4, 6, 30].

According to existing research in health communication, individuals' active participation in their inquiries is influenced not only by their own health literacy and communication skills [6, 12, 30], but also the clinicians' communication styles and the nature of the inquiry topics [9, 19, 23, 30]. For example, when clinicians express empathy through verbal acknowledgment (e.g., “*I understand how pain it is*” [25, 37]) or supportive facial expressions (e.g., attentive nodding, a comforting smile [2]), patients feel validated and emotionally safe, and thus are willing to ask more questions and openly express worries [6, 9, 30]. On the other hand, prior work has shown that patients' communicative acts and needs vary across clinical contexts. For example, sensitive health topics may make patients hesitant to disclose information because of stigma, embarrassment, and privacy concerns [23]. In contrast, patients with mental health concerns may want to seek opportunities for disclosing their emotional distress and gaining validation and supports [19].

Extending these findings from human-human interaction to human-AI interaction, we see the opportunities to encourage communicative acts through optimizing the design of AI systems for health inquiry [14, 29, 34, 38]. In this regard, we explore how the AI's **modality of empathetic expression** and the **conversational**

context influence individuals' communicative acts. Specifically, we examine three common health inquiry contexts: *General*, *Sensitive*, and *Mental Health*, and design four conditions where the chatbot displayed empathy through different modality configurations: *Baseline* (no explicit empathy cues), *Verbal empathy only*, *Visual empathy only*, and *Multimodal empathy* (both verbal and visual cues) conditions. By comparing participants' communicative acts across these conditions and contexts, we found that empathetic verbal expressions increased participants' reply length, whereas the specific communicative acts users performed were more strongly shaped by the conversational context.

This work contributes to HCI and the IMWUT community by providing an empirical understanding of user engagement with health inquiry AI through the lens of communicative acts. Our findings offer design implications for building context-sensitive AI systems for health inquiry to better elicit users' active information sharing and seeking behaviors that can result in higher quality of retrieved information.

2 Method

2.1 Data Source

We analyzed conversation logs collected from a previously conducted 2 (*without* vs. *with verbal empathy*) $\times$ 2 (*without* vs. *with visual empathy*) $\times$ 3 (*general*, *sensitive*, *mental health context*) within-subjects experiment ($N = 48$), followed by optional debriefing interviews ($n = 16$). The original experiment examined the effects of empathetic expression of health inquiry chatbot on people's empathy and authenticity perception, as well as their trust, satisfaction, and dependency toward the chatbot [38]. This paper reports a complementary analysis of the same dataset, shifting the analytical focus to users' interactions with the chatbot. Particularly, we compare their communicative acts across study conditions.

To maintain topic consistency and protect participant privacy, the health contexts were presented as hypothetical scenarios with visual illustrations and text descriptions rather than prompting users to discuss their actual medical conditions. This approach has been widely used in prior work to examine user perceptions in sensitive settings [7, 14, 34].

Through a custom-built interface, participants first read a hypothetical scenario, adopted the scenario character's perspective, chatted with the chatbot about the described health issue, completed post-task ratings, and then proceeded to the next scenario. The chatbot's responses were generated to manipulate whether empathy was expressed through verbal cues, visual cues, both channels, or neither. The presentation of these conditions were counterbalanced and the design of empathetic expression passed manipulation check. Full details regarding the experimental setup, procedure, and the design of the empathetic cues are available in [38]. The experiment was conducted in mainland China and was approved by the ethics review committee in the authors' institution. Figure 1 illustrates the conversation examples collected from the experiment.

2.2 Data Analysis

To understand participants' reactions to different empathy modalities, we analyzed only the dialogue occurring after the chatbot's first reply (after participants were exposed to the condition). We first coded each of the messages based on Street's definition of

communicative acts [31], and then compared their likelihood to appear across conditions. Next, we triangulate these quantitative data with participants' qualitative pre- and post-task responses and interviews to contextualize the findings.

2.2.1 Communicative Acts Coding. According to Street et al., communicative acts—the process of active patient participation during encounters with clinicians—include *asking questions*, *expressing concerns*, and *making assertive responses*, which are essential for effectively addressing one's informational needs [31]. Additionally, we extended this framework by incorporating *providing unprompted contextual information*, to capture the proactive sharing of personal background or symptom history that complements the picture of their health condition and concerns [6, 36].

Two researchers independently coded the same one-third of the dataset based on the four dimensions of communicative acts, and then compared their coding results to discuss and resolve any discrepancies. Next, they independently coded an additional one-third of the dataset to achieve an inter-coder reliability of Cohen's $\kappa = .94$, then the first author finished coding the remaining data. In the following, we describe each of these dimensions in detail with example quotes from participants' message. Note that one message may contain multiple communicative acts.

- **Asking questions:** seeking information, clarification, or guidance in an interrogative form (e.g., “*sometimes I feel a little dizzy. Does that count as another symptom?*” and “*I haven't showered for several days. Could that be the main reason, or is it something related to my body itself?*”).
- **Expressing concerns:** expressing worry, anxiety, fear, frustration, anger, or other negative emotions (e.g., “*What could be causing this? I'm actually quite scared right now.*” and “*Fatigue has been constant. You can probably tell that I really have too much on my plate.*”).
- **Making assertive responses:** stating one's beliefs, preferences, interpretations, or expectations of the condition, which may involve disagreement and suggestions (e.g., “*I did not use medication to help me sleep, because I think that if I can fall asleep on my own, I should avoid relying on medication.*” and “*I don't really think it's likely to be an infection? It was just contact with unclean underwear.*”).
- **Providing unprompted contextual information:** mentioning contexts about one's situations that are not included in the scenario descriptions of the study materials or asked by the chatbot, which sometimes involved participants' own experience or expectations (e.g., “*Recently the weather has been gloomy, and I have been feeling very down. I do not want to do anything, my chest feels tight, and I keep feeling anxious.*” and “*No, I don't want to exercise because I feel like I don't have much energy in my body.*”).

2.2.2 Quantitative Analysis. We first compared participants' reply length (total number of Chinese characters) per message using a linear mixed-effects model [20]. The model included the experimental condition and their interaction as fixed effects, with each participant as a random effect. Next, we analyzed the above coded communicative acts separately using mixed-effects logistic regression, with the same fixed- and random-effects structure. Finally, we

conducted robustness checks to assess the stability of the results. To validate the trends observed about the communicative acts across conversational contexts, we conducted an independent analysis of participants' initial messages sent to the chatbot (upon seeing the conversational context and before exposing to the empathy modality) and verified the results were consistent.

3 Results

3.1 Reply Length

Compared to the *Baseline* condition, participants typed significantly longer replies (5–6 more characters on average) in the *Verbal-only* ($\beta = 5.54, p = .024$) and *Multimodal* conditions ($\beta = 5.60, p = .024$). No significant effect was found between the *Visual-only* and the *Baseline* condition ($\beta = 2.01, p = .598$). Regarding conversational contexts, replies were significantly longer in the *Mental Health* context than in the *General* ($\beta = 7.27, p < .001$) and *Sensitive* contexts ($\beta = 6.06, p < .001$). The *General* and *Sensitive* contexts did not significantly differ from each other ($\beta = 1.21, p = .476$). No significant interaction effect was identified between empathy cues and health contexts on reply length ($\chi^2(6) = 0.50, p = .998$).

3.2 Communicative Acts

Across the four communicative acts—asking questions, expressing concerns, making assertive responses, and providing unprompted contextual information—most measures significantly differed across health contexts, while only a few minor differences were identified across empathy modality conditions. Specifically, none of the contrasts in empathy modality reached statistical significance, despite that the *Verbal-only* ($\Delta = 0.96, p = .5407$) and *Multimodal* conditions ($\Delta = 1.10, p = .4918$) showed higher rates of question-asking than the *Baseline* condition and the *Visual-only* condition showed a higher likelihood of providing unprompted contextual information than *Baseline* ($\Delta = 0.70, p = .3325$). Below, we focused on describing the varied interaction patterns across conversational contexts.

3.2.1 More Question Asking in Sensitive Context. Participants were more likely to ask questions in *Sensitive* context than in both the *General* context ($\Delta = 2.9, p = .0002$) and the *Mental Health* context ($\Delta = 1.5, p = .0068$). From their conversation logs and interpretations noted in the scenarios, we found that in the *Sensitive* context, participants were navigating a tense intersection of social stigma and acute anxiety, noting that “*This symptom feels really embarrassing*” and sometimes raising multiple questions at once: “*Could this have been caused by not changing my underwear, or cleaning my private area over the past few days?*” As a result, they exhibited an urgent, protective need to immediately identify the exact cause and resolve the issue (e.g., “*What is happening? What exactly are these small bumps? How can they go away?*”).

3.2.2 Heightened Concerns, More Assertive Responses and Disclosure in Mental Health Inquiries. In *Mental Health* context, participants expressed higher intensity of concerns than in *General* ($\Delta = 3.07, p < .0001$) and *Sensitive* contexts ($\Delta = 4.08, p < .0001$); they were also more likely to input assertive responses ($\Delta = 2.80$ vs. *General*, $p = .0018$) and provide unprompted contextual information (*General*: $\Delta = 0.79, p = .0284$; *Sensitive*: $\Delta = 1.45, p = .0003$). The qualitative data reveals that because mental health struggles lack visible biomarkers, participants relied heavily on narrative disclosure, situating their distress in relation to its duration, triggers,

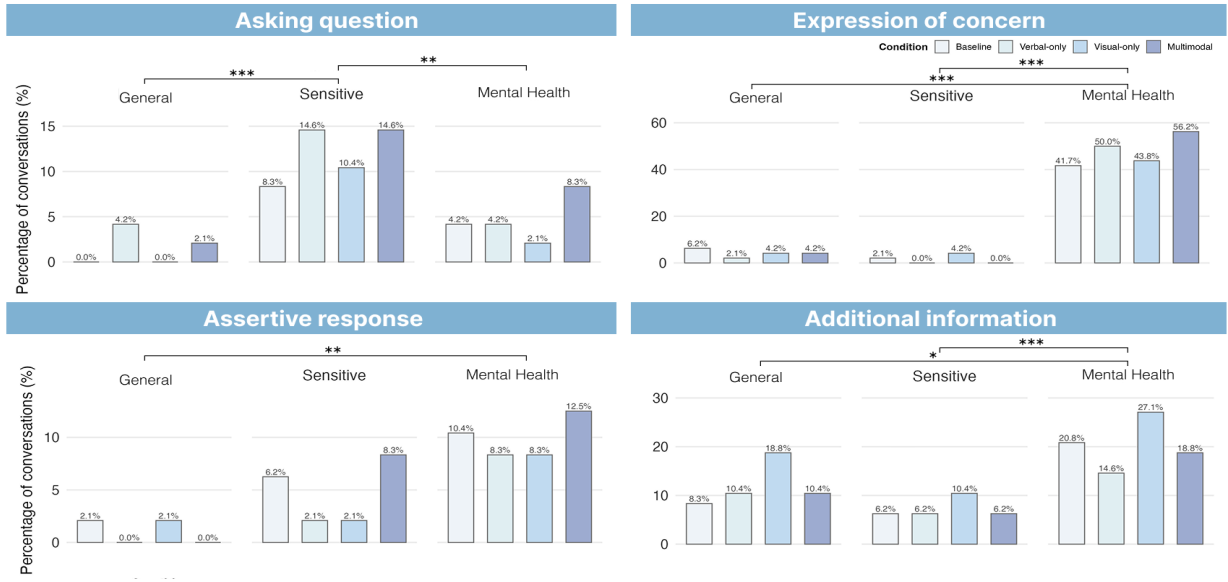


Figure 2: Main result visualization of users' communicative acts across empathy conditions and health contexts. Statistically significant differences across conversational contexts are marked. Note. * $p < .05$, ** $p < .01$, * $p < .001$.**

and daily impacts to make their internal states understandable to the AI. This included keyword labeling of their conditions such as “anxious,” “frustrated,” “helpless,” and “exhausted,” and sometimes “self-diagnosis” that attributed their conditions to “stress,” “depression,” “ADHD,” or “difficulties with emotion regulation.” Moreover, to convey the severity of their concerns, they shared unprompted details on how their sleep, work performance, relationships, and daily functioning are impacted.

4 Discussion

Our analysis showed that while the modality of empathetic expression increased participants reply length, such increase did not necessarily contribute to communicative acts. This finding suggested that while empathetic cues (particularly verbal and multimodal ones) are effective at keeping users engaged and encouraging them to type more, the additional text may consist of conversational politeness or social reciprocity [3, 8] rather than proactive information-seeking or assertive behaviors. On the other hand, we found that participants' communicative acts were mainly driven by the conversational context—the complexity and urgency of their health conditions. This indicates that active participation, such as intensive information-seeking, sharing one's diagnostic interpretations, and volunteering unprompted context, is fundamentally a problem-driven coping mechanism rather than a reaction to the AI's social behaviors [16, 21, 30].

Taken together, we see the opportunities to tailor the AI's communication strategy and inquiry interface based on the conversational contexts, so as to maximize the benefits of the communicative acts that users exhibit. For example, when confronted with the acute anxiety and stigma of a *Sensitive* condition, the system can deploy structural UI scaffolding, such as interactive FAQ accordion blocks [24, 32] or etiology-tracing visualizations to rapidly provide cognitive closure [33]. When navigating *Mental Health* burdens,

the AI may prioritize understanding and compassion to alleviate their concerns [19]. This sheds light on combining empathetic expression of AI; although in our findings, the effects of empathy modality on communicative acts were not significant, it was still essential for keeping participants engaged. At the same time, how to dynamically process the information from users' self-disclosure and handle their self-diagnosis without over-stepping into clinical diagnosis or dismissively shutting down their intuition remains an open challenge for future research [11].

5 Limitations and Future Work

To maintain consistent experimental conditions, we held the AI responses constant; as a result, our study does not establish whether the characterized “communicative acts” improve the model's response quality [5, 27]. Future work should involve clinical professionals in assessing AI responses to these communicative acts with real-world conversation logs. In addition, the results of this study may not generalize across cultures [17, 26].

Acknowledgments

We thank our participants for their time and contributions. We also appreciate the thoughtful suggestions provided by the anonymous reviewers. This work is supported by City University of Hong Kong (#7005997).

References

- [1] Mahyar Abbasian, Iman Azimi, Amir M. Rahmani, and Ramesh Jain. 2024. Conversational Health Agents: A Personalized LLM-Powered Agent Framework. arXiv:2310.02374 [cs.CL] <https://arxiv.org/abs/2310.02374>
- [2] Janet B. Bavelas and Nicole Chovil. 2000. Visible acts of meaning: An integrated message model of language in face-to-face dialogue. *Journal of Language and Social Psychology* 19, 2 (2000), 163–194. doi:10.1177/0261927X00019002001
- [3] Robert Bowman, Orla Cooney, Joseph W. Newbold, Anja Thieme, Leigh Clark, Gavin Doherty, and Benjamin Cowan. 2024. Exploring how politeness impacts

- the user experience of chatbots for mental health support. *International Journal of Human-Computer Studies* 184 (2024), 103181. doi:10.1016/j.ijhcs.2023.103181
- [4] Donald J Cegala and Douglas M Post. 2009. The impact of patients' participation on physicians' patient-centered communication. *Patient education and counseling* 77, 2 (2009), 202–208.
 - [5] Yufeng Du, Minyang Tian, Srikanth Ronanki, Subendhu Rongali, Sravan Babu Bodapati, Aram Galstyan, Azton Wells, Roy Schwartz, Eliu A Huerta, and Hao Peng. 2025. Context Length Alone Hurts LLM Performance Despite Perfect Retrieval. In *Findings of the Association for Computational Linguistics: EMNLP 2025*, Christos Christodoulopoulos, Tanmoy Chakraborty, Carolyn Rose, and Violet Peng (Eds.). Association for Computational Linguistics, Suzhou, China, 23281–23298. doi:10.18653/v1/2025.findings-emnlp.1264
 - [6] Thomas A D'Agostino, Thomas M Atkinson, Lauren E Latella, Madeline Rogers, Dana Morrissey, Antonio P DeRosa, and Patricia A Parker. 2017. Promoting patient participation in healthcare interactions through communication skills training: a systematic review. *Patient education and counseling* 100, 7 (2017), 1247–1257.
 - [7] Uğur Genç and Himanshu Verma. 2024. Situating Empathy in HCI/CSCW: A Scoping Review. *Proc. ACM Hum.-Comput. Interact.* 8, CSCW2, Article 513 (Nov. 2024), 37 pages. doi:10.1145/3687052
 - [8] Yuanyuan Guo, Linlin Xu, and Chaoyou Wang. 2025. Exploring the effect of empathic response and its boundaries in artificial intelligence service recovery. *Journal of Retailing and Consumer Services* 82 (2025), 104065. doi:10.1016/j.jretconser.2024.104065
 - [9] M Jawad Hashim. 2017. Patient-centered communication: basic skills. *American family physician* 95, 1 (2017), 29–34.
 - [10] Zhe He, Balu Bhasuran, Qiao Jin, Shubo Tian, Karim Hanna, Cindy Shavor, Lisbeth Garcia Arguello, Patrick Murray, and Zhiyong Lu. 2024. Quality of answers of generative large language models versus peer users for interpreting laboratory test results for lay patients: evaluation study. *Journal of medical Internet research* 26 (2024), e56655.
 - [11] Bright Huo, Amy Boyle, Nana Marfo, Wimonchat Tangamornsuksan, Jeremy P Steen, Tyler McKechnie, Yung Lee, Julio Mayol, Stavros A Antoniou, Arun James Thirunavukarasu, et al. 2025. Large language models for chatbot health advice studies: a systematic review. *JAMA Network Open* 8, 2 (2025), e2457879.
 - [12] Marra G Katz, Terry A Jacobson, Emir Veledar, and Sunil Kripalani. 2007. Patient literacy and question-asking behavior during the medical encounter: a mixed-methods analysis. *Journal of general internal medicine* 22, 6 (2007), 782–786.
 - [13] Wojciech Kusa, Edoardo Mosca, and Aldo Lipani. 2023. “Dr LLM, what do I have?”: The Impact of User Beliefs and Prompt Formulation on Health Diagnoses. In *Proceedings of the Third Workshop on NLP for Medical Conversations*, Sopan Khosla (Ed.). Association for Computational Linguistics, Bali, Indonesia, 13–19. doi:10.18653/v1/2023.nlpmc-1.2
 - [14] Bingjie Liu and S Shyam Sundar. 2018. Should Machines Express Sympathy and Empathy? Experiments with a Health Advice Chatbot. *Cyberpsychology, Behavior, and Social Networking* 21, 10 (Oct. 2018), 625–636. PMID: 30334655. doi:10.1089/cyber.2018.0110
 - [15] Siru Liu, Aileen P Wright, Allison B McCoy, Sean S Huang, Julian Z Jenkins, Josh F Peterson, Yaa A Kumah-Crystal, William Martinez, Babatunde Carew, Dara Mize, et al. 2024. Using large language model to guide patients to create efficient and comprehensive clinical care message. *Journal of the American Medical Informatics Association* 31, 8 (2024), 1665–1670.
 - [16] Rebecca D. McMullan, David Berle, Sandra Arnáez, and Vladan Starcevic. 2019. The relationships between health anxiety, online health information seeking, and cyberchondria: Systematic review and meta-analysis. *Journal of Affective Disorders* 245 (Feb. 2019), 270–278. doi:10.1016/j.jad.2018.11.037
 - [17] T. Mojaverian, T. Hashimoto, and H. S. Kim. 2013. Cultural differences in professional help seeking: a comparison of Japan and the U.S. *Frontiers in Psychology* 3 (2013), 615. Published 2013 Jan 11. doi:10.3389/fpsyg.2012.00615
 - [18] Jeremy Y Ng. 2025. Prompt engineering for generative artificial intelligence chatbots in health research: a practical guide for traditional, complementary, and integrative medicine researchers. *Integrative Medicine Research* (2025), 101222.
 - [19] Daisy Parker, Richard Byng, Chris Dickens, and Rose McCabe. 2020. Patients' experiences of seeking help for emotional concerns in primary care: doctor as drug, detective and collaborator. *BMC family practice* 21, 1 (2020), 35.
 - [20] R Core Team. 2024. *R: A Language and Environment for Statistical Computing*. R Foundation for Statistical Computing, Vienna, Austria. <https://www.R-project.org/>
 - [21] Stephen A Rains and Riva Tukachinsky. 2015. Information seeking in uncertainty management theory: exposure to information about medical uncertainty and information-processing orientation as predictors of uncertainty management success. *Journal of Health Communication* 20, 11 (2015), 1275–1286.
 - [22] Felix G Rebitschek, Alessandra Carella, Silja Kohlrausch-Pazin, Michael Zitzmann, Anke Steckelberg, and Christoph Wilhelm. 2025. Evaluating evidence-based health information from generative AI using a cross-sectional study with laypeople seeking screening information. *npj Digital Medicine* 8, 1 (2025), 343.
 - [23] Pamela Sankar and Nora L Jones. 2005. To tell or not to tell: primary care patients' disclosure deliberations. *Archives of internal medicine* 165, 20 (2005), 2378–2383.
 - [24] Janet E Sansoni, Pam Grootemaat, and Cathy Duncan. 2015. Question prompt lists in health consultations: a review. *Patient education and counseling* 98, 12 (2015), 1454–1464.
 - [25] Lennart Seitz. 2024. Artificial empathy in healthcare chatbots: Does it feel authentic? *Computers in Human Behavior: Artificial Humans* 2, 1 (2024), 100067. doi:10.1016/j.chbah.2024.100067
 - [26] N. Sheeran, L. Jones, R. Pines, B. Jin, A. Pamoso, J. Eigeland, and M. Benedetti. 2023. How culture influences patient preferences for patient-centered care with their doctors. *Journal of Communication in Healthcare* 16, 2 (2023), 186–196. Epub 2022 Jul 13. doi:10.1080/17538068.2022.2095098
 - [27] Freda Shi, Xinyun Chen, Kanishka Misra, Nathan Scales, David Dohan, Ed Chi, Nathanael Schärli, and Denny Zhou. 2023. Large language models can be easily distracted by irrelevant context. In *Proceedings of the 40th International Conference on Machine Learning (Honolulu, Hawaii, USA) (ICML '23)*. JMLR.org, Article 1291, 18 pages.
 - [28] Karan Singhal, Shekoofeh Azizi, Tao Tu, S Sara Mahdavi, Jason Wei, Hyung Won Chung, Nathan Scales, Ajay Tanwani, Heather Cole-Lewis, Stephen Pfohl, et al. 2023. Large language models encode clinical knowledge. *Nature* 620, 7972 (2023), 172–180.
 - [29] Silke ter Stal, Gerbrich Jongbloed, and Monique Tabak. 2021. Embodied Conversational Agents in eHealth: How Facial and Textual Expressions of Positive and Neutral Emotions Influence Perceptions of Mutual Understanding. *Interacting with Computers* 33, 2 (07 2021), 167–176. doi:10.1093/iwc/iwab019
 - [30] Richard L. Jr. Street, Howard S. Gordon, Martha M. Ward, Edward Krupat, and Richard L. Kravitz. 2005. Patient participation in medical consultations: why some patients are more involved than others. *Medical Care* 43, 10 (Oct. 2005), 960–969. doi:10.1097/01.mlr.0000178172.40344.70
 - [31] Richard L. Jr. Street and Barbara Millay. 2001. Analyzing patient participation in medical encounters. *Health Communication* 13, 1 (2001), 61–73. doi:10.1207/S15327027HC1301_06
 - [32] Marguerite Tracy, Julie Ayre, Olivia Mac, Tessa Copp, Emerita Lyndal Trevena, and Heather Shepherd. 2022. Question prompt lists and endorsement of question-asking support patients to get the information they seek—A longitudinal qualitative study. *Health Expectations* 25, 4 (2022), 1652–1663.
 - [33] Claire Woodcock, Brent Mittelstadt, Dan Busbridge, and Grant Blank. 2021. The impact of explanations on layperson trust in artificial intelligence-driven symptom checker apps: experimental study. *Journal of medical Internet research* 23, 11 (2021), e29386.
 - [34] Yue You, Chun-Hua Tsai, Yao Li, Fenglong Ma, Christopher Heron, and Xinning Gui. 2023. Beyond Self-diagnosis: How a Chatbot-based Symptom Checker Should Respond. *ACM Trans. Comput.-Hum. Interact.* 30, 4, Article 64 (Sept. 2023), 44 pages. doi:10.1145/3589959
 - [35] Hye Sun Yun and Timothy Bickmore. 2025. Online health information-seeking in the era of large language models: cross-sectional web-based survey study. *Journal of medical Internet research* 27 (2025), e68560.
 - [36] J.D. Zamfirescu-Pereira, Richmond Y. Wong, Bjoern Hartmann, and Qian Yang. 2023. Why Johnny Can't Prompt: How Non-AI Experts Try (and Fail) to Design LLM Prompts. In *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems (Hamburg, Germany) (CHI '23)*. Association for Computing Machinery, New York, NY, USA, Article 437, 21 pages. doi:10.1145/3544548.3581388
 - [37] Junbo Zhang, Qi Chen, Jiandong Lu, Xiaolei Wang, Luning Liu, and Yuqiang Feng. 2024. Emotional expression by artificial intelligence chatbots to improve customer satisfaction: Underlying mechanism and boundary conditions. *Tourism Management* 100 (2024), 104835. doi:10.1016/j.tourman.2023.104835
 - [38] Xi Zheng, Xuyu Yang, Can Liu, and Yuhua Luo. 2026. Disentangling the Effects of Verbal and Visual Cues for Expressing Empathy on Conversational User Experience with Health Inquiry Chatbots. Available at SSRN (2026). SSRN 6278334.
 - [39] Tao Zhou and Songtao Li. 2026. Understanding user switch of information seeking: From search engines to generative AI. *Journal of librarianship and information science* 58, 1 (2026), 696–708.