

UC Merced

Proceedings of the Annual Meeting of the Cognitive Science Society

Title

Analyzing Human Heuristics and Strategies in Everyday Decision-Making Conversations for Conversational AI Design

Permalink

<https://escholarship.org/uc/item/9jp7m44d>

Journal

Proceedings of the Annual Meeting of the Cognitive Science Society, 48(0)

Authors

Kang, Sora

Jeon, Soyun

Eun, Jinsu

et al.

Publication Date

2026

Copyright Information

This work is made available under the terms of a Creative Commons Attribution License, available at <https://creativecommons.org/licenses/by/4.0/>

Peer reviewed

Analyzing Human Heuristics and Strategies in Everyday Decision-Making Conversations for Conversational AI Design

Sora Kang (sorakang@snu.ac.kr), Soyun Jeon, Jinsu Eun, Kwangwon Lee
Chaerin Song, Minyoung Joo, Joonhwan Lee

Human-Computer Interaction+Design Lab, Seoul National University, Seoul, Republic of Korea

Abstract

Conversational AI increasingly supports everyday decision-making, yet most systems rely on data-centric reasoning rather than the heuristic and interactional strategies people use in natural conversation. To ground design in actual human practice, we analyze 955 real-world Korean conversations (15,476 utterances) involving food and travel decisions, applying a decision-making codebook through an LLM-assisted coding pipeline. Our findings reveal that people prioritize satisficing over optimization, relying heavily on internal knowledge and interactional strategies to manage cognitive load. Critically, we identify a frequency–efficiency mismatch: the most prevalent heuristics sustain conversational flow during exploration, whereas infrequent, rule-based strategies are highly effective at driving resolution during exploitation. By mapping how these patterns transfer across the spectrum of human-AI interaction, this work provides empirical grounding consistent with cognitive theories of decision-making and offers design implications that align AI systems with human heuristic processes.

Keywords: Conversational AI; Human-AI Interaction; Decision Support Systems; Heuristics; Conversation Analysis

Introduction

Recent advances in large language models (LLMs) have positioned conversational AI as an increasingly capable interaction partner, not merely a tool for information retrieval. Yet most contemporary conversational systems still operate primarily as answer providers: they respond to well-formed questions, summarize relevant information, and produce recommendations given explicit constraints. This framing fits tasks with clear correctness criteria (e.g., factual queries or constrained lookups). In contrast, much of everyday life consists of decisions without a single objective “right” answer—what to eat tonight, where to travel next, or which hobby to start—where choices are shaped by personal values, affect, context, and limited time.

Decision-making in such settings is classically explained through bounded rationality and satisficing (i.e., seeking an option that is “good enough” rather than globally optimal) (Simon, 1955, 1956). Under bounded resources, people rely on heuristics—fast, experience-based shortcuts—to reduce cognitive effort and reach closure. These heuristics can be adaptive and efficient, but they also introduce systematic distortions and biases. For a conversational AI meant to support everyday decision-making, this creates a central design tension: effective support should feel lightweight and conversation-friendly

(as in human dialogue), while also helping users structure reasoning, reduce bias, and maintain agency. Yet we still lack clear empirical answers to a key question: How do humans actually talk their way toward everyday decisions, and what heuristics and interactional strategies structure this process?

Moreover, conversational decision-making encompasses a range of participant structures. Two friends may co-decide where to eat, one person may recommend options while the other chooses, or a knowledgeable party may guide an uncertain one toward a decision. These configurations—co-decider, recommender, advisor—reflect the fact that language use in decision-making is fundamentally a coordinated activity whose structure varies with the roles, stakes, and relationships of the participants (Clark, 1996). Understanding how cognitive heuristics and interactional strategies operate across these varied structures is essential not only for theories of human decision-making but also for designing conversational AI systems, which may occupy any of these roles.

To address this gap, we propose a human-grounded approach that treats everyday decision-making as a conversational process rather than a one-shot recommendation problem. We analyze real human–human decision dialogues to identify recurring strategies and heuristics as they appear in conversation, and we compile these patterns into a reusable codebook. This paper investigates the following research questions:

- RQ1. How are decision-making strategies and heuristics manifested and distributed in everyday human–human decision-making conversations?
- RQ2. Which conversational strategies and heuristics are associated with whether a decision is ultimately reached?

Our contributions are: (1) a conversation-grounded codebook of decision-making heuristics and strategies; (2) an LLM-assisted coding pipeline for large-scale analysis; and (3) an empirical characterization of how conversational strategies relate to decision outcomes, with implications for conversational AI design.

Related Work

Decision-Making Theories: Heuristics and Biases

Human decision-making is defined by cognitive and environmental constraints, contrasting the rational agent model with Simon’s theory of bounded rationality (Simon, 1955). Individuals navigate these limits through satisficing—selecting

"good enough" options rather than exhaustive optimization (Simon, 1955, 1956). Under bounded resources, people rely on heuristics: fast, intuitive shortcuts that reduce cognitive effort (Tversky & Kahneman, 1974). These include judgment heuristics such as availability (Tversky & Kahneman, 1973), representativeness (Tversky & Kahneman, 1974), affect (Slovic et al., 2002), and anchoring (Tversky & Kahneman, 1974). Furthermore, choices are sensitive to framing effects and loss aversion, which skew decisions based on presentation (Kahneman & Tversky, 1979; Tversky & Kahneman, 1981). Decision-making is also a social process influenced by heuristics like consistency, reciprocity, and social proof (Cialdini, 2001). For selection tasks, adaptive agents utilize choice heuristics—such as lexicographic rules or elimination by aspects—flexibly switching between them based on task complexity (Gigerenzer et al., 1999; Payne et al., 1993). While functional, these processes can result in systematic biases (Kahneman, 2011).

Conversational AI and Decision Support

AI-assisted decision-making aims for complementary performance between human and machine (Bansal et al., 2021; Lai et al., 2023; Reverberi et al., 2022; Zhang et al., 2024). However, recommendation-centric systems often ignore the underlying reasoning process, leading to inappropriate reliance where users either over-rely on or disregard suggestions (Lai et al., 2023; Ma et al., 2023). While AI explanations are intended to foster trust, they can paradoxically increase over-reliance by reducing active cognitive engagement (Bućinca et al., 2021; Miller, 2023). Recent research thus advocates for AI as a cognitive assistant that encourages independent judgment (Ma et al., 2025; Park & Kulkarni, 2023). Systems like Extend AI and frameworks focusing on causal rather than probabilistic reasoning aim to align AI with human cognition and mitigate over-reliance (Miller, 2019; Reber et al., 2025). While analyses of user–chatbot interaction logs have informed such designs (Li et al., 2020; Zhang et al., 2021), these logs often reflect existing system limitations rather than natural interaction; human–human dialogues reveal richer relational cues and patterns for designing more empathic and effective agents (Salman et al., 2021). We build on this by empirically characterizing naturalistic decision-making talk to ground future AI that aligns with human heuristics.

Method

This paper presents a conversation-analytic study of everyday decision-making interactions. We analyze real-world conversations to uncover the strategies individuals use in decision-making and to characterize how heuristics surface in naturalistic dialogue.

Human Conversation Dataset

To examine how individuals interact and address the inherent challenges of decision-making, we employed the Subject-Specific Text Data for Everyday Conversation dataset from

AI-Hub (National Information Society Agency (NIA), 2021), a Korean national platform providing open AI resources. The dataset consists of free-form conversations in Korean, spanning 20 different subjects. For this study, we focused on two domains—Food & Drink (5,176 conversations) and Travel (5,550 conversations)—as they frequently involve collaborative decision-making tasks, such as choosing meals or selecting travel destinations. From these subsets, we filtered for conversations containing explicit decision-making sequences, excluding those lacking a clear decision goal (e.g., simple reporting of past events) or containing insufficient interactional turns. After this filtering, a set of 955 conversations (15,476 utterances) remained, which served as the basis for both codebook validation and analysis.

Conversation Codebook

The codebook was developed through an iterative process to adapt theoretical concepts to systematic conversation analysis. The initial schema was grounded in established behavioral economics and decision theory, drawing on Simon's satisficing framework (Simon, 1955), Tversky and Kahneman's heuristics (Tversky & Kahneman, 1974), and Slovic's affect heuristic (Slovic et al., 2002). We conducted a pilot coding process on a random sample of 50 conversations from the initial AI-Hub pool (distinct from the final analysis set), followed by reconciliation sessions among the research team to review discrepancies, identify ambiguities, and clarify category definitions. Through this consensus-building process, we refined the codebook iteratively until the coders reached a shared understanding.

Through this refinement, the codebook evolved to capture specific conversational nuances. New dimensions were added, such as distinguishing the *source* of information (internal preferences vs. external facts) and specifying the *type* of information exchanged. The Suggestion and Response codes were expanded to include finer-grained response types (e.g., implicit agreement or soft deferral), moving beyond simple binary classifications. Overlapping or unclear categories were consolidated or removed to improve coding consistency.

Conversation-Level Codes. We coded the conversations on two dimensions. *Outcome* captures whether a decision was reached or if the dialogue ended unresolved. *Decision Method* identifies the strategy used: *Satisficing* refers to selecting a "good enough" option without extensive search, whereas *Maximizing* involves seeking the optimal choice by comparing multiple variables.

Utterance-Level Codes At the utterance level, our coding framework is organized into five dimensions.

Information Collection categorizes utterances where participants request or provide decision-relevant information from external sources (e.g., online reviews, third-party opinions) or internal knowledge (e.g., personal preferences, prior experiences), covering types such as price, time, or logistics.

Table 1: The Decision-Making Conversation Codebook. The framework consists of conversation-level outcomes and methods, utterance-level interaction structures, and cognitive heuristics.

Category	Definitions and Indicators
I. Conversation Level Analysis	
Outcome	Decision Reached: Interaction concludes with a mutual decision. / No Decision: Unresolved.
Method	Satisficing: Stopping the search once a “good enough” option is found. Prioritizes sufficiency. Maximizing: Considering all variables and alternatives to find the “optimal” choice.
II. Utterance Level Analysis (Structure)	
Info. Collection	Source: <i>External</i> (Objective facts, reviews) vs. <i>Internal</i> (User preferences, situations). Type: Covers <i>Schedule, Cost, Activity, Place, Time, Reputation, Other</i> .
Suggestion	Suggestion: Proposing clear options or recommendations. – <i>Context-based:</i> Grounded in previous dialogue context. – <i>Explanation-based:</i> Accompanied by reasoning or justification.
Response	Explicit: Agreement (Positive preference) vs. Rejection (Negative preference). Implicit Agreement: Positive cues or feedback without explicit confirmation. Implicit Rejection: – <i>Alternative (Type 1):</i> Indirect refusal by suggesting a different option. – <i>Avoidance (Type 2):</i> Indirect refusal by evading the decision. – <i>Reasoning (Type 3):</i> Indirect refusal by providing reasons for dislike. Other: <i>Indifference (NM)</i> (“No matter” but hidden preference), <i>Suspension</i> (Deferring decision).
Strategies	Choice Overload Mitigation (COM): Narrowing the option set to reduce cognitive load. Open-Ended Question (OEQ): Broadening the scope to invite free-form opinions.
III. Cognitive Heuristics Analysis	
Judgment	<i>Availability</i> (Recall ease), <i>Representativeness</i> (Similarity), <i>Anchoring</i> (Reference point), <i>Affect</i> (Emotion), <i>Framing</i> (Gain/Loss presentation).
Social Influence	<i>Consistency, Reciprocity, Scarcity, Social Proof, Likability, Authority</i> .
Choice Rules	<i>Lexicographic</i> (Most important attribute), <i>Attribute Elimination</i> (Sequential rejection), <i>Additive Difference, Conjunctive/Disjunctive</i> (Thresholds).

It applies whenever information necessary for a decision or for proposing an option is revealed; simple exclamations or reactions are not coded.

Suggestion includes utterances proposing specific options or recommendations. Suggestions may be grounded in previous dialogue context or accompanied by explanations justifying why a particular option should be considered.

Response identifies utterances in which a speaker reacts to a previous suggestion or piece of information. Responses may express agreement or rejection directly, hint at acceptance or rejection indirectly, indicate indifference, or signal that the decision should be put on hold. This dimension captures the dynamics of acceptance, negotiation, resistance, and deferral in decision-making conversations.

Choice Overload Mitigation (COM) refers to utterances that reduce cognitive load by narrowing available alternatives, often presenting a limited set of options (e.g., “How about pizza or pasta?”) rather than an open-ended array.

Open-Ended Question identifies utterances that invite the other party to respond freely by sharing opinions, preferences, or alternative suggestions. Unlike COM, open-ended questions deliberately broaden the conversation scope (e.g., “What do you feel like eating?”), creating opportunities for participants to express unarticulated preferences.

Heuristics The heuristic level comprises three dimensions (Table 1). *Judgment Heuristics* code intuitive shortcuts in evaluating options or probabilities: availability (recall ease), representativeness (similarity to a familiar category), anchoring (initial reference point), affect (immediate emotional reactions), and framing (gain/loss presentation). *Social Influence Heuristics* code appeals to social cues: consistency, reciprocity, scarcity, social proof, likability, and authority. *Choice Heuristics* code simplified selection rules: lexicographic (best on the most important attribute), attribute elimination (sequential exclusion), additive difference (weighted comparison), and conjunctive/disjunctive (minimum thresholds).

LLM-Assisted Coding Procedure

Following the development of the codebook through human coding, we applied it to the full dataset using a human-in-the-loop validation pipeline to ensure reliability. We utilized the GPT-4 model for automatic coding. To bridge the gap between human nuance and machine processing, we implemented an iterative refinement process. Initial trials revealed discrepancies, particularly in subtle categories such as implicit rejection and choice overload mitigation. We addressed this by incorporating concrete examples derived from the pilot study directly into the system prompt to guide the model’s reasoning, effectively utilizing few-shot prompting. Furthermore, we conducted manual reviews on random batches of 20

conversations after each prompt adjustment. We compared the output against human consensus and refined the definitions until the model demonstrated consistent alignment with human annotators. We proceeded to full-scale coding of the 955 conversations only after three consecutive batches showed no major deviations from human judgment.

Results

Decision Outcomes and Methods

As summarized in Table 2, among the 955 conversations, 693 (72.6%) resulted in a decision. Across all conversations, satisficing (41.9%) was the most frequent strategy, significantly exceeding maximizing (15.7%) ($\chi^2 = 113.6, p < .001$), suggesting a general preference for cognitive economy over exhaustive comparison.

Utterance-Level Patterns

Information Collection Information collection was the most frequent code (45.2%). A Chi-square test reveals a significant reliance on internal knowledge (55.1%) over external information (14.4%) ($\chi^2 = 1668.5, p < .001$), indicating that everyday decision-making operates primarily through memory retrieval rather than active information discovery.

Suggestions and Responses Suggestions occurred in 1,992 utterances, with 63.1% relying solely on context without justification. In responses ($N = 5,785$), explicit agreement was common (33.2%), whereas explicit rejection was rare (1.9%). Instead, participants utilized indirect disagreement strategies (combined 38.1%), such as implicit deferral or rejection.

Table 2: Descriptive Statistics: Conversations and Utterances

Category	Count	Percentage
Conversation Level ($N = 955$)		
Decision Reached	693	72.6%
No Decision	262	27.4%
<i>Method used (full sample):</i>		
Satisficing	400	41.9%
Maximizing	150	15.7%
Unclear/Mixed	405	42.4%
Utterance Level ($N = 15,476$)		
1. Information Collection	6,993	45.2%
Internal Knowledge	3,853	(55.1%)*
External Information	1,006	(14.4%)*
Source Unspecified/Other	2,134	(30.5%)*
2. Response	5,785	37.4%
Explicit Agreement	1,922	(33.2%)*
Implicit Agreement	1,303	(22.5%)*
Implicit Rejection	1,299	(22.5%)*
Implicit Deferral/Suspension	902	(15.6%)*
Explicit Rejection	112	(1.9%)*
Other (NM, N/A)	247	(4.3%)*
3. Suggestion	1,992	12.9%
4. Choice Overload Mitigation	1,433	9.3%
5. Open-Ended Question	988	6.4%

*Percentage within the specific category

Conversational Strategies

Choice Overload Mitigation (COM) Choice overload mitigation appeared in 1,433 utterances (9.3%). Crucially, almost half of all suggestions (48.8%) were framed using this strategy, with speakers frequently narrowing the field by presenting binary contrasts (e.g., “Pizza vs. Pasta”) rather than open lists.

Heuristic Usage and Decision Success

We analyzed heuristic usage ($N = 3,416$) to examine the relationship between heuristic frequency and decision outcomes.

Prevalence vs. Efficiency As illustrated in Figure 1 and detailed in Table 3, the data revealed a mismatch between the usage frequency of heuristics and their effectiveness in resolving decisions.

- **High-Frequency, Low-Efficiency Heuristics:** Affect (46.2%) and Availability (27.5%) were the most frequently observed heuristics. However, their efficiency was relatively low; Availability showed a success ratio of 1.53:1.
- **Low-Frequency, High-Efficiency Heuristics:** Conversely, the most decisive heuristics were rare. Attribute Elimination (1.6%) demonstrated the highest success ratio (6.86:1), significantly outperforming the baseline average ($p < .001$).

Social Heuristics Strategies relying on external validation, such as Authority (1.38:1), were associated with lower decision success, whereas relational heuristics like Reciprocity (4.10:1) were highly effective.

Table 3: Frequency and Success Ratio of Key Heuristics

Heuristic	Freq. (%)	Ratio (T:F)
<i>Choice Rules</i>		
Attribute Elimination	1.6%	6.86 : 1
Conjunctive	2.4%	3.26 : 1
Disjunctive	10.3%	2.73 : 1
Lexicographic	18.2%	2.75 : 1
<i>Social Influence</i>		
Reciprocity	3.1%	4.10 : 1
Likability	11.5%	3.08 : 1
Social Proof	7.1%	2.17 : 1
Authority	0.9%	1.38 : 1
<i>Judgment</i>		
Affect	46.2%	2.59 : 1
Availability	27.5%	1.53 : 1
Baseline Average	-	2.43 : 1

Discussion

Everyday decision-making in conversation deviates significantly from classical optimization, operating instead as a satisficing process driven by heuristics and interactional strategies. In this section, we unpack the cognitive mechanisms underlying these conversational patterns, examine how this bounded rationality transfers to human–AI contexts, and outline design implications that align conversational AI with human cognitive processes.

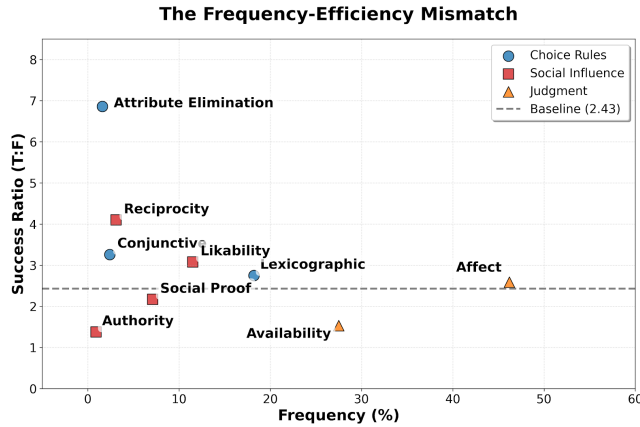


Figure 1: Frequency–Efficiency Mismatch across heuristic types. High-frequency heuristics (e.g., Affect) show lower decision success ratios, whereas low-frequency rule-based strategies (e.g., Attribute Elimination) show the highest.

Cognitive Heuristics in Dialogue

Frequency–Efficiency Mismatch in Heuristics One of the central findings of this study is the inverse relationship between heuristic frequency and decision success (Figure 1). Affective heuristics were the most prevalent strategy but were not strongly associated with decision resolution, whereas rule-based strategies, such as Attribute Elimination, were rare yet exhibited the highest success ratio.

We propose that this mismatch reflects a functional division within conversational decision-making. High-frequency heuristics such as affect and availability serve a dual function: they operate as cognitive shortcuts for evaluating options, but they simultaneously perform a conversational function—maintaining dialogue flow, building shared context, and signaling engagement with the partner’s perspective (Slovic et al., 2002). Their prevalence may therefore reflect not their decision-closing power but their role in sustaining the interactional substrate on which decisions are built.

By contrast, rule-based heuristics such as Attribute Elimination carry a higher cognitive cost: they require the speaker to articulate explicit criteria, which demands more effortful processing (Payne et al., 1993). They may also carry a social cost, as eliminating an option can implicitly reject a partner’s suggestion. These dual costs help explain why such strategies are rarely initiated spontaneously, despite their effectiveness. When they do appear, they function as convergence mechanisms that move the conversation from open exploration to closure.

This pattern suggests that conversational decision-making proceeds through two functionally distinct phases: an *exploration phase* dominated by low-cost, socially facilitative heuristics (affect, availability) that expand and maintain the decision space, and an *exploitation phase* in which higher-cost, rule-based heuristics narrow the options and drive toward resolution. This two-phase structure parallels the exploration–

exploitation trade-off described in computational models of decision-making (Cohen et al., 2007), but here it emerges from the interactional dynamics of conversation rather than from individual optimization. Future work could examine whether the timing of the transition between phases predicts decision success.

Satisficing as a Conversational Achievement The predominance of satisficing over maximizing in our data is consistent with decades of work on bounded rationality (Simon, 1955, 1956), but our findings suggest that in conversational contexts, satisficing is not merely a default mode of individual reasoning but a joint accomplishment enabled by specific interactional mechanisms.

Choice Overload Mitigation (COM) and implicit agreement together scaffold the satisficing process. COM constrains the option space—typically through binary contrasts—so that a “good enough” threshold can be reached more quickly. Implicit agreement provides a low-cost pathway to acceptance without the exhaustive justification that maximizing would demand. In this sense, satisficing in conversation is a property of the dialogue system, not only of individual cognition.

Internal Knowledge and Ecological Rationality The predominance of internal knowledge over external information retrieval supports the ecological rationality framework (Gigerenzer et al., 1999), which argues that heuristic strategies are adapted to the structure of the decision environment. In everyday, low-stakes domains such as food and travel, the relevant information is often already available in memory—past experiences, personal preferences, situational constraints—making external search unnecessary or even counterproductive by adding cognitive load without proportionate benefit.

This finding extends the ecological rationality account from individual decision-making to conversational settings. In dialogue, the reliance on internal knowledge takes a specific form: participants retrieve and share personal experiences, articulate latent preferences, and use affective reactions as information. The conversational process thus functions less as a joint information search and more as a mutual preference clarification process.

Transferability to Human–AI Interaction

Applying findings from human–human conversation to human–AI design requires addressing a structural difference: human–human decision-making often involves joint commitment, where both parties must agree to act on the outcome, whereas human–AI interaction is typically framed as individual decision support. We address this at two levels.

First, the core cognitive patterns identified in this study—satisficing, reliance on internal knowledge, the frequency–efficiency mismatch, and choice overload mitigation—are properties of individual bounded rationality that surface in conversation but are not dependent on joint commitment. These patterns arise from the decision-maker’s own cognitive constraints and would be expected to operate regardless

of whether the conversational partner is human or artificial.

Second, the social patterns in our data are likely to appear in human–AI interaction as well, both because users apply social heuristics to machines exhibiting conversational cues (Gamino et al., 2020; Nass & Moon, 2000) and because AI systems themselves increasingly occupy relational roles—from neutral information provider to opinionated advisor to persistent conversational partner with accumulated user knowledge. The role a conversational AI occupies modulates which patterns from our codebook apply most directly, but across this range, strategies such as implicit rejection, conversational deferral, and suspension remain relevant. A design-critical implication is that the absence of explicit negation should not be interpreted as confirmation, precisely because users bring social cognition to machine interlocutors.

Rather than treating human–human and human–AI decision-making as categorically distinct, we suggest that they occupy a shared spectrum of conversational decision structures, differentiated by the degree of commitment, social obligation, and role asymmetry between participants. Our codebook is designed as a flexible analytical resource that can be selectively applied across this spectrum.

Design Implications

Building on the theoretical analysis and transferability arguments above, we propose the following design guidelines grounded in the observed cognitive and interactional patterns.

1. Memory-Based Preference Elicitation. Given the predominance of internal knowledge and the interpretation of everyday decisions as preference clarification rather than information search, systems should minimize the cognitive cost of external retrieval. Instead of querying attribute parameters (e.g., price, location), the dialogue manager should employ associative prompting—triggering episodic memory to surface latent preferences (e.g., “Do you recall a similar experience you enjoyed?”). This shifts the system function from database querying to memory support.

2. Operationalizing Implicit Rejection. Standard intent recognition models often classify non-response or topic shifts as “neutral” or “continuation.” However, given the prevalence of implicit strategies and the rarity of explicit rejection in our data, and given that the CASA literature predicts users will employ these same face-saving strategies with AI interlocutors (Nass & Moon, 2000), dialogue policies should be calibrated to interpret ambiguous user turns (e.g., hesitation, unrelated inquiries) as potential negative feedback signals, triggering exploration of alternative branches rather than persisting with the current candidate.

3. Structured Pairwise Presentation. To address choice overload, systems should avoid presenting unstructured recommendation lists. Instead, the interface should adopt a pairwise comparison framework, presenting options in pairs (e.g., A vs. B). This aligns with the conversational satisficing patterns observed in our data, reducing the cognitive load of global optimization to a simpler task of local preference judg-

ment (Iyengar & Lepper, 2000).

4. Phase-Adaptive Heuristic Support. The two-phase structure identified in our theoretical analysis provides a concrete policy for conversational agents. During the *exploration phase*, the system should process and mirror affective and availability cues to maintain engagement and build shared context. As the conversation progresses to the *exploitation phase*—signaled by turn count, heuristic cycling, or user signs of decision fatigue—the system should actively supply the rule-based mechanisms that users rarely initiate spontaneously, introducing constraint-based proposals (e.g., “Excluding options over distance X ”) to drive toward closure. This phase-adaptive strategy addresses the frequency–efficiency mismatch by providing the right kind of support at each stage rather than defaulting to a single mode.

Limitations and Future Work

This study is limited by the scope of its data and setting. The study focuses on Korean everyday decision scenarios, which may limit generalizability to other cultural contexts or domains (e.g., medicine, finance). Second, the conversations are text-based and may differ from multimodal or face-to-face interactions. Third, our LLM-assisted coding pipeline, while calibrated against human coders, may still contain residual annotation errors. Future work can extend this analysis to additional domains and languages, investigate sequential patterns (e.g., how specific strategies unfold turn by turn), and embed these findings into the design and evaluation of conversational AI agents for decision support.

Conclusion

We presented a conversation-analytic study of everyday decision-making dialogue, using Korean conversation datasets to develop and apply a codebook of decision strategies, conversational moves, and heuristics. Our findings show that satisficing, internal knowledge, affective judgments, and choice-overload mitigation are central to how people talk their way toward decisions, and that specific heuristic patterns are associated with whether conversations reach resolution. These insights provide a grounded foundation for designing conversational AI that complements bounded rationality and human heuristics, supporting everyday decision-making in a way that aligns with how people actually think and converse.

Acknowledgments

This work was partly supported by Institute of Information & Communications Technology Planning & Evaluation (IITP) grant funded by the Korea government (MSIT) [NO.RS-2021-II211343, Artificial Intelligence Graduate School Program (Seoul National University)], SNU-Global Excellence Research Center establishment project, and Basic Science Research Program through the National Research Foundation of Korea (NRF) funded by the Ministry of Education (No. RS-2025-25421701).

References

- Bansal, G., Wu, T., Zhou, J., Fok, R., Nushi, B., Kamar, E., Ribeiro, M. T., & Weld, D. (2021). Does the whole exceed its parts? the effect of ai explanations on complementary team performance. *Proceedings of the 2021 CHI Conference on Human Factors in Computing Systems*, Article 81, 1–16. <https://doi.org/10.1145/3411764.3445569>
- Buçinca, Z., Malaya, M. B., & Gajos, K. Z. (2021). To trust or to think: Cognitive forcing functions can reduce over-reliance on ai in ai-assisted decision-making. *Proceedings of the ACM on Human-Computer Interaction*, 5(CSCW1), Article 188, 21. <https://doi.org/10.1145/3449287>
- Cialdini, R. B. (2001). *Influence: Science and practice* (4th ed.). Allyn & Bacon.
- Clark, H. H. (1996). *Using language*. Cambridge University Press. <https://doi.org/10.1017/CBO9780511620539>
- Cohen, J. D., McClure, S. M., & Yu, A. J. (2007). Should I stay or should I go? How the human brain manages the trade-off between exploitation and exploration. *Philosophical Transactions of the Royal Society B: Biological Sciences*, 362(1481), 933–942. <https://doi.org/10.1098/rstb.2007.2098>
- Gambino, A., Fox, J., & Ratan, R. A. (2020). Building a stronger CASA: Extending the computers are social actors paradigm. *Human-Machine Communication*, 1, 71–86. <https://doi.org/10.30658/hmc.1.5>
- Gigerenzer, G., Todd, P. M., & The ABC Research Group. (1999). *Simple heuristics that make us smart*. Oxford University Press.
- Iyengar, S. S., & Lepper, M. R. (2000). When choice is demotivating: Can one desire too much of a good thing? *Journal of Personality and Social Psychology*, 79(6), 995–1006.
- Kahneman, D. (2011). *Thinking, fast and slow*. Farrar, Straus; Giroux.
- Kahneman, D., & Tversky, A. (1979). Prospect theory: An analysis of decision under risk. *Econometrica*, 47(2), 263–291. <https://doi.org/10.2307/1914185>
- Lai, V., Chen, C., Smith-Renner, A., Liao, Q. V., & Tan, C. (2023). Towards a science of human-ai decision making: An overview of design space in empirical human-subject studies. *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*, 1369–1385. <https://doi.org/10.1145/3593013.3594087>
- Li, C.-H., Yeh, S.-F., Chang, T.-J., Tsai, M.-H., Chen, K., & Chang, Y.-J. (2020). A conversation analysis of non-progress and coping strategies with a banking task-oriented chatbot. *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*, 1–12. <https://doi.org/10.1145/3313831.3376209>
- Ma, S., Chen, Q., Wang, X., Zheng, C., Peng, Z., Yin, M., & Ma, X. (2025). Towards human-ai deliberation: Design and evaluation of llm-empowered deliberative ai for ai-assisted decision-making. *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*, Article 261, 1–23. <https://doi.org/10.1145/3706598.3713423>
- Ma, S., Lei, Y., Wang, X., Zheng, C., Shi, C., Yin, M., & Ma, X. (2023). Who should i trust: Ai or myself? leveraging human and ai correctness likelihood to promote appropriate trust in ai-assisted decision-making. *Proceedings of the 2023 CHI Conference on Human Factors in Computing Systems*, Article 759, 1–19. <https://doi.org/10.1145/3544548.3581058>
- Miller, T. (2019). Explanation in artificial intelligence: Insights from the social sciences. *Artificial Intelligence*, 267, 1–38. <https://doi.org/10.1016/j.artint.2018.07.007>
- Miller, T. (2023). Explainable ai is dead, long live explainable ai! hypothesis-driven decision support using evaluative ai. *Proceedings of the 2023 ACM Conference on Fairness, Accountability, and Transparency*, 333–342. <https://doi.org/10.1145/3593013.3594001>
- Nass, C., & Moon, Y. (2000). Machines and mindlessness: Social responses to computers. *Journal of Social Issues*, 56(1), 81–103. <https://doi.org/10.1111/0022-4537.00153>
- National Information Society Agency (NIA). (2021). Subject-Specific Text Data for Everyday Conversation [Accessed: 2026-01-29].
- Park, S., & Kulkarni, C. (2023). Thinking assistants: Llm-based conversational assistants that help users think by asking rather than answering. *arXiv preprint arXiv:2312.06024*. <https://arxiv.org/abs/2312.06024>
- Payne, J. W., Bettman, J. R., & Johnson, E. J. (1993). *The adaptive decision maker*. Cambridge University Press. <https://doi.org/10.1017/CBO9781139173933>
- Reber, U., et al. (2025). Ai, help me think—but for myself: Assisting people in complex decision-making by providing different kinds of cognitive support. *Proceedings of the 2025 CHI Conference on Human Factors in Computing Systems*. <https://doi.org/10.1145/3706598.3713295>
- Reverberi, C., Rigon, T., Solari, A., et al. (2022). Experimental evidence of effective human-ai collaboration in medical decision-making. *Scientific Reports*, 12(14952). <https://doi.org/10.1038/s41598-022-18751-2>
- Salman, S., Richards, D., & Caldwell, P. (2021). Analysis of empathic dialogue in actual doctor-patient calls and implications for design of embodied conversational agents. *Italian Journal of Computational Linguistics (IJCoL)*, 7(1-2), 91–112. <https://doi.org/10.4000/ijcol.862>
- Simon, H. A. (1955). A behavioral model of rational choice. *The Quarterly Journal of Economics*, 69(1), 99–118. <https://doi.org/10.2307/1884852>
- Simon, H. A. (1956). Rational choice and the structure of the environment. *Psychological Review*, 63(2), 129–138. <https://doi.org/10.1037/h0042769>
- Slovic, P., Finucane, M. L., Peters, E., & MacGregor, D. G. (2002). The affect heuristic. In T. Gilovich, D. Griffin, & D. Kahneman (Eds.), *Heuristics and biases: The psychology of intuitive judgment* (pp. 397–420). Cambridge University Press. <https://doi.org/10.1017/CBO9780511808098.025>
- Tversky, A., & Kahneman, D. (1973). Availability: A heuristic for judging frequency and probability. *Cognitive Psy-*

- chology*, 5(2), 207–232. [https://doi.org/10.1016/0010-0285\(73\)90033-9](https://doi.org/10.1016/0010-0285(73)90033-9)
- Tversky, A., & Kahneman, D. (1974). Judgment under uncertainty: Heuristics and biases. *Science*, 185(4157), 1124–1131. <http://www.jstor.org/stable/1738360>
- Tversky, A., & Kahneman, D. (1981). The framing of decisions and the psychology of choice. *Science*, 211(4481), 453–458. <https://doi.org/10.1126/science.7455683>
- Zhang, Z. T., Feger, S. S., Dullenkopf, L., Liao, R., Süßlin, L., Liu, Y., & Butz, A. (2024). Beyond recommendations: From backward to forward ai support of pilots' decision-making process. *Proceedings of the ACM on Human-Computer Interaction*, 8(CSCW2), Article 485, 32. <https://doi.org/10.1145/3687024>
- Zhang, Z. T., Liu, Y., & Hussmann, H. (2021). Forward reasoning decision support: Toward a more complete view of the human-ai interaction design space. *Proceedings of the 14th Biannual Conference of the Italian SIGCHI Chapter*, Article 18, 1–5. <https://doi.org/10.1145/3464385.3464696>