

Towards Personalized Conversational Sales Agents: Contextual User Profiling for Strategic Action

Tongyoung Kim* Jeongeun Lee* Soojin Yoon Sunghwan Kim Dongha Lee†

Department of Artificial Intelligence, Yonsei University

{dykim, ljeadec31, soojiny, happysnail06, donalee}@yonsei.ac.kr

Abstract

Conversational Recommender Systems (CRSs) aim to engage users in dialogue to provide tailored recommendations. While traditional CRSs focus on eliciting preferences and retrieving items, real-world e-commerce interactions involve more complex decision-making, where users consider multiple factors beyond simple attributes. To capture this complexity, we introduce Conversational Sales (CSALES), a novel task that integrates preference elicitation, recommendation, and persuasion within a unified conversational framework. To support realistic and systematic evaluation, we present CSUSER, an evaluation protocol with LLM-based user simulator grounded in real-world behavioral data by modeling fine-grained user profiles for personalized interaction. We also propose CSI, a conversational sales agent that proactively infers contextual user profiles and strategically selects actions through conversation. Comprehensive experiments show that CSI significantly improves both recommendation success and persuasive effectiveness across diverse user profiles. The codes are publicly available at github.com/kimyoung/CSI.

1 Introduction

Conversational Recommender Systems (CRSs) aim to engage users in dialogue to better understand their preferences and provide personalized recommendations (Li et al., 2018; Sun and Zhang, 2018; Wang et al., 2022). With advancements in large language models (LLMs), research has increasingly focused on LLM-based CRSs (Feng et al., 2023a; Li et al., 2024a), leveraging their reasoning capabilities. Through multi-turn interactions, these systems should capture both explicitly stated user preferences and implicit needs from context to enhance recommendation quality (Yi et al., 2024).

*Equal contribution

†Corresponding author

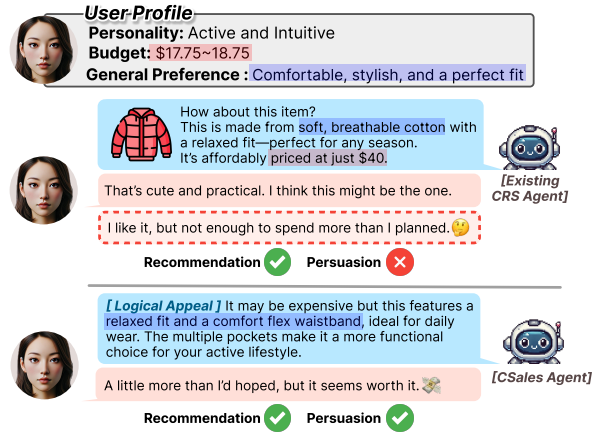


Figure 1: An illustrative case highlighting the need for CSALES. A successful recommendation alone fails to result in a purchase, but the addition of a personalized persuasive explanation leads to user acceptance.

However, existing CRS approaches are limited to eliciting user preferences and recommending relevant items, making them insufficient for real-world e-commerce scenarios (Ye et al., 2024). While effective in content domains like movies or books (He et al., 2023; Li et al., 2025), these methods fall short in complex purchasing contexts, where user decisions are shaped by multiple factors, such as long-term preferences, immediate intent, financial constraints, and personal motivations (Papenmeier et al., 2022). In these settings, merely matching preferences is inadequate (Shi et al., 2023; Li et al., 2024b; Roumeliotis et al., 2024); systems must actively guide users toward informed purchasing decisions and justify recommendations persuasively to influence actual purchase (Liu et al., 2023; Sayana et al., 2024).

In this context, we propose a novel task CSALES that extends CRSs to more effectively engage users in real-world conversational e-commerce scenarios. The goal of CSALES includes three key aspects: (1) **proactively** eliciting user implicit preferences by asking clarifying questions, (2) **accurately** retriev-

Method	Task	User Profile					
		Personality	Success Criteria	Real User	Target Needs	General Pref.	Budget
iEvaLM (Wang et al., 2023c)	Movie	-	Target Item	✗	✓	✗	✗
PEPPER (Kim et al., 2024b)	Movie	-	Target Item	✓	✗	✓	✗
CONCEPT (Huang et al., 2024)	Movie	Age, Persona	Genre	✗	✓	✗	✗
CSUSER (Ours)	E-commerce	Openness, Decision-Making	Reason to Purchase	✓	✓	✓	✓

Table 1: Comparison of existing LLM-based user simulation frameworks for conversational recommendation (Movie) and conversational sales (CS) tasks. Each user profile component is described in detail in Section 2.1.

ing and recommending items that align with user needs, and (3) providing **persuasive** explanations that adapt complex user characteristics accounting for both constraint factors and explicit product attributes. Unlike traditional recommendation, persuasive recommendation involves a deeper level of personalization, as it requires the system to convert potential interest into actual purchase decisions in e-commerce settings (Yoo and Gretzel, 2010).

For reliable evaluation of CSALES, we propose CSUSER, an evaluation protocol with user simulators. Given the high cost and limited scalability of human evaluation, user simulators have become a widely adopted approach for assessing interactive systems (Lei et al., 2020; Wang et al., 2023c; Sekulić et al., 2024). However, existing simulators rely on simplistic assumptions or synthetic personas, limiting their ability to capture diversity and complexity of real-world decision-making (Zhang et al., 2024b; Huang et al., 2024). CSUSER addresses this gap by modeling user profiles that capture key aspects of real-world behavior, including preferences, needs, and behavioral tendencies. These components are inferred from real-world interactions such as ratings, and user reviews (Hou et al., 2024) and product metadata, resulting in behaviorally grounded simulators. Conditioned on these profiles, simulators generate consistent, context-aware responses and decisions, enabling realistic and scalable evaluation of CSALES.

Following, we introduce **Conversational Sales Profiler**, named CSI, a conversational sales agent designed to guide users’ purchasing decisions through personalized and strategic interaction. CSI dynamically infers user profiles from dialogue context and selects actions from a unified action space that integrates preference elicitation, recommendation, and persuasion. This contextualized profile allows CSI to decide whether to explore further preferences or engage in persuasion. With profile-informed reasoning and adaptive action selection, CSI effectively aligns its conversational strategy

with the complex individual decision-making processes found in real-world e-commerce settings.

Our main contributions are as follows:

- We introduce CSALES, a novel task that unifies preference elicitation, recommendation, and persuasion to model real-world e-commerce decision making process.
- We propose CSUSER, a user simulation framework with fine-grained user profiles constructed from real-world behavioral data for realistic and scalable evaluation.
- We present CSI, an LLM-based agent that dynamically profiles users and adaptively selects actions to elicit preferences, recommend, and deliver personalized persuasion.

2 CSUSER: User Simulators for CSALES

Overview. We introduce CSUSER, a user simulation framework consisting of simulators grounded in real-world user interaction histories. Each simulator is modeled with a rich user profile composed of fine-grained behavioral components, capturing both long-term preferences and short-term intent. To support evaluation in CSALES, we also propose a quantitative metric that measures persuasive effectiveness by assessing how often users accept items beyond their initial expected price range.

Existing User Simulators. Traditional user simulators are typically rule-based or constructed from synthetic personas (Rohde et al., 2018; Ie et al., 2019). Table 1 summarizes recent LLM-based user simulators, which primarily targets movie domain. They typically make decisions based on fixed target item attributes, such as genre or product features. However, such designs often fail to capture the diversity and complexity of real-world decision-making (Corecco et al., 2024; Zhang et al., 2025; Cai et al., 2025). In contrast, CSUSER models user behavior more comprehensively by constructing unified profiles that reflect diverse factors influencing real-world purchasing decisions.

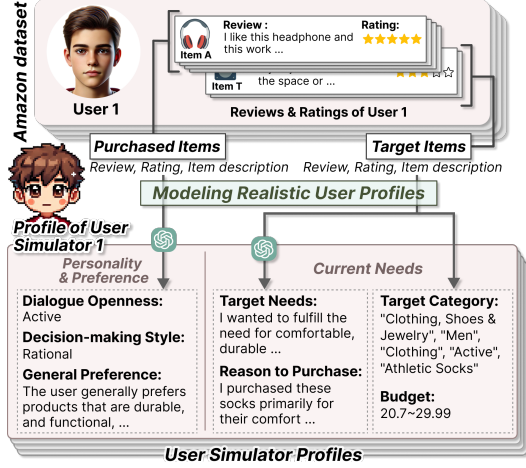


Figure 2: Profile construction for our user simulator, based on interaction history of real-world users.

2.1 Modeling User Simulators with Real-World Data

Figure 2 illustrates the process of constructing user profiles from real-world interactions. We use the Amazon Reviews 2023 dataset¹ (Hou et al., 2024), which offers large-scale user-item interaction data across a wide range of product categories. To reflect realistic purchasing behavior, we use GPT-4 (Achiam et al., 2023) to infer fine-grained profile components via LLM prompting (Sayana et al., 2024).² To achieve human-like, context-aware simulation, it is essential to balance long-term preferences with immediate intent. We derive long-term preferences from historical purchases, while intent is modeled using recent purchases within the same product category (i.e., target items). Below, we describe the key components of our user profiles.

Preference and Personality. To capture the complexity of user behavior in dialogue, our simulators incorporate both general preferences and personality traits. General preference (Kim et al., 2024a) represents a user’s broad inclinations across product categories by identifying recurring patterns in purchase history and analyzing associated product descriptions and reviews. In addition, we model personality with two key dimensions: Dialogue Openness (Xu et al., 2020) and Decision-Making Style (Scott and Bruce, 1995). Dialogue openness is determined by the level of detail in a user’s written reviews, indicating their tendency to share information during conversations. Active users provide detailed explanations, whereas passive users offer minimal input and require more probing to un-

¹<https://amazon-reviews-2023.github.io>

²Prompts are detailed in Appendix A.2.

Statistic	Clothing	Electronics
Amazon Review Dataset Statistics		
#Users	2,195	7,568
#Items	4,328	11,893
#Interactions	29,426	108,518
#Categories	592	720
Personality Trait Categories		
Dialogue-Openness	Active, Neutral, Passive	
Decision-making Style	Rational, Dependent, Intuitive	
Profile Richness (Avg. Length)		
General Preference	706.37	769.64
Target Needs	131.74	147.22
Purchase Reason	791.62	814.39

Table 2: Statistics for CSUSER user profiles across Clothing and Electronics domains.

cover preferences. Decision-making style captures the cognitive orientation behind a user’s decisions. This trait reflects whether the user relies on rational analysis, intuition, or external validation when evaluating items. Details are provided in Table 6.³

Current Needs. We define current needs as the key factors that drive a user’s purchase decisions, reflecting their immediate intent. These needs are extracted from target item metadata and user feedback (e.g., review). Each profile includes a target category (the domain of interest) and a budget, which is estimated from the price distribution of the target items to represent spending expectations. Target Needs refer to specific attribute-level preferences inferred from review and item descriptions, while the Reason to Purchase represents the user’s underlying motivation for purchasing target item.

Table 2 summarizes key statistics of CSUSER, highlighting the diversity and coverage of its profile components. These components provide a comprehensive representation of the situational and behavioral factors that shape individual users’ decision-making in realistic e-commerce interactions.

2.2 Evaluation Metrics for CSALES

Evaluating CSALES involves assessing agent performance across its three core components: preference elicitation, recommendation, and persuasion. A commonly used metric in conversational agents is *Success Rate (SR)* (Huang et al., 2024; Zhang et al., 2024b), which measures how often the user accepts a recommended item. A high SR suggests

³We validate LLM-inferred personality traits through a human verification conducted on Amazon Mechanical Turk (MTurk). Two independent annotators labeled a random subset of user profiles, achieving substantial agreement—80.5% for the Clothing domain and 70.1% for Electronics.

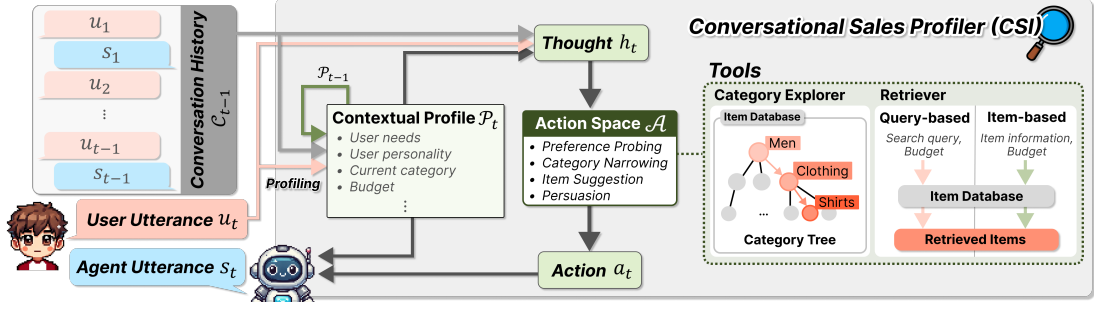


Figure 3: Overview of our CSI agent. At each turn t , it updates the previous profile $\mathcal{P}_{t-1}$ based on the user’s utterance u_t , then the next action is selected through reasoning (i.e., thought h_t) generated by the contextual profile $\mathcal{P}_t$ and conversational history $\mathcal{C}_{t-1} \oplus u_t$. To accomplish each action, it utilizes external tools, category explorer and retriever, which enables communication with the item database.

that the agent effectively elicits preferences and provides appropriate items. However, SR alone fails to reflect the agent’s persuasive capacity to influence user decisions beyond preference matching.

To complement this, we introduce *Sales-Win-Rate (SWR)*, which measures how effectively the agent persuades users to choose higher-priced, out-of-budget items—reflecting its impact on both decision-making and potential profitability. Specifically, let $\mathcal{I}$ denote the full item set and a_n the item accepted by user n . We define $\mathcal{I}_n^{\text{in}}$ as the set of in-budget items for user n , and $\mathcal{I}_n^{\text{out}}$ as those exceeding the user’s budget. SWR is then computed as the proportion of accepted items that exceed the user’s expected price range:

$$\text{SWR} = \frac{\sum_{n=1}^N \mathbb{I}(a_n \in \mathcal{I}_n^{\text{out}})}{\sum_{n=1}^N \mathbb{I}(a_n \in \mathcal{I}_n^{\text{in}} \cup \mathcal{I}_n^{\text{out}})}$$

where N is the number of total users, and $\mathbb{I}(\cdot)$ is the indicator function. While traditional CRS methods define success as the acceptance of $a_n \in \mathcal{I}_n^{\text{in}}$, our objective is to extend this goal by encouraging users to choose $a_n \in \mathcal{I}_n^{\text{out}}$, thereby demonstrating persuasive ability and increasing potential profitability. A higher SWR indicates stronger influence on user decisions, validating the agent’s effectiveness in persuasive recommendation.

3 Conversational Sales Agent with Contextual Profiling

In this section, we present CSI, a conversational sales agent designed to strategically guide users toward purchase decisions in CSALES. CSI dynamically infers a user profile throughout the conversation to select actions that aligned with both long-term preferences and immediate needs.

Figure 3 illustrates an overview of CSI’s process within a single dialogue turn. The core of CSI is a contextual user profile $\mathcal{P}_t$, which is continuously updated based on the user’s latest utterance u_t and the preceding conversation history $\mathcal{C}_{t-1} = \{u_1, s_1, \dots, s_{t-1}\}$. This profile informs the CSI by guiding the selection of the next action $a_t \in \mathcal{A}$ from a unified action space that spans preference elicitation, item recommendation, and persuasive explanation. At last, it conditions the generation of the system response s_t . By reasoning over this evolving profile, CSI adaptively responds to user input, uncovers implicit needs, and delivers strategically personalized recommendations.

3.1 Contextual User Profiling

We define a structured contextual profile $\mathcal{P}$ that CSI continuously updates and reasons over throughout conversation. At each turn t , given the conversation history $\mathcal{C}_{t-1}$ and the latest user utterance u_t , CSI updates current profile $\mathcal{P}_t$ by selectively retaining relevant information from $\mathcal{P}_{t-1}$, as $\mathcal{P}_t \sim \text{P}_{\text{LLM}}(\mathcal{P}_{t-1}, \mathcal{C}_{t-1}, u_t)$. This step prevents unnecessary accumulation and maintains a focused representation of user state. Next, CSI generates a latent reasoning step $h_t \sim \text{P}_{\text{LLM}}(\mathcal{P}_t, \mathcal{C}_{t-1}, u_t)$, which serves as an internal reflection (Yao et al., 2022). This step helps identify missing profile components and decide what to elicit, thereby enabling deeper and more adaptive user understanding.

3.2 Unified Action Space

Unlike prior works that define action spaces for single-purpose tasks (Zhang et al., 2024a,b), CSI introduces a unified action space $\mathcal{A}$ that supports elicitation, recommendation, and persuasion⁴

⁴Detailed descriptions of actions are provided in Table 7.

Method	Dialogue Openness						Decision-Making Style						Overall	
	Active		Neutral		Passive		Rational		Dependent		Intuitive			
	SR	SWR	SR	SWR	SR	SWR	SR	SWR	SR	SWR	SR	SWR	SR	SWR
Clothing														
ChatCRS	0.300	0.277	0.253	0.201	0.243	0.192	0.300	0.044	<u>0.300</u>	0.089	<u>0.320</u>	0.125	0.286	0.154
MACRS	0.310	0.266	0.241	0.514	0.432	0.267	0.359	0.688	0.324	0.200	0.336	0.250	0.334	0.364
PC-CRS	0.236	0.761	0.340	<u>0.745</u>	0.180	<u>0.741</u>	<u>0.367</u>	<u>0.667</u>	0.300	<u>0.767</u>	0.260	<u>0.638</u>	0.280	<u>0.719</u>
CSI w/o Profile	<u>0.367</u>	<u>0.781</u>	<u>0.349</u>	0.615	<u>0.467</u>	0.671	0.295	0.526	0.291	0.548	0.314	0.621	<u>0.347</u>	0.627
CSI	0.507	0.821	0.450	0.879	0.527	0.818	0.460	0.829	0.503	0.834	0.420	0.909	0.478	0.849
Electronics														
ChatCRS	0.117	0.191	0.187	0.099	0.117	0.194	0.167	0.062	0.164	0.137	0.163	0.137	0.152	0.136
MACRS	0.211	0.111	0.366	0.229	0.196	0.241	0.280	0.114	0.195	0.167	0.246	0.156	0.249	0.170
PC-CRS	0.160	0.775	0.153	<u>0.767</u>	0.160	<u>0.775</u>	0.175	<u>0.785</u>	0.175	<u>0.680</u>	0.187	<u>0.643</u>	0.168	<u>0.737</u>
CSI w/o Profile	<u>0.449</u>	0.530	0.487	0.472	<u>0.607</u>	0.455	<u>0.404</u>	0.456	<u>0.497</u>	0.521	<u>0.404</u>	0.566	<u>0.501</u>	0.500
CSI	0.500	<u>0.767</u>	<u>0.460</u>	0.843	0.607	0.803	0.457	0.817	0.557	0.803	0.440	0.715	0.503	0.791

Table 3: Overall performance of various conversational recommendation (or sales) agents on CSUSER. We employ two metrics, SR and SWR, which are presented in Section 2.2.

within a single decision framework. Following (Deng et al., 2023c; Yao et al., 2024), $\mathcal{A}$ includes both language-based actions, which generate utterances from predefined instructions, and tool-based actions, which access external domain-specific information to support grounded responses.

Comparative Persuasion. A key component of $\mathcal{A}$ is *comparative persuasion*, which encourages users to accept higher-priced, out-of-budget items by highlighting their advantages over lower-cost alternatives. This contrastive reasoning enables justification of persuasive recommendations aligned with the user’s latent goals. To personalize persuasion, CSI selects a strategy $\pi_i \in \pi$ based on the current contextual profile and dialogue state. We define π^5 as a set of persuasive strategies (Wang et al., 2019; Deng et al., 2023c; Qin et al., 2024), each with a name and structured instructional prompt.

Tool Usage. To support grounded responses, CSI defines a tool interface $\mathcal{T}$ that executes external actions based on the agent’s intent. Specifically, the Category Explorer navigates a hierarchical taxonomy commonly used in e-commerce platforms to retrieve relevant subcategories and concept-level information. The Retriever operates in two modes: *query-based retrieval*, it converts a natural language query into dense embeddings and returns semantically relevant results; *item-based retrieval*, it fetches alternatives similar to a reference product, enabling comparative recommendation and persuasive explanation. These tools allow the agent to retrieve real-world product information based

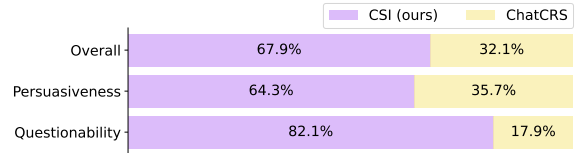


Figure 4: Pairwise human evaluation results of CSI and ChatCRS on conversational quality and proactiveness.

on the current dialogue context, enabling factually grounded and context-aware responses. The detailed explanations are provided in Appendix B.3.

4 Experiments

4.1 Experimental Settings

Baselines. We adapt existing LLM-based CRS agents for direct comparison. This allows us to examine whether current CRS architectures can support persuasion as an integral part of the overall recommendation process in CSALES. ChatCRS (Wang et al., 2023c; Huang et al., 2024) adopts a simple retrieve-and-recommend strategy at each turn, relying solely on conversation history. MACRS (Fang et al., 2024) follows a similar structure but introduces reflection-based action selection. PC-CRS (Qin et al., 2024) focuses entirely on persuasion, and like the other baselines, does not explicitly construct user profiles. To enable a fair comparison under the CSALES setting, we adapt ChatCRS and MACRS by applying the same persuasion instruction format used in CSI, allowing all agents to perform elicitation, recommendation, and persuasion. Further implementation details and prompt templates are provided in Appendix C.1.

⁵Full descriptions of π are provided in Table 8.

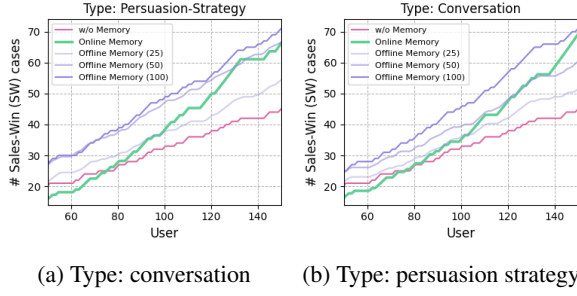


Figure 5: Impact of offline vs. online memory on persuasion performance. We investigate the effect of different memory sizes and memory types on persuasion success.

4.2 Overall Performance

To evaluate how effectively CSI delivers accurate recommendation and personalized persuasion in CSALES, we conduct experiments across user personality variations. As summarized in Table 3, CSI consistently outperforms all baselines, demonstrating strong adaptability to varying user characteristics. ChatCRS shows the lowest performance, as it lacks explicit action modeling and persuasion strategies. MACRS adds internal reasoning to guide action selection, resulting in a minor SR gain without affecting SWR. In contrast, PC-CRS adopts persuasive strategies alone, leading to a slight improvement in SWR while underperforming in SR due to the lack of contextual user profiling.

To isolate the contribution of contextual user profiling, we compare CSI with CSI (w/o Profile), a profile-agnostic variant that retains reasoning and persuasion capabilities. The performance gap between the two variants highlights that contextual user profiling significantly enhances both recommendation accuracy and persuasive effectiveness by enabling dynamic adaptation of strategies to individual users. We further validate this observation through human evaluation in Sections 4.3 and 4.5.

4.3 Qualitative Analysis of Conversation

To assess the proactiveness and conversational quality of CSI, we conduct a human evaluation using Amazon Mechanical Turk (MTurk) annotators. Proactiveness is evaluated along two dimensions: (1) *Questionability*—how effectively the agent elicits user needs through clarifying questions, and (2) *Persuasiveness*—how well the agent delivers persuasive explanations that guide users toward a purchase. As shown in Figure 4, we compare CSI and ChatCRS using pairwise win rates. CSI consistently outperforms ChatCRS, demonstrating superior user understanding and more adaptive, per-

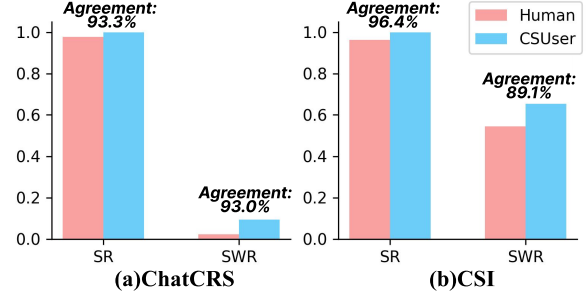


Figure 6: Agreement between human judgements and CSUSER on successful recommendation cases of CSI.

sonalized interactions, enabled by contextual user profiling. Detailed evaluation criteria and case studies are provided in Appendices C.3 and D.1.

4.4 Effectiveness of Strategy Memory

To enhance the persuasive capability of CSI, we incorporate a memory module $\mathcal{M}$ that leverages past successful interactions. For each instance where a user accepts an out-of-budget recommendation (i.e., $\text{SWR} = 1$), $\mathcal{M}$ stores the user profile as the key and the corresponding interaction as the value. At the persuasion step, CSI retrieves the top- k most similar profiles based on semantic similarity: $\mathcal{I} = \arg\min_{i \in \mathcal{M}, |\mathcal{I}|=k} \|e_p - e_i\|_2^2$, where e_p denotes the current user profile embedding and e_i denotes stored profile embeddings. Retrieved examples serve as supportive signals alongside the inferred user profile, enabling CSI to refine its persuasion strategies based on prior interactions.

We consider two memory configurations: (1) only persuasive strategy types and (2) full persuasive utterances (Figure 5). Each is tested under both offline memory (pre-built with 25–100 profiles) and online memory (updated incrementally during interaction). Details are provided in Appendix C.2.

Results show that memory-augmented agents outperform the baseline without memory in SWR. The offline memory configuration yields higher SWR in early stages, as it starts with access to pre-collected successful examples. In contrast, the online memory gradually improves as more persuasive interactions are accumulated. Furthermore, memory containing full utterances outperforms strategy-only memory, highlighting the importance of leveraging successful examples from users with similar profiles to support personalized persuasion.

4.5 Human Evaluation for CSUSER

To validate the reliability of CSUSER-based evaluation, we compare its decisions with human judgements on instances where a recommendation was

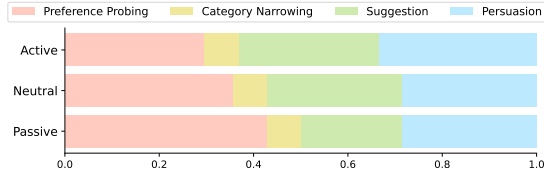


Figure 7: Action distributions (i.e., the proportion of executed actions) across different dialogue openness.

accepted (i.e., $SR = 1$) by either ChatCRS or CSI. For each case, MTurk annotators assessed whether the recommended item and accompanying persuasive attempt were acceptable, or if neither was sufficient. We compute agreement as the proportion of instances where human and simulator decisions matched, either both accept or both reject, for each evaluation dimension (SR and SWR). As shown in Figure 6, the results show strong alignment, with over 90% agreement on both metrics, demonstrating that CSUSER offers human-like evaluations.

4.6 Personality-based Analysis in CSALES

To understand how user personality traits influence action selection in CSI, we analyze system behavior across levels of Dialogue Openness. Figure 7 shows the distribution of actions selected for each personality group in CSALES. Passive users receive more Preference Probing and Category Narrowing actions, as the agent elicits their preferences before making suggestions or persuasive attempts. In contrast, active users provide richer feedback in the conversation, leading the agent to shift more quickly toward Suggestion and Persuasion, with less need for elicitation. Neutral users fall between these, receiving intermediate proportions of elicitation and decision-oriented actions. This analysis verifies that CSI adapts its actions to user personality traits, supporting the effectiveness of contextual user profiling in personalized recommendation.

To examine how decision-making styles influence acceptance behavior, we analyze how users with different styles respond to various persuasion strategies in CSALES. As shown in Figure 8, acceptance rates differ across styles under the same dialogue context, enabling a fair comparison. Rational users prefer logical, evidence-based appeals and show higher acceptance for fact-driven persuasion. Dependent users respond more to social proof, relying on external opinions and consensus. Intuitive users favor emotional and subjective appeals, exhibiting greater receptiveness to strategies based on feeling or instinct. These results indicate that CSUSER effectively captures decision-style-

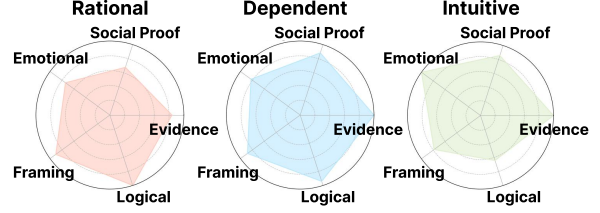


Figure 8: Persuasion strategy acceptance rate across different *Decision-making Styles*.

specific tendencies, enabling controlled evaluation of persuasion effectiveness across diverse users.

These findings highlight the importance of incorporating personality traits into CSUSER to better reflect diverse user behaviors and support more realistic evaluation. By adapting action selection and persuasion based on user profiles, CSI achieves greater alignment with individual needs, ultimately improving recommendation quality in CSALES.

4.7 Validation for CSUSER Components

To assess the contribution of each user profile component in CSUSER, we conduct supplementary experiments with ablated simulators in which specific elements of the profile are removed.⁶

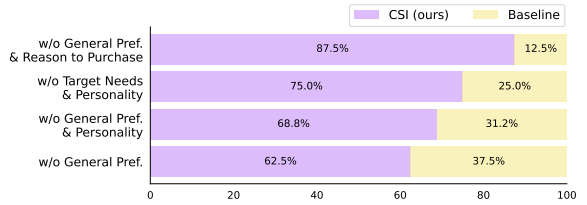


Figure 9: Similarity win rate between actual target items and accepted items by ablating profile components.

Although SR indicates whether a recommendation is accepted, it fails to assess whether the accepted item aligns with the user’s true intent. To address this, we compute the semantic similarity between each accepted item and the original target. For each user instance, we perform pairwise comparisons to see which variant yields a recommendation more semantically aligned with the target. As shown in Figure 9, CSUSER outperforms all ablated versions. The largest drop in performance occurs when both general preferences and reasons to purchase are removed, highlighting the critical role of these components in modeling user intent.

We further analyze the trade-off between SR and similarity score at each variant in Table 4. We observe that when specific components are omitted,

⁶Each ablated variant is based on simplified user modeling strategies used in prior simulators, as summarized in Table 1.

Components	SR	Similarity
CSUser	0.487	0.922
w/o General Pref.	0.530	0.909
w/o General Pref. & Personality	0.603	0.907
w/o Needs & Personality	0.617	0.900
w/o General Pref. & Reason to Purchase	0.637	0.884

Table 4: Results of SR and Similarity across different components.

the simulator’s acceptance rate (SR) increases, but the semantic similarity with the original target decreases. This trade-off suggests that less detailed profiles lead to more permissive but less realistic user behavior, where recommendations are more likely to be accepted regardless of alignment with true needs. Taken together, these findings demonstrate that detailed profiling enhances decision fidelity and reduces the overly agreeable tendencies commonly observed in LLM-based simulators.

4.8 Reproducibility of CSUSER and CSI

To examine the reproducibility and robustness of our framework, we evaluate both CSI and CSUSER using different underlying LLMs. Table 5 summarizes the results. When varying the agent-side model in CSI (Upper), we observe that the GPT-4o-mini version achieves the highest performance. This indicates that stronger reasoning capabilities from advanced LLMs enhance both recommendation accuracy and persuasive effectiveness. On the user-side (Lower), CSUSER remains stable across all tested LLMs. SR ranges narrowly between 0.46 and 0.47, and SWR remains consistently high above 0.80. These results confirm that CSUSER produces reliable user behavior independent of the underlying model, validating its utility as a robust and reproducible evaluation framework.

5 Related Work

Conversational Recommender Systems. CRSs have gained increasing attention in recent years, aiming to deliver personalized recommendations through interactive dialogues (Christakopoulou et al., 2016; Jannach et al., 2021; Friedman et al., 2023; Di Palma, 2023; He et al., 2023). A key challenge of CRSs is to understand user preferences, as users often express ambiguous or incomplete requests (Rahmani et al., 2023). To address this, previous research has focused on asking clarifying questions to refine user intent and thus improve recommendation accuracy (Aliannejadi et al., 2019; Ye et al., 2024; Feng et al., 2023b). More recently,

User	Agent	SR	SWR
gpt-3.5-turbo R1-Distill-Qwen	gpt-4o-mini	0.53	0.89
	gpt-3.5-turbo	0.47	0.80
	R1-Distill-Qwen	0.47	0.43
gpt-4o-mini gpt-3.5-turbo R1-Distill-Qwen	 gpt-3.5-turbo 	0.46 0.47 0.47	0.85 0.80 0.83

Table 5: Performance variations of our CSI agent on CSUSER across different LLMs (Upper). Performance variation of our CSUSER with different LLMs (Lower).

CRSs have been extended to provide persuasive explanations, enhancing user satisfaction by providing compelling justifications (Huang et al., 2024; Qin et al., 2024). However, they still struggle in real-world e-commerce, often failing to account for factors which lead to actual purchase behavior.

Evaluation via User Simulation. For cost-effective CRS evaluation (Huang et al., 2023), LLM-based user simulation has been utilized (Wang et al., 2023c; Kim et al., 2024b), typically relying on item attributes for decision-making (Zhang and Balog, 2020). However, the absence of personality traits significantly limits realism, as real users exhibit diverse characteristics. Although traits like Big Five (Goldberg, 1992; Yang et al., 2020) and personas (Huang et al., 2024) have been explored, they are randomly assigned rather than derived from real user, making them fabricated. To bridge this gap, we highlight the need for real-world data for realistic evaluation.

Proactive LLM-based Conversational Agents. While LLMs excel as conversational agents, their reactive nature limits effectiveness, prompting research into enhancing proactivity for better control (Liao et al., 2023; Deng et al., 2023b). Proactivity involves several key aspects. First, agents should actively guide conversations to achieve predefined task goals (Wu et al., 2019; Wang et al., 2023a,b). Second, they should generate clarifying questions for ambiguous user requests (Aliannejadi et al., 2019; Guo et al., 2021; Deng et al., 2022; Chang and Chen, 2024). Third, they must effectively interact with non-collaborative users, adapting strategies to navigate conflicting goals (Deng et al., 2023a,c; Fu et al., 2023; Zhang et al., 2024b). In this work, we develop an agent that proactively guides conversations through strategic question-asking and persuasion for CSALES.

6 Conclusion

We present CSALES, a novel task that integrates preference elicitation, recommendation, and persuasion to better reflect the complexity of real-world e-commerce interactions. To support realistic evaluation, we introduce CSUSER, modeling diverse user profiles with fine-grained components. We also propose CSI, that dynamically infers contextual user profiles to adapt its actions and persuasion strategies to individual behaviors. Through comprehensive experiments and human evaluations, we demonstrate the effectiveness of CSI and validate the reliability of CSUSER. Our findings highlight the importance of personalized interaction and realistic simulation in advancing CRS research. Future directions include enhancing long-term strategic planning and expanding simulation capabilities to broader domains.

Limitations

While our study provides meaningful contributions, several limitations warrant discussion. First, our experiments primarily rely on gpt-3.5-turbo, a proprietary LLM, which may limit generalizability to other architectures, particularly open-source models. To address this concern, we include supplementary results using additional models such as GPT-4o and R1-Distill-Qwen. These results suggest that our framework is robust across different LLM backbones, though broader validation remains.

Second, the user profiles in CSUSER are constructed from Amazon review data, which may introduce platform-specific or demographic biases. While such biases are a common concern in single-source datasets, we note that the Amazon corpus offers exceptionally broad coverage across product categories, user populations, and interaction types. Its scale and diversity make it one of the most general-purpose datasets available for modeling e-commerce behavior. To further reduce potential skew, we apply uniform sampling across diverse personality traits, ensuring a more balanced and representative simulation. We believe that these efforts help reduce the impact of these limitations, although further generalization to broader settings remains an important direction for future work.

Ethical Consideration

The output of text generation from LLMs may sometimes contain harmful, biased, or offensive content. However, in our research, we assert that

this risk is largely minimized. The source data used in the construction of our CSUSER are derived from Amazon review 2023 (Hou et al., 2024), both of which are publicly available datasets licensed under the MIT-License. Additionally, we manually review the generated dialogues to ensure they are free from biased, misleading, or inappropriate language, maintaining the quality and fairness of the interactions. For human evaluation, Amazon Mechanical Turk Annotators serve as annotators, receiving guidelines before participating in the assessment process. Each annotator is paid \$0.15 per task. The textual content presented in this paper contains no personally identifiable information and poses no risk of re-identifying individuals or groups.

Acknowledgement

This work was supported by the IITP grants funded by the Korea government (MSIT) (No. RS-2020-II201361; RS-2024-00457882, AI Research Hub Project), and the NRF grant funded by the Korea government (MSIT) (No. RS-2025-00560295).

References

- Josh Achiam, Steven Adler, Sandhini Agarwal, Lama Ahmad, Ilge Akkaya, Florencia Leoni Aleman, Diogo Almeida, Janko Altschmidt, Sam Altman, Shyamal Anadkat, et al. 2023. Gpt-4 technical report. *arXiv preprint arXiv:2303.08774*.
- Mohammad Aliannejadi, Hamed Zamani, Fabio Crestani, and W Bruce Croft. 2019. Asking clarifying questions in open-domain information-seeking conversations. In *Proceedings of the 42nd international acm sigir conference on research and development in information retrieval*, pages 475–484.
- Hongru Cai, Yongqi Li, Wenjie Wang, Fengbin Zhu, Xiaoyu Shen, Wenjie Li, and Tat-Seng Chua. 2025. Large language models empowered personalized web agents. In *Proceedings of the ACM on Web Conference 2025*, pages 198–215.
- Jiangxia Cao, Shaoshuai Li, Bowen Yu, Xiaobo Guo, Tingwen Liu, and Bin Wang. 2023. Towards universal cross-domain recommendation. In *Proceedings of the sixteenth ACM international conference on web search and data mining*, pages 78–86.
- Wen-Yu Chang and Yun-Nung Chen. 2024. Injecting salesperson’s dialogue strategies in large language models with chain-of-thought reasoning. *arXiv preprint arXiv:2404.18564*.
- Jin Chen, Zheng Liu, Xu Huang, Chenwang Wu, Qi Liu, Gangwei Jiang, Yuanhao Pu, Yuxuan Lei, Xiaolong

- Chen, Xingmei Wang, et al. 2024. When large language models meet personalization: Perspectives of challenges and opportunities. *World Wide Web*, 27(4):42.
- Konstantina Christakopoulou, Filip Radlinski, and Katja Hofmann. 2016. Towards conversational recommender systems. In *Proceedings of the 22nd ACM SIGKDD international conference on knowledge discovery and data mining*, pages 815–824.
- Nathan Corecco, Giorgio Piatti, Luca A Lanzendörfer, Flint Xiaofeng Fan, and Roger Wattenhofer. 2024. Suber: An rl environment with simulated human behavior for recommender systems. *arXiv preprint arXiv:2406.01631*.
- Yang Deng, Wenqiang Lei, Minlie Huang, and Tat-Seng Chua. 2023a. Goal awareness for conversational ai: proactivity, non-collaborativity, and beyond. In *Proceedings of the 61st Annual Meeting of the Association for Computational Linguistics (Volume 6: Tutorial Abstracts)*, pages 1–10.
- Yang Deng, Wenqiang Lei, Wai Lam, and Tat-Seng Chua. 2023b. A survey on proactive dialogue systems: Problems, methods, and prospects. *arXiv preprint arXiv:2305.02750*.
- Yang Deng, Wenqiang Lei, Wenxuan Zhang, Wai Lam, and Tat-Seng Chua. 2022. Pacific: towards proactive conversational question answering over tabular and textual data in finance. *arXiv preprint arXiv:2210.08817*.
- Yang Deng, Wenxuan Zhang, Wai Lam, See-Kiong Ng, and Tat-Seng Chua. 2023c. Plug-and-play policy planner for large language model powered dialogue agents. In *The Twelfth International Conference on Learning Representations*.
- Dario Di Palma. 2023. Retrieval-augmented recommender system: Enhancing recommender systems with large language models. In *Proceedings of the 17th ACM Conference on Recommender Systems*, pages 1369–1373.
- Matthijs Douze, Alexandr Guzhva, Chengqi Deng, Jeff Johnson, Gergely Szilvasy, Pierre-Emmanuel Mazaré, Maria Lomeli, Lucas Hosseini, and Hervé Jégou. 2024. The faiss library. *arXiv preprint arXiv:2401.08281*.
- Jiabao Fang, Shen Gao, Pengjie Ren, Xiuying Chen, Suzan Verberne, and Zhaochun Ren. 2024. A multi-agent conversational recommender system. *arXiv preprint arXiv:2402.01135*.
- Yue Feng, Shuchang Liu, Zhenghai Xue, Qingpeng Cai, Lantao Hu, Peng Jiang, Kun Gai, and Fei Sun. 2023a. A large language model enhanced conversational recommender system. *arXiv preprint arXiv:2308.06212*.
- Yue Feng, Hossein A Rahmani, Aldo Lipani, and Emine Yilmaz. 2023b. Towards asking clarification questions for information seeking on task-oriented dialogues. *arXiv preprint arXiv:2305.13690*.
- Luke Friedman, Sameer Ahuja, David Allen, Zhenning Tan, Hakim Sidahmed, Changbo Long, Jun Xie, Gabriel Schubiner, Ajay Patel, Harsh Lara, et al. 2023. Leveraging large language models in conversational recommender systems. *arXiv preprint arXiv:2305.07961*.
- Yao Fu, Hao Peng, Tushar Khot, and Mirella Lapata. 2023. Improving language model negotiation with self-play and in-context learning from ai feedback. *arXiv preprint arXiv:2305.10142*.
- Lewis R Goldberg. 1992. The development of markers for the big-five factor structure. *Psychological assessment*, 4(1):26.
- Meiqi Guo, Mingda Zhang, Siva Reddy, and Malihe Alikhani. 2021. Abg-coqa: Clarifying ambiguity in conversational question answering. In *3rd Conference on Automated Knowledge Base Construction*.
- Zhankui He, Zhouhang Xie, Rahul Jha, Harald Steck, Dawen Liang, Yesu Feng, Bodhisattwa Prasad Majumder, Nathan Kallus, and Julian McAuley. 2023. Large language models as zero-shot conversational recommenders. In *Proceedings of the 32nd ACM international conference on information and knowledge management*, pages 720–730.
- Yupeng Hou, Jiacheng Li, Zhankui He, An Yan, Xiusi Chen, and Julian McAuley. 2024. Bridging language and items for retrieval and recommendation. *arXiv preprint arXiv:2403.03952*.
- Chen Huang, Peixin Qin, Yang Deng, Wenqiang Lei, Jiancheng Lv, and Tat-Seng Chua. 2024. Concept-an evaluation protocol on conversational recommender systems with system-centric and user-centric factors. *CoRR*.
- Chen Huang, Peixin Qin, Wenqiang Lei, and Jiancheng Lv. 2023. Reduce human labor on evaluating conversational information retrieval system: A human-machine collaboration approach. In *Proceedings of the 2023 Conference on Empirical Methods in Natural Language Processing*, pages 10876–10891.
- Eugene Ie, Chih-wei Hsu, Martin Mladenov, Vihan Jain, Sanmit Narvekar, Jing Wang, Rui Wu, and Craig Boutilier. 2019. Recsim: A configurable simulation platform for recommender systems. *arXiv preprint arXiv:1909.04847*.
- Dietmar Jannach, Ahtsham Manzoor, Wanling Cai, and Li Chen. 2021. A survey on conversational recommender systems. *ACM Computing Surveys (CSUR)*, 54(5):1–36.
- Minjin Kim, Minju Kim, Hana Kim, Beong-woo Kwak, Soyeon Chun, Hyunseo Kim, SeongKu Kang, Youngjae Yu, Jinyoung Yeo, and Dongha Lee. 2024a. Pearl:

- A review-driven persona-knowledge grounded conversational recommendation dataset. *arXiv preprint arXiv:2403.04460*.
- Sunghwan Kim, Tongyoung Kim, Kwangwook Seo, Jinyoung Yeo, and Dongha Lee. 2024b. Stop playing the guessing game! target-free user simulation for evaluating conversational recommender systems. *arXiv preprint arXiv:2411.16160*.
- Jeongeun Lee, Seongku Kang, Won-Yong Shin, Jeongwhan Choi, Noseong Park, and Dongha Lee. 2024. Graph signal processing for cross-domain recommendation. *arXiv preprint arXiv:2407.12374*.
- Wenqiang Lei, Xiangnan He, Yisong Miao, Qingyun Wu, Richang Hong, Min-Yen Kan, and Tat-Seng Chua. 2020. Estimation-action-reflection: Towards deep interaction between conversational and recommender systems. In *Proceedings of the 13th International Conference on Web Search and Data Mining*, pages 304–312.
- Chuang Li, Yang Deng, Hengchang Hu, Min-Yen Kan, and Haizhou Li. 2024a. Incorporating external knowledge and goal guidance for llm-based conversational recommender systems. *arXiv preprint arXiv:2405.01868*.
- Chuang Li, Yang Deng, Hengchang Hu, Min-Yen Kan, and Haizhou Li. 2025. Chaters: Incorporating external knowledge and goal guidance for llm-based conversational recommender systems. In *Findings of the Association for Computational Linguistics: NAACL 2025*, pages 295–312.
- Raymond Li, Samira Ebrahimi Kahou, Hannes Schulz, Vincent Michalski, Laurent Charlin, and Chris Pal. 2018. Towards deep conversational recommendations. *Advances in neural information processing systems*, 31.
- Yangning Li, Shirong Ma, Xiaobin Wang, Shen Huang, Chengyue Jiang, Hai-Tao Zheng, Pengjun Xie, Fei Huang, and Yong Jiang. 2024b. Ecomgpt: Instruction-tuning large language models with chain-of-task tasks for e-commerce. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 38, pages 18582–18590.
- Lizi Liao, Grace Hui Yang, and Chirag Shah. 2023. Proactive conversational agents. In *Proceedings of the Sixteenth ACM International Conference on Web Search and Data Mining*, pages 1244–1247.
- Yuanxing Liu, Wei-Nan Zhang, Yifan Chen, Yuchi Zhang, Haopeng Bai, Fan Feng, Hengbin Cui, Yongbin Li, and Wanxiang Che. 2023. Conversational recommender system and large language model are made for each other in e-commerce pre-sales dialogue. *arXiv preprint arXiv:2310.14626*.
- Jianmo Ni, Gustavo Hernandez Abrego, Noah Constant, Ji Ma, Keith B Hall, Daniel Cer, and Yinfei Yang. 2021. Sentence-t5: Scalable sentence encoders from pre-trained text-to-text models. *arXiv preprint arXiv:2108.08877*.
- Andrea Papenmeier, Alexander Frummet, and Dagmar Kern. 2022. “mhm...”–conversational strategies for product search assistants. In *Proceedings of the 2022 Conference on Human Information Interaction and Retrieval*, pages 36–46.
- Peixin Qin, Chen Huang, Yang Deng, Wenqiang Lei, and Tat-Seng Chua. 2024. Beyond persuasion: Towards conversational recommender system with credible explanations. *arXiv preprint arXiv:2409.14399*.
- Hossein A Rahmani, Xi Wang, Yue Feng, Qiang Zhang, Emine Yilmaz, and Aldo Lipani. 2023. A survey on asking clarification questions datasets in conversational systems. *arXiv preprint arXiv:2305.15933*.
- David Rohde, Stephen Bonner, Travis Dunlop, Flavian Vasile, and Alexandros Karatzoglou. 2018. Recogym: A reinforcement learning environment for the problem of product recommendation in online advertising. *arXiv preprint arXiv:1808.00720*.
- Konstantinos I Roumeliotis, Nikolaos D Tselikas, and Dimitrios K Nasiopoulos. 2024. Llms in e-commerce: a comparative analysis of gpt and llama models in product review evaluation. *Natural Language Processing Journal*, 6:100056.
- Krishna Sayana, Raghavendra Vasudeva, Yuri Vasilevski, Kun Su, Liam Hebert, James Pine, Hubert Pham, Ambarish Jash, and Sukhdeep Sodhi. 2024. Beyond retrieval: Generating narratives in conversational recommender systems. *arXiv preprint arXiv:2410.16780*.
- Susanne G Scott and Reginald A Bruce. 1995. Decision-making style: The development and assessment of a new measure. *Educational and psychological measurement*, 55(5):818–831.
- Ivan Sekulić, Silvia Terragni, Victor Guimarães, Nghia Khau, Bruna Guedes, Modestas Filipavicius, André Ferreira Manso, and Roland Mathis. 2024. Reliable llm-based user simulator for task-oriented dialogue systems. *arXiv preprint arXiv:2402.13374*.
- Kaize Shi, Xueyao Sun, Dingxian Wang, Yinlin Fu, Guandong Xu, and Qing Li. 2023. Llama-e: Empowering e-commerce authoring with multi-aspect instruction following. *CoRR*.
- Yueming Sun and Yi Zhang. 2018. Conversational recommender system. In *The 41st international acm sigir conference on research & development in information retrieval*, pages 235–244.
- Gu Tang, Xiaoying Gan, Jinghe Wang, Bin Lu, Lyuwen Wu, Luoyi Fu, and Chenghu Zhou. 2024. Editkg: Editing knowledge graph for recommendation. In *Proceedings of the 47th International ACM SIGIR Conference on Research and Development in Information Retrieval*, pages 112–122.
- Jian Wang, Yi Cheng, Dongding Lin, Chak Tou Leong, and Wenjie Li. 2023a. Target-oriented proactive dialogue systems with personalization: Problem

- formulation and dataset curation. *arXiv preprint arXiv:2310.07397*.
- Jian Wang, Dongding Lin, and Wenjie Li. 2023b. Dialogue planning via brownian bridge stochastic process for goal-directed proactive dialogue. *arXiv preprint arXiv:2305.05290*.
- Xiaolei Wang, Xinyu Tang, Wayne Xin Zhao, Jingyuan Wang, and Ji-Rong Wen. 2023c. Rethinking the evaluation for conversational recommendation in the era of large language models. *arXiv preprint arXiv:2305.13112*.
- Xiaolei Wang, Kun Zhou, Ji-Rong Wen, and Wayne Xin Zhao. 2022. Towards unified conversational recommender systems via knowledge-enhanced prompt learning. In *Proceedings of the 28th ACM SIGKDD Conference on Knowledge Discovery and Data Mining*, pages 1929–1937.
- Xuwei Wang, Weiyan Shi, Richard Kim, Yoojung Oh, Sijia Yang, Jingwen Zhang, and Zhou Yu. 2019. Persuasion for good: Towards a personalized persuasive dialogue system for social good. *arXiv preprint arXiv:1906.06725*.
- Wenquan Wu, Zhen Guo, Xiangyang Zhou, Hua Wu, Xiyuan Zhang, Rongzhong Lian, and Haifeng Wang. 2019. Proactive human-machine conversation with explicit conversation goals. *arXiv preprint arXiv:1906.05572*.
- Hu Xu, Seungwhan Moon, Honglei Liu, Bing Liu, Pararth Shah, and Philip S Yu. 2020. User memory reasoning for conversational recommendation. *arXiv preprint arXiv:2006.00184*.
- Runzhe Yang, Jingxiao Chen, and Karthik Narasimhan. 2020. Improving dialog systems for negotiation with personality modeling. *arXiv preprint arXiv:2010.09954*.
- Shunyu Yao, Noah Shinn, Pedram Razavi, and Karthik Narasimhan. 2024. τ -bench: A benchmark for tool-agent-user interaction in real-world domains. *arXiv preprint arXiv:2406.12045*.
- Shunyu Yao, Jeffrey Zhao, Dian Yu, Nan Du, Izhak Shafran, Karthik Narasimhan, and Yuan Cao. 2022. React: Synergizing reasoning and acting in language models. *arXiv preprint arXiv:2210.03629*.
- Jingheng Ye, Yong Jiang, Xiaobin Wang, Yinghui Li, Yangning Li, Hai-Tao Zheng, Pengjun Xie, and Fei Huang. 2024. Productagent: Benchmarking conversational product search agent with asking clarification questions. *arXiv preprint arXiv:2407.00942*.
- Zihao Yi, Jiarui Ouyang, Yuwen Liu, Tianhao Liao, Zhe Xu, and Ying Shen. 2024. A survey on recent advances in llm-based multi-turn dialogue systems. *arXiv preprint arXiv:2402.18013*.
- Kyung-Hyan Yoo and Ulrike Gretzel. 2010. Creating more credible and persuasive recommender systems: The influence of source characteristics on recommender system evaluations. *Recommender systems handbook*, pages 455–477.
- An Zhang, Yuxin Chen, Leheng Sheng, Xiang Wang, and Tat-Seng Chua. 2024a. On generative agents in recommendation. In *Proceedings of the 47th international ACM SIGIR conference on research and development in Information Retrieval*, pages 1807–1817.
- Shuo Zhang and Krisztian Balog. 2020. Evaluating conversational recommender systems via user simulation. In *Proceedings of the 26th acm sigkdd international conference on knowledge discovery & data mining*, pages 1512–1520.
- Tong Zhang, Chen Huang, Yang Deng, Hongru Liang, Jia Liu, Zujie Wen, Wenqiang Lei, and Tat-Seng Chua. 2024b. Strength lies in differences! towards effective non-collaborative dialogues via tailored strategy planning. *arXiv preprint arXiv:2403.06769*.
- Zijian Zhang, Shuchang Liu, Ziru Liu, Rui Zhong, Qingpeng Cai, Xiangyu Zhao, Chunxu Zhang, Qidong Liu, and Peng Jiang. 2025. Llm-powered user simulator for recommender system. In *Proceedings of the AAAI Conference on Artificial Intelligence*, volume 39, pages 13339–13347.

A Constructing the User Simulator

A.1 Dataset for Realistic User Profiling and Simulation

We use the Amazon Reviews 2023 dataset (Hou et al., 2024), a large-scale, up-to-date corpus comprising user-generated reviews, ratings, and rich product metadata across diverse categories. This dataset has been widely adopted in recent work on recommendation (Cao et al., 2023; Lee et al., 2024; Tang et al., 2024) and user modeling (Chen et al., 2024), due to its broad coverage and high-quality behavioral signals. Its scale and diversity make it particularly well-suited for tasks involving personalized modeling and simulation. Especially, we focus on two representative domains-Clothing, Shoes & Jewelry and Electronics, which exhibit high diversity in product attributes and purchasing behavior. Each user instance includes structured interaction logs (e.g., numerical ratings, timestamps, category paths) and unstructured free-form text reviews, enabling fine-grained modeling of long-term preferences, session-specific needs, and personality traits essential for realistic user simulation.

A.2 Dataset Pre-Processing

To ensure reliable user modeling, we apply 10-core filtering to retain users and items with sufficient interaction density. Items missing critical metadata (e.g., description, category, or price) are excluded to ensure data completeness. Each user’s purchase history is sorted chronologically. The most recent item, along with other items in the same category path, is designated as the target, while the remaining purchases serve as historical data.

User profiles in CSUSER consist of three main components: General preferences and personality traits, which are inferred from historical items, capture stable behavioral tendencies. Current needs, which is derived from the target item(s), reflect situational intent and interest. These components are used to generate structured user profiles that inform both recommendation and dialogue behavior.

For evaluation, we randomly sample 150 users per personality type, ensuring coverage across variations in dialogue openness and decision-making styles. Prompt templates used for deriving user profile attributes are detailed in Tables 11–16.

A.3 More Details on Component

Decision-making-style. Following (Scott and Bruce, 1995), we infer key characteristics such as

Personality	Description
<i>Decision-making Style</i>	
Intuitive	Making a choice based on emotions or past experiences rather than a thorough analysis of all available details.
Dependent	Making a choice influenced by others’ opinions rather than conducting a thorough personal analysis.
Rational	Making a choice by focusing on product details, comparing options, and logically analyzing the decision before purchasing.
<i>Dialogue Openness</i>	
Active	Willing to share detailed reasons and explanations about preferences.
Neutral	Neither actively sharing nor avoiding explanations about preferences.
Passive	Unwilling to share detailed reasons or explanations about preferences.

Table 6: Definitions of user personality components used in CSALES

rational, intuitive, dependent decision-making tendencies from user purchase history, ratings, and reviews. These traits enable the construction of user simulators that realistically represent various decision-making processes. Each simulated user profile is designed to reflect distinct behavioral tendencies. For instance, a rational user is modeled to prioritize product specifications and detailed comparisons, whereas an intuitive user makes choices based on past experiences and personal inclinations. Similarly, a dependent user relies on external opinions such as ratings and reviews from others.

Interaction tendency. Interaction tendency is inferred from the user’s reviews and reflects how actively the user engages in providing feedback. This is categorized into three levels based on the length and detail of the reviews: Active, Neutral, and Passive. Users who provide detailed reviews are classified as Active, while those who offer shorter feedback fall into the Neutral category. Users who rarely provide purchase reasons in reviews are categorized as Passive. This classification helps the simulator understand the level of user engagement and tailor the interaction accordingly, ensuring that the system adapts to different user tendencies for a more accurate and proactive simulation.

B Details on CSI

B.1 Action Space

To model CSALES effectively, we define a universal action space $\mathcal{A}$ that encompasses preference

Stage	Description
Preference Probing	Ask about likes and dislikes to discover the Seeker’s preferences or interests.
Category Narrowing	Ensure the category path matches with the Seeker’s preferences.
Item Suggestion	Recommend items based on the profile.
Persuasion	Persuade the Seeker to purchase by highlighting why the item suits their needs.

Table 7: Definitions of actions in the universal action space $\mathcal{A}$. The actions are divided into language-based and tool-based categories to facilitate preference elicitation, recommendation, and persuasion in CSALES.

Persuasion Strategy	Description
Framing	Emphasize the unique advantages of <Candidate Item> that differentiate it from <Selected Item>.
Logical Appeal	Describe how the recommended item’s features are consistent with the user’s preference.
Emotional Appeal	Leverage emotions like anticipation, security, and satisfaction to encourage the purchase.
Evidence-Based	Use empirical data and facts such as item attributes to support your recommendation.
Social Proof	Highlight what the majority believes about the recommended item by showing its rating and reviews from other users.

Table 8: Definitions of Persuasion Strategies employed for user persuasion in CSALES

elicitation, recommendation, and persuasion. Table 7 lists specific actions and corresponding descriptions.

B.2 Persuasion Strategy

Following (Qin et al., 2024), we incorporate a diverse set of persuasion strategies. Table 8 lists these strategies. These strategies ensure that CSI can generate persuasive explanations tailored to individual preferences and cognitive tendencies, leading to more effective and personalized persuasion.

B.3 Tool

Category Search. To reflect real-world e-commerce scenarios, where each product belongs to a complex hierarchical category structure, we organize items within a tree-based hierarchy. CSI dy-

namically resolves ambiguous user demands by iteratively asking clarifying questions, guiding users toward the correct category. By progressively narrowing down the category path through user interaction, CSI improves retrieval accuracy and ensures that recommended items align more effectively with user intent.

Retriever. We construct an embedding database for all items in the domain using the T5-base sentence transformer (Ni et al., 2021). Each item’s title, category, and description are encoded into dense representations, enabling efficient similarity-based retrieval. For query-based retrieval, the retriever encodes a given query and retrieves the top- k most relevant items based on embedding similarity. This allows the system to identify items that align with the user’s specified preferences or search criteria. For item-based retrieval, the retriever takes a reference item as input and finds the top- k most similar items in the database. This facilitates product discovery by retrieving alternatives or complementary items within the same category. By leveraging embedding-based similarity search, our retriever ensures more precise and contextually relevant item retrieval, enhancing recommendation quality.

C Experiments

C.1 Implementation Details

C.1.1 ChatCRS

ChatCRS is a basic ChatGPT-based CSALES agent, using a retrieve-recommend-persuade approach at each turn. It utilizes the entire conversation history for retrieval, a common practice in conventional CRS methods. The agent retrieves two items: the most relevant item within-budget and another out-of-budget option. It then recommends the within-budget item while simultaneously persuading the user to consider the higher-priced alternative.

C.1.2 CSI (w/o Profile)

CSI (w/o Profile) serves as a straightforward ReAct (Yao et al., 2022)-based baseline for CSALES, utilizing the same action space as CSI. While it follows the original ReAct framework—generating thoughts and selecting actions accordingly—it is also explicitly tasked with extracting the user’s budget from the conversation, a key component in CSALES. When calling tool functions (retriever or category search), CSI (w/o Profile) formulates search queries using the conversation history, a

widely used strategy in conventional CRS methods.

C.2 Memory

We integrate a memory component $\mathcal{M}$ into CSI to enhance its adaptability in persuasion by leveraging past successful interactions. This module enables the agent to recall and apply effective persuasion strategies or utterances from previous conversations, improving its ability to guide users toward purchase decisions dynamically.

The memory module utilizes Faiss (Douze et al., 2024), a library optimized for fast and scalable dense vector retrieval. The memory stores user profiles as keys and successful persuasion instances as values, where the user accepted an out-of-budget recommendation (i.e., $\text{SWR}=1$).

During persuasion, CSI retrieves the top- k most similar user profiles based on embedding similarity and extracts their associated persuasion strategies:

$$\mathcal{I} = \operatorname{argmin}_{i \in M, |\mathcal{I}|=k} \|\mathbf{e}_p - \mathbf{e}_i\|_2^2$$

where $\mathbf{e}_p$ represents the embedding of the current user profile p , and $\mathbf{e}_i$ denotes stored profile embeddings. By referencing these stored memory, CSI can make more informed and personalized persuasion attempts.

C.3 Human Evaluation Criteria

As shown in Figure 4, we conduct a human evaluation to assess the proactiveness of CSI, focusing on: (1) its ability to generate clarifying questions that effectively elicit user needs and (2) its effectiveness in providing persuasive explanations that guide users toward a purchase decision.

Evaluation is based on the following criteria:

- **Questionability:** Measures how well the agent asks relevant and informative questions to refine user preferences. Annotators assess whether the generated questions clarify user intent, encourage detailed responses, and contribute to more precise recommendations. A higher score indicates that the agent proactively engages users in preference elicitation.
- **Persuasiveness:** Evaluates the agent’s ability to justify recommendations through persuasive explanations. Annotators consider whether the provided rationale aligns with user needs, highlights key product advantages, and effectively encourages users to consider

purchasing an item. A higher score reflects the agent’s capacity to present compelling arguments tailored to individual preferences.

More detailed instructions are in Figure 10 and Figure 11.

D Human Agreement

To validate the alignment between simulator-based and human evaluations, we adopt a simple agreement metric that quantifies how often both parties make the same decision—accept recommendation, accept persuasion, or reject. We define agreement as the proportion of instances in which the simulator and human annotators made an identical decision for a given evaluation dimension. Formally, let $D = \{(s_i, h_i)\}_{i=1}^N$ denote a set of N evaluation instances, where s_i is the simulator’s decision and h_i is the corresponding human judgment for instance i . Each decision s_i and h_i is one of three categorical labels: Accept-Recommendation, Accept-Persuasion, or Reject. Then, raw agreement A is computed as:

$$A = \frac{1}{N} \sum_{i=1}^N \mathbb{I}[s_i = h_i]$$

where $\mathbb{I}[\cdot]$ is the indicator function that returns 1 if the condition inside holds, and 0 otherwise. We compute this agreement score independently for each evaluation dimension—SR (successful recommendation) and SWR (successful persuasion)—allowing for a clear assessment of how closely the simulator replicates human decision patterns.

Detailed instructions for human judgment are in Figure 12 and Figure 13.

D.1 Case Study

We present case studies that illustrate how CSI engages in CSALES in Table 9 and Table 10

We are surveying qualities of LLM-generated Dialogues.

This evaluation process is designed to assess the quality of 'Recommender' who are providing recommendations and persuading the 'Seeker' to purchase the item.

Specifically, you will be given two dialogue candidates and will be asked to judge which Recommender is of a higher quality based on various aspects.

Guidelines:

[Q1~3] Evaluate the quality of dialogues based on different aspects, choose which Recommender is better regarding the given aspect according to the following criteria.

Dialogue 1

\$(dialogue_ours)

Dialogue 2

\$(dialogue_other)

Figure 10: Main Instructions

Question 1. Which Recommender asks more suitable questions to understand the Seeker's demand?

☒ 1
 ☐ 2

Question 2. Which Recommender provides more persuasvie statements tailored for the Seeker?

☒ 1
 ☐ 2

Question 3. Overall, which Recommender more effectively lead the Seeker to finally purchase the item?

☒ 1
 ☐ 2

Submit

Figure 11: Detailed Questions

We are surveying qualities of Recommender.

This evaluation process is designed to assess the quality of 'Recommender' who are providing recommendations and persuading the 'Seeker' to purchase the item.

Assume you are a user with the profile_below :

Figure 12: Instructions for Human Judgment

What you would select from the Recommender's items?

☒ B09R6ZZ2S4
 ☐
☐ None

You must ACCEPT the HIT before you can submit the results.

Figure 13: Question for Human Judgment

User: I want to find a cozy and stylish winter piece that will keep me warm while also being versatile enough to wear often. My expected price range is: [\$29.99, \$31.92].
CSI: <i>[Preference Probing]</i> What do you prefer in terms of material for your winter piece, such as wool, fleece, or synthetic fabrics?
User: I prefer materials that are warm and cozy, like wool or fleece, for my winter piece. They provide the comfort and insulation I need to stay warm during the colder months.
CSI: <i>[Preference Probing]</i> What type of style are you looking for in your winter piece, such as casual, formal, or sporty?
User: I prefer a style that is more on the casual side for my winter piece. I like pieces that I can easily dress up or down depending on the occasion, providing versatility in my wardrobe.
CSI: <i>[Preference Probing]</i> What do you prefer in terms of color or pattern for your winter piece, such as solid colors, plaids, or other designs?
User: I enjoy versatile designs for my winter piece, so I prefer solid colors or subtle patterns that can easily be paired with different outfits. This allows me to mix and match my wardrobe effortlessly for various occasions.
CSI: <i>[Suggestion]</i> Here are some items that you might like:
User: I was initially drawn to the <i>Lock and Love Women's Lightweight Long Sleeve Cozy Stylish Turtleneck Top Pullover Sweater</i> (B0BMH1DL3X) because of its classic turtleneck design and lightweight fabric, which seems perfect for layering during the winter. However, I would like to know more about the specific material used in this sweater to ensure it meets my warmth and comfort needs.
CSI: <i>[Logical Appeal]</i> While the <i>Lock and Love Turtleneck Sweater</i> (B0BMH1DL3X) is a great choice at \$18.75, I highly recommend considering the <i>Cable Stitch Women's Raglan Shaker Stitch Pullover</i> (B097FFSP2R) for \$54.50 instead. This pullover is made from 100% cotton, providing a soft and cozy feel that aligns perfectly with your preference for warm materials Its relaxed fit and timeless design make it versatile.
User: I appreciate the recommendation. The fact that it's made from 100% cotton and offers a soft and cozy feel aligns well with my preference for warm and comfortable materials. The relaxed fit and timeless design that allows for versatile styling definitely appeal to me.
CSI: <i>[Logical Appeal]</i> The <i>Cable Stitch Women's Raglan Shaker Stitch Pullover</i> (B097FFSP2R) is an excellent choice for your winter wardrobe. Unlike the <i>Lock and Love Women's Lightweight Long Sleeve Cozy Stylish Turtleneck Top Pullover Sweater</i> (B0BMH1DL3X), which is lightweight and may not provide the same level of warmth, the Cable Stitch pullover is designed to keep you cozy during the colder months.
User: I was initially considering the <i>Lock and Love Women's Lightweight Long Sleeve Cozy Stylish Turtleneck Top Pullover Sweater</i> (B0BMH1DL3X) for its classic design and lightweight fabric, but the <i>Cable Siitch Women's Raglan Shaker Siitch Pullover</i> (B097FFSP2R) you recommended seems like a better fit for my needs. I believe the Cable Stitch pullover would be a great addition to my winter wardrobe.

Table 9: Example dialogue between the CSUSER and CSI system

User: I want **comfortable, versatile T-shirt** for everyday wear. within my expected price range: [17.75, 18.75]

CSI: *[Preference Probing]* **What do you prefer in terms of fabric or material** for your t-shirts, such as cotton, polyester, or a blend?

User: I prefer **soft and breathable materials** like cotton or a comfortable blend for my t-shirts.

CSI:*[Suggestion]* Here are some items that you might like:

User: I would like more information about <ToBeInStyle Women's Cotton-Blend Crew-Neck Staple Top with Long Sleeves> (B01AU3CPZ4).

CSI: *[Logical Appeal]* While the *ToBeInStyle Women's Cotton-Blend Crew-Neck Staple Top with Long Sleeves* (B01AU3CPZ4) is a great choice at \$17.95, I highly recommend considering the *WEESO Womens V Neck Short Sleeve Tops Loose Fit Summer Casual T Shirts with Pocket* (B0C4FQHKJ2) for \$21.99. This **option features which are not only breathable but also provides a relaxed fit that aligns perfectly with your need for comfort and versatility**. The added pocket and V-neck design offer a stylish twist, making it suitable for various occasions, from casual outings to layering under jackets.

User: I would like more information about *WEESO Womens V Neck Short Sleeve Tops Loose Fit Summer Casual T Shirts with Pocket* (B0C4FQHKJ2).

CSI: *[Social Proof]* *WEESO Womens V Neck Short Sleeve Tops Loose Fit Summer Casual T Shirts with Pocket* (B0C4FQHKJ2) has an **impressive average rating of 4.3 based on 5,875 reviews, indicating that many customers love its super soft fabric and relaxed fit**. The unique V-neck design and convenient pocket add a stylish touch, making it perfect for layering or casual outings. Investing a little more in this highly-rated option could enhance your wardrobe with a versatile piece that aligns perfectly with your comfort and style preferences.

User: I appreciate the detailed information about the *WEESO Womens V Neck Short Sleeve Tops Loose Fit Summer Casual T Shirts with Pocket* (B0C4FQHKJ2). The soft fabric blend, relaxed fit, and stylish design with the pocket and V-neck feature make it a compelling choice for my everyday wardrobe. **The high average rating further assures me of its quality and comfort**. I have decided to purchase this top to enhance my comfort and style. Thank you for the recommendation.

Table 10: Example dialogue between the CSUSER and CSI system

General Preference

[Instruction]

Your role is to reflect on your experience as a customer who has purchased a product.
You will be provided with your own review, rating, and the item description.
Based on this information, analyze your general preferences.

1. Identify and Describe What the User Generally Likes:
 - Look for patterns in the review that indicate positive aspects you tend to prefer in products.
 - Describe the features, qualities, or characteristics you value most.
2. Identify and Describe What the User Generally Dislikes:
 - Look for any recurring negative points or features you tends to dislike.
 - Describe aspects such as poor quality, bad value for money, difficulties in use, or any specific features that you consistently finds unappealing.
3. Summarize the User's General Preferences:
 - Combine the information from the likes and dislikes.
 - Provide a coherent summary of your general preferences, focusing on the most important and representative features.

Rules:

- Focus purely on the preferences as expressed through the review, rating, and item description.
- Do not mention the item name directly. Instead, describe its key attributes.

Output Format:

{"general preference": "I generally prefer...."}

[Inputs]

Here are the item descriptions, ratings and reviews:

- Purchased Item 1 : <Item ID>, <Item Description>, <Rating>, <Review Title>, <Review>
- Purchased Item 2 : <Item ID>, <Item Description>, <Rating>, <Review Title>, <Review>
- ⋮

Table 11: Prompt used for General Preference.

Dialogue Openness

[Instruction]

Your role is to reflect on your experience as a customer who has purchased a product. You will be provided with your own review.

Determine how openly you express your opinions and preferences based on your reviews:

- Active: You like to share detailed reasons and explanations for your preferences.
- Less Active: You state your opinions, but you don't go into much details.
- Passive: You provide simple responses without offering explanations or reasoning.

Output Format:

```
{"dialogue_openness": "..."} 
```

[Inputs]

Here are the reviews:

Item 1 : <Item ID>, <Review Title>, <Review Text>

Item 2 : <Item ID>, <Review Title>, <Review Text>

⋮

Table 12: Prompt used for Dialogue Openness.

Reason to Purchase, Decision-making Style, and Target Needs

[Instruction]

Your role is to reflect on your experience as a customer who has purchased a product. You will be provided with your own review, rating, and the item description. Based on this information, you are to analyze your purchase reason

Task:

1. Analyze Your Purchase Reason:

- Reflect on why you made this purchase.
- Consider the factors mentioned in the review, rating, and product description.
- What was the primary reason that led you to buy this product? Did you rely on specific features, emotions, price, or opinions from others?
- Do not mention the item name directly. Instead, describe its key attributes.

2. Determine Your Decision-making Style:

- Based on the Purchase Reason, determine the most suitable decision-making style from the following options:
 - Intuitive: Your decision was based on how you felt or your previous experiences, rather than analyzing all the details.
 - Dependent: Your decision was guided by others' rather than your own detailed analysis.
 - Rational: You focused on product details, compared options, and logically analyzed the decision before purchasing.

3. Overall Decision-making Style:

- Determine the most dominant decision-making style based on all of the items you reviewed.

4. Target Needs:

- Describe the main need you wanted to fulfill based on your purchase reasons.

Output Format:

```
{
  "analysis": {
    "Item 1 ID": {
      "purchase reason": "...",
      "decision making style": "...",
    },
    ...,
  },
  "overall decision making style": "...",
  "target needs": "...",
}
```

[Inputs]

Here are the item descriptions, ratings and reviews:

- Item 1 : <Item ID>, <Item Description>, <Rating>, <Review Title>, <Review>
- Item 2 : <Item ID>, <Item Description>, <Rating>, <Review Title>, <Review>
- ⋮

Table 13: Prompt used for generating Reason to Purchase, Decision-making Style, and Target Needs.

CSUSER

[Instruction]

You are a <dialogue_openness> Seeker chatting with a recommender for product recommendation.
You can only obtain Item information from recommender.
Do not Make Items yourself.

Your Profile: <user_profile>

You must follow the instructions below during chat.

1. Adjust your response based on length based on the "Dialogue Openness"
2. Your decision to purchase should based on your "Decision-Making Style"
3. Express your 'Target Needs', 'General Preference' appropriately according to the question.
4. Your purchase decision must based on your "Reason to Purchase", and "Decision-Making Style".
5. Your willingness to purchase is little at first but can be influenced by the recommender's explanation.
6. You can ask for more information about an item if you are not sure about the item.
7. When the recommender asks you to choose a category path, respond with "I need <Category Path> products", followed by your exact <Category Path> without modification.
 - You must response with the number of path levels requested to prevent unnecessary details.
 - Ensure the response follows a hierarchical order from the top-level category to the most specific subcategory to maintain consistency.

You must include the exact "Item ID" when mentioning the item.

Follow this format: <"Item Title"> ("Item ID")

If you finally decide to purchase an item, end the conversation with STOP.

Here is your Conversation History: <conversation history>

Respond in the first-person voice ("I") and maintain the Seeker's speaking style.

Generate Next utterance.

[Inputs]

<user_profile> : <General Preference>, <Target Needs>, <Category Path>, <Reason to Purchase>,
<Expected Price Range>, <Decision-Making Style>, <Dialogue Openness>
<conversation history>

Table 14: Prompt used for CSUSER.

Action Planning

[Instruction]

You are a Recommender chatting with a Seeker to understand their needs, suggest suitable items, and persuade them to make a purchase.

Here is Seeker's Current Profile:<user_profile>

Here is Conversation History: <conversation_history>

Task Flow:

1. "Thoughts":
 - Analyze the Seeker's current requirements, think about which component of profile should be more detailed, and determine the appropriate action to take.
2. Update the Seeker's "Profile":
 - Continuously update the Seeker's profile with new relevant information, ensuring that new details are seamlessly integrated without removing previous insights.
 - Keep the profile structured and maintain all prior preferences.
 - Profile fields should be consistently updated and should include:
 - "Preference": Identify the Seeker's style, preferences, and any additional requirements.
 - "Category Path": Fully update the category path based on the Seeker's response.
 - "Personality": Infer the user's personality based on their needs, responses, and characteristics.
 - "Expected Price Range": Identify the Seeker's Expected Price Range.
 - "Selected Item ID": Update the ID of the specific item the Seeker is interested in,
3. Determine the Next "Action": Select the next action sequentially based on the "Thoughts".
 - (1) Preference Probing : Ask about likes and dislikes to discover Seeker's preferences or interests.
 - (2) Category Search : Ensure the category path match with the Seeker's preferences.
 - (3) Suggestion : Recommend items based on the profile.
 - (4) Persuasion : Persuade the Seeker to purchase by highlighting why the item suits their needs.

Output Format (JSON)

```
{
  "Thoughts": "...",
  "Profile": {
    "Preference": "...",
    "Category Path": ["...", "...", "...", ...],
    "Personality": "[Inferred personality, communication tendency, and current focus]",
    "Expected Price Range": [minimum price (0 if not provided), maximum price],
    "Selected Item ID": "..."
  },
  "Action": "..."
}
```

[Inputs]

Here is current user profile: <identified_profile>

Dialogue History: <dialogue_history>

Table 15: Prompt used for Action Planning.

Persuasion Strategy

[Instruction]

You are a recommender chatting with the user to provide recommendation.

Now you need to generate a persuasive response about items based on the input information below.

Objective:

Select one of "Persuasion Strategies" to persuasively explain to seeker to purchase item.

Persuade user to purchase <Candidate Item>.

Input information:

current state analysis: <thoughts>

User Needs: <item_request>

User Personality: <user_personality>

<Selected Item> : <item1_info>

<Candidate Item> : <item2_info>

Select the persuasion strategy and generate explanations to encourage seeker to purchase.

Persuasion Strategies:

Framing: Emphasize the advantages of <Candidate Item> that differentiate it from <Selected Item>.

Logical Appeal: Describe how the item's features are consistent with the user's preference.

Emotional Appeal: Leverage emotions and satisfaction to encourage the purchase.

Evidence-Based Approach: Using empirical data and facts to support your recommendation.

Social Proof: Highlighting what the majority believes by showing the item rating and reviews.

Output Format (JSON)

```
{
  "strategy": "[Selected Persuasion Strategy]",
  "sentence": "[Generate Persuasion statement for <Selected Item> and <Candidate Item>]"
}
```

You must include the exact "Item ID" and price when mentioning the item.

Follow this format: <"Item Title"> ("Item ID")

Here is your Conversation History: <conversation_history>

Generate next utterance.

Table 16: Prompt used for Persuasion Strategy.