

Document Version

Final published version

Licence

CC BY

Citation (APA)

Tarvirdians, M., Chandrasegaran, S., Hung, H., Jonker, C. M., & Oertel, C. (2026). Reflecti-Mate: A Conversational Agent for Adaptive Decision-Making Support Through System 1 and System 2 Thinking. In *UMAP 2026 - Proceedings of the 34th ACM International Conference on User Modeling, Adaptation and Personalization* (pp. 213-222). Association for Computing Machinery (ACM). <https://doi.org/10.1145/3774935.3806176>

Important note

To cite this publication, please use the final published version (if applicable).
Please check the document version above.

Copyright

In case the licence states "Dutch Copyright Act (Article 25fa)", this publication was made available Green Open Access via the TU Delft Institutional Repository pursuant to Dutch Copyright Act (Article 25fa, the Taverne amendment). This provision does not affect copyright ownership.
Unless copyright is transferred by contract or statute, it remains with the copyright holder.

Sharing and reuse

Other than for strictly personal use, it is not permitted to download, forward or distribute the text or part of it, without the consent of the author(s) and/or copyright holder(s), unless the work is under an open content license such as Creative Commons.

Takedown policy

Please contact us and provide details if you believe this document breaches copyrights.
We will remove access to the work immediately and investigate your claim.

implications, while rarely considering emotional concerns such as distance from family or an intuitive sense of discomfort shaped by prior experiences.

Despite the documented benefits of such integration, people often have difficulty to do so in practice. A key challenge is limited metacognitive awareness of one's own thinking processes [5, 43], which makes it difficult to recognize how different sources of judgment influence their decision. For instance, people frequently overestimate the extent to which their decisions are guided by cognition while underestimating the influence of intuition and emotion [35], or fail to accurately identify which types of considerations dominate their reflection [49]. This lack of awareness makes it difficult for individuals to deliberately broaden or deepen their reflection in ways that integrate diverse considerations when making important decisions.

Receiving personalized support is one potential way to help individuals overcome low-awareness cognitive patterns [14]. However, existing systems (e.g. [31, 39]) often equate reflection support with promoting cognitive reasoning. Drawing on dual-process theory [17, 26, 48], emotional and intuitive considerations are frequently treated as heuristic or bias-prone (System 1) and are therefore de-emphasized or left unsupported. As a result, adaptation in such systems is typically reactive, relying only on conversational history to generate follow-up questions that steer users toward further cognitive elaboration, regardless of individual reflection patterns.

In this work, we address this gap by introducing a conversational reflection assistant that personalizes its dialog strategy to support integrative pre-decisional reflection (Fig. 1). Rather than aiming to alter users' dominant reasoning mode (i.e. if a person prefers cognitive reasoning we respect this rather than changing it to an emotional dominant reasoning style), our approach aims to support reflection by helping users surface potentially overlooked considerations and further elaborate underdeveloped thoughts. We present a user model that captures how individuals distribute attention across different thought categories and how their ideas evolve over the course of a reflective interaction. Building on this model, the agent dynamically balances exploratory questions that engage underrepresented types of thoughts with exploitative, Socratic follow-up questions [11, 37] that deepen existing ideas. This balance enables deliberate reflection—applying System 2-level elaboration to all thoughts, including emotional and intuitive ones typically associated with System 1—without privileging cognitive reasoning over other considerations. This implies that if a person indicates that not considering a certain aspect is deliberate, we respect the choice and do not try to change this.

We evaluate our approach through a user study in the context of high-stakes life decisions, examining its effects on users' reflective language use and perceived effectiveness of the support agent.

We are answering the following research questions in this study:

- **RQ1:** What differences emerge in reflective language during assisted reflection when participants interact with the experimental agent compared to a baseline agent?
- **RQ2:** How do participants' perceptions of the interaction differ between the experimental agent and the baseline agent?
- **RQ3:** To what extent do the experimental and baseline agents preserve users' dominant reasoning style?

2 Related Work

Humans are prone to a range of well-documented biases and cognitive limitations during decision making process, such as confirmation bias, framing effects, and bounded rationality [6, 25, 29, 53]. These limitations can distort how individuals reason and form opinions. Therefore, a growing body of research has explored technological approaches to support people in the pre-decisional phase, when they form opinions that will guide their subsequent decision. For example, these approaches aim to help users overcome specific biases that can distort their reasoning [24, 46, 56].

Some studies have explored personalized interventions on information seeking behavior, as information seeking is an important aspect of the pre-decisional phase [16]. These systems often model how users interact with external information to provide adaptive support that facilitates less biased behavior. This could entail engagement with both sides (pro and con) of an argument [55] and various topics to disrupt self-imposed filter-bubbles [2, 3]. However, these studies do not support people's own thinking process but rather focus on the process of exploring external information and opinions.

In contrast, other works have explored interventions that build directly upon individuals' own thoughts and rationales regarding their decisions. For instance, Reicherts et al. [40] probe users' investment rationales and, tailored to their individual portfolio profile, provide an augmented version of their thoughts that highlights overlooked facts in order to mitigate anchoring bias. However, this work focuses on a one-shot response rather than a reflective dialog that fosters the human's own reasoning.

Beyond one-shot interactions, some approaches employ multi-turn dialog to support decision making. Hadoux et al. [22] model users' beliefs and concerns throughout the conversation to steer them towards particular choices. This work focuses therefore more on steering the user's decision rather than fostering integrative reflection. Similarly, Chiang et al. [12] work with users' expressed thoughts by modeling the intent and stance underlying their reasoning. This work supports users in reflecting on the correctness of their stance, but does not aim to foster reflection that integrates cognitive, emotional, and intuitive perspectives.

Collectively, these works leave a gap: no existing system models users' thought patterns over the course of reflection to adaptively support integrative pre-decisional reasoning. Our work addresses this gap by introducing a conversational reflection assistant that adapts to users' evolving thought patterns to broaden and deepen reflection, while respecting individual reasoning preferences.

3 Methodology

This section presents our approach and the user study conducted to evaluate it.

3.0.1 Reflection Phases. In this study, the decision maker undergoes a two-phase reflection process without making a decision in either phase, focusing instead on externalizing thoughts related to the decision.

- (1) **Unaided Reflection Phase:** The decision-maker first reflects on the decision topic independently to establish the first baseline. They proceed to the next phase once they feel they have shared sufficiently.

- (2) **Assisted Reflection Phase:** In the second phase, the decision maker interacts with a conversational agent that acts as a reflection assistant, posing questions grounded in the user's prior unaided reflection. The interaction is process-oriented [57], with the agent refraining from offering its own insights or solutions [36].

3.1 User Profile

3.1.1 Reflection Profile. Once the decision maker enters the “assisted reflection phase”, our agent first constructs a computational representation of the user's previously shared reflection (in the “unaided reflection phase”). This representation serves as the structured knowledge base guiding the agent's subsequent decisions.

Each reflection is represented as a collection of individual thoughts, where each thought is characterized by a category, indicating the type of consideration it expresses, and by its elaborations, which serve as an indicator of its reasoning depth. An elaboration is a justification or supporting detail that deepens a main thought rather than introducing a new one.

Formally, a reflection is represented as:

$$R = (T, G) \quad (1)$$

where

- $T = \{t_1, t_2, \dots, t_n\}$ is the set of main thoughts,
- $G = \{(t_i, E_i) \mid t_i \in T\}$ is the set of parent-child relations, where each E_i is a set of elaborations associated with t_i (a main thought).

Each main thought t_i is represented as a tuple:

$$t_i = (\text{text}_i, \text{category}_i) \quad (2)$$

where text_i is the content of the main thought and $\text{category}_i \in C$ is the category of the thought, with

$$C = \{\text{internal}, \text{external}, \text{experiential}, \text{other}\}. \quad (3)$$

We generate this representation using a locally hosted Phi-4 14B [1] model (Q4_K_M quantization, temperature = 0), chosen after pilot testing several models that could be deployed on local servers.

3.1.2 Thought categories. The category of a thought reflects the type of consideration expressed in its content. These categories operationalize the head-heart-gut framework [15, 49], in combination with findings from Klein [27]. Thoughts labeled as “internal” express internally oriented considerations such as feelings or personal values. “External” thoughts refer to considerations shaped by external factors, such as situational constraints, other stakeholders, or factual information. “Experiential” thoughts draw on past experiences or prior outcomes, which have been shown to play a central role in human decision making [27]. The “other” category captures thoughts that do not clearly fall into the above categories.

We deliberately make no claims about underlying cognitive processes, focusing instead on their behavioral correlates as expressed in language. While these correlates relate to different sources of judgment, the relationship is not one-to-one. Internal considerations, for example, can emerge from heart or gut (emotional or intuitive processes); external considerations from head (cognitive process); and experiential considerations from gut or head (intuitive or cognitive processes).

3.2 Adaptation Algorithm

Once the agent has the user reflection profile, it aims to identify patterns in it. In particular, it attempts to detect: first, fixation on specific categories of thought, which can cause the user's decision to be based on narrow thinking; and second, surface-level reasoning, which can cause the user's decision to be based on fast, insufficiently considered thoughts.

To this end, the agent constructs a user's reflective pattern model based on the profile. For such modeling, two dimensions are defined:

- (1) **Breadth** – reflecting how many distinct main thoughts the user expresses within each thought category.
- (2) **Depth** – reflecting how elaborated those thoughts are through associated explanations or expansions.

The depth of a main thought t_i is defined based on its elaborative richness:

$$D_i = 1 + |E_i| \quad (4)$$

where $|E_i|$ is the number of elaborations associated with t_i . The constant term 1 ensures that a thought without any elaborations is assigned a depth of 1, representing the thought itself.

The breadth of reflection within a thought category k is defined as:

$$B_k = |\{t_i \in T \mid C(t_i) = k\}| \quad (5)$$

which represents the number of distinct thoughts the user produced from that perspective.

Using these two measures, we derive a cumulative indicator of the user's fixation on a thought category:

$$S_k = \sum_{t_i: C(t_i)=k} D_i = \sum_{t_i: C(t_i)=k} (1 + |E_i|) = B_k + \sum_{t_i: C(t_i)=k} |E_i| \quad (6)$$

It increases both with the number of thoughts in a thought category (breadth) and with the number of elaborations attached to them (depth), thus reflecting the overall focus of reflection within that category.

In order to model the user's reasoning depth patterns, the following indicator is defined:

$$\bar{D}_k = \begin{cases} \frac{S_k}{B_k}, & \text{if } B_k > 0 \\ 0, & \text{otherwise} \end{cases} \quad (7)$$

This indicator provides an average measure of elaboration within each thought category, independent of the number of thoughts expressed, thus solely reflecting the average degree of reasoning depth within each thought category.

The reflection profile and pattern model (the indicators) are initially constructed at the start of the “assisted reflection phase” and dynamically updated after each turn to guide the agent based on the user's current reflective state.

The agent has two main objectives in this interaction:

- (1) To disrupt fixation patterns, thereby broadening the user's reflection.
- (2) To disrupt shallow reasoning patterns, thereby deepening the user's reflection.

3.2.1 Disrupt disadvantageous thought patterns. (1) To disrupt fixation: The agent uses the fixation indicator (Equation 6) to identify the categories receiving the most attention. Rather than reinforcing this focus, it redirects the user to the least explored category-analogous to uncertainty-based sampling in adaptive learning [32, 45]-on the

assumption that these categories offer the greatest opportunity to broaden reflection.

(2) To disrupt shallow reasoning: The agent uses depth indicators (Equation 7) to identify underdeveloped thoughts. Categories with lower depth scores contain less elaborated reflections, and within them, thoughts with the lowest depth scores (Equation 4) are prioritized for follow-up, maximizing potential gains in reflection depth.

3.2.2 Reflection Enhancing Strategies. To satisfy both objectives throughout the interaction, the agent formulates its decision process as an *exploration-exploitation tradeoff* [50]. At each turn t , it decides whether to elicit new thoughts to broaden reflection (*exploration*) or deepen existing thoughts (*exploitation*).

This decision is guided by a continuously updated internal state s_t , which includes the structured reflection profile, the user's reflective pattern model, the history of asked questions, and the dialog context. Several signals derived from s_t inform action selection: the agent favors novel prompts to reduce repetition, prioritizes underdeveloped thoughts for exploitation, and adjusts probing based on implicit user feedback, reducing the utility of thoughts that previously elicited minimal elaboration, signaling low chance of further development.

The agent estimates the relative utility of exploration and exploitation and applies an ϵ -greedy policy [50, 54], selecting exploration with probability ϵ and exploitation with probability $1 - \epsilon$:

$$a_t = \begin{cases} \text{explore,} & \text{with probability } \epsilon \\ \text{exploit,} & \text{with probability } 1 - \epsilon \end{cases} \quad (8)$$

During exploration, the agent draws questions targeting underrepresented categories from a curated set of category-specific prompts (Appendix B, peer-reviewed by three researchers for clarity), to eliminate phrasing as a potential confounding variable. During exploitation, it generates Socratic follow-ups [11, 37] using a locally hosted Phi-4 14B model (Q4_K_M quantization, temperature = 0; Appendix-C), fostering elaboration without introducing solutions.

After each response, the agent updates the reflection profile: exploration responses are first categorized and added as new thought nodes, and exploitation responses are appended as child elaborations. The adaptation algorithm is then updated accordingly, ensuring subsequent decisions reflect the current reflective state.

The agent maintains conversational coherence by leveraging dialog history and favoring questions within the same category across adjacent turns, minimizing unnecessary topic shifts and supporting a natural interaction.

3.3 User Study

We conducted a between-subjects study comparing two versions of a conversational reflection assistant during the “assisted reflection phase.” (See Figure 2)

In the experimental condition, the agent's behavior was guided by the reflection profile and adaptive strategies described above, which tracked users' reflective patterns to decide when to broaden reflection or prompt deeper elaboration.

The baseline agent was instructed to pursue the same high-level objectives without maintaining an explicit profile of the user's reflection (Appendix-D). Both agents used the same underlying language model (Phi-4 14B) and shared interaction goals to isolate the

effect of profile-driven adaptive dialog management while controlling for language generation, task, and interaction context.

3.3.1 Procedure. The study was conducted on a web-based platform. After providing informed consent describing the study purpose, data use, and participant rights, participants completed a pre-questionnaire collecting demographic information, self-reported decision clarity (5-point Likert scale), and the short form of the Self-Reflection and Insight Scale (SRIS) [47]. Participants then selected a decision topic and entered the “unaided reflection phase”, freely articulating their thoughts in text without external intervention. Once they indicated they had no further thoughts, they proceeded to the “assisted reflection phase”. Here, participants interacted with a conversational agent and were randomly assigned to the experimental or baseline condition, without being informed of the assignment. Interaction continued until participants chose to end the session, with a minimum of ten conversational turns to ensure sufficient engagement. After the interaction, participants rated post-statements assessing the agent's perceived effectiveness. All statements were reviewed by three researchers for ambiguity, bias, and ethical considerations prior to the study, in line with established content validity and scale development procedures [8, 33].

3.3.2 Participants. We recruited 128 participants (64 per condition) via the crowd-sourcing platform Prolific¹, with sample size determined a priori using G*Power [18, 19] to achieve 80% statistical power. Inclusion criteria were fluency in English and currently being in the process of making one of the big life decisions [10] listed in Appendix-A. These decisions were chosen because they represent moments when reflection is particularly crucial [34], given their potential long-term impact on people's lives. A chi-square test revealed no significant association between topic and condition ($\chi^2(13) = 21.73$, $p = .060$), suggesting balanced topic distribution across conditions.

Participants were diverse in age (mean = 41.2, SD = 13.3, median = 40.1, range 18–65+), gender (66 female, 62 male), education (ranging from no diploma to doctorate), self-reflection tendencies measured with the short form of Self-Reflection and Insight Scale [47] (mean = 46.08, SD = 7.26, median = 46), and decision clarity, self-reported on a 5-point Likert scale (median = 4).

3.3.3 Measures. The effectiveness of the interaction was evaluated across two dimensions:

Reflective language. We analyzed participants' reflections using linguistic markers associated with cognitive, emotional, and intuitive processing, quantified with the Linguistic Inquiry and Word Count dictionary (LIWC-22) [38] and operationalized with established category groupings [21, 23, 44]. Cognitive processing included insight, causation, discrepancy, and tentative categories; emotional processing included positive and negative emotion; intuitive processing included past-tense verbs and perceptual processes (seeing, hearing, feeling).

Perceived support. We evaluated participants' perceptions of the agent's effectiveness in integrating head, heart, and gut perspectives during reflection and in deepening their reflections (see Table 3). Ratings were collected using 5-point Likert scales, and participants provided open-ended responses to elaborate on their evaluations.

3.3.4 Ethical Considerations. We obtained ethical approval for running this study from the university's Human Research Ethics

¹<https://www.prolific.com>

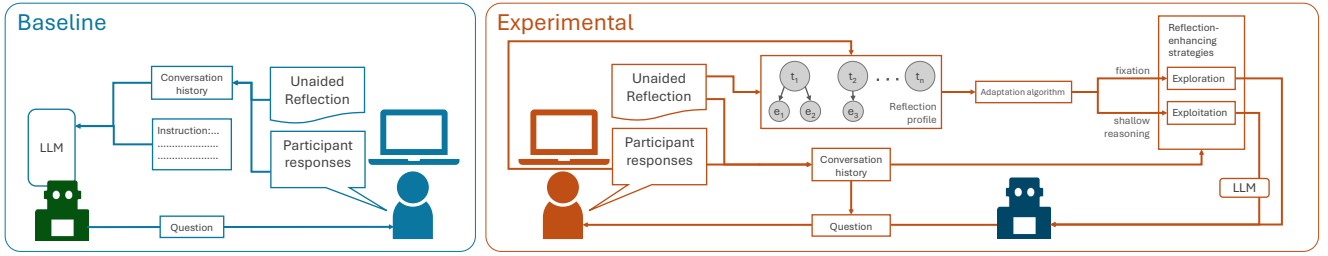


Figure 2: Overview of the Conditions. The baseline agent selects the next prompt based on the conversational history—including unaided user reflection—and a fixed instruction specifying objectives and constraints, which are passed to the LLM for prompt generation. The experimental agent additionally uses an adaptation algorithm informed by the user’s reflection profile to identify thought patterns, and pose either an exploration or exploitation question accordingly.

Committee (HREC). All participants provided informed consent and were compensated in accordance with the crowd-sourcing platform’s policies. Given the sensitivity of personal decision making, the AI agent was designed to avoid directive advice and instead support reflective inquiry. When participants requested opinions or recommendations, the agent clearly stated that it could not provide advice, and participants were not asked to reach or justify a final decision. All LLMs were deployed locally to ensure that reflection data remained on secure servers and were not shared with third parties. Data were anonymized, securely stored, and personally identifiable information was separated from reflection content; participants were informed of data handling and retention practices during consent.

4 Results

4.1 Reflection Behavior

To examine whether dialog management strategy (baseline vs experimental) influenced how users expressed and developed their reflections, we analyzed linguistic markers in users’ conversational turns during the assisted reflection phase. These markers were aggregated into three composite dimensions, cognitive, emotional, and intuitive using LIWC (see User Study–Measures).

A multivariate analysis of variance (MANOVA) conducted on users’ conversational reflections revealed a significant effect of condition on the joint distribution of reflective language dimensions, Wilks’ $\Lambda = 0.878$, $F(3, 124) = 5.74$, $p = .001$.² This result indicates that users interacting with the experimental and baseline agents differed in how cognitive, emotional, and intuitive linguistic markers were jointly expressed during assisted reflection.

4.1.1 Distribution and Differences in Dimensions. To characterize the nature of this effect, we examined within-condition differences among the three dimensions and between-condition differences during assisted reflection.

In the interaction with the baseline agent, users’ conversational reflections exhibited significant differences across dimensions, $F = 7.53$, $p < .001$. This pattern was characterized by a pronounced dominance of cognitive linguistic markers relative to emotional and intuitive

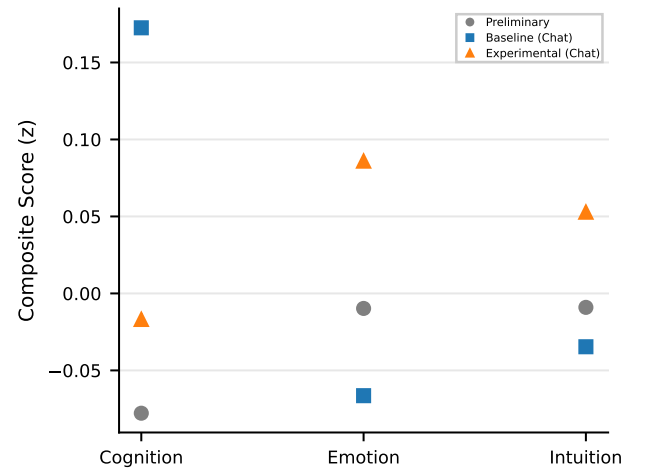


Figure 3: Transformation of Linguistic Markers by Condition. The figure illustrates mean composite z-scores for cognitive, emotional, and intuitive dimensions across Preliminary (unaided) reflections and both conditions: Baseline and Experimental. While the Baseline condition results in a pronounced dominance of cognitive markers at the expense of other dimensions, the experimental agent facilitates multi-modal participation.

markers (see Figure 3). In contrast, no significant differences across dimensions were observed in the experimental condition, $F = 1.12$, $p = .33$, with users expressing cognitive, emotional, and intuitive markers at more comparable levels during the assisted reflection phase (see Figure 3).

To quantify between-condition differences in reflective expression, we computed Cohen’s d for each composite dimension. Relative to the baseline condition, users interacting with the experimental agent exhibited lower cognitive composite scores ($d = -0.51$) and higher emotional ($d = 0.35$) and intuitive ($d = 0.24$) composite scores.

These differences reflect shifts in relative salience rather than suppression of cognitive reflection. Cognitive markers remained robustly present in the experimental condition; however, emotional and intuitive markers were more strongly expressed alongside cognitive markers.

²Assumption checks indicated violations of covariance homogeneity (Box’s M , $p = .022$) and multivariate normality (Henze–Zirkler). With equal sample sizes, MANOVA is robust to such violations [51]. All four test statistics yielded consistent results ($p = .001$), including Pillai’s trace, and no multicollinearity was detected (all $r < .20$).

4.2 Perceived Support

Participants rated perceived holistic integration. As shown in Fig. 4, ratings were significantly higher in the experimental condition: 72% agreed or strongly agreed, compared to 44% in the baseline condition. A chi-square test on collapsed categories (Disagree/Neutral/Agree) confirmed a significant association between condition and response, $\chi^2(2, N = 128) = 15.95, p = .002$ (Holm-corrected), with a moderate effect size ($V = .35$). Baseline participants also exhibited greater uncertainty, with more neutral responses (36% vs. 8%).

In contrast, ratings for thought elaboration did not differ between conditions, with high agreement in both groups (experimental: 72%, baseline: 79%), $\chi^2(2, N = 128) = 0.69, p = 1.0, V = .07$.

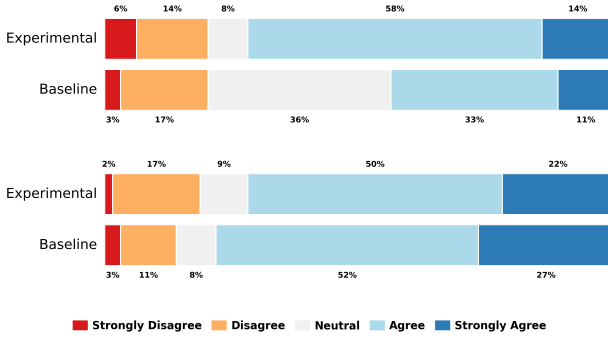


Figure 4: Distribution of participant responses for the two perceived support statements (see Appendix-E). Top: holistic integration. Bottom: elaboration depth.

4.3 Heterogeneous vs. Homogeneous Effects

The experimental agent adapted its questioning strategy online based on the evolving reflection profile of its users. To explore whether this adaptive strategy produced heterogeneous downstream effects on reflective language use across participants, we conducted a post-hoc cluster analysis. Users were clustered based on their pre-interaction (unaided reflection phase) linguistic profiles using k-means, yielding three interpretable archetypes:

- **Cluster 0 – Reserved** ($n = 73$): Low expression across all modes (Cog: -0.11 , Emo: -0.43 , Int: -0.21).
- **Cluster 1 – Intuition-dominant** ($n = 24$): High intuitive markers (Int: 0.91).
- **Cluster 2 – Emotion-dominant** ($n = 31$): High emotional markers (Emo: 1.25).

These clusters were used to investigate whether the effect of the dialogue strategy varied as a function of users' initial reflective language patterns, addressing RQ3.

4.3.1 Condition Effects Varied by Initial Profile. Figure 5 reveals that the baseline and experimental agents produced different effects depending on users' initial reflective profiles.

Across all three clusters, the baseline agent amplified cognitive markers while reducing the relative prominence of users' initially dominant modes. In Cluster 0 (Reserved), all dimensions increased but cognition dominated ($\Delta_{\text{Cog}} = 0.43, d = 0.74, p < .001$). In

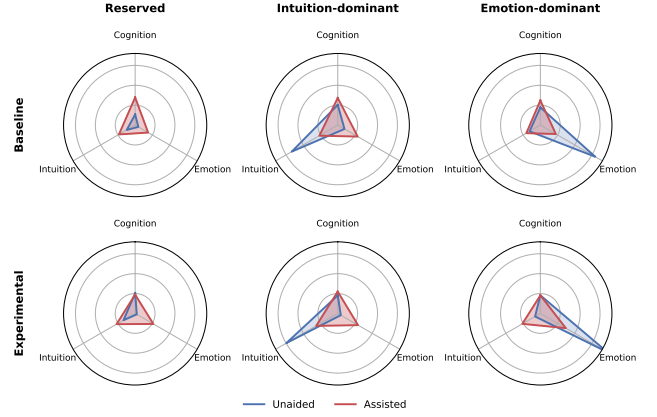


Figure 5: Reflection trajectories across clusters and conditions. Radar plots visualize the transition from unaided (blue) to assisted (red) reflection across three user archetypes. Axis scales represent standardized reflection scores (z-scores), where distance from the center indicates the relative intensity and reliance on a specific reflective language mode

Cluster 1 (Intuition-dominant), the initially prominent intuition collapsed ($\Delta_{\text{Int}} = -0.80, d = -1.60, p < .001$) as cognition and emotion increased. In Cluster 2 (Emotion-dominant), emotional expression was substantially reduced ($\Delta_{\text{Emo}} = -1.14, d = -1.14, p < .001$). Despite starting from markedly different baseline profiles, all three clusters moved toward a similar post-interaction configuration characterized by elevated cognition, indicating a homogenizing effect independent of initial reflective patterns.

In contrast, the experimental agent selectively expanded underrepresented modes while the initially dominant mode typically remained the most prominent. In Cluster 0, emotional ($d = 0.90, p < .001$) and intuitive markers ($d = 0.49, p = .005$) increased. In Cluster 1, intuition remained the dominant mode ($d = -1.86, p < .001$) while emotion was scaffolded ($d = 0.90, p = .010$). In Cluster 2, emotional markers remained dominant ($d = -1.14, p < .001$) while intuition was scaffolded ($d = 0.67, p = .026$). Rather than converging toward a single profile, users in the experimental condition retained cluster-specific structure while exhibiting broader multi-modal participation.

4.3.2 Differential Impact on Cognitive Markers. The most striking difference between conditions emerged in cognitive language use. The baseline agent produced uniformly large cognitive increases across all clusters (Δ_{Cog} range: 0.17 – $0.43, d = 0.27$ – $0.74, \text{all } p < .05$). In contrast, the experimental agent showed non-significant cognitive changes ($\text{all } |\Delta_{\text{Cog}}| < 0.10, \text{all } p > .31$), suggesting that adaptive scaffolding focused on expanding all modes rather than uniformly amplifying cognition.

5 Discussion

We investigated whether an adaptive conversational agent could support more integrative pre-decisional reflection than instruction-based approaches. In this section, we interpret our findings and consider their implications for designing decision support agents.

In RQ1, we investigated differences in reflective language use between our experimental and baseline agent. While prior work [7, 49]

has argued for an integrative approach to decision making, and [15] demonstrated the benefits of such integration in real-world leadership and business scenarios, our agent was able to enhance these patterns compared to the baseline. Although assessing task outcomes would require longitudinal studies, which are beyond the scope of the current work, our findings indicate the potential of conversational agents to support longer-term decision making.

The baseline agent's failure illuminates why computational user modeling matters. Despite identical goals (support breadth and depth across various thoughts) and access to the same conversational context, it could not achieve integration because it lacked the structured representation needed to detect when reflection was imbalanced or shallow. Instructions alone cannot substitute for this computational awareness. This suggests that effective reflection support requires not just good prompts, but principled adaptation grounded in real-time user modeling.

In **RQ2**, we explored the participants' perceptions of the interaction, and its results further suggest that participants not only found the agent helpful in supporting integration but also in elaborating on their thoughts. While prior work (e.g., [2]) shows increased elaborateness when participants are presented with diverse arguments, our results indicate that depth of reasoning can also be increased by engaging with the underlying driving forces of participants' own thoughts.

In **RQ3**, we investigated whether our agent could respect users' dominant reasoning style while supporting integration and elaboration at the same time. We found that it could, as it produced divergent effects—maintaining the prominence of each user's dominant mode while selectively expanding underrepresented dimensions. Importantly, these differentiated outcomes emerged even though the experimental agent was not explicitly conditioned on cluster membership, suggesting that real-time adaptive scaffolding can support personalized trajectories without requiring predefined user clusters.

This finding is particularly encouraging given prior work [4], which suggests that although psychological support extends beyond purely cognitive nudging toward more holistic approaches, LLM-based agents often fall short in preserving individual reasoning styles, focusing narrowly on users' initial thoughts and asking follow-up questions that deepen but do not broaden reflection. Our baseline agent exemplified this limitation. Its dominance of cognitive reflective markers indicates that it primarily probed users to further rationalize their thoughts, echoing prior work equating reflection support with promoting cognitive reasoning alone [39]. This imposed a one-size-fits-all reasoning style, driving all users toward cognitive dominance regardless of their initial reflective tendencies. This demonstrates the potential of our approach in decision making support, where the agent respects diverse types of reasoning as valuable and avoids enforcing convergence toward a single "ideal" profile.

While our findings demonstrate the effectiveness of adaptive inquiry-based support, comments from participants who rated the agent as less effective revealed a persistent tension inherent to reflection support systems: the gap between users' expectations for advice and the system's deliberate focus on facilitative inquiry [13]. Some participants ($n=10$ across both conditions) expressed frustration that the agent "is only trained to ask questions" or explicitly stated they expected to "get advice also."

On the one hand, this observation points to a promising direction: systems could offer users flexible control over interaction modes,

allowing them to switch between inquiry-based reflection and directive advice as their needs evolve—an approach explored in conversational recommendation systems [30]. Such flexibility could accommodate diverse user preferences and support needs. On the other hand, caution is needed; prioritizing users' desire for advice should not come at an ethical cost. Prior work suggests that inquiry-driven interactions better support user agency and autonomy, preventing premature advice in high-stakes situations [28]. From this perspective, the absence of advice represents a deliberate pedagogical trade-off rather than a system limitation. Nevertheless, users' expectations highlight the need for expectation management, suggesting that agents may benefit from explicitly negotiating their role and clearly communicating system boundaries to maintain alignment with system objectives.

Ultimately, in an age where people increasingly rely on LLMs for many tasks, including high-stakes decision making [9], there is a risk that these systems may overly influence or bias users. Rather than adapting to the LLM, we aim to contribute to work that fosters and preserves human autonomy without steering decisions or homogenizing reasoning across individuals. That said, we emphasize that real-world deployment of such systems requires further research and additional ethical review, particularly to better understand long-term effects and to address the needs of vulnerable users or crisis situations.

6 Future Work and Limitations

While the experimental agent tailored its exploitation questions to the user's answers, the exploration questions were scripted, primarily for experimental purposes. This suggests an interesting future direction: exploring how these questions can be personalized to individual users without compromising their ability to facilitate broad exploration. Leveraging human expert input and domain knowledge could guide this adaptation.

A limitation of our evaluation is that LIWC captures linguistic expressions of reflection rather than the underlying sources of those expressions. While LIWC has been widely used as a proxy for psychological and cognitive-affective processes, the markers associated with different processes should not be interpreted as direct evidence of engagement of a specific "system." Accordingly, our analysis focused on shifts in expressed reflective patterns rather than latent states, future work could triangulate these findings with behavioral or physiological measures.

7 Conclusion

Grounded in theories of human decision making, we explored whether conversational agents can support an integrative pre-decision reflection process. Our approach, built directly upon users' observable thought patterns, successfully elicited more integrative reflective language and was perceived as supporting a holistic reflection process, all while respecting the user's reasoning tendencies. In contrast, the baseline LLM-based agent followed a "one-size-fits-all" approach, pushing all users toward cognition. Our novel reflection modeling and adaptation approach, supported by these findings, opens a promising direction for personally meaningful decision support systems, where emotional and intuitive aspects are also treated as complementary forms of reasoning worthy of personalized computational support.

Appendix

A Big Life Decision Topics

Decision Category	Decision Topic
Career	Start/Quit a job/position
Career	Start a new business
Education	Pursue a degree
Family	Have/adopt a child
Family	Care for a family member
Family	Get pet
Finances	Buy home
Relationships	Begin/End romantic relationship
Relationships	Get married/divorced
Relocation	Move to a new city/country

Table 1: Decision Topics and Categories, fourteen topics in total

B Exploration Questions by Category

Category	Question
Internal	What are your gut feelings about {decision}? When you think about {decision}, what does your heart want? What emotions come up when you think about making this decision?
Experiential	What personal (first-hand) experiences have you had that relate to {decision}? What stories and experiences from your network (second-hand experiences) can you think of in relation to {decision}? What lessons or insights from your past experiences might help you in the process of making this decision?
External	What external factors are supporting this decision? What external constraints are posing challenges in {decision}?

Table 2: Exploration questions used by the experimental agent, grouped by targeted category.

C Exploitation Prompt (Experimental Agent)

You are an AI assistant designed to generate follow-up questions that encourage deeper elaboration on a user's argument about {topic}.

Problem:

- The user's argument: {span} lacks depth and requires further exploration.

Your Task:

- Generate a single, concise follow-up question for the argument.

Question Generation Rules:

- This question can be a **Socratic** question.
- Ensure that the question is in relation to the decision to be made: {topic} not just a general question.
- Ensure that the question fits within the current context by considering the conversation history.
- If they are earlier relevant points raised in the conversation history, you should acknowledge them in your question to show that you are aware of the context and ask if the user has more to say.
- Using the conversation history, avoid asking about information the user has already elaborated upon.
- The question should be neutral, non-leading, and avoid introducing new viewpoints.
- The question should be **short**, **clear**, and **directly** reference the original argument: {span}.
- Use a conversational tone
- Ensure the question helps the user **deepen** their reflection, not simply repeat previous information.

Constraints:

- Focus the question on **gathering new insights**, not repeating what has already been said.

- Do not suggest new perspectives or opinions.
- Ask only **one** question per turn.
- Avoid compound or double-barreled questions.
- Stay focused on one idea at a time.
- Keep the question objective and concise.
- Ensure to explicitly include the original text span within the follow-up question.
- The question must fit logically within the flow of the conversation.

D Baseline Agent Instruction

You are an AI assistant designed to support users in reflecting before making a big decision.

Problem:

- The user wants to enhance the **breadth** and **depth** of their reflection by integrating their head, heart, and gut, and by exploring aspects they may have under-considered or overlooked.

Your Task:

- At each turn, generate a single, concise question based on the conversation history to support this process.

Question Generation Rules:

- Ensure that the question is in relation to the decision to be made: {topic}, not just a general question.
- Ensure that the question fits within the context by considering the conversation history.
- If they are earlier relevant points raised in the conversation history, you should acknowledge them in your question to show that you are aware of the context
- Using the conversation history, avoid asking about information the user has already said.
- The question should be neutral, non-leading, and avoid introducing new viewpoints.
- The question should be **short**, **clear**.
- Use a conversational tone
- Ensure the question helps the user **deepen** or **broaden** their reflection, not simply repeating previous information.
- This question can be a **Socratic** question.

Constraints:

- Do not suggest new perspectives or opinions.
- Ask only **one** question per turn.
- Avoid compound or double-barreled questions.
- Stay focused on one idea at a time.
- The question must fit logically within the flow of the conversation.

E Post-Questionnaire Items

Statement	Construct
The agent posed questions that helped me integrate head, heart, and gut.	Holistic integration
The agent posed questions that helped me elaborate further on my thoughts.	Elaboration depth

Table 3: Post-experiment statements used to evaluate perceived support on a 5-point Likert scale (from strongly disagree to strongly agree).

Acknowledgments

We gratefully acknowledge the support of Delft AI Initiative and Designing Intelligence (DI) lab. This research was (partially) funded by the Hybrid Intelligence Center, a 10-year programme funded by the Dutch Ministry of Education, Culture and Science through the Netherlands Organisation for Scientific Research, Hybrid Intelligence Centre, grant number 024.004.022. In addition, it has been partially supported by the BOLD cities initiative. It was also partially funded by the project RISE: Relapse Intervention and Self-Management (Technology Assisted Self-Management: Preventing Relapse and Crisis by the Severe Mentally Ill Themselves) with file number KICH1.GZ03.21.002 of the research programme KIC - MISSIE Zorg in eigen leefomgeving 2021 which is (partly) financed by the Dutch Research Council (NWO).

Gen AI Declaration. Generative AI tools were used to assist with copy-editing, but the authors are fully accountable for the content.

References

- [1] Marah I. Abdin, Jyoti Aneja, Harkirat Behl, Sébastien Bubeck, Ronen Eldan, Suriya Gunasekar, Michael Harrison, Russell J. Hewett, Mojan Javaheripi, Piero Kauffmann, James R. Lee, Yin Tat Lee, Yuanzhi Li, Weishung Liu, Caio C.T. Mendes, Anh Nguyen, Eric Price, Gustavo de Rosa, Olli Saarikivi, Adil Salim, Shital Shah, Xin Wang, Rachel Ward, Yue Wu, Dingli Yu, Cyril Zhang, and Yi Zhang. 2024. *Phi-4: Technical Report*. Technical Report MSR-TR-2024-57. Microsoft Research. <https://arxiv.org/abs/2412.08905> Technical Report.
- [2] Annalena Aicher, Daniel Kornmüller, Yuki Matsuda, Stefan Ultes, Wolfgang Minker, and Keichi Yasumoto. 2023. Towards breaking the self-imposed filter bubble in argumentative dialogues. (2023).
- [3] Annalena Aicher, Wolfgang Minker, and Stefan Ultes. 2022. Towards modelling self-imposed filter bubbles in argumentative dialogue systems. (2022).
- [4] Riku Arakawa and Hiromu Yakura. [n. d.]. Coaching Copilot: Blended Form of an LLM-Powered Chatbot and a Human Coach to Effectively Support Self-Reflection for Leadership Growth. In *Proceedings of the 6th ACM Conference on Conversational User Interfaces*. 1–14.
- [5] Ruben T Azevedo, Salvatore Maria Aglioti, and Bigna Lenggenhager. 2016. Participants' above-chance recognition of own-heart sound combined with poor metacognitive awareness suggests implicit knowledge of own heart cardiodynamics. *Scientific reports* 6, 1 (2016), 26545.
- [6] Max H Bazerman and Dolly Chugh. 2006. Bounded awareness: Focusing failures in negotiation. In *Negotiation theory and research*. Psychology Press, 7–26.
- [7] Antoine Bechara and Antonio R Damasio. 2005. The somatic marker hypothesis: A neural theory of economic decision. *Games and economic behavior* 52, 2 (2005), 336–372.
- [8] Godfred O Boateng, Torsten B Neilands, Edward A Frongillo, Hugo R Melgar-Quinonez, and Sera L Young. 2018. Best practices for developing and validating scales for health, social, and behavioral research: a primer. *Frontiers in public health* 6 (2018), 149.
- [9] Becca Caddy. 2025. *Altman says Gen Z uses ChatGPT for life decisions, here's why that's both smart and risky*. TechRadar. <https://www.techradar.com/computing/artificial-intelligence/altman-says-gen-z-uses-chatgpt-for-life-decisions-heres-why-thats-both-smart-and-risky> Accessed: 2025-12-05.
- [10] Adrian R Camilleri. 2023. An investigation of big life decisions. *Judgment and Decision Making* 18 (2023), e32.
- [11] Timothy A Carey and Richard J Mullan. 2004. What is Socratic questioning? *Psychotherapy: theory, research, practice, training* 41, 3 (2004), 217.
- [12] Chun-Wei Chiang, Zhuoran Lu, Zhuoyan Li, and Ming Yin. 2024. Enhancing ai-assisted group decision making through llm-powered devil's advocate. In *Proceedings of the 29th International Conference on Intelligent User Interfaces*. 103–119.
- [13] NAJ Cornelissen, RJM Van Eerd, HK Schraffenberger, and Willem FG Haselager. 2022. Reflection machines: increasing meaningful human control over Decision Support Systems. *Ethics and Information Technology* 24, 2 (2022), 19.
- [14] Patricia G Devine, Patrick S Forscher, Anthony J Austin, and William TL Cox. 2012. Long-term reduction in implicit race bias: A prejudice habit-breaking intervention. *Journal of experimental social psychology* 48, 6 (2012), 1267–1278.
- [15] David L Dotlich, Peter C Cairo, and Stephen H Rhinesmith. 2006. *Head, Heart and Guts: How the world's best companies develop complete leaders*. John Wiley & Sons.
- [16] Glyn Elwyn and Talya Miron-Shatz. 2010. Deliberation before determination: the definition and evaluation of good decision making. *Health Expectations* 13, 2 (2010), 139–147.
- [17] Jonathan St B T Evans. 2008. Dual-processing accounts of reasoning, judgment, and social cognition. *Annual Review of Psychology* 59 (2008), 255–278.
- [18] Franz Faul, Edgar Erdfelder, Axel Buchner, and Albert-Georg Lang. 2009. Statistical power analyses using G* Power 3.1: Tests for correlation and regression analyses. *Behavior research methods* 41, 4 (2009), 1149–1160.
- [19] Franz Faul, Edgar Erdfelder, Albert-Georg Lang, and Axel Buchner. 2007. G* Power 3: A flexible statistical power analysis program for the social, behavioral, and biomedical sciences. *Behavior research methods* 39, 2 (2007), 175–191.
- [20] Robert E Goodin and Simon J Niemeyer. 2003. When does deliberation begin? Internal reflection versus public discussion in deliberative democracy. *Political Studies* 51, 4 (2003), 627–649.
- [21] Adam J Guastella and Mark R Dadds. 2006. Cognitive-behavioral models of emotional writing: A validation study. *Cognitive Therapy and Research* 30, 3 (2006), 397–414.
- [22] Emmanuel Hadoux, Anthony Hunter, and Sylwia Polberg. 2023. Strategic argumentation dialogues for persuasion: Framework and experiments based on modelling the beliefs and concerns of the persuadee. *Argument & Computation* 14, 2 (2023), 109–161.
- [23] Kate E Hamilton-West and Lyn Quine. 2007. Effects of written emotional disclosure on health outcomes in patients with ankylosing spondylitis. *Psychology and Health* 22, 6 (2007), 637–657.
- [24] Hsieh-Hong Huang, Jack Shih-Chieh Hsu, and Cheng-Yuan Ku. 2012. Understanding the role of computer-mediated counter-argument in countering confirmation bias. *Decision Support Systems* 53, 3 (2012), 438–447.
- [25] Bryan D Jones. 1999. Bounded rationality. *Annual review of political science* 2, 1 (1999), 297–321.
- [26] Daniel Kahneman and Shane Frederick. 2002. Representativeness revisited: Attribute substitution in intuitive judgment. In *Heuristics and Biases: The Psychology of Intuitive Judgment*, Thomas Gilovich, Dale Griffin, and Daniel Kahneman (Eds.). Cambridge University Press, 49–81.
- [27] G Klein. 1998. *Sources of Power: How People Make decisions* MIT Press Cambridge MA. (1998).
- [28] Philipp Koralus. 2025. The philosophic turn for AI agents: replacing centralized digital rhetoric with decentralized truth-seeking: P. Koralus. *Mind & Society* (2025), 1–24.
- [29] Russell F Korte. 2003. Biases in decision making and implications for human resource development. *Advances in Developing Human Resources* 5, 4 (2003), 440–457.
- [30] Ivica Kostic, Krisztian Balog, and Ujwal Gadriaju. 2025. Should We Tailor the Talk? Understanding the Impact of Conversational Styles on Preference Elicitation in Conversational Recommender Systems. In *Proceedings of the 33rd ACM Conference on User Modeling, Adaptation and Personalization*. 164–173.
- [31] Nguyen-Thinh Le and Laura Wartschinski. 2018. A cognitive assistant for improving human reasoning skills. *International Journal of Human-Computer Studies* 117 (2018), 45–54.
- [32] David D. Lewis and William A. Gale. 1994. A sequential algorithm for training text classifiers. *Proceedings of the 17th Annual International ACM SIGIR Conference on Research and Development in Information Retrieval* (1994), 3–12.
- [33] Mary R Lynn. 1986. Determination and quantification of content validity. *Nursing research* 35, 6 (1986), 382–386.
- [34] Ine Mols, Elise Van den Hoven, and Berry Eggen. 2016. Informing design for reflection: An overview of current everyday practices. In *Proceedings of the 9th Nordic Conference on Human-Computer Interaction*. 1–10.
- [35] Paul Norris and Seymour Epstein. 2011. An experiential thinking style: Its facets and relations with objective and subjective criterion measures. *Journal of personality* 79, 5 (2011), 1043–1080.
- [36] Soya Park and Chinmay Kulkarni. 2023. Thinking assistants: Llm-based conversational assistants that help users think by asking rather than answering. *arXiv preprint arXiv:2312.06024* (2023).
- [37] Richard Paul and Linda Elder. 2007. Critical thinking: The art of Socratic questioning. *Journal of developmental education* 31, 1 (2007), 36.
- [38] James W. Pennebaker, Ryan L. Boyd, Richard J. Booth, A. Ashokkumar, and M. E. Francis. 2022. *Linguistic Inquiry and Word Count: LIWC-22*. Pennebaker Conglomerates, Austin, TX. <https://www.liwc.app>
- [39] Leon Reicherts, Gun Woo Park, and Yvonne Rogers. 2022. Extending Chatbots to probe users: Enhancing complex decision-making through probing conversations. In *Proceedings of the 4th Conference on Conversational User Interfaces*. 1–10.
- [40] Leon Reicherts, Zelin Tony Zhang, Elisabeth von Oswald, Yunting Liu, Yvonne Rogers, and Mariam Hassib. 2025. AI, help me think—but for myself: Assisting people in complex decision-making by providing different kinds of cognitive support. In *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*. 1–19.
- [41] Troy D Sadler and Dana L Zeidler. 2005. Patterns of informal reasoning in the context of socioscientific decision making. *Journal of Research in Science Teaching: The Official Journal of the National Association for Research in Science Teaching* 42, 1 (2005), 112–138.
- [42] Lucrezia Savioni, Stefano Triberti, Ilaria Durosini, and Gabriella Pravettoni. 2023. How to make big decisions: A cross-sectional study on the decision making process in life choices. *Current Psychology* 42, 18 (2023), 15223–15236.
- [43] Irene Scopelliti, Carey K Morewedge, Erin McCormick, H Lauren Min, Sophie Lebrecht, and Karim S Kassam. 2015. Bias blind spot: Structure, measurement, and consequences. *Management Science* 61, 10 (2015), 2468–2486.

- [44] Sarah Seraj, Kate G Blackburn, and James W Pennebaker. 2021. Language left behind on social media exposes the emotional and cognitive costs of a romantic breakup. *Proceedings of the National Academy of Sciences* 118, 7 (2021), e2017154118.
- [45] Burr Settles. 2009. *Active learning literature survey*. Ph. D. Dissertation. University of Wisconsin-Madison.
- [46] Li Shi, Houjiang Liu, Yian Wong, Utkarsh Mujumdar, Dan Zhang, Jacek Gwizdka, and Matthew Lease. 2024. Argumentative experience: Reducing confirmation bias on controversial issues through llm-generated multi-persona debates. *arXiv preprint arXiv:2412.04629* (2024).
- [47] Paul J Silvia. 2022. The self-reflection and insight scale: Applying item response theory to craft an efficient short form. *Current Psychology* 41, 12 (2022), 8635–8645.
- [48] Steven A. Sloman. 1996. The empirical case for two systems of reasoning. *Psychological Bulletin* 119, 1 (1996), 3–22.
- [49] Grant Soosalu, Suzanne Henwood, and Arun Deo. 2019. Head, heart, and gut in decision making: Development of a multiple brain preference questionnaire. *Sage Open* 9, 1 (2019), 2158244019837439.
- [50] Richard S. Sutton and Andrew G. Barto. 2018. *Reinforcement Learning: An Introduction*. MIT Press.
- [51] Barbara G. Tabachnick and Linda S. Fidell. 2019. *Using Multivariate Statistics* (7 ed.). Pearson.
- [52] Morita Tarvirdians, Senthil Chandrasegaran, Hayley Hung, Catholijn M Jonker, and Catharine Oertel. 2025. Reflection Before Action: Designing a Framework for Quantifying Thought Patterns for Increased Self-awareness in Personal Decision Making. *arXiv preprint arXiv:2510.04364* (2025).
- [53] Amos Tversky and Daniel Kahneman. 1974. Judgment under Uncertainty: Heuristics and Biases: Biases in judgments reveal some heuristics of thinking under uncertainty. *Science* 185, 4157 (1974), 1124–1131.
- [54] Christopher J. C. H. Watkins and Peter Dayan. 1992. Q-learning. *Machine Learning* 8, 3-4 (1992), 279–292.
- [55] Klaus Weber, Annalena Aicher, Wolfgang Minker, Stefan Ultes, and Elisabeth André. 2023. Fostering user engagement in the critical reflection of arguments. *arXiv preprint arXiv:2308.09061* (2023).
- [56] Yu Zhang, Jingwei Sun, Li Feng, Cen Yao, Mingming Fan, Liuxin Zhang, Qianying Wang, Xin Geng, and Yong Rui. 2024. See widely, think wisely: Toward designing a generative multi-agent system to burst filter bubbles. In *Proceedings of the 2024 CHI Conference on Human Factors in Computing Systems*. 1–24.
- [57] Zelun Tony Zhang and Leon Reicherts. 2025. Augmenting Human Cognition With Generative AI: Lessons From AI-Assisted Decision-Making. *arXiv preprint arXiv:2504.03207* (2025).