

Rapid Judgment in the Age of AI: A Thin-Slice Study of Virtual Humans to Support Caregivers

Sharon Mozgai¹, Lila Rabinovich², Caroline P. Nguyen²,
Katie Seymour³, Sujeet Rao², Todd Keitz³, and Marco Angrisani²

¹University of Southern California ICT, Los Angeles, CA 90094, USA

²University of Southern California, Los Angeles, CA 90089, USA

³My Care Friends, Boynton Beach, FL, 33472 USA

ABSTRACT

Caregivers often face barriers to accessing timely and sustained support, motivating the development of scalable digital interventions such as virtual humans. As exposure to artificial intelligence (AI) systems becomes increasingly widespread, users may bring different expectations to their initial encounters with such systems. This is particularly consequential in caregiving contexts, where first impressions can shape trust, perceived appropriateness, and willingness to engage with supportive digital interventions. This study examines how individuals evaluate a caregiver-support virtual human under thin-slice conditions, based on a 10-second exposure. Participants (N = 354) viewed a brief video and provided open-ended impressions, which were analyzed using inductive thematic coding and aggregated into higher-level constructs. Responses were compared across groups defined by prior AI exposure. Results indicate that capability-related concerns were significantly more prevalent among participants with no exposure, whereas relational concerns increased among those with prior use. These findings suggest that initial evaluations vary systematically with prior experience, with implications for designing caregiver-facing systems as AI exposure increases.

Keywords: Virtual human, Thin-slice perception, Caregiving support, Embodied conversational agents, Qualitative analysis

INTRODUCTION

Caregivers provide essential support to individuals with chronic illness, disability, and aging-related needs, often while navigating substantial emotional, cognitive, and logistical demands. Despite the availability of informational and psychosocial resources, many caregivers experience barriers to accessing timely, personalized, and sustained support, particularly outside of formal care settings (Adelman et al., 2014; Eden and Schulz, 2016). Digital interventions have been proposed as a scalable means of extending support to caregivers; however, their effectiveness depends not only on content, but also on whether users are willing to engage with them at all.

Recent advances in artificial intelligence and real-time rendering have enabled the development of virtual humans capable of delivering interactive, socially grounded support. These systems are increasingly positioned as

front-facing interfaces for health and wellness interventions, including those targeting stress, coping, and behavior change. In this context, the initial encounter between a user and a virtual human may be consequential. Unlike traditional digital tools, virtual humans present as social actors, and users may rapidly form impressions about their credibility, warmth, and appropriateness for the caregiving context. These early impressions may shape expectations for interaction, perceived usefulness, and willingness to return.

At the same time, public exposure to AI systems has increased substantially, particularly through the proliferation of conversational agents and generative models. This shift raises an open question: do individuals with prior experience interacting with AI interpret virtual humans differently at first encounter compared to those without such exposure? Differences in familiarity may influence expectations about system capabilities, perceived intentionality, and tolerance for ambiguity or imperfection. Understanding how these factors shape initial perceptions is particularly important for caregiver-facing systems, where trust and perceived alignment with user needs are critical.

To examine these issues, the present study adopts a thin-slice paradigm, presenting participants with a brief (10-second) clip of a virtual human designed to support caregivers. Thin-slice methodology is well-suited to this context, as prior work demonstrates that individuals form rapid, socially meaningful judgments about others based on minimal behavioral cues (Ambady and Rosenthal, 1992). This approach is particularly appropriate for virtual humans, which are designed to function as social actors and are evaluated using many of the same perceptual and interpersonal heuristics applied in human–human interaction. As a result, even brief exposure may be sufficient for users to infer qualities such as warmth, credibility, and intent, which are central to engagement in caregiving contexts. The present study examines how such initial impressions vary as a function of prior exposure to AI. Participants provided qualitative judgments of the agent, enabling analysis of how users interpret its role, intent, and interpersonal qualities from minimal exposure.

This work makes three contributions. First, it extends prior research on virtual humans in health contexts by focusing specifically on caregiver support, a domain with distinct relational and emotional demands. Second, it applies a thin-slice methodology to examine how users form impressions of virtual humans under conditions of minimal exposure. Third, it introduces AI exposure as a lens for interpreting variation in these judgments, contributing to a growing body of work on how user experience with AI shapes interaction with emerging systems.

RELATED WORK

Caregivers play a critical role in supporting individuals with chronic illness, disability, and aging-related needs, often under conditions of sustained psychological and logistical burden. Digital interventions have emerged as a scalable means of providing education, coping support, and behavioral

guidance. Systematic reviews indicate that such interventions can improve caregiver outcomes including psychological well-being, self-efficacy, skills, quality of life, and perceived social support, while also highlighting challenges related to usability, personalization, and sustained engagement (Zhai et al., 2023). These findings underscore the importance of designing caregiver-facing systems that are not only functional but also immediately interpretable and engaging at first encounter.

Embodied conversational agents (ECAs) and virtual humans have been increasingly deployed in healthcare contexts to deliver coaching, assessment, and psychosocial support. Prior reviews identify a range of design features, including embodiment, verbal and nonverbal behavior, appearance, and interaction modality, as shaping user experience and engagement (Laranjo et al., 2018; Ter Stal et al., 2020b). More recent work continues to emphasize that while ECAs show promise across chronic disease management and mental health applications, the field lacks clear guidance on how specific design features influence user perceptions and downstream outcomes (Jiang et al., 2024). In caregiver contexts specifically, emerging qualitative work suggests that trust, perceived empathy, and clarity of role are central to acceptance of virtual assistants (Sezgin et al., 2025).

This literature motivates closer examination of how users form initial impressions of such systems. A substantial body of psychological research demonstrates that individuals form meaningful social judgments from brief samples of behavior, often referred to as “thin slices.” In a seminal meta-analysis, Ambady and Rosenthal (1992) showed that judgments based on exposures as short as a few seconds can predict interpersonal outcomes with above-chance accuracy. Subsequent work has examined the conditions under which such judgments are reliable, highlighting the roles of slice length, behavioral channel (e.g., visual vs. auditory), and construct being assessed (Carney et al., 2007). Reviews of thin-slice methodology further indicate that brief exposures can yield stable impressions even when observers have limited contextual information, making thin slices a useful paradigm for studying rapid social perception (Murphy and Hall, 2021). This literature provides a theoretical and methodological foundation for examining how users interpret virtual humans from short clips.

Research in human-agent interaction suggests that thin-slice processes extend to virtual characters. Users rapidly form impressions of virtual humans based on nonverbal cues such as gaze, facial expression, posture, and interpersonal distance (Cafaro et al., 2012; Cafaro et al., 2016). Experimental studies demonstrate that even brief exposures (e.g., 10–15 seconds) are sufficient for users to infer personality traits, social attitudes, and relational stance. In eHealth contexts, first impressions are further shaped by agent characteristics such as apparent age, gender, and professional role, which influence perceived credibility and relatability (Ter Stal et al., 2020a; Ter Stal et al., 2020b).

Taken together, this work highlights not only how such impressions are formed, but also raises the question of their consequences for user engagement and interaction in healthcare contexts where trust and perceived safety influence engagement and disclosure. Work with virtual humans has shown

that users may be more willing to disclose sensitive information to agents than to human interviewers under certain conditions, particularly when the interaction is perceived as nonjudgmental (Lucas et al., 2014). At the same time, broader research on AI in healthcare indicates that trust is shaped by both technical factors and human-facing design elements, including perceived empathy, transparency, and alignment with user expectations (Shevtsova et al., 2024). For caregiver-facing systems, early impressions may therefore play a critical role in determining whether users engage with or disengage from the intervention.

Additionally, user familiarity with AI systems has been identified as a potential moderator of trust and acceptance. Studies on AI literacy suggest that individuals with greater exposure to AI technologies tend to exhibit higher baseline trust, while also demonstrating more calibrated expectations depending on context and perceived risk (Huang and Ball, 2024). Generational and experiential differences in interaction with AI assistants have similarly been associated with variation in perceived usefulness, ease of use, and willingness to adopt such systems (Alanzi et al., 2023). These findings suggest that prior exposure to AI may influence how users interpret the behavior and intent of virtual humans, particularly during initial encounters.

Prior work establishes that (1) caregivers can benefit from digital support, (2) virtual humans are a promising modality for delivering such support, (3) first impressions can be formed rapidly from thin slices of behavior, and (4) these impressions influence trust, engagement, and disclosure. Additionally, emerging evidence suggests that user familiarity with AI may shape expectations and interpretations of AI-driven systems. However, there is limited research examining how individuals form qualitative judgments of caregiver-support virtual humans based on very brief exposures, or how these judgments differ as a function of prior AI exposure. The present study addresses this gap by analyzing thin-slice impressions derived from a 10-second clip of a virtual human designed to support caregivers.

METHODS

Participants and Recruitment

Participants ($N = 384$; 56% female; 63% Non-Hispanic White, 23% Non-Hispanic Non-White, 14% Hispanic; M age = 59 years) were recruited were drawn from the Understanding America Study (UAS), a probability-based online panel of approximately 14,000 U.S. residents administered by the University of Southern California's Center for Economic and Social Research. Panel members are recruited using address-based sampling, with provisions to include households without prior internet access. The panel maintains publicly documented recruitment, participation, and attrition metrics. Post-stratification weights are applied to align the sample with national population benchmarks. Evaluations of both weighted and unweighted distributions indicate strong correspondence with U.S. demographic characteristics across multiple variables. Further details on sampling, recruitment, and weighting procedures are provided in the Appendix. All participants provided informed consent in accordance with panel procedures.

Stimulus and Procedure

Participants completed the study online via the panel platform. Prior to viewing the stimulus, participants were presented with a brief introduction framing the concept of a virtual human as a digital assistant capable of providing personalized support, guidance, and resources for caregiving. This description was intended to establish a general understanding of the system without providing detailed information about its capabilities or limitations. Participants were then shown a brief (10-second) video clip depicting a virtual human designed for caregiving contexts. In the clip, the agent introduces itself and describes its role as a source of support for caregivers, using verbal and nonverbal behaviors typical of supportive, health-oriented interactions. The stimulus was designed to approximate an initial encounter scenario, including conversational tone, facial expression, and posture consistent with a socially interactive agent (see Figure 1). Following exposure, participants were asked to provide an open-ended explanation of their likelihood to use a virtual human for caregiving support. The open-ended prompt elicited participants' immediate impressions of the agent, including perceptions of its role, intent, interpersonal qualities, and overall appropriateness for caregiving support. No additional contextual information about the agent was provided beyond the initial framing, in order to capture first impressions under conditions of minimal exposure.



Figure 1: Virtual human used in the study. Participants viewed a brief (10-second) introductory clip in which the agent presented itself as a source of support for caregivers.

The study employed a thin-slice paradigm, in which participants form judgments based on brief exposures to behavioral stimuli. Prior research demonstrates that individuals can generate rapid and stable social judgments from short observational windows (Ambady and Rosenthal, 1992). This approach is particularly appropriate for virtual human evaluation, as users often form initial impressions within the first moments of interaction, which may shape subsequent engagement, trust, and willingness to continue use. The thin-slice design, therefore, enables examination of early perceptual processes under conditions that approximate the first encounter with a virtual agent.

Open-ended responses were analyzed using an inductive thematic analysis approach (Braun & Clarke, 2006). Two independent coders first coded approximately one-third of the dataset ($n = 118$) to identify emergent themes and collaboratively develop a preliminary codebook. Inter-coder agreement on this double-coded subset was 83.4%, and coding discrepancies were discussed and resolved through consensus. The finalized codebook was then applied to the remaining responses, which were divided between the two coders. The coding framework was iteratively refined throughout the process to incorporate newly identified themes, with refinements discussed and agreed upon by both coders. The full codebook, including definitions, is provided in the Appendix.

Construct Development

Codes were grouped into higher-level constructs reflecting broader conceptual themes, including trust and risk concerns, relational and human-centered limitations, capability and usability barriers, and overall attitudinal responses toward AI. This process of aggregating codes into analytically meaningful categories is consistent with established approaches to qualitative data reduction and abstraction (Miles et al., 2014). Table 1 illustrates the mapping of individual codes to higher-level constructs. To enable comparison across participants, constructs were operationalized as binary indicators at the participant level, such that a value of 1 indicated the presence of at least one corresponding code in a participant's response, and 0 indicated its absence. This approach allows qualitative data to be systematically transformed for quantitative analysis while preserving the interpretive basis of the original coding (Creswell and Plano Clark, 2017). Exposure to AI and text-generation tools was self-reported and grouped into three categories: no prior exposure, awareness without use, and prior use.

Table 1: Code aggregation and exposure grouping: hierarchical organization of qualitative codes into higher-level constructs for analysis. Individual codes derived from inductive thematic analysis were grouped into four overarching categories.

Construct	Codes and Code Names	Operational Definition
Trust / Risk / Societal Concern	Codes 1, 2, 9, 10, 11: Distrust of AI; Privacy/Security Concerns; Environmental Impact; Conditional Trust; Broader Social Concern	Participant expresses concern related to risk, safety, trust, or societal consequences of AI use
Relational / Human-Centered Concerns	Codes 3, 5, 8: Preference for Human Connection; Lack of Empathy/ Authenticity; Anti-Embodiment/Uncanny Valley	Participant expresses concern about AI's ability to function as a social or relational agent
Capability / Usability Barriers	Codes 6, 12, 16, 19, 20: Limited VH Capability; Caregiving Target Limitations; Misunderstanding of VH Role; Limited Own Capability; Generational Resistance	Participant expresses limitations related to usability, understanding, or ability to engage with AI/VH systems

(Continued)

Table 1: Continued.

Construct	Codes and Code Names	Operational Definition
Attitudes / Orientation	Codes 7, 13, 14, 15, 17, 18, 22: Preference for ChatGPT; Willing/Need to Learn; Positive/Optimistic; General Negative Response; No Perceived Need; Preference for Self-Reliance; Not Sure	Participant expresses an overall stance, orientation, or positioning toward AI use
Residual / Contextual	Codes 4, 21: Negative Past Chatbot Experiences; Other – Uncodable	Contextual or uncategorized responses

RESULTS

Distribution of Thematic Content

At the level of coded themes, attitudinal content was most prevalent (42.4% of code applications), followed by relational concerns (26.9%), trust and risk concerns (12.3%), and capability barriers (9.4%). The distribution of thematic content varied by exposure group. Responses from participants with no exposure more frequently included capability-related themes (18.3% of code applications) than those from participants with awareness (7.5%) or prior use (7.2%). In contrast, relational concerns were more prevalent among participants with prior use (33.0%) compared to those with no exposure (25.6%) or awareness without use (25.1%). Trust and risk themes also increased across exposure groups, from 6.1% among participants with no exposure to 15.5% among users. Because multiple themes could be assigned to a single response, these distributions are descriptive and are not treated as independent observations for statistical testing.

Participant-Level Differences in Thematic Categories

To examine whether thematic prevalence differed across exposure groups while preserving independence of observations, analyses were conducted at the participant level. Each participant was coded as expressing a thematic category if at least one instance of that category was identified in their response.

Capability Barriers. A significant association was observed between exposure and capability-related themes, $\chi^2(2, N = 354) = 7.96, p = .019$, Cramér's $V = .15$. Capability-related themes were most prevalent among participants with no exposure (21.1%) and substantially less common among those with awareness (8.9%) or prior use (10.0%).

Relational Concerns. A marginal association was observed for relational themes, $\chi^2(2, N = 354) = 5.88, p = .053$, Cramér's $V = .13$. The proportion of participants expressing relational concerns increased from 25.4% among those with no exposure to 41.4% among those with prior use.

Trust and Risk Concerns. No statistically significant association was observed for trust and risk themes, $\chi^2(2, N = 354) = 5.15, p = .076$, Cramér's $V = .12$. However, a monotonic increase across exposure groups was observed descriptively (7.0%, 16.4%, and 20.0%, respectively).

Attitudinal Responses. Attitudinal themes were prevalent across all exposure groups (42.3%–51.6%) but did not vary meaningfully. No inferential test was conducted for this category due to minimal observed differences.

Decomposition of Relational Concerns. To better understand the marginal association observed for relational themes, follow-up analyses examined individual codes within this category. Among the three relational codes—preference for human connection, lack of empathy or authenticity, and anti-embodiment or uncanny valley—only uncanny valley concerns demonstrated a meaningful pattern across exposure groups. The proportion of participants expressing discomfort with human-like agents increased from 8.5% among participants with no exposure to 17.1% among those with prior use, $\chi^2(2, N = 354) = 5.77, p = .056$, Cramér's $V = .13$. In contrast, preference for human connection ($\chi^2(2) = 0.32, p = .85$) and lack of empathy or authenticity ($\chi^2(2) = 1.76, p = .41$) did not vary significantly across groups.

Summary of Results. Three primary patterns emerged. First, themes related to capability were significantly more prevalent among participants with no prior exposure, indicating that early-stage concerns center on perceived ability to use or understand AI systems. Second, relational concerns increased with exposure, driven specifically by embodiment-related discomfort rather than general dissatisfaction with AI interaction. Third, trust and risk themes increased descriptively but were not statistically significant, and attitudinal themes remained stable across groups.

DISCUSSION

AI Literacy as a Mechanism of Changing Perception

The findings suggest that differences in how individuals evaluate virtual humans reflect a process of increasing AI literacy, rather than simple variation in preference. Participants with no prior exposure primarily expressed self-oriented concerns, focusing on their ability to use or understand the technology. In contrast, participants with prior use shifted toward system-oriented critique, evaluating how virtual humans behave and where they fall short. Here, exposure enabled more specific and differentiated critiques, particularly in the domains of relational and embodiment-related features. These findings indicate that familiarity with AI systems relates to more critical evaluation of physical embodiment and the role of virtual humans in caregiving support.

From Capability Barriers to Design Critique

The observed reduction in capability-related concerns with increasing exposure suggests that early barriers to engagement are primarily associated with perceived usability and self-efficacy. As these barriers diminish, users appear to reorient their evaluations toward the design and behavior of the system itself. The emergence of relational concerns among more experienced users reflects this shift. Rather than rejecting AI systems outright, participants

with prior use identified specific limitations in how virtual humans attempt to simulate human interaction. In particular, the increase in embodiment-related discomfort suggests that human-like features may introduce new points of friction as users become more familiar with AI systems.

Implications for Design in a Post-AI Novelty Context

These findings point to a transition in human–agent interaction, in which engagement is less driven by novelty and increasingly shaped by user expectations formed through repeated exposure to AI systems. First, systems must continue to support early-stage engagement by reducing perceived barriers to use. This includes clear onboarding, intuitive interaction design, and the minimization of assumptions about prior experience. Second, systems must also be designed to withstand more critical evaluation over time. As users gain experience, they appear less tolerant of mismatches between expected and observed behavior, particularly in domains involving social and relational interaction. Third, the observed increase in trust and risk concerns suggests that trust should be treated as an ongoing process rather than a static design feature. Systems may benefit from mechanisms that support transparency, expectation setting, and calibration of capabilities over time. Finally, the findings suggest that anthropomorphic design should be applied selectively. While embodiment may initially attract attention, it may also introduce new expectations that are difficult to meet, particularly for more experienced users.

LIMITATIONS AND FUTURE DIRECTIONS

Several limitations should be considered. First, exposure was used as a proxy for AI literacy and does not capture variation in depth or quality of experience. Future work should incorporate more granular measures, such as frequency of use or familiarity with specific systems. Second, the distribution of participants across exposure groups was uneven, which may limit statistical power for detecting differences involving smaller groups. Third, the coding framework reflects interpretive decisions made during qualitative analysis. Although systematically applied, alternative coding approaches may yield different thematic structures. Fourth, the use of binary indicators limits the ability to capture variation in the intensity or nuance of participant responses. Future work could explore more fine-grained modeling approaches. Fifth, the study examined responses to a single virtual human stimulus, limiting the generalizability of findings to other agent designs, appearances, or interaction styles. Future work should examine a broader range of virtual human characteristics and fidelity levels to assess the robustness of these patterns. Finally, the cross-sectional design limits causal interpretation. Longitudinal studies are needed to better understand how perceptions evolve as users gain experience with AI systems.

ACKNOWLEDGMENT

The authors acknowledge the support of Public Exchange at the USC Dornsife College of Letters, Arts and Sciences, which also provided funding for this study.

REFERENCES

- Adelman, R. D., et al. (2014). Caregiver burden: a clinical review. *JAMA*, 311(10), pp. 1052–1060.
- Alanzi, T., et al. (2023). AI-powered mental health virtual assistants' acceptance: an empirical study on influencing factors among generations X, Y, and Z. *Cureus*, 15(11).
- Ambady, N., Rosenthal, R. (1992). Thin slices of expressive behavior as predictors of interpersonal consequences: A meta-analysis. *Psychological Bulletin*, 111(2), p. 256.
- Braun, V., Clarke, V. (2006). Using thematic analysis in psychology. *Qualitative Research in Psychology*, 3(2), pp. 77–101.
- Cafaro, A., et al. (2016). First impressions in human–agent virtual encounters. *ACM Transactions on Computer-Human Interaction (TOCHI)*, 23(4), pp. 1–40.
- Cafaro, A., et al. (2012). “First impressions: Users’ judgments of virtual agents’ personality and interpersonal attitude in first encounters”, in: *International Conference on Intelligent Virtual Agents*. Springer, pp. 67–80.
- Carney, D. R., et al. (2007). A thin slice perspective on the accuracy of first impressions. *Journal of Research in Personality*, 41(5), pp. 1054–1072.
- Creswell, J. W., Plano Clark, V. L. (2017). *Designing and Conducting Mixed Methods Research*. Sage Publications.
- Eden, J., Schulz, R. (2016). *Families Caring for an Aging America*. National Academies of Sciences, Engineering, and Medicine.
- Huang, K.-T., Ball, C. (2024). The influence of AI literacy on user’s trust in AI in practical scenarios: A digital divide pilot study. *Proceedings of the Association for Information Science and Technology*, 61(1), pp. 937–939.
- Jiang, Z., et al. (2024). Embodied conversational agents for chronic diseases: Scoping review. *Journal of Medical Internet Research*, 26, e47134.
- Laranjo, L., et al. (2018). Conversational agents in healthcare: A systematic review. *Journal of the American Medical Informatics Association*, 25(9), pp. 1248–1258.
- Lucas, G. M., et al. (2014). It’s only a computer: Virtual humans increase willingness to disclose. *Computers in Human Behavior*, 37, pp. 94–100.
- Miles, M. B., et al. (2014). *Qualitative Data Analysis*. Sage.
- Murphy, N. A., Hall, J. A. (2021). Capturing behavior in small doses: A review of comparative research in evaluating thin slices for behavioral measurement. *Frontiers in Psychology*, 12, 667326.
- Sezgin, E., et al. (2025). Perceptions about the use of virtual assistants for seeking health information among caregivers of young childhood cancer survivors. *Digital Health*, 11, 20552076251326160.
- Shevtsova, D., et al. (2024). Trust in and acceptance of artificial intelligence applications in medicine: mixed methods study. *JMIR Human Factors*, 11(1), e47031.
- Ter Stal, S., et al. (2020a). Design features of embodied conversational agents in eHealth: a literature review. *International Journal of Human-Computer Studies*, 138, 102409.

- Ter Stal, S., et al. (2020b). Who do you prefer? The effect of age, gender and role on users' first impressions of embodied conversational agents in eHealth. *International Journal of Human-Computer Interaction*, 36(9), pp. 881–892.
- Zhai, S., et al. (2023). Digital health interventions to support family caregivers: An updated systematic review. *Digital Health*, 9, 20552076231171967.

APPENDIX A: PANEL RECRUITMENT AND WEIGHTING

Details about the panel's recruitment, sampling, and attrition rates, along with a full description of the weighting procedure and evaluations of demographic representativeness, are publicly available and regularly updated on the panel website <https://uasdata.usc.edu/page/Weights>.

APPENDIX B: CODEBOOK BARRIERS AND ATTITUDES TOWARD AI AND VH TECHNOLOGY

1. **Distrust of AI.** General skepticism or mistrust toward AI systems' reliability, accuracy, or intentions.
2. **Privacy/Security Concerns.** Worries about data protection, surveillance, or misuse of personal information.
3. **Preference for Human Connection.** Desire for interaction with real people over artificial agents.
4. **Negative Past Chatbot Experiences.** Prior negative chatbot encounters shaping current attitudes toward AI/VH systems.
5. **Lack of Empathy/Authenticity.** Perception that AI cannot genuinely understand or respond to human emotions.
6. **Limited VH Capability.** Belief that chatbots/VHs lack depth or nuance for complex, sensitive topics.
7. **Preference for ChatGPT.** Loyalty to ChatGPT; no interest in alternative chatbot or VH services.
8. **Anti-Embodiment/Uncanny Valley.** Discomfort with avatars perceived as 'too human' or artificial.
9. **Environmental Impact.** Concern about AI's energy consumption and carbon footprint.
10. **Conditional Trust.** Willingness to use AI for limited purposes but not emotional or sensitive interactions.
11. **Broader Social Concern.** Worries about AI's societal effects (job displacement, dependency, misinformation, existential risk).
12. **Caregiving Target Limitations.** Belief that the intended caregiving population would struggle to use chatbot/VH technology.
13. **Willing/Need to Learn.** Limited AI knowledge coupled with openness to learn more.
14. **Positive/Optimistic About AI.** Generally favorable attitudes toward AI (enthusiasm, curiosity, confidence).
15. **General Negative Response.** Broadly unfavorable reactions not fitting more specific codes.
16. **Misunderstanding of VH Role.** Incorrect or incomplete understanding of what virtual humans are or how they function.

17. **No Perceived Need.** No felt need for chatbot or VH assistance; no personal use case.
18. **Preference for Self-Reliance.** Preference for handling tasks independently rather than relying on AI/VH systems.
19. **Limited Own Capability to Use AI/VH.** Self-reported lack of skill, confidence, or access to use AI/VH effectively.
20. **Generational Resistance.** Resistance attributed to generational differences in technology adoption.
21. **Other – Uncodable.** Responses that do not align with any established code.
22. **Not Sure / Don't Know.** Uncertainty, indecision, or insufficient knowledge to form a clear opinion.