

Melina Mehrmam

A Thesis in The John Molson School of
Business Department of Marketing

Presented in Partial Fulfillment of the
Requirements for the Degree of
Master of Science (Option Marketing)
at Concordia University
Montreal, Quebec, Canada

July 2025

© Melina Mehrmam, 2025

CONCORDIA UNIVERSITY

School of Graduate Studies

This is to certify that the thesis

Prepared By: Melina Mehrmam

Entitled: Consumers' perceptions of AI-based recommendations

and submitted in partial fulfillment of the requirements for the degree of

Master of Science (Marketing)

Complies with the regulations of the University and meets the accepted standards with respect to originality and quality.

Signed by the final Examining Committee:

_____ Chair / Examiner

Dr. Kamila Sobol

_____ Examiner

Dr. Darlene Walsh

_____ Thesis Supervisor

Dr. Caroline Roux

Approved by:

_____ Chair of Department

Dr. Tieshan Li

_____ Dean, John Molson School of Business

Dr. Anne-Marie Croteau

Abstract

Consumers' perceptions of AI-based recommendations

Melina Mehrmam

As digital platforms increasingly rely on recommendation systems, understanding how these systems influence consumer behavior is essential. While AI-based recommendations offer strong personalization, research suggests that consumers may respond negatively when artificial intelligence is explicitly mentioned, particularly in emotional or pleasure-driven consumption contexts. This thesis examines how the framing of recommendations (AI-based, behavior-based, or similarity-based), product type (utilitarian or hedonic), and individual characteristics shape consumer responses.

Three studies were conducted. A pilot study with 184 participants assessed consumers' familiarity with recommendation types and measured AI-related traits such as literacy, perceived capability, magical beliefs, fear, and usage. Study 1, with 490 participants, tested how recommendation framing and product type influenced trust, perceived relevance, and behavioral intentions. Study 2, with 361 participants, used a new shopping context and introduced novelty seeking as an additional moderator.

Results showed no consistent advantage for AI-based recommendations. However, consumers with higher perceived AI capability responded more positively to AI framing, while those with higher fear or lower AI usage were less receptive. Novelty-seeking individuals were more favorable toward similarity-based recommendations, particularly when evaluating hedonic products.

This research contributes to the literature on AI in marketing by showing that consumer responses to AI-based recommendations depend on both context and individual characteristics. Marketers should consider aligning recommendation framing with product category and customer traits to build trust and increase engagement.

Acknowledgments

This thesis was not a difficult burden but rather a fascinating journey. I truly enjoyed every step from start to finish. What made it especially meaningful was the opportunity to grow alongside the incredible people who supported and inspired me throughout the process.

I feel truly lucky to have had Professor Caroline Roux as my supervisor. Your guidance, encouragement, and thoughtful feedback helped me stay focused and motivated throughout the process. You created a space where I felt supported, challenged, and inspired — and for that, I'm very grateful.

I would also like to sincerely thank my committee members, Professor Kamila Sobol and Professor Darlene Walsh, for their time, insightful feedback, and valuable suggestions. Your perspectives enriched this thesis and pushed me to think more critically about my research.

To my parents and sister, thank you for always believing in me. Your constant support, both emotional and practical, gave me the strength and confidence to move forward. Your faith in my abilities means more than I can express.

To my friends, thank you for bringing balance and lightness to this process. Your support, encouragement, and timely distractions made this journey more enjoyable and reminded me not to lose sight of the bigger picture.

In reflecting on this experience, I am reminded of Rumi's words:

من ز مقصد نگذرم، اما نظر بر راه دارم / که به ره خوشتر بود دل را ز منزل‌های دور

“I do not dismiss the goal, but my heart is in the road — for it is sweeter than far-off destinations”.

This work marks an important milestone, but it is the learning, growth, and relationships formed along the way that I will carry with me long after the final submission.

Table of Contents

List of Figures	vii
1 Introduction.....	1
2 Literature Review.....	2
2.1 Consumers' Responses to AI	2
2.2 The Role of Recommendation Type	3
2.3 The Moderating Role of Product Type	4
2.4 The Moderating Role of AI-Related Individuals Differences.....	4
3 Overview of the Studies.....	6
4 Pilot Study.....	6
4.1 Methods.....	6
4.2 Results and Discussion.....	7
5 Study 1	8
5.1 Methods.....	8
5.2 Results and Discussion.....	9
6 Study 2	11
6.1 Methods.....	12
6.2 Results and Discussion.....	14
7 General Discussion	18
7.1 Theoretical Contributions.....	19
7.2 Managerial Implications.....	20
7.3 Limitations and Future Research Directions.....	20
7.4 Conclusion	21
References.....	22
Appendices.....	26
A. Pilot Study Questions.....	26
B. Study 1 Stimuli.....	31
C. Study 1 Questions	31
D. Study 1 ANOVA Results	34
E. Study 1 Correlation Table	38
F. Study 2 Stimuli.....	40

G.	Study 2 Questions	40
H.	Study 2 Correlation Table	43
I.	Study 2 PROCESS Model 1 Outputs	44
J.	Study 2 PROCESS Model 3 Outputs	51

List of Figures

Figure 1 Conceptual framework	5
Figure 2 Interaction effect of recommendation and product type on the utilitarian/hedonic manipulation checks – Study 1	10
Figure 3 Additional conceptual framework – Study 2.....	12
Figure 4 Interaction effect of recommendation type and AI capability on recommendation confidence – Study 2.....	Error! Bookmark not defined.
Figure 5 Interaction effect of recommendation type and AI fear on recommendation confidence – Study 2	15
Figure 6 Interaction effect of recommendation type and AI use on recommendation confidence – Study 2	16
Figure 7 Interaction effect of recommendation type, product type, and novelty seeking on profile creation intentions – Study 2.....	17
Figure 8 Interaction effect of recommendation and product type on the utilitarian/hedonic manipulation checks – Study 2	18

1 Introduction

Consumers frequently receive personalized recommendations while shopping or consuming content online. Platforms like Amazon recommending products to buy, Netflix suggesting the next movie to watch, or Spotify curating a music playlist. Such marketing strategy has proven highly effective as, for instance, Amazon has reported a 29% increase in sales thanks to its recommendation system (Rejoinder, 2021). Similarly, Spotify's algorithmic recommendations account for 33% of music discovery on the platform (Spotify, 2023). These examples highlight how recommendations shape consumer choices and personalize experiences across various digital platforms.

The technologies and approaches underlying recommendation systems are in constant evolution, and various types of recommendations are being used by digital platforms. Given the important role that recommendations play in consumption decisions, it is crucial to better understand which type(s) most effectively influence such decisions. I thus aim to compare the effects of more "common" versus more novel recommendation approaches in my thesis. More "common" approaches include recommending products based on the interests of similar customers (e.g., "customers like you also like") or on a consumer's online behavior (e.g., "based on your history"). More novel approaches tend to rely on more advanced algorithms or artificial intelligence (e.g., "powered by AI"), as they leverage advanced data analytics and deep learning to predict what a consumer might want based on a variety of data sources and recommendation approaches (including more "common" ones).

Although consumers have become quite familiar with similarity- and behavior-based recommendations, AI-based recommendations are fairly new. A survey found that, although 65% of consumers trust businesses that use AI technology, only 37% of consumers are receptive to AI-based product recommendations, and 55% of consumers are concerned with AI being used for music and movie/TV show recommendations (Haan & Watts, 2023). Therefore, even if AI-driven recommendation systems can provide retailers with a competitive advantage and create value for consumers (Habil et al., 2023), the latter may be less receptive to AI-based recommendations as compared to other recommendation approaches.

My thesis thus aims to investigate if the way recommendations are framed (i.e., similarity-, behavior-, or AI-based) influences how consumers perceive and respond to them. I will first review prior research on recommendation types, product types (i.e., utilitarian vs. hedonic), and AI-related individual differences (e.g., literacy, magical beliefs, fear) to provide my rationale for my hypotheses. I will then present the results of a pilot and two studies, in which participants were asked to evaluate shopping scenarios featuring different types of recommendations. The studies also explored the roles of product type and AI-related individual differences in the effects of recommendation types on consumers' responses. Lastly, I will discuss the implications of my research and possible future research directions that build on its limitations.

2 Literature Review

2.1 Consumers' Responses to AI

As the technologies underlying recommendation systems continue to evolve, it is important to compare the effectiveness of more traditional methods with newer AI-driven approaches. Prior research suggests that, while consumers may appreciate the efficiency and personalization enabled by AI, their behavioral responses to AI-powered recommendations and technologies are often mixed. A growing body of literature reveals a persistent skepticism among consumers toward AI-based technologies, driven by concerns over privacy and security, and a lack of understanding of how AI functions (Cheatham et al., 2019; Gursoy & Cai, 2024). For instance, although many consumers recognize the benefits of AI, such as enhanced personalization, the explicit mention of “Artificial Intelligence” in product or service descriptions can decrease purchase intentions. This is due to reduced emotional trust, particularly in contexts where perceived risk is high (Cicek, Gursoy, & Lu, 2025). Specifically, the study by Cicek et al. (2025) found that revealing the presence of AI in a product description can significantly lower consumers' willingness to buy, with emotional trust serving as a key mediating factor. This effect is further amplified in high-risk contexts, suggesting that despite the potential value AI offers to retailers, consumers remain hesitant, particularly when it is made salient during the decision-making process.

Further, with the rapid advancement of AI, recommendation systems have evolved to become increasingly efficient, often surpassing human-curated suggestions in scalability and data-driven precision. AI-driven recommenders, particularly those powered by deep learning and reinforcement learning, analyze vast datasets to generate personalized recommendations, as exemplified by the YouTube recommendation system developed by Covington et al. (2016). These systems can even outperform human recommendations in terms of accurately reflecting individual preferences (Yeomans et al., 2019). Despite these technical achievements, empirical evidence reveals a persistent phenomenon known as “algorithm aversion” (Dietvorst et al., 2015), where consumers show reluctance to follow AI-generated suggestions and instead prefer human-based recommendations (Longoni et al., 2019; Önköl et al., 2009). One reason for this is that human-based recommendations often provide rich narrative descriptions that help recipients relate the product to their own needs (Simonson & Rosen, 2014). This ability to engage consumers through storytelling and shared experience elicits stronger mentalizing and self-referencing processes, which are cognitive mechanisms that enhance perceived relevance and persuasiveness (Wien & Peluso, 2021).

Although prior research has explored comparisons between AI- and human-generated recommendations (Longoni et al., 2019; Yeomans et al., 2019), and investigated consumers' reactions to certain types of recommendations (Boerman et al., 2017; Zhang et al., 2020), to my knowledge, there is a lack of research comparing more “common” approaches (i.e., similarity- or behavior-based) to newer ones (i.e., AI-based). My thesis will therefore attempt to address this research gap. Understanding how various recommendation approaches affect consumer choices can help improve the design and use of recommendation systems. By identifying how consumers interpret different sources of recommendations, my thesis will provide insights into which frames are more effective at building trust and motivating action in digital environments.

2.2 The Role of Recommendation Type

Given the growing influence of recommendation systems on consumer behavior, it becomes crucial to understand how consumers perceive and respond to different types of recommendations. Among the more established techniques are user-based collaborative filtering (e.g., “customers like you also like”) and behavioral targeting. (e.g., “based on your history”). On the one hand, user-based collaborative filtering recommends products by identifying other users with similar preferences and suggesting items those similar users (but not the target consumer) have liked, consumed, or purchased (Zhang et al., 2020; Thakkar et al., 2018). For example, Netflix might recommend a show because people with similar viewing patterns enjoyed it. This method relies on user feedback, such as ratings, to calculate similarity scores and generate recommendations (Pu et al., 2012). Behavioral targeting, on the other hand, looks at a user’s own activity such as browsing history, searches, and clicks to predict what they might want and recommend related products or content. It uses tools like data mining and machine learning to build a profile of the user and personalize suggestions (Boerman et al., 2017).

While both methods aim to deliver personalized experiences, they differ in how they work and how consumers perceive them. User-based collaborative filtering builds on shared preferences among users, while behavioral targeting is based on each person’s unique digital behavior. Therefore, unlike collaborative filtering, which uses data from other users, behavioral targeting is fully based on individual behavior, making it very specific to each user, which can raise more concerns about privacy (Summers, Smith, & Reczek, 2016).

More novel methods firms can use to make recommendations tend to rely on more advanced algorithms or artificial intelligence (e.g., “powered by AI”). Such methods leverage advanced data analytics and deep learning to predict what a consumer might want based on a variety of data sources and recommendation methods (including more “common” approaches). AI-driven recommendation systems can thus enable retailers to better understand consumers’ needs and predict their future behaviors than commonly used methods (Covington et al., 2016). However, as previously discussed, the potential benefits of AI-based recommendations for consumers are undermined by various concerns (e.g., privacy, lack of trust; Cheatham et al., 2019; Cicek, Gursoy, & Lu, 2025; Gursoy & Cai, 2024; Haan & Watts, 2023) they have toward such technology.

To deepen our understanding of how consumers interpret and respond to different types of recommendations, I will investigate the role of framing, which refers to how information is presented, and which affects how it is perceived and understood, thus shaping individuals’ perceptions and decisions (Hardisty, 2024). In the context of recommendation systems, the same recommendation can produce different reactions depending on how it is framed, such as being presented as having been generated based on similar users, one’s browsing history, or AI. These distinct labels are likely to trigger various cognitive and emotional responses, including trust, perceived relevance, and behavioral intentions.

Prior research suggests that consumers generally seem reluctant about the use of AI for generating product recommendations (Dietvorst et al., 2015; Haan & Watts, 2023). I therefore hypothesize:

H1: Product recommendations labeled as AI-based will prompt more negative consumer responses (e.g., trust, perceived relevance, behavioral intentions) than those labeled as similarity- or behavior-based.

However, given consumers' mixed responses to AI (as previously discussed), I will explore conditions under which highlighting the use of AI (vs. more common approaches) for generating product recommendations may produce positive (vs. negative) responses. Next, I will discuss two potential moderators of the effects of recommendation type: product type and AI-related individual differences.

2.3 The Moderating Role of Product Type

Prior research has consistently shown that the type of product, whether utilitarian or hedonic, plays a crucial role in how consumers respond to product recommendations. Utilitarian products are defined as those that are practical, functional, and goal oriented, serving instrumental purposes (e.g., appliances, office supplies), whereas hedonic products are associated with pleasure, enjoyment, and experiential consumption (e.g., fashion, entertainment; Dhar & Wertenbroch, 2000). This fundamental distinction can influence how consumers evaluate product recommendations. For instance, Longoni and Cian (2020) introduced the “word-of-machine” effect, demonstrating that consumers tend to prefer AI-based recommendations for utilitarian products due to perceptions of analytical competence, but resist them in hedonic contexts where emotional resonance is more important (Longoni & Cian, 2020). Similarly, Wien and Peluso (2021) found that human recommenders were more persuasive for hedonic products, as they triggered stronger emotional and self-referential engagement. Additionally, Choi et al. (2011) observed that social presence in recommender systems had a greater positive effect on trust and reuse intentions for hedonic products than for utilitarian ones. Together, these findings suggest that the effectiveness of recommendation systems is contingent not only on the source of the recommendation but also on the nature of the product being recommended. Since AI-based recommendations are typically seen as more analytical and data-driven, and behavior- and similarity-based recommendations may retain a sense of humanity by reflecting actual human preferences and actions, which can enhance emotional resonance especially in hedonic contexts, it is expected that:

H2: Product recommendations labeled as AI-based (vs. similarity- or behavior-based) will prompt more positive consumer responses (e.g., trust, perceived relevance, behavioral intentions) for utilitarian products, and more negative responses for hedonic products.

2.4 The Moderating Role of AI-Related Individual Differences

Beyond contextual variables such as product type or recommendation framing, individual differences in how consumers understand, interact with, and emotionally respond to artificial intelligence could play a role in shaping their reactions to AI-based product recommendations. One important factor is AI literacy, defined as a consumer's objective knowledge about how AI systems operate. Interestingly, recent research by Tully, Longoni, and Appel (2024) find that lower AI literacy is associated with greater AI receptivity. This challenges earlier assumptions that greater knowledge automatically enhances trust (Campolo & Crawford, 2020). Therefore:

H3a: Consumers with lower AI literacy will respond more positively to AI-based recommendations.

Tully et al. (2024) also found that this was largely because consumers with less understanding of AI are more likely to perceive it as being magical or superhuman, which elicits feelings of awe that increase trust and engagement (Khare, Kautish, & Khare, 2023).

H3b: Consumers who perceive AI as magical will respond more positively to AI-based recommendations.

Conversely, fear of AI, including concerns over autonomy loss, job displacement, or algorithmic bias, has been shown to suppress AI adoption (Cave & Dihal, 2019; Schepman & Rodway, 2020).

H3c: Consumers with greater fear of AI will respond less favorably to AI-based recommendations.

Another key factor is perceived AI capability, or the belief in an AI system's competence and reliability, which has been shown to predict consumer trust and acceptance (Tully et al., 2024).

H3d: Consumers with higher perceived AI capability will respond more positively to AI-based recommendations.

Finally, AI usage, defined as the frequency and breadth of consumers' interactions with AI tools in daily life, significantly influences responses (Tang et al., 2022). Liu et al. (2024) demonstrate that AI usage can simultaneously empower users by enhancing their technological self-efficacy and induce anxiety due to perceived threats such as job insecurity or resource depletion.

H3e: Greater AI usage will lead to more positive responses to AI-based recommendations.

Together, these individual differences offer a nuanced framework for understanding variation in consumer responses to AI recommendations.

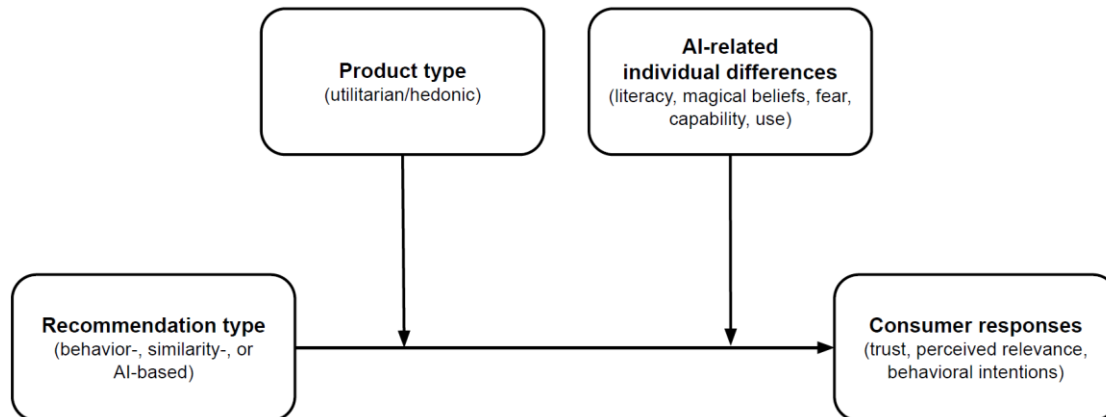


Figure 1. Conceptual framework

3 Overview of the Studies

My thesis consists of a pilot study followed by two experimental studies. The pilot study explores consumers' perceptions of different recommendation types and assesses relevant individual differences. Study 1 examines how recommendation type influences consumer perceptions (e.g., trust and perceived relevance) and behavioral intentions (e.g., learn more, add to cart). It also investigates how product type (utilitarian vs. hedonic) and AI-related individual differences moderate this relationship. Study 2 introduces a new variable – novelty seeking – to further explore the boundaries of the effects. The pilot and two experimental studies also allow for a comparison of consumers' perceptions (pilot study) with their behavioral intentions (experimental studies), providing a more comprehensive understanding of how recommendation framing, product type, and individual characteristics shape consumer responses.

4 Pilot Study

I first conducted a pilot study to understand online participants' experiences with different types of product recommendations and their levels of trust in such recommendations. The pilot study also aimed to identify which products they most often purchased on Amazon, and which streaming services they most frequently used for listening to music or watching movies, as I intended to use such platforms as part of my experimental stimuli. The pilot study also measured participants' literacy, perceived capability, magical beliefs, trust, fear, and usage of AI.

4.1 Methods

Two hundred and five participants were recruited from Amazon Mechanical Turk via CloudResearch and were compensated \$1.00 USD for a 7-minute study. Participants who did not pass the attention checks, did not agree to the consent form, experienced technical issues and/or distractions, or did not take the survey seriously (based on data quality questions) were excluded from the analyses. The final sample consisted of 184 participants ($M_{\text{age}} = 43.61$, $SD = 11.61$; 41.3% female).

Participants first provided informed consent and answered attention checks (e.g., “A chicken is a type of insect” [True/False]). They were then asked to select the products they most often buy on Amazon and the streaming services they most frequently use for listening to music or watching movies. This step helped identify popular product categories and streaming services to inform the design of future studies.

Next, participants answered questions regarding how frequently they encountered different types of recommendations, such as “recommendations based on customers like you,” “recommendations based on your browsing history,” and “recommendations powered by AI” (scale: 1 = Never to 7 = Very often). They also rated their trust in AI-based recommendations for various products and services, including music playlists, movies/TV shows, restaurants, books, news articles, and online courses (scale: 1 = Strongly disagree to 7 = Strongly agree; Malhotra et al., 2004; Bol et al., 2018).

Subsequently, participants answered questions measuring individual differences related to AI. They rated AI's capability to recommend products, movies/TV shows, playlists, or books (scale: 1 = Definitely incapable to 7 = Definitely capable; Tully et al., 2023). AI literacy was assessed through multiple-choice questions, selected from the AI literacy scale developed by Tully et al.

(2024; e.g., “How do supervised machine learning algorithms learn?”). Participants received a score of 1 for correct answers and 0 for incorrect ones. The total score, ranging from 0 to 4, was calculated by summing the correct responses. Trust in AI was measured using items adapted from existing scales (Malhotra et al., 2004; Bol et al., 2018; e.g., “To what extent do you believe that AI-based recommendations are trustworthy?”; scale: 1 = Strongly disagree to 7 = Strongly agree). Participants also rated their beliefs about the “magical” nature of AI using the scale developed by Khare et al. (2023; e.g., “If I think about artificial intelligence that powers products and services, I experience feelings of awe”; scale: 1 = Strongly disagree to 7 = Strongly agree). Their fear of AI was measured with the scale from Cave and Dihal (2019; e.g., “I fear that artificial intelligence will make us lose our humanity”; scale: 1 = Not at all to 7 = Very much). Additionally, participants reported how frequently they used AI in their daily lives (scale: 0 = Never to 7 = Extremely often). Finally, participants completed standard demographic questions and data quality checks (e.g., distractions, technological issues). The full list of questions is available in appendix A.

4.2 Results and Discussion

First, indices were created for the scales related to AI capability ($\alpha = .93$), literacy ($\alpha = .38$), trust ($\alpha = .94$), magic ($\alpha = .95$), and fear ($\alpha = .90$) by averaging the items from their respective scales.

The most popular types of products purchased by participants from Amazon were electronics (16.98%) and health and personal care (16.47%), while the most frequently used music and movie streaming services were Spotify (30.63%) and Netflix (21.41%), respectively. Results also indicated that behavior-based recommendations (i.e., based on purchase/browsing history) were the most frequently encountered type of recommendation ($M = 5.61$, $SD = 1.35$), followed by similarity-based recommendations (i.e., “based on consumers like you”; $M = 5.03$, $SD = 1.54$).

Participants reported varying levels of trust in AI-based recommendations depending on the type of product or service. For product types, AI-based recommendations were more trusted for small appliances ($M = 4.34$; $SD = 1.69$) followed by electronics ($M = 4.24$; $SD = 1.79$). For streaming services, AI-based recommendations were more trusted for movies ($M = 4.54$; $SD = 1.68$) and music ($M = 4.43$; $SD = 1.70$).

Regarding individual differences related to AI, participants’ perceived AI capability was moderately high ($M = 4.85$, $SD = 1.46$). Participants’ level of trust in AI was around the midpoint of the scale ($M = 4.10$, $SD = 1.44$), while the average perception of AI as magical was lower ($M = 2.74$, $SD = 1.72$). Finally, participants’ AI literacy scores had a mean of 2.53 ($SD = 1.00$).

In terms of correlations, AI literacy was negatively correlated with perceiving AI as magical ($r = -0.26$), fear of AI ($r = -0.19$), and AI usage ($r = -0.20$). Trust in AI was positively correlated with perceiving AI as magical ($r = 0.57$) and AI usage ($r = 0.53$). Fear of AI was positively correlated with perceiving AI as magical ($r = 0.24$). AI usage was also positively correlated with perceiving it as magical ($r = 0.57$). Additionally, age was negatively correlated with perceiving AI as magical ($r = -.13$), such that younger participants perceived AI as more magical than older ones.

In sum, the pilot study provided valuable insights that guided the development of Study 1. The findings helped identify electronics, as well as health and personal care, as the most commonly purchased products on Amazon, informing the selection of stimuli for the main experiment.

Additionally, the results showed that participants' trust in AI-based recommendations seemed relatively higher for entertainment-related services (e.g., movie and music streaming) compared to utilitarian product categories (e.g., small appliances, electronics). This observation informed the decision to include product type (hedonic vs. utilitarian) as a moderator in the main study. Furthermore, the pilot study confirmed that participants were generally familiar with behavior-based recommendations, supporting the use of such recommendation types in the experimental stimuli.

5 Study 1

The goal of Study 1 was to test the effects of different recommendation types on consumer responses, with product type and individual differences related to the use of AI as moderators. Specifically, the study investigated consumer responses to AI-based and behavior-based product recommendations across two scenarios involving different product types: a utilitarian product (power bank) and a hedonic product (hoodie). The study aimed to evaluate how the type of recommendation – whether AI-based or based on one's browsing history – impacted consumers' trust, perceived relevance, and behavioral intentions, such as the likelihood of wanting to learn more about the recommended product or adding it to their cart.

5.1 Methods

Five hundred and twenty-seven participants were recruited via CloudResearch Connect and were compensated \$1.00 USD for a 7-minute study. The same data exclusions used in the pilot study were also applied here. Participants who failed the attention checks, did not agree to the consent form, experienced technical issues and/or distractions, or did not take the survey seriously (based on data quality questions) were excluded from the analyses. The final sample consisted of 490 participants ($M_{\text{age}} = 39.27$, $SD = 12.26$; 48.6% female).

Participants first provided informed consent and answered attention checks (e.g., “A shark is a type of bird” [True/False]). Next, participants were randomly assigned to one of four conditions in a 2 (product type: utilitarian vs. hedonic) $\times$ 2 (recommendation type: behavior-based vs. AI-based) design. Depending on the condition, participants were presented with the following scenarios:

Utilitarian product: *“Imagine you're planning a busy day out, and your phone's battery tends to drain quickly. You think it would be helpful to have a new power bank to stay connected throughout the day. As you browse online, you see a power bank that is being recommended for you.”*

Hedonic product: *“Imagine the weather is getting colder, and you're thinking about refreshing your wardrobe. You think it would be helpful to have a new hoodie to keep you warm. As you browse online, you see a hoodie that is being recommended for you.”*

Again, depending on the condition, participants were told that the product was recommended either based on their browsing history or by an AI-powered algorithm. They were also presented with an image of the product alongside the scenario's text. See appendix B for the complete stimuli.

Following the manipulation, participants rated how much they trusted the product recommendation (scale: 1 = Not at all to 7 = Extremely) and its relevance (scale: 1 = Not at all to 7 = Extremely). A correlation analysis revealed that these two items were correlated ($r = 0.55, p < 0.001$), so they were averaged to form an index of recommendation confidence. Next, participants rated their likelihood of clicking on the recommended product to learn more about it, and of adding the recommended product to their cart (scale: 1 = Extremely unlikely to 7 = Extremely likely). A correlation analysis indicated these two items were correlated ($r = 0.77, p < 0.001$), so they were averaged to form an index of behavioral intentions.

Participants were then presented with manipulation checks asking them whether they believed purchasing electronics (clothes) was an entirely objective or subjective decision and a rational or emotional purchase (bipolar scales: 1-7). These two items were correlated ($r = 0.54, p < 0.001$) and were averaged to create an index of the product type manipulation checks. Following this, participants were asked about their preferences when buying the focal product type (i.e., electronics or clothes). Specifically, they rated how much they agreed with the following statements: “I know exactly what I want to buy,” “I have very specific preferences,” and “I only buy a (certain) specific brand(s)” (scale: 1 = Strongly disagree to 7 = Strongly agree). A reliability analysis showed that these three items loaded onto one factor, yielding a Cronbach’s alpha of 0.75. These items were averaged to form an index of preference strength. AI capability (for electronics and for clothes), literacy ($\alpha = .10$), perceptions of being magical ($\alpha = 0.92$), and fear ($\alpha = 0.89$) were assessed using the same items and scales as in the pilot study. Participants then completed additional attention checks regarding the product shown in their scenario and the type of recommendation they saw, to verify the effectiveness of the manipulations. Finally, participants answered standard demographic questions, including a question related to participants’ frequency with which they used AI in their daily lives (scale: 0 = Never to 7 = Extremely often), and data quality checks (e.g., distractions, technological issues). The complete list of questions is provided in appendix C.

5.2 Results and Discussion

I conducted a series of two-way ANOVAs to examine the effects of product recommendation type (i.e., AI-based vs. behavior-based) on consumer responses, with the dependent variables being trust and perceived relevance and behavioral intentions such as likelihood of learning more and adding the product to cart. The results showed no significant interaction between recommendation and product types on either dependent variable (all $ps > .05$). Additionally, the main effects of recommendation type were non-significant across all outcomes (all $ps > .05$). However, product type significantly influenced trust and perceived relevance ratings, with the utilitarian product ($M = 5.53, SD = 1.01$) receiving higher scores than the hedonic one ($M = 5.07, SD = 1.31; F(1486) = 18.68, p < .001$). This suggests that participants evaluated the utilitarian product recommendation as more relevant and trustworthy than the hedonic one. See appendix D for detailed results. These results were consistent when controlling for individual differences such as preference strength, perceived AI capability, and age.

The correlation matrix of AI Literacy, AI magical beliefs, AI fear, and AI use shows that these variables are not highly correlated with one another (all $rs < .44$; see appendix E). I thus analyzed the moderation effects of these individual differences independently using PROCESS Model 2 (see Figure 1). In separate analyses, trust and perceived relevance of the recommendation, and

behavioral intentions, such as the likelihood of exploring more details or adding the product to the cart were the dependent variables and recommendation type (i.e., AI-based vs. behavior-based) was the independent variable. For the moderators, product type (i.e., utilitarian vs. hedonic) and individual differences (i.e., AI literacy, AI magical beliefs, AI fear, and AI use) were included as separate moderators.

For the model with trust and perceived relevance of the recommendation as the dependent variable, recommendation type as the independent variable, product type as the first moderator, and AI fear the second moderator, the analyses showed no significant interactions. Specifically, neither the interaction between recommendation type and product type ($\beta = -0.17$, $SE = 0.21$, $t = -0.83$, $p > 0.40$) nor the interaction between recommendation type and AI fear ($\beta = 0.05$, $SE = 0.06$, $t = 0.80$, $p > 0.42$) significantly influenced trust and perceived relevance of the recommendation. All other combinations of moderators and dependent variables were analyzed, and none of the interactions or main effects were significant.

Finally, to further explore the effect of product type, I analyzed the related manipulation checks. I first conducted a two-way ANOVA to determine whether the product manipulation was successful. The results showed a significant interaction between recommendation and product types ($F(1, 486) = 3.90$, $p = .04$), no main effect of recommendation type ($F(1, 486) = .59$, $p = .44$), and a significant main effect of product type ($F(1, 486) = 32.37$, $p = .00$). The hoodie was seen as more hedonic ($M = 3.26$, $SD = 1.39$) than the power bank ($M = 2.56$, $SD = 1.34$), which suggests that the product type manipulation worked as intended. Although the interaction was marginal, I also looked at the pairwise comparisons. Recommendation type did not impact the perceived utilitarian/hedonic nature of the hoodie ($F(1, 486) = 0.71$, $p > .39$), but it did for the power bank ($M_{\text{Behavior-based}} = 2.72$, $SD = 1.44$, $M_{\text{AI-based}} = 2.39$, $SD = 1.21$, $F(1, 486) = 3.82$, $p < .05$). These exploratory results suggest that the AI-based product recommendation prompted participants to perceive the power bank as being even more utilitarian than for the behavior-based recommendation.

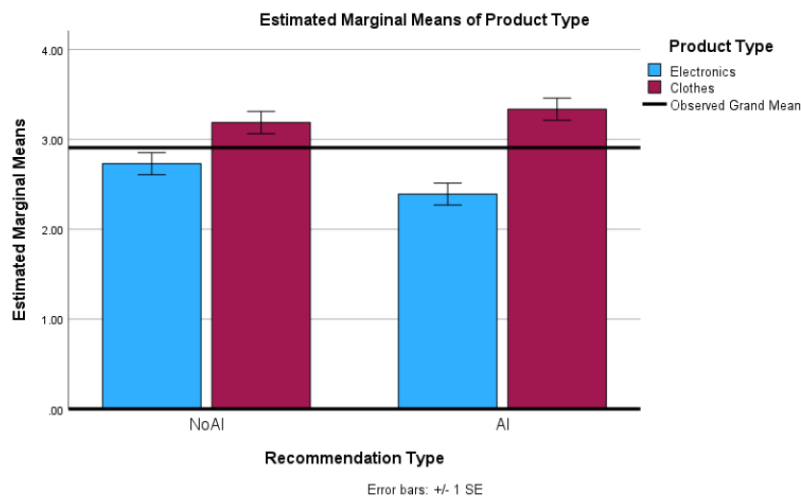


Figure 2. Interaction effect of recommendation and product type on the utilitarian/hedonic manipulation checks – Study 1

At the end of the experiment, I conducted attention checks to assess participants' recall of the product and recommendation types they encountered. Only 9 participants incorrectly identified the product type, but 186 participants incorrectly identified the recommendation type. Despite these

issues, I did not exclude these participants from the analyses, as doing so would have greatly reduced the sample size. However, these results underscore the need to improve the stimuli for Study 2. In the next study, I will attempt to make the recommendation type manipulation more explicit to enhance clarity for participants. Additionally, a potential limitation of Study 1 was the use of graphics from Amazon, a platform participants may have been overly familiar with. This familiarity might have reduced the effectiveness of the manipulation, as participants may not have paid close attention to the recommendation type. To address this, I will use a different platform in Study 2 to try to minimize this potential familiarity bias and enhance the manipulation's impact.

6 Study 2

The lack of effect of recommendation type in Study 1 may have been due to the fact that the online shopping context used involved relatively common shopping decisions on platforms most American consumers are familiar with. Recommendation type may be more relevant in contexts that provide opportunities for discovery and exploration, as digital technologies tend to promote novelty seeking (Shanmugasundaram & Tamilarasu, 2023).

Therefore, to further explore the role of recommendation type, I will use a less familiar shopping context in study 2 (i.e., meal planner website). I will also assess participants' novelty seeking tendencies as a potential moderator. Novelty seeking refers to consumers' tendency to seek new, varied, and stimulating experiences, often showing a preference for trying unfamiliar products or brands (Hirschman, 1980; Katz & Lazarsfeld, 1955). It also seems related to product type, as consumers with strong hedonic values are more receptive to novel product suggestions (Wang et al., 2000).

Based on this, the following hypotheses are proposed:

H4a: Consumers higher in novelty seeking will respond more negatively to an AI-based (vs. similarity-based) recommendation for a hedonic product, while there will be no difference for those lower in novelty seeking.

H4b: Consumers higher in novelty seeking will respond more positively to an AI-based (vs. similarity-based) recommendation for a utilitarian product, while there will be no difference for those lower in novelty seeking.

Therefore, in addition to the main conceptual model (see Figure 1), I will test an additional conceptual model in Study 2 (see Figure 2).

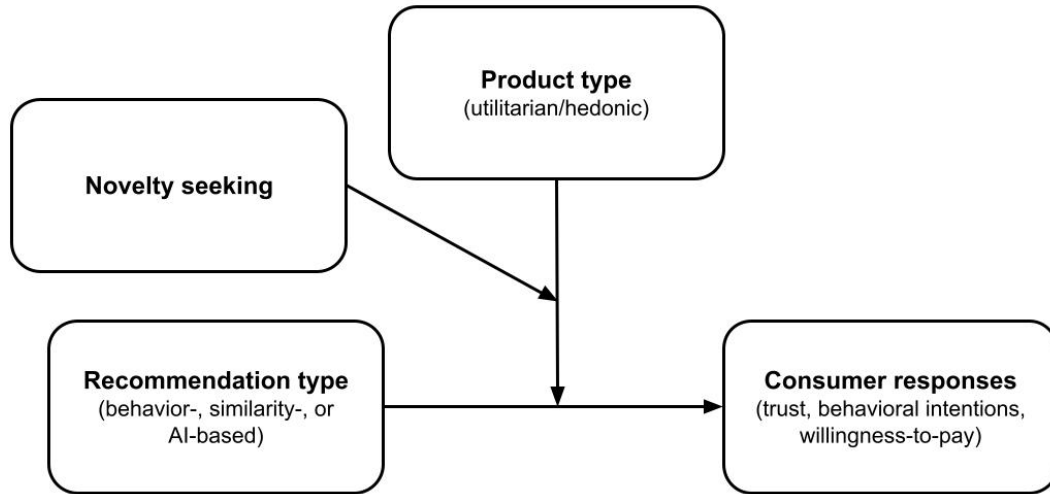


Figure 3. Additional conceptual framework – Study 2

The goal of Study 2 was to further investigate consumers’ responses to different types of product recommendations. In this study, I will compare the effects of AI-based recommendations to similarity-based ones, as the novel shopping context made the use of behavior-based recommendations (used in Study 1) less credible. Additionally, instead of changing the product to manipulate product type, I will use one shopping context (i.e., meal planning website) and manipulate how the service is described: (1) *“quick and easy recipe recommendations with minimal steps and ingredients designed to reduce mental load and simplify meal planning and cooking,”* which focuses on the utilitarian aspects of the service, as it highlights its more functional and practical characteristics; versus (2) *“fun and yummy recipe recommendations, made with delicious ingredients designed to add pleasure and excitement to meal planning and cooking,”* which focuses on the hedonic aspects of the service, by highlighting its more enjoyable and experiential characteristics. The study aimed to evaluate how the type of recommendation – whether AI-based or based on “users like you” – impacted consumers’ trust, perceived relevance, and behavioral intentions, including their willingness to create a profile on the website, activate a free trial, and pay for a monthly subscription.

6.1 Methods

Five hundred twenty-six participants were recruited via CloudResearch Connect and were compensated \$1.00 USD for a 7-minute study. The data exclusions used in the pilot study and Study 1 were also applied here. Participants who failed the attention checks, did not agree to the consent form, experienced technical issues and/or distractions, or did not take the survey seriously (based on data quality questions) were excluded from the analyses. Additionally, in this study, participants were asked to indicate their willingness to pay by entering a value between \$0 and \$20 in a text box. Those who did not comply with these instructions (e.g., by indicating a value greater than \$20) were excluded from the analyses. Recommendation and product types attention checks were also used in the study and, this time around, participants who failed these checks were excluded from the analyses ($N = 140$ because, unlike in Study 1, these manipulations were more obvious and the stimuli were presented multiple times throughout the study. The final sample consisted of 361 participants ($M_{\text{age}} = 38.82$, $SD = 11.65$; 50.7% female).

Participants first provided informed consent and answered attention checks (e.g., “A crow is a type of fish” [True/False]). Next, participants were randomly assigned to one of four conditions in a 2 (product type: utilitarian vs. hedonic) $\times$ 2 (recommendation type: AI-based vs. similarity-based) design. All participants were asked to imagine that they had been struggling to come up with meal ideas. After searching the internet for inspiration, they found a subscription service offering personalized recipe recommendations. This service provided either quick and easy recipes (utilitarian condition) or fun and yummy recipes (hedonic condition). The recipes were recommended by an AI assistant that processed data to generate personalized recommendations (AI-based condition) or by an algorithm that matched their profile with those of other users like them to generate personalized recommendations (similarity-based condition). To receive recommendations, participants needed to create a profile indicating their food preferences, allergies, dietary restrictions, budget, and cooking habits. The service allowed them to preview their recommended recipes after creating their profile, and they could gain full access to their personalized recommendations by activating a 14-day free trial. After the free trial expired, participants needed to purchase a monthly subscription, which was cancelable at any time with one click, to continue receiving recommendations. See appendix F for the complete stimuli.

Following the manipulation, participants rated how much they trusted and liked the product recommendations (scale: 1 = Not at all, 7 = Extremely). A correlation analysis revealed that these two items were strongly correlated ($r = 0.71$, $p < 0.001$), so they were averaged to form an index of recommendation confidence. Next, participants rated their likelihood of creating a profile, activating a free trial, and paying for a monthly subscription (scale: 1 = Extremely unlikely, 7 = Extremely likely). They were also asked about their willingness to pay for the monthly subscription (WTP), where they had to enter an amount between \$0 and \$20 USD.

Participants were then presented with manipulation checks asking them whether the recipe recommendation service was mostly functional versus enjoyable, practical versus pleasurable, and utilitarian versus hedonic (bipolar scales: 1-7). A reliability analysis yielded a Cronbach’s alpha of 0.90, so these items were averaged to create an index of the product type manipulation checks. Next, AI capability was assessed using one item (“To what extent do you believe AI is capable of recommending recipes?” scale: 1 = Definitely incapable, 7 = Definitely capable). AI literacy ($\alpha = .07$), perceptions of being magical ($\alpha = 0.93$), and fear ($\alpha = 0.87$) were assessed using the same items and scales as in the pilot study and Study 1.

Additionally, participants’ novelty seeking was assessed using a scale from Steenkamp and Baumgartner (1995; e.g., “I like to continue doing the same old things rather than trying new and different things,” scale: -3 = Completely false, 3 = Completely true). A reliability analysis yielded a Cronbach’s alpha of 0.88, so these items were averaged into a novelty seeking index. Participants then completed additional attention checks regarding the type of recommendation they saw, to verify the effectiveness of the manipulations. Finally, participants answered standard demographic questions, including a question related to participants’ frequency with which they used AI in their daily lives (scale: 0 = Never to 7 = Extremely often), and data quality checks (e.g., distractions, technological issues). The complete list of questions is provided in appendix G.

6.2 Results and Discussion

I first conducted a series of two-way ANOVAs to examine the effects of product recommendation type (i.e., AI-based vs. similarity-based) on consumer responses, with the dependent variables being recommendation confidence, behavioral intentions (including willingness to create a profile, activate a free trial, and pay for a monthly subscription), and willingness to pay for the recipe recommendation service. The results showed no significant interaction between recommendation and product types on all dependent variables (all $ps > .05$). Additionally, the main effects of recommendation type and product type were non-significant across all outcomes (all $ps > .05$). These results were consistent when controlling for individual differences such as age.

The correlation matrix of AI capability, literacy, magical beliefs, fear, and use shows that these variables are not highly correlated with one another (all $rs < .39$; see appendix H). I thus analyzed the moderation effects of these individual differences independently using PROCESS Model 2 (see Figure 1). In separate analyses, recommendation confidence, behavioral intentions (including willingness to create a profile, activate a free trial, and pay for a monthly subscription), and willingness to pay for the recipe recommendation service were the dependent variables, and recommendation type (i.e., AI-based vs. similarity-based) was the independent variable. For the moderators, product type (i.e., utilitarian vs. hedonic) and individual differences (i.e., AI capability, literacy, magical beliefs, fear, and use) were included as separate moderators. Across all the analyses, there was no significant interaction between recommendation and product type (all $ps > .05$). However, there were significant interactions between recommendation type and AI capability, fear, and use for some of the dependent variables, while the interactions with AI literacy and magical beliefs were non-significant across all the analyses.

I re-ran the analyses that produced significant interactions between recommendation type and AI-related individual differences using PROCESS Model 1, thus collapsing the analyses across the product type condition, since it did not produce any significant main or interaction effects in the previous analyses. For the model with recommendation confidence as the dependent variable, recommendation type as the independent variable, and AI capability as the moderator, the analyses showed a significant interaction effect ($\beta = .43$, $SE = .08$, $t = 5.54$, $p < .00$). There were also significant main effects of recommendation type ($\beta = -2.62$, $SE = .44$, $t = -5.98$, $p < 0.00$) and AI capability ($\beta = .31$, $SE = .06$, $t = 5.17$, $p < 0.00$). Specifically, participants with lower AI capability ratings (-1 SD) were less confident in the AI-based (vs. similarity-based) recommendation ($\beta = -.53$, $SE = .12$, $t = -4.58$, $p < 0.00$). Conversely, participants with higher AI capability ratings (+1 SD) were more confident in AI-based (vs. similarity-based) recommendation ($\beta = .35$, $SE = .15$, $t = 2.28$, $p = 0.02$). Similar significant moderation effects by AI capability were observed for behavioral intentions related to creating a profile and using the free trial. See appendix I for detailed results.

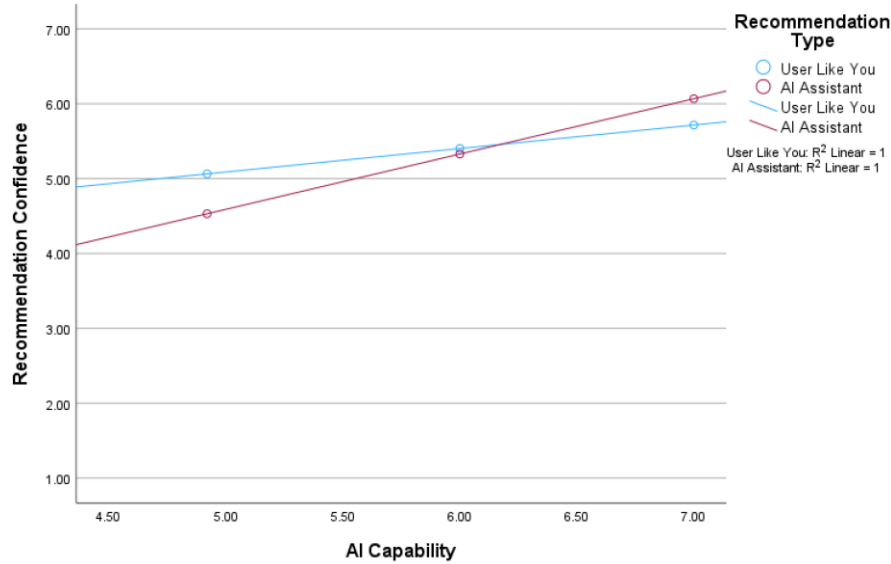


Figure 4. Interaction effect of recommendation type and AI capability on recommendation confidence – Study 2

For the model with recommendation confidence as the dependent variable, recommendation type as the independent variable, and AI fear as the moderator, the analyses showed a significant interaction effect ($\beta = -0.21$, $SE = 0.08$, $t = -2.49$, $p < 0.01$). There also was a marginal main effect of recommendation type ($\beta = .52$, $SE = .31$, $t = 1.64$, $p < .09$), and the main effect of AI fear was not significant ($\beta = -0.04$, $SE = 0.06$, $t = -0.68$, $p > .49$). Specifically, for participants with lower AI fear ratings (-1 SD), there was no significant difference in terms of recommendation confidence ($\beta = .20$, $SE = .20$, $t = 0.96$, $p > .33$). However, participants with higher AI fear ratings (+1 SD) were less confident in AI-based (vs. similarity-based) recommendation ($\beta = -0.55$, $SE = .19$, $t = -2.78$, $p < .00$). Similar significant moderation effects by AI fear were observed for behavioral intentions related to creating a profile. See appendix I for detailed results.

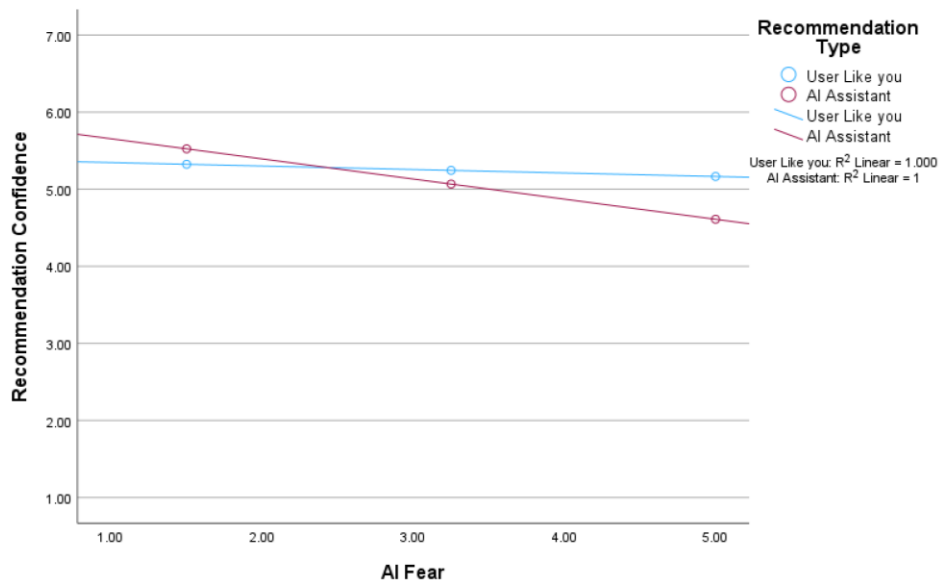


Figure 5. Interaction effect of recommendation type and AI fear on recommendation confidence – Study 2

For the model with recommendation confidence as the dependent variable, recommendation type as the independent variable, and AI use as the moderator, the analyses showed a significant interaction effect ($\beta = .15$, $SE = .06$, $t = 2.37$, $p < 0.01$). There also were significant main effects of recommendation type ($\beta = -.94$, $SE = .32$, $t = -2.88$, $p < .00$) and AI use ($\beta = .10$, $SE = .05$, $t = 2.10$, $p > .03$). Specifically, participants with lower AI use ratings (-1 SD) were less confident in AI-based (vs. similarity-based) recommendation ($\beta = -.63$, $SE = .21$, $t = -2.95$, $p < .003$). Conversely, for participants with higher AI fear ratings (+1 SD), there was no significant difference in terms of recommendation confidence ($\beta = .15$, $SE = .21$, $t = 0.72$, $p > .46$). See appendix I for detailed results.

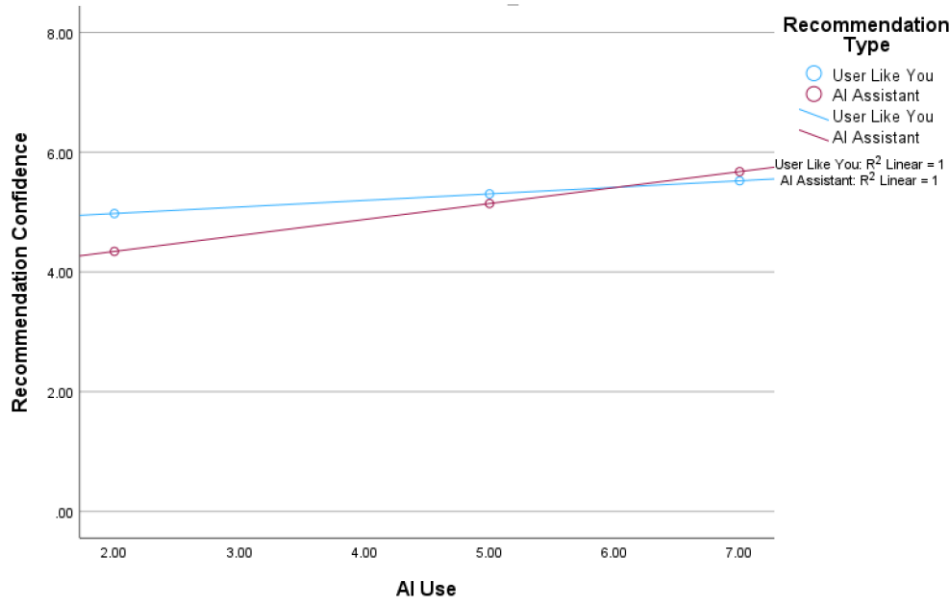


Figure 6. Interaction effect of recommendation type and AI use on recommendation confidence – Study 2

In sum, AI literacy and magical beliefs did not impact participants' responses to AI-based (vs. similarity-based) recommendations. Participants' confidence in the recommendations was impacted by AI capability, fear, and use. Their intentions of creating a profile on the website were impacted by AI capability and fear, and their intentions to activate the free trial were impacted by AI capability. None of the individual differences impacted participants' intentions to pay for a monthly subscription, or their willingness to pay for the recipe recommendation service. AI capability and, to a lesser extent, fear and use, thus seem to play an important role toward the top of the marketing funnel (De Haan, Wiesel, & Pauwels, 2016), but AI-related individual differences seem to have less impact toward the bottom of the funnel.

Next, I analyzed the moderation effects of novelty seeking using PROCESS Model 3 (see Figure 3). In separate analyses, recommendation confidence, behavioral intentions (including willingness to create a profile, activate a free trial, and pay for a monthly subscription), and willingness to pay for the recipe recommendation service were the dependent variables and recommendation type (i.e., AI-based vs. similarity-based) was the independent variable. Product type (i.e., utilitarian vs. hedonic) and novelty seeking were included as moderators. The analyses produced no significant results for recommendation confidence, subscription intentions, and willingness-to-pay. However, novelty seeking played a role in participants' intentions to create a profile and activate the free

trial, which are behaviors related to discovery. For the model with intentions to create a profile as the dependent variable, the analyses showed a significant 3-way interaction effect ($\beta = -.63$, $SE = .31$, $t = -2.08$, $p = .04$). However, there were no significant 2-way interaction or main effects (all $ps > .05$). Specifically, participants with higher novelty seeking tendencies (+1 SD) had marginally lower intentions to create a profile when the recipes were provided by AI (vs. a similarity-based algorithm) when the hedonic characteristics of the service were highlighted ($\beta = -.69$, $SE = .37$, $t = -1.86$, $p = 0.06$). Similar effects were observed for participants' intentions of activating the free trial, where those with higher novelty seeking tendencies (+1 SD) had significant (marginal) higher (lower) intentions to activate the free trial when the recipes were provided by AI (vs. a similarity-based algorithm) when the utilitarian (hedonic) characteristics of the service were highlighted. See appendix J for detailed results.

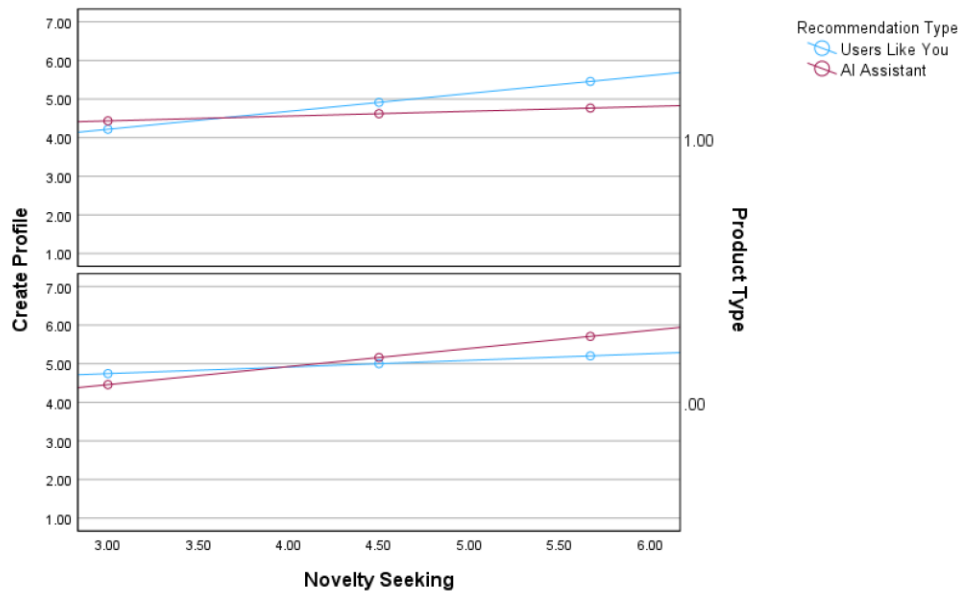


Figure 7. Interaction effect of recommendation type, product type, and novelty seeking on profile creation intentions – Study 2

Finally, to further explore the effect of product type as in Study 1, I analyzed the related manipulation checks. I first conducted a two-way ANOVA to determine whether the product manipulation was successful. The results showed a significant interaction between recommendation and product types ($F(1, 357) = 15.05$, $p < .001$), a marginal main effect of recommendation type ($F(1, 357) = 3.58$, $p = .06$), and a significant main effect of product type ($F(1, 357) = 79.09$, $p < .001$). The hedonic description was seen as more hedonic ($M = 4.63$, $SD = 1.58$) than the utilitarian one ($M = 3.29$, $SD = 1.60$), which suggests that the product type manipulation worked as intended. Because the interaction was significant, I also looked at the pairwise comparisons. Recommendation type did not impact the perceived utilitarian/hedonic nature of the utilitarian description ($F(1, 357) = 2.04$, $p > .1$), but it did for the hedonic description ($M_{\text{Similarity-based}} = 5.25$, $SD = 1.14$, $M_{\text{AI-based}} = 4.27$, $SD = 1.68$, $F(1, 357) = 16.13$, $p < .001$). These exploratory results suggest that the AI-based product recommendation prompted participants to perceive the hedonic description of the service as being less hedonic than for the similarity-based recommendation. The finding that the hedonic product was perceived as more utilitarian in the AI condition could also be attributed to a typo in the stimuli, as this condition included utilitarian language.

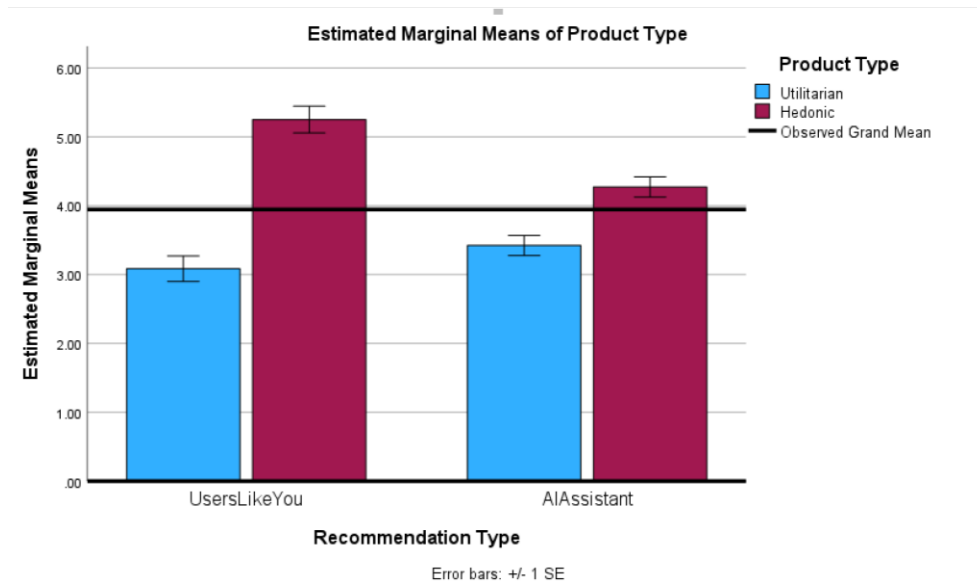


Figure 8. Interaction effect of recommendation and product type on the utilitarian/hedonic manipulation checks – Study 2

7 General Discussion

Recommendation systems have become central to digital platforms, influencing decisions and driving sales, as seen with Amazon’s 29% increase in sales and Spotify’s 33% music discovery through such algorithms (Rejoinder, 2021; Spotify, 2023). While AI-powered recommendation systems offer strong personalization capabilities (Covington et al., 2016), prior research shows that consumers often react negatively when AI is explicitly mentioned, due to lower emotional trust, especially in hedonic contexts (Cicek et al., 2025; Longoni & Cian, 2020; Wien & Peluso, 2021). Building on the concept of framing (Hardisty, 2024), I tested whether how a recommendation is labeled – whether AI-generated, based on one’s browsing history, or based on similar users’ preferences – influences consumers’ evaluations and behavioral intentions. I also investigated whether product type (i.e., utilitarian vs. hedonic; Longoni & Cian, 2020; Wien & Peluso, 2021; Choi et al., 2011), AI-related individual differences – literacy (Tully et al., 2024), perceived capability, magical beliefs (Khare et al., 2023), fear (Cave & Dihal, 2019), and usage (Liu et al., 2024) – and novelty seeking (Wang et al., 2000), shape consumers’ responses. These effects were tested across a pilot and two experimental studies.

The pilot study found that consumers were most familiar with behavior-based recommendations, followed by similarity- and AI-based ones. Trust in AI-based recommendations varied by context, with higher trust reported for entertainment-related services, like music and movie streaming, compared to utilitarian products, like electronics or small appliances. Based on this difference, product type (i.e., utilitarian vs. hedonic) was selected as a key moderator in the experimental studies. The pilot also revealed interesting patterns for AI-related individual differences: correlation analyses indicated that AI literacy was negatively related to magical beliefs, AI fear, and AI usage, while magical beliefs were positively related to AI fear and AI usage. These results guided the inclusion of individual difference variables in the experiments.

Study 1 tested the effects of AI-based versus behavior-based Amazon-style product recommendations across hedonic (i.e., hoodie) and utilitarian (i.e., power bank) products. Although the main effects of recommendation type were not significant, product type significantly influenced participants' confidence in the recommendations, as the utilitarian product (i.e., power bank) was rated more favorably than the hedonic one (i.e., hoodie). Moreover, AI-related individual differences, especially magical beliefs and usage, were positively associated with recommendation confidence. In hindsight, study 1 had several limitations that may have hindered its results, which I attempted to address in study 2.

Study 2 explored AI-based versus similarity-based recommendations in a novel context (i.e., recipe suggestions) and introduced novelty seeking as an additional moderator. Although the main and interaction effects of recommendation and product types were not significant, significant interactions were found with some of the AI-related individual differences and novelty seeking. Overall, AI capability and, to a lesser extent, fear and use, impacted responses toward the top of the marketing funnel (i.e., confidence in the recommendations, intentions to create a profile, intentions to activate a free trial), but they did not impact those toward the bottom of the funnel (i.e., intentions to pay for a monthly subscription, willingness-to-pay). Further, novelty seeking impacted responses related to discovery (i.e., intentions to create a profile and activate the free trial), but not other outcomes. Participants with higher novelty seeking tendencies had lower (higher) intentions to engage in these behaviors when the recipes were provided by AI (vs. a similarity-based algorithm) when the hedonic (utilitarian) characteristics of the service were highlighted.

In addition, both studies revealed intriguing preliminary findings regarding the potential impact of recommendation type on consumers' perceptions. Study 1 showed that AI-based (vs. behavior-based) recommendations prompted participants to evaluate the utilitarian product as being more utilitarian. Conversely, study 2 showed that AI-based (vs. similarity-based) recommendations prompted participants to evaluate the hedonic description as being less hedonic. These exploratory results thus suggest that AI-based recommendations may be more appropriate for utilitarian products, but future research is needed to further unpack these findings.

7.1 Theoretical Contributions

My research contributes to the growing literature on consumer responses to AI by demonstrating that reactions to AI-based recommendations are nuanced and context-dependent (Zehnle, Hildebrand, & Valenzuela, 2025). While prior studies often report either positive or negative attitudes toward AI-based technologies (Cheatham et al., 2019; Gursoy & Cai, 2024; Cicek, Gursoy, & Lu, 2025; Dietvorst et al., 2015; Longoni et al., 2019), my findings reveal more nuanced consumer responses, which also vary across different product types (i.e., utilitarian vs. hedonic) and individual characteristics (i.e., AI-related, novelty seeking). For example, AI-based (vs. similarity-based) recommendations were negatively perceived for a hedonic (vs. utilitarian) product by consumers higher in novelty seeking. Importantly, not all individual differences predicted consumer responses consistently. Specifically, while AI perceived capability and, to a lesser extent, fear and usage (depending on the context and outcome) significantly shaped consumer responses, other characteristics like AI literacy and magical beliefs had no effect, which is inconsistent with prior research on AI receptivity (Tully et al., 2024). My findings thus offer a more nuanced understanding of how both contextual and psychological variables can influence the

effectiveness of how recommendation systems are framed. They also suggest that consumers' receptivity to AI should be studied in specific marketing (e.g., product recommendations, chatbots, search) and decision (e.g., utilitarian vs. hedonic, material vs. experiential, self-control related) contexts, and consider relevant individual differences, rather than being assessed more generally.

7.2 Managerial Implications

My findings highlight the importance of aligning the framing of product recommendations with both the type of product being recommended and individual consumer traits. For instance, my results show that consumers with higher novelty-seeking tendencies respond more favorably to similarity-based (vs. AI-based) recommendations when evaluating a hedonic product. This suggests that marketers promoting hedonic offerings such as fashion, entertainment, or gourmet food should frame their recommendations as “based on users like you” rather than emphasizing AI for these consumers.

More broadly, my findings caution against relying on one-size-fits-all approaches to the design of recommendations. Context matters, what works for utilitarian products may not be effective for hedonic ones, and individual consumer characteristics are critical. Firms should move beyond demographic segmentation and incorporate psychographic traits, such as novelty seeking, into their segmentation and personalization strategies. By aligning the framing of recommendations with both product type and individual characteristics, marketers can enhance their perceived relevance, build trust, and increase behavioral engagement.

7.3 Limitations and Future Research Directions

This research has several limitations that open important avenues for future investigation. First, a high failure rate for the recommendation type manipulation check, particularly in Study 1 where 182 out of the final sample of 490 participants failed to correctly recall the type of recommendation they were exposed to, suggests that the stimuli may not have been memorable or noticeable enough. This could reflect participants' inattention or confusion between AI-based and behavior/similarity-based framing. Alternatively, since most participants who failed the manipulation checks in both studies were in the “no AI” condition, it may be that participants believe any type of recommendation is an AI-based recommendation. Future research should conduct more rigorous pre-tests to ensure the clarity and salience of the manipulations, and could also investigate whether consumers tend to perceive all types of product recommendations as being AI-based, even when they are not labeled as such.

Second, Study 1 did not yield significant effects, which may be attributed to the highly familiar context used (i.e., Amazon) and the use of a single product recommendation, while consumers are typically exposed to multiple recommendations on such platforms. Familiarity may thus have reduced the extent to which participants paid attention to the framing of the recommendation, and the single recommendation may have reduced the credibility of the stimuli. By contrast, Study 2 used a less familiar context (i.e., novel recipe recommendation service) and yielded some significant effects. Future research could investigate how familiarity (i.e., known vs. unknown brand/platform) and the number of recommendations (i.e., one vs. multiple) influence participants' attention and response to recommendations' framing. Future research could also investigate whether AI-based recommendations are better received in contexts of discovery (e.g., finding new

recipes, music, or movies), which anecdotal evidence from market research seems to suggest (Haan & Watts, 2023).

Third, another possible limitation relates to participants' general familiarity with the different types of recommendation sources. As observed in the pilot study, participants were most familiar with behavior-based recommendations, followed by similarity-based, and least familiar with AI-based ones. Since people tend to prefer what they know, this lack of familiarity may have contributed to the relatively less favorable responses toward AI-based recommendations across studies. Future research should consider measuring and controlling for perceived familiarity, or experimentally manipulating it, to better isolate the framing effects.

In addition, this thesis examined each recommendation type in a separate study. Future research should consider testing AI-based, behavior-based, and similarity-based recommendations within the same study. This would enable a direct comparison of their effects and provide a more complete understanding of how consumers evaluate different types of recommendation sources.

Fourth, individual differences that have been shown to predict responses to AI in prior work, such as AI literacy (Tully et al., 2024), did not influence participants' responses in my thesis. Future research could delve deeper into which individual characteristics matter most for different AI usage contexts (e.g., product recommendations, chatbots, search), and under what conditions they exert influence. Lastly, although not a primary focus of my thesis, my results revealed that consumers' perceptions of utilitarian and hedonic products might be impacted by recommendation type (i.e., AI-based vs. behavior/similarity-based). Future research could further explore the effects of recommendation type on product evaluations to better understand their role in shaping consumers' perceptions.

7.4 Conclusion

In sum, across one pilot and two experimental studies, I found that consumer responses to recommendation systems are shaped by a combination of framing, product type, and individual characteristics, with perceived AI capability and novelty seeking emerging as key moderators. While effects found in prior work, such as those related to AI literacy (Tully et al., 2024), were not observed in my thesis, my findings underscore the importance of context and personalization in the design of recommendations. Although more research is needed to better understand the mechanisms and boundary conditions of the effects, my thesis contributes to advancing our understanding of how AI-based recommendations are perceived, and how they can be better tailored to improve consumers' receptivity and engagement in digital environments.

References

- Boerman, S. C., Kruikemeier, S., & Zuiderveen Borgesius, F. J. (2017). Online behavioral advertising: A literature review and research agenda. *Journal of Advertising*, 46(3), 363–376. <https://doi.org/10.1080/00913367.2017.1339368>
- Campolo, Alexander, and Kate Crawford (2020), “Enchanted determinism: Power without responsibility in artificial intelligence,” *Engaging Science, Technology, & Society*, 6, 1-19.
- Cave, S., & Dihal, K. (2019). Hopes and fears for intelligent machines in fiction and reality. *Nature Machine Intelligence*, 1(2), 74–78. <https://doi.org/10.1038/s42256-019-0020-9>
- Cheatham, B., Javanmardian, K., & Samandari, H. (2019). Confronting the risks of artificial intelligence. McKinsey & Company. <https://www.mckinsey.com/~/media/McKinsey/Business%20Functions/McKinsey%20Analytics/Our%20Insights/Confronting%20the%20risks%20of%20artificial%20intelligence/Confronting-the-risks-of-artificial-intelligence-vF.pdf>
- Choi, J., Lee, H. J., & Kim, Y. C. (2011). The influence of social presence on customer intention to reuse online recommender systems: The roles of personalization and product type. *International Journal of Electronic Commerce*, 16(1), 129–154. <https://doi.org/10.2753/JEC1086-4415160105>
- Cicek, M., Gursoy, D., & Lu, L. (2025). Adverse impacts of revealing the presence of “Artificial Intelligence (AI)” technology in product and service descriptions on purchase intentions: The mediating role of emotional trust and the moderating role of perceived risk. *Journal of Hospitality Marketing & Management*, 34(1), 1–23. <https://doi.org/10.1080/19368623.2024.2368040>
- Covington, P., Adams, J., & Sargin, E. (2016). Deep neural networks for YouTube recommendations. In *Proceedings of the 10th ACM Conference on Recommender Systems* (pp. 191–198). <https://doi.org/10.1145/2959100.2959190>
- De Haan, E., Wiesel, T., & Pauwels, K. (2016). The effectiveness of different forms of online advertising for purchase conversion in a multiple-channel attribution framework. *International journal of research in marketing*, 33(3), 491-507.
- Dhar, Ravi and Klaus Wertenbroch (2000), “Consumer Choice Between Hedonic and Utilitarian Goods,” *Journal of Marketing Research*, 37 (1), 60–71.
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), 114–126. <https://doi.org/10.1037/xge0000033>
- Gursoy, D., & Cai, R. (2024). Artificial intelligence: An overview of research trends and future directions. *International Journal of Contemporary Hospitality Management*. <https://doi.org/10.1108/IJCHM-03-2024-032>

Haan, K., & Watts, R. (2023, July 20). Over 75% of consumers are concerned about misinformation from artificial intelligence. *Forbes*. Retrieved from <https://www.forbes.com/advisor/business/artificial-intelligence-consumer-sentiment/>

Habil, S., El-Deeb, S., & El-Bassiouny, N. (2023). AI-based recommendation systems: The ultimate solution for market prediction and targeting. In D. R. Fortin & S. D. Rodgers (Eds.), *The Palgrave Handbook of Interactive Marketing* (pp. 683–704). Springer. https://doi.org/10.1007/978-3-031-14961-0_30

Hardisty, D. J. (2024). Framing. In A. R. Heerde & K. N. Lemon (Eds.), *Elgar encyclopedia of consumer behavior* (Chapter 53). Edward Elgar Publishing. <https://doi.org/10.4337/9781803926278.ch53>

Hirschman, E. C. (1980). Innovativeness, novelty seeking, and consumer creativity. *Journal of Consumer Research*, 7 (3), 288-295

Katz, E., & Lazarsfeld, P. F. (1955). *Personal Influence: The Part Played by People in the Flow of Mass Communications*. New York: The Free Press

Khare, A., Kautish, P., & Khare, A. (2023). The online flow and its influence on awe experience: an AI-enabled e-tail service exploration. *International Journal of Retail & Distribution Management*, 51(6), 713-735.

Liu, Y., Li, Y., Song, K., & Chu, F. (2024). The two faces of Artificial Intelligence (AI): Analyzing how AI usage shapes employee behaviors in the hospitality industry. *International Journal of Hospitality Management*, 122, 103875. <https://doi.org/10.1016/j.ijhm.2024.103875>

Longoni, C., Bonezzi, A., & Morewedge, C. K. (2019). Resistance to medical artificial intelligence. *Journal of Consumer Research*, 46(4), 629–650. <https://doi.org/10.1093/jcr/ucz013>

Longoni, C., & Cian, L. (2020). Artificial intelligence in utilitarian vs. hedonic contexts: The “word-of-machine” effect. *Journal of Marketing*, 86(1), 91–108. <https://doi.org/10.1177/0022242920957347>

Önkal, D., Goodwin, P., Thomson, M., Gönül, M. S., & Pollock, A. (2009). The relative influence of advice from human experts and statistical methods on forecast adjustments. *Journal of Behavioral Decision Making*, 22(4), 390–409. <https://doi.org/10.1002/bdm.637>

Pu, P., Chen, L., & Hu, R. (2012). Evaluating recommender systems from the user’s perspective: Survey of the state of the art. *User Modeling and User-Adapted Interaction*, 22(4–5), 317–355. <https://doi.org/10.1007/s11257-011-9115-7>

Rejoinder. (2021). Amazon’s recommendation secret: The single most important thing Amazon does to sell online. <https://www.rejoinder.com/resources/amazon-recommendations-secret-selling-online>

Schepman, A., & Rodway, P. (2020). Initial validation of the general attitudes towards artificial intelligence scale. *Computers in Human Behavior Reports*, 1, 100014. <https://doi.org/10.1016/j.chbr.2020.100014>

Shanmugasundaram, M., & Tamilarasu, A. (2023). The impact of digital technology, social media, and artificial intelligence on cognitive functions: a review. *Frontiers in Cognition*, 2, 1203077.

Simonson, I., & Rosen, E. (2014). *Absolute value: What really influences customers in the age of (nearly) perfect information*. New York: Harper Collins.

Spotify for Artists. (2023). Discovery Mode. <https://artists.spotify.com/discovery-mode>

Steenkamp, J.-B. E. M., & Baumgartner, H. (1995). Development and cross-cultural validation of a short form of CSI as a measure of optimum stimulation level. *International Journal of Research in Marketing*, 12(2), 97–104. [https://doi.org/10.1016/0167-8116\(93\)E0035-8](https://doi.org/10.1016/0167-8116(93)E0035-8)

Summers, C. A., Smith, R. W., & Reczek, R. W. (2016). An audience of one: Behaviorally targeted ads as implied social labels. *Journal of Consumer Research*, 43(1), 156–178. <https://doi.org/10.1093/jcr/ucw012>

Tang, P. M., Koopman, J., McClean, S. T., Zhang, J. H., Li, C. H., De Cremer, D., & Ng, C. T. S. (2022). When conscientious employees meet intelligent machines: An integrative approach inspired by complementarity theory and role theory. *Academy of Management Journal*, 65(3), 1019–1054.

Thakkar, P., Varma, K., Ukani, V., Mankad, S., & Tanwar, S. (2019). Combining user-based and item-based collaborative filtering using machine learning. In S. C. Satapathy & A. Joshi (Eds.), *Information and communication technology for intelligent systems* (Vol. 107, pp. 173–180). Springer. https://doi.org/10.1007/978-981-13-1747-7_17

Tully, S. M., Longoni, C., & Appel, G. (2024). Lower artificial intelligence literacy predicts greater AI receptivity. *Marketing Science Institute Working Paper Series*, Report No. 24-132.

Wang, C.-L., Chen, Z.-X., Chan, A. K. K., & Zheng, Z.-C. (2000). The influence of hedonic values on consumer behaviors: An empirical investigation in China. *Journal of Global Marketing*, 14(1-2), 169–186. https://doi.org/10.1300/J042v14n01_09

Wien, A. H., & Peluso, A. M. (2021). Influence of human versus AI recommenders: The roles of product type and cognitive processes. *Journal of Business Research*, 137, 13–27. <https://doi.org/10.1016/j.jbusres.2021.08.016>

Yeomans, M., Shah, A., Mullainathan, S., & Kleinberg, J. (2019). Making sense of recommendations. *Journal of Behavioral Decision Making*, 32(4), 403–414. <https://doi.org/10.1002/bdm.2118>

Zehnle, M., Hildebrand, C., & Valenzuela, A. (2025). Not all AI is created equal: A meta-analysis revealing drivers of AI resistance across markets, methods, and time. *International Journal of Research in Marketing*

Zhang, Z., Zhang, Y., & Ren, Y. (2020). Employing neighborhood reduction for alleviating sparsity and cold start problems in user-based collaborative filtering. *Information Retrieval Journal*, 23(5), 449–472. <https://doi.org/10.1007/s10791-020-09378-w>

Appendices

A. Pilot Study Questions

[Attention checks]

The following questions help us ensure that you are paying attention to the study's instructions.

	True (1)	False (0)
A chicken is a type of insect.	☐	☐
A butterfly is a type of mammal.	☐	☐
A dog is a type of plant.	☐	☐

[Shopping habits]

What types of products do you most often buy from Amazon?

Select all that apply.

- ☐ I do not buy products on Amazon
- ☐ Electronics
- ☐ Books (e.g., physical and/or digital format)
- ☐ Entertainment (e.g., music, movies, box sets)
- ☐ Home & Kitchen
- ☐ Clothing & Accessories
- ☐ Health & Personal Care
- ☐ Groceries
- ☐ Other (Please specify): _____

Which streaming services do you most often use to listen to music?

Select all that apply.

- ☐ I do not use streaming services for music
- ☐ Spotify
- ☐ Apple Music
- ☐ Amazon Music
- ☐ YouTube Music
- ☐ SoundCloud
- ☐ Deezer
- ☐ Tidal
- ☐ Other (Please specify): _____

What streaming services do you most often use to watch movies and/or TV shows?

Select all that apply.

- ☐ I do not use streaming services for movies/tv shows
- ☐ Netflix
- ☐ Amazon Prime video
- ☐ Disney
- ☐ Hulu
- ☐ HBO Max

- ☐ Apple TV+
☐ Paramount+
☐ Peacock
☐ Other (Please specify): _____

How often do you encounter the following types of recommendations?

	Never 1	2	3	4	5	6	Very often 7
Consumers like you also like this song / movie / product.	o	o	o	o	o	o	o
Based on your purchase history / your listening habits / the last shows you watched we recommend you...	o	o	o	o	o	o	o
This recommendation is powered by AI (artificial intelligence).	o	o	o	o	o	o	o

[Product-specific AI trust]

I (would) trust AI-based product recommendations when purchasing...

	Strongly disagree 1	2	3	4	5	6	Strongly agree 7
Electronics	o	o	o	o	o	o	o
Fashion	o	o	o	o	o	o	o
Health	o	o	o	o	o	o	o
Beauty	o	o	o	o	o	o	o
Fitness/sports	o	o	o	o	o	o	o
Small appliances (e.g., vacuum cleaner, mixer)	o	o	o	o	o	o	o
Home decor	o	o	o	o	o	o	o
Groceries	o	o	o	o	o	o	o

I (would) trust AI-based recommendations for...

	Strongly disagree 1	2	3	4	5	6	Strongly agree 7
Music playlists	☐	☐	☐	☐	☐	☐	☐
Movies/TV shows	☐	☐	☐	☐	☐	☐	☐
Restaurants	☐	☐	☐	☐	☐	☐	☐
Books	☐	☐	☐	☐	☐	☐	☐
News articles	☐	☐	☐	☐	☐	☐	☐
Online courses	☐	☐	☐	☐	☐	☐	☐

[Perceived AI capability]

To what extent do you believe that AI (artificial intelligence) is capable of completing each of the following:

	Definitely incapable 1	2	3	4	5	6	Definitely capable 7
Recommending a product you would want to purchase	☐	☐	☐	☐	☐	☐	☐
Recommending a movie or TV show you would enjoy watching	☐	☐	☐	☐	☐	☐	☐
Recommending a playlist you would enjoy listening to	☐	☐	☐	☐	☐	☐	☐
Recommending a book you would enjoy reading	☐	☐	☐	☐	☐	☐	☐

[AI literacy]

The following questions aim to test your understanding of artificial intelligence. Please answer the questions to the best of your abilities.

In which area does AI (artificial intelligence) typically excel?

- ☐ Emotional understanding
- ☐ Pattern recognition
- ☐ Moral reasoning
- ☐ Creativity

What is a common form of knowledge representation in AI (artificial intelligence)?

- ☐ Neural networks
- ☐ Waterfall model
- ☐ Agile methodology
- ☐ SWOT analysis

How do supervised machine learning algorithms learn?

- ☐ From labeled data
- ☐ From rewards and punishments
- ☐ By observing human behavior
- ☐ From intrinsic motivation

Which statement best describes the programmability of AI (artificial intelligence) systems?

- ☐ They cannot be programmed by humans
- ☐ They program themselves
- ☐ They are programmed using data
- ☐ They are programmed by computer code

[General AI trust]

To what extent do you believe that AI-based recommendations...

	Strongly disagree 1	2	3	4	5	6	Strongly agree 7
are trustworthy.	☐	☐	☐	☐	☐	☐	☐
tell the truth.	☐	☐	☐	☐	☐	☐	☐
fulfill the promises made in the information provided.	☐	☐	☐	☐	☐	☐	☐
keep my best interests in mind when making recommendations.	☐	☐	☐	☐	☐	☐	☐
are generally predictable and consistent.	☐	☐	☐	☐	☐	☐	☐
are always honest.	☐	☐	☐	☐	☐	☐	☐

[AI magical beliefs]

If I think about artificial intelligence that powers products and services...

	Strongly disagree 1	2	3	4	5	6	Strongly agree 7
I experience feelings of awe.	☐	☐	☐	☐	☐	☐	☐

I feel that I am witness to something grand.	☐	☐	☐	☐	☐	☐	☐
I perceive vastness and I feel my jaw drop.	☐	☐	☐	☐	☐	☐	☐
I have goosebumps and gasp.	☐	☐	☐	☐	☐	☐	☐

[AI fear]

I fear that artificial intelligence will...

	Not at all 1	2	3	4	5	6	Very much 7
make us lose our humanity.	☐	☐	☐	☐	☐	☐	☐
turn against humans.	☐	☐	☐	☐	☐	☐	☐
make humans become obsolete.	☐	☐	☐	☐	☐	☐	☐
make humans lose the ability to connect with each other.	☐	☐	☐	☐	☐	☐	☐

[AI use; included as part of the demographics]

How often do you use artificial intelligence in your everyday life?

☐ Never ☐ 1 ☐ 2 ☐ 3 ☐ 4 ☐ 5 ☐ 6 ☐ Extremely often 7

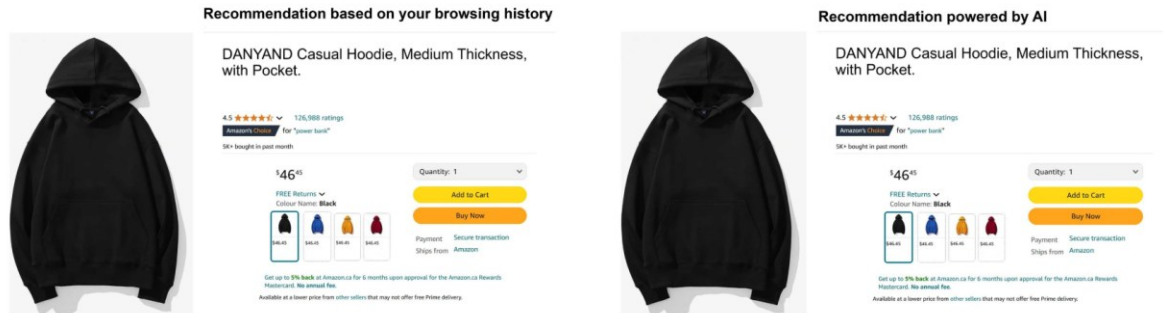
B. Study 1 Stimuli

[Utilitarian product]

Imagine you're planning a busy day out, and your phone's battery tends to drain quickly. You think it would be helpful to have a new power bank to stay connected throughout the day. As you browse online, you see a power bank that is being recommended for you. [Please carefully look at the image below.]

[Behavior-based recommendation]

[AI-based recommendation]

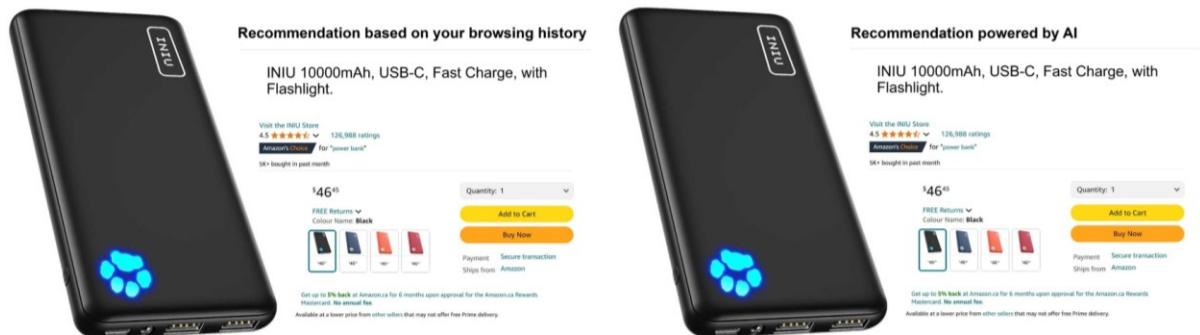


[Hedonic product]

Imagine the weather is getting colder, and you're thinking about refreshing your wardrobe. You think it would be helpful to have a new hoodie to keep you warm. As you browse online, you see a hoodie that is being recommended for you. [Please carefully look at the image below.]

[Behavior-based recommendation]

[AI-based recommendation]



[Please refer back to the image when answered the following questions.]

C. Study 1 Questions

[Attention checks]

Similar to the one shown in appendix A

[Recommendation confidence]

To what extent...

	Not at all 1	2	3	4	5	6	Extremely 7
do you trust this product recommendation?	o	o	o	o	o	o	o
do you find this product recommendation relevant?	o	o	o	o	o	o	o

[Behavioral intentions]

What is the likelihood that you would...

	Extremely unlikely 1	2	3	4	5	6	Extremely likely 7
click on the recommended product to learn more about it?	o	o	o	o	o	o	o
put the recommended product in your cart?	o	o	o	o	o	o	o

[Product type manipulation checks]

Purchasing electronics, such as a power bank [hoodie] is an...

	1	2	3	4	5	6	7	
Entirely objective purchase	o	o	o	o	o	o	o	Entirely subjective purchase
Entirely rational purchase	o	o	o	o	o	o	o	Entirely emotional purchase

[Preference strength]

When purchasing electronics [clothes], such as a power bank [hoodie]...

	Strongly disagree 1	2	3	4	5	6	Strongly agree 7
I know exactly what I want to buy.	o	o	o	o	o	o	o

I have very specific preferences.	☐	☐	☐	☐	☐	☐	☐
I only buy a(certain) specific brand(s).	☐	☐	☐	☐	☐	☐	☐

[AI capability]

To what extent do you believe that AI (artificial intelligence is capable of) recommending the following products:

	Definitely incapable 1	2	3	4	5	6	Definitely capable 7
Electronics	☐	☐	☐	☐	☐	☐	☐
Clothes	☐	☐	☐	☐	☐	☐	☐

[AI literacy, magical beliefs, fear]

See appendix A

[Attention checks]

In this survey, I was asked to evaluate a...

- ☐ Power bank (electronics)
- ☐ Hoodie (clothes)

In the scenario, the product recommendation was...

- ☐ Based on your browsing history
- ☐ Powered by AI

[AI use; included as part of the demographics]

See appendix A

D. Study 1 ANOVA Results

→ Univariate Analysis of Variance

Between-Subjects Factors

		Value Label	N
ElectroClothes	0	Electronics	248
	1	Clothes	242
AIReco	0	NoAI	244
	1	AI	246

Descriptive Statistics

Dependent Variable: mTrustRelevant

ElectroClothes	AIReco	Mean	Std. Deviation	N
Electronics	NoAI	5.4556	1.06689	124
	AI	5.6089	.95543	124
	Total	5.5323	1.01356	248
Clothes	NoAI	5.0875	1.32631	120
	AI	5.0615	1.30571	122
	Total	5.0744	1.31329	242
Total	NoAI	5.2746	1.21310	244
	AI	5.3374	1.17281	246
	Total	5.3061	1.19224	490

Tests of Between-Subjects Effects

Dependent Variable: mTrustRelevant

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	27.175 ^a	3	9.058	6.591	<.001
Intercept	13779.082	1	13779.082	10026.304	<.001
ElectroClothes	25.666	1	25.666	18.675	<.001
AIReco	.495	1	.495	.360	.549
ElectroClothes * AIReco	.984	1	.984	.716	.398
Error	667.906	486	1.374		
Total	14491.000	490			
Corrected Total	695.082	489			

a. R Squared = .039 (Adjusted R Squared = .033)

Univariate Analysis of Variance

Between-Subjects Factors

		Value Label	N
ElectroClothes	0	Electronics	248
	1	Clothes	242
AIReco	0	NoAI	244
	1	AI	246

Descriptive Statistics

Dependent Variable: mSubjEmo

ElectroClothes	AIReco	Mean	Std. Deviation	N
Electronics	NoAI	2.7298	1.44323	124
	AI	2.3911	1.21902	124
	Total	2.5605	1.34389	248
Clothes	NoAI	3.1875	1.27197	120
	AI	3.3361	1.50130	122
	Total	3.2624	1.39142	242
Total	NoAI	2.9549	1.37811	244
	AI	2.8598	1.44337	246
	Total	2.9071	1.41061	490

Tests of Between-Subjects Effects

Dependent Variable: mSubjEmo

Source	Type III Sum of Squares	df	Mean Square	F	Sig.
Corrected Model	68.793 ^a	3	22.931	12.325	<.001
Intercept	4151.835	1	4151.835	2231.497	<.001
ElectroClothes	60.237	1	60.237	32.376	<.001
AIReco	1.107	1	1.107	.595	.441
ElectroClothes * AIReco	7.270	1	7.270	3.908	.049
Error	904.232	486	1.861		
Total	5114.250	490			
Corrected Total	973.025	489			

a. R Squared = .071 (Adjusted R Squared = .065)

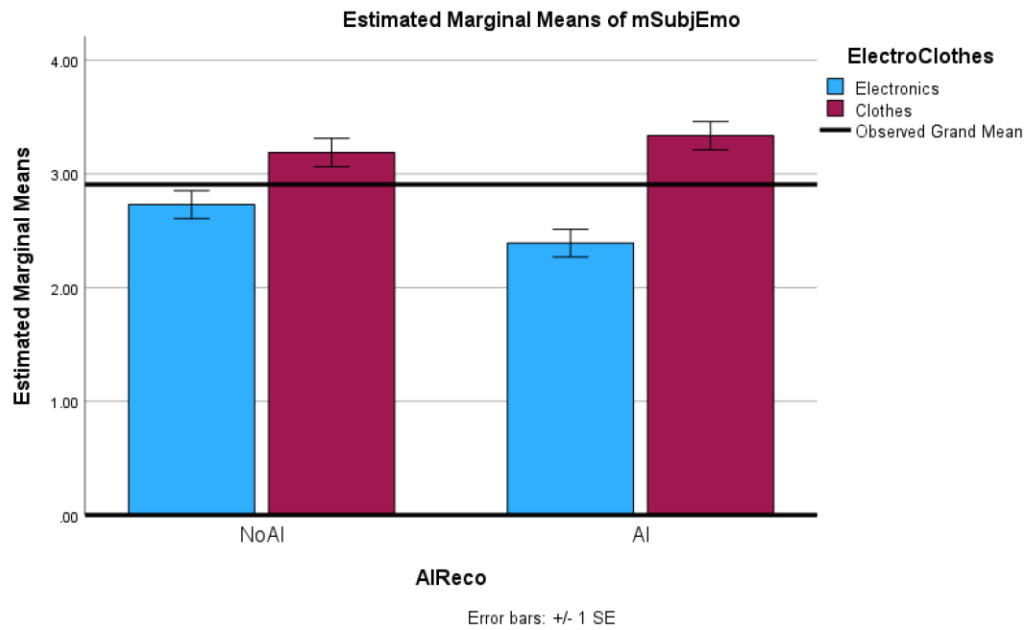
Estimated Marginal Means

ElectroClothes * AIReco

Dependent Variable: mSubjEmo

ElectroClothes	AIReco	Mean	Std. Error	95% Confidence Interval	
				Lower Bound	Upper Bound
Electronics	NoAI	2.730	.122	2.489	2.971
	AI	2.391	.122	2.150	2.632
Clothes	NoAI	3.188	.125	2.943	3.432
	AI	3.336	.123	3.093	3.579

Profile Plots



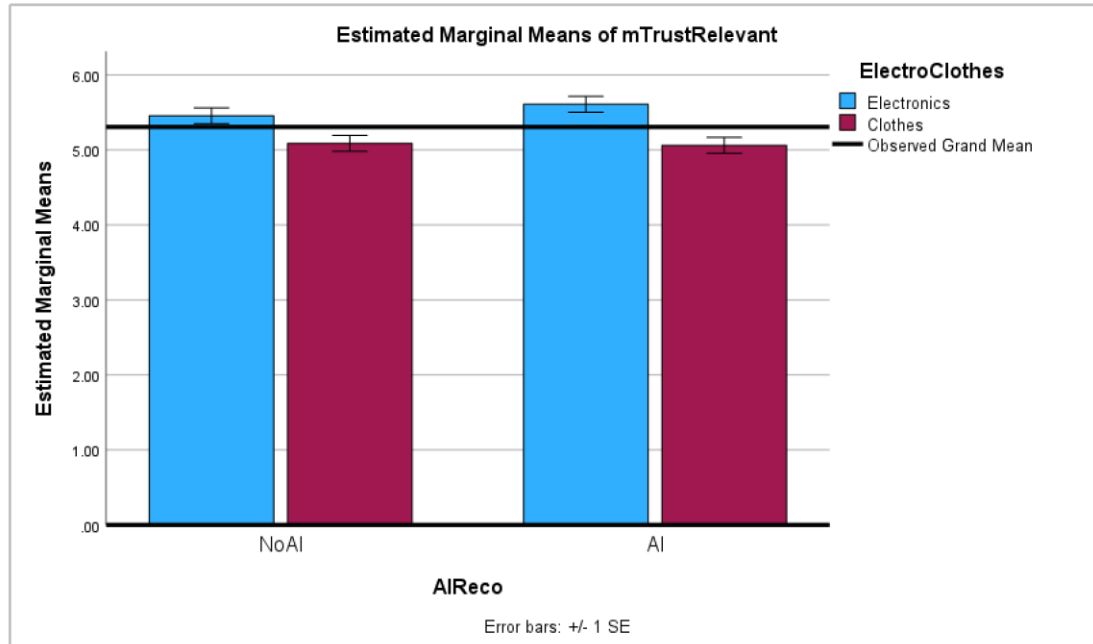
Estimated Marginal Means

ElectroClothes * AIReco

Dependent Variable: mTrustRelevant

ElectroClothes	AIReco	Mean	Std. Error	95% Confidence Interval	
				Lower Bound	Upper Bound
Electronics	NoAI	5.456	.105	5.249	5.662
	AI	5.609	.105	5.402	5.816
Clothes	NoAI	5.088	.107	4.877	5.298
	AI	5.061	.106	4.853	5.270

Profile Plots



E. Study 1 Correlation Table

Correlations

		Correlations				
		SumLiteracy	mStrenghtPrefs	mAlMagical	mAlFear	AlUse
SumLiteracy	Pearson Correlation	1	-.080	-.189**	-.027	-.056
	Sig. (2-tailed)		.078	<.001	.558	.215
	N	490	490	490	490	490
mStrenghtPrefs	Pearson Correlation	-.080	1	.260**	.063	.167**
	Sig. (2-tailed)	.078		<.001	.165	<.001
	N	490	490	490	490	490
mAlMagical	Pearson Correlation	-.189**	.260**	1	.098*	.447**
	Sig. (2-tailed)	<.001	<.001		.029	<.001
	N	490	490	490	490	490
mAlFear	Pearson Correlation	-.027	.063	.098*	1	-.196**
	Sig. (2-tailed)	.558	.165	.029		<.001
	N	490	490	490	490	490
AlUse	Pearson Correlation	-.056	.167**	.447**	-.196**	1
	Sig. (2-tailed)	.215	<.001	<.001	<.001	
	N	490	490	490	490	490

** . Correlation is significant at the 0.01 level (2-tailed).

* . Correlation is significant at the 0.05 level (2-tailed).

Model : 2

Y : mTrustRe

X : AIReco

W : ElectroC

Z : mAlFear

Sample

Size: 490

OUTCOME VARIABLE:

mTrustRe

Model Summary

R	R-sq	MSE	F	df1	df2	p
.2132	.0454	1.3709	4.6077	5.0000	484.0000	.0004

Model

	coeff	se	t	p	LLCI	ULCI
constant	5.7127	.1850	30.8717	.0000	5.3491	6.0762
AIReco	-.0156	.2580	-.0603	.9519	-.5224	.4913
ElectroC	-.3692	.1499	-2.4624	.0141	-.6638	-.0746
Int_1	-.1757	.2116	-.8303	.4068	-.5915	.2401
mAlFear	-.0783	.0464	-1.6878	.0921	-.1694	.0128
Int_2	.0515	.0641	.8032	.4223	-.0745	.1774

Product terms key:

Int_1 : AIReco x ElectroC

Int_2 : AIReco x mAlFear

Test(s) of highest order unconditional interaction(s):

	R2-chng	F	df1	df2	p
X*W	.0014	.6893	1.0000	484.0000	.4068
X*Z	.0013	.6451	1.0000	484.0000	.4223
BOTH	.0026	.6598	2.0000	484.0000	.5174

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:

95.0000

WARNING: Variables names longer than eight characters can produce incorrect output when some variables in the data file have the same first eight characters. Shorter variable names are recommended. By using this output, you are accepting all risk and consequences of interpreting or reporting results that may be incorrect.

—— END MATRIX ——

F. Study 2 Stimuli

Imagine you have been struggling with coming up with meal ideas. After doing an internet search for inspiration, you find a subscription service that offers personalized recipe recommendations.

[Utilitarian product]

The service offers quick and easy recipes, with minimal steps and ingredients, that are specifically designed to reduce your mental load and simplify your meal planning and cooking.

[Hedonic product]

The service offers fun and yummy recipes, made with delicious ingredients, that are specifically designed to add pleasure and excitement to your meal planning and cooking.

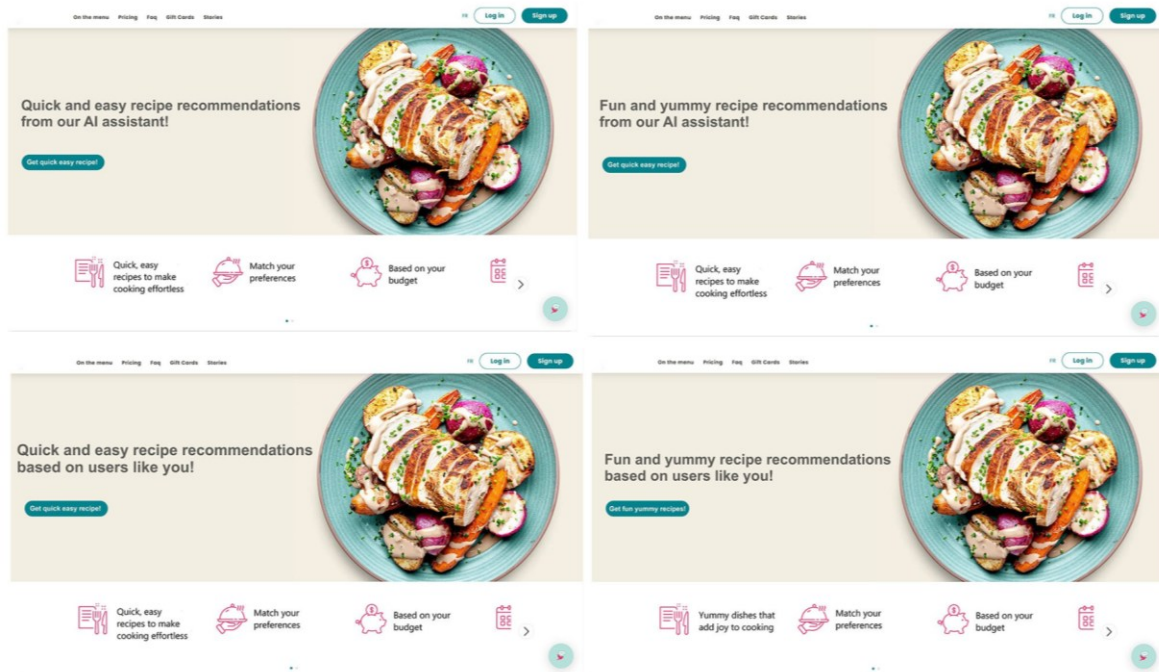
To get recipe recommendations, you need to create a profile where you indicate your food preferences and allergies, diet restrictions, budget, and cooking habits.

[AI-based recommendation]

An AI assistant then processes your data to generate your personalized recipe recommendations.

[Similarity-based recommendation]

An algorithm then matches your profile with those of other users like you to generate your personalized recipe recommendations.



G. Study 2 Questions

[Attention checks]

Similar to the one shown in appendix A

[Recommendation confidence]

To what extent would you...

	Not at all 1	2	3	4	5	6	Extremely 7
trust the recommended recipes	☐	☐	☐	☐	☐	☐	☐
like getting recipe recommendations based on users like you	☐	☐	☐	☐	☐	☐	☐

[Behavioral intentions]

The service allows you to preview your recipe recommendations after creating a profile, and get full access to the recommended recipes by activating a 14-days free trial. You then need to purchase a monthly subscription (cancelable at any time with one click) to continue getting recipe recommendations after the free trial expires. What is the likelihood that you would...

	Extremely unlikely 1	2	3	4	5	6	Extremely likely 7
create a profile to see your recommended recipes	☐	☐	☐	☐	☐	☐	☐
activate a free trial of the service	☐	☐	☐	☐	☐	☐	☐
pay for a monthly subscription	☐	☐	☐	☐	☐	☐	☐

[Willingness-to-pay]

How much would you be willing to pay for the monthly subscription? Enter a dollar amount between US\$0 and US\$20. _____

[Product type manipulation checks]

The recipe recommendation service mostly emphasizes the _____ aspects of cooking.

	1	2	3	4	5	6	7	
Functional	☐	☐	☐	☐	☐	☐	☐	Enjoyable
Practical	☐	☐	☐	☐	☐	☐	☐	Pleasurable
Utilitarian	☐	☐	☐	☐	☐	☐	☐	Hedonic

[Ai literacy, magical beliefs, and fear]
See appendix A

[Novelty seeking]

How well do the following statements describe your personality?

	Completely false -3	-2	-1	0	1	2	Completely true 3
I like to continue doing the same old things rather than trying new and different things.	☐	☐	☐	☐	☐	☐	☐
I like to experience novelty and change in my daily routine.	☐	☐	☐	☐	☐	☐	☐
I am continually seeking new ideas and experiences.	☐	☐	☐	☐	☐	☐	☐
I like continually changing activities.	☐	☐	☐	☐	☐	☐	☐
When things get boring, I like to find some new and unfamiliar experience.	☐	☐	☐	☐	☐	☐	☐
I prefer a routine way of life to an unpredictable one full of change.	☐	☐	☐	☐	☐	☐	☐

[Attention checks]

In the scenario about the recipe recommendation service, the recipes were...

- ☐ recommended based on users like you
- ☐ recommended by an AI assistant

[AI use; included as part of the demographics]

See appendix A

H. Study 2 Correlation Table

Correlations

		Correlations				
		AIUse	mAIFear	mAIMagic	SumLiteracy	AICapability
AIUse	Pearson Correlation	1	-.206**	.390**	.020	.279**
	Sig. (2-tailed)		<.001	<.001	.702	<.001
	N	361	361	361	361	361
mAIFear	Pearson Correlation	-.206**	1	-.111*	-.081	-.210**
	Sig. (2-tailed)	<.001		.034	.124	<.001
	N	361	361	361	361	361
mAIMagic	Pearson Correlation	.390**	-.111*	1	-.081	.136**
	Sig. (2-tailed)	<.001	.034		.125	.010
	N	361	361	361	361	361
SumLiteracy	Pearson Correlation	.020	-.081	-.081	1	-.065
	Sig. (2-tailed)	.702	.124	.125		.220
	N	361	361	361	361	361
AICapability	Pearson Correlation	.279**	-.210**	.136**	-.065	1
	Sig. (2-tailed)	<.001	<.001	.010	.220	
	N	361	361	361	361	361

** . Correlation is significant at the 0.01 level (2-tailed).

* . Correlation is significant at the 0.05 level (2-tailed).

I. Study 2 PROCESS Model 1 Outputs

Model : 1
Y : mTrustLi
X : Alreco
W : AICapabi

Sample
Size: 361

OUTCOME VARIABLE:
mTrustLi

Model Summary

R	R-sq	MSE	F	df1	df2	p
.6602	.4359	.9475	91.9429	3.0000	357.0000	.0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	3.5198	.3424	10.2811	.0000	2.8465	4.1931
Alreco	-2.6247	.4386	-5.9849	.0000	-3.4872	-1.7623
AICapabi	.3137	.0606	5.1733	.0000	.1945	.4330
Int_1	.4251	.0767	5.5402	.0000	.2742	.5761

Product terms key:

Int_1 : Alreco x AICapabi

Test(s) of highest order unconditional interaction(s):

	R2-chng	F	df1	df2	p
X*W	.0485	30.6940	1.0000	357.0000	.0000

Focal predict: Alreco (X)
Mod var: AICapabi (W)

Conditional effects of the focal predictor at values of the moderator(s):

AICapabi	Effect	se	t	p	LLCI	ULCI
4.9200	-.5331	.1165	-4.5773	.0000	-.7621	-.3040
6.0000	-.0739	.1117	-.6617	.5086	-.2936	.1458
7.0000	.3512	.1540	2.2804	.0232	.0483	.6541

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
95.0000

W values in conditional tables are the 16th, 50th, and 84th percentiles.

Model : 1

Y : CreatePr

X : Alreco

W : AICapabi

Sample

Size: 361

OUTCOME VARIABLE:

CreatePr

Model Summary

R	R-sq	MSE	F	df1	df2	p
.4068	.1655	2.6810	23.5944	3.0000	357.0000	.0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	4.2029	.5759	7.2981	.0000	3.0703	5.3355
Alreco	-3.0755	.7377	-4.1689	.0000	-4.5263	-1.6246
AICapabi	.1267	.1020	1.2418	.2151	-.0739	.3273
Int_1	.5313	.1291	4.1160	.0000	.2774	.7852

Product terms key:

Int_1 : Alreco x AICapabi

Test(s) of highest order unconditional interaction(s):

	R2-chng	F	df1	df2	p
X*W	.0396	16.9411	1.0000	357.0000	.0000

Focal predict: Alreco (X)

Mod var: AICapabi (W)

Conditional effects of the focal predictor at values of the moderator(s):

AICapabi	Effect	se	t	p	LLCI	ULCI
4.9200	-.4615	.1959	-2.3556	.0190	-.8467	-.0762
6.0000	.1123	.1879	.5978	.5503	-.2572	.4819
7.0000	.6436	.2591	2.4843	.0134	.1341	1.1532

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:

95.0000

W values in conditional tables are the 16th, 50th, and 84th percentiles.

Model : 1
Y : FreeTria
X : Alreco
W : AICapabi

Sample
Size: 361

OUTCOME VARIABLE:
FreeTria

Model Summary

R	R-sq	MSE	F	df1	df2	p
.3241	.1050	3.2866	13.9627	3.0000	357.0000	.0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	4.5787	.6376	7.1808	.0000	3.3247	5.8326
Alreco	-3.2120	.8168	-3.9324	.0001	-4.8183	-1.6056
AICapabi	-.0083	.1129	-.0733	.9416	-.2304	.2138
Int_1	.5748	.1429	4.0217	.0001	.2937	.8559

Product terms key:

Int_1 : Alreco x AICapabi

Test(s) of highest order unconditional interaction(s):

R2-chng	F	df1	df2	p	
X*W	.0405	16.1743	1.0000	357.0000	.0001

Focal predict: Alreco (X)

Mod var: AICapabi (W)

Conditional effects of the focal predictor at values of the moderator(s):

AICapabi	Effect	se	t	p	LLCI	ULCI
4.9200	-.3840	.2169	-1.7705	.0775	-.8106	.0425
6.0000	.2368	.2080	1.1380	.2559	-.1724	.6459
7.0000	.8115	.2869	2.8291	.0049	.2474	1.3757

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
95.0000

W values in conditional tables are the 16th, 50th, and 84th percentiles.

Model : 1

Y : mTrustLi

X : Alreco

W : mAlFear

Sample

Size: 361

OUTCOME VARIABLE:

mTrustLi

Model Summary

R	R-sq	MSE	F	df1	df2	p
.2490	.0620	1.5754	7.8657	3.0000	357.0000	.0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	5.3908	.2493	21.6282	.0000	4.9006	5.8810
Alreco	.5247	.3180	1.6498	.0999	-.1008	1.1502
mAlFear	-.0452	.0660	-.6847	.4940	-.1749	.0846
Int_1	-.2160	.0865	-2.4979	.0129	-.3861	-.0459

Product terms key:

Int_1 : Alreco x mAlFear

Test(s) of highest order unconditional interaction(s):

	R2-chng	F	df1	df2	p
X*W	.0164	6.2393	1.0000	357.0000	.0129

Focal predict: Alreco (X)

Mod var: mAlFear (W)

Conditional effects of the focal predictor at values of the moderator(s):

mAlFear	Effect	se	t	p	LLCI	ULCI
1.5000	.2006	.2085	.9622	.3366	-.2095	.6107
3.2500	-.1774	.1369	-1.2956	.1960	-.4467	.0919
5.0000	-.5555	.1996	-2.7831	.0057	-.9480	-.1630

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:

95.0000

W values in conditional tables are the 16th, 50th, and 84th percentiles.

Model : 1
 Y : CreatePr
 X : Alreco
 W : mAlFear

Sample
 Size: 361

OUTCOME VARIABLE:
 CreatePr

Model Summary

R	R-sq	MSE	F	df1	df2	p
.1216	.0148	3.1650	1.7860	3.0000	357.0000	.1494

Model

	coeff	se	t	p	LLCI	ULCI
constant	4.5702	.3533	12.9358	.0000	3.8754	5.2650
Alreco	.8076	.4508	1.7916	.0740	-.0789	1.6942
mAlFear	.0958	.0935	1.0243	.3064	-.0881	.2796
Int_1	-.2594	.1226	-2.1164	.0350	-.5005	-.0184

Product terms key:

Int_1 : Alreco x mAlFear

Test(s) of highest order unconditional interaction(s):

	R2-chng	F	df1	df2	p
X*W	.0124	4.4791	1.0000	357.0000	.0350

Focal predict: Alreco (X)
 Mod var: mAlFear (W)

Conditional effects of the focal predictor at values of the moderator(s):

mAlFear	Effect	se	t	p	LLCI	ULCI
1.5000	.4185	.2956	1.4158	.1577	-.1628	.9998
3.2500	-.0356	.1941	-.1832	.8547	-.4173	.3462
5.0000	-.4896	.2829	-1.7306	.0844	-1.0460	.0668

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:
 95.0000

W values in conditional tables are the 16th, 50th, and 84th percentiles.

Model : 1

Y : mTrustLi

X : Alreco

W : AIUse

Sample

Size: 361

OUTCOME VARIABLE:

mTrustLi

Model Summary

R	R-sq	MSE	F	df1	df2	p
.3449	.1190	1.4797	16.0708	3.0000	357.0000	.0000

Model

	coeff	se	t	p	LLCI	ULCI
constant	4.7568	.2511	18.9454	.0000	4.2630	5.2505
Alreco	-.9472	.3281	-2.8868	.0041	-1.5925	-.3019
AIUse	.1093	.0520	2.1045	.0360	.0072	.2115
Int_1	.1572	.0661	2.3762	.0180	.0271	.2873

Product terms key:

Int_1 : Alreco x AIUse

Test(s) of highest order unconditional interaction(s):

	R2-chng	F	df1	df2	p
X*W	.0139	5.6463	1.0000	357.0000	.0180

Focal predict: Alreco (X)

Mod var: AIUse (W)

Conditional effects of the focal predictor at values of the moderator(s):

AIUse	Effect	se	t	p	LLCI	ULCI
2.0000	-.6328	.2140	-2.9577	.0033	-1.0536	-.2121
5.0000	-.1613	.1364	-1.1829	.2377	-.4295	.1069
7.0000	.1531	.2103	.7278	.4672	-.2605	.5667

***** ANALYSIS NOTES AND ERRORS *****

Level of confidence for all confidence intervals in output:

95.0000

W values in conditional tables are the 16th, 50th, and 84th percentiles.

J. Study 2 PROCESS Model 3 Outputs

Model : 3
 Y : CreatePr
 X : Alreco
 W : UtilHed
 Z : mNovelSe

Sample
 Size: 361

OUTCOME VARIABLE:

CreatePr

Model Summary

R	R-sq	MSE	F	df1	df2	p
.2561	.0656	3.0359	3.5395	7.0000	353.0000	.0011

Model

	coeff	se	t	p	LLCI	ULCI
constant	4.2265	.7329	5.7664	.0000	2.7850	5.6680
Alreco	-1.1762	.9620	-1.2226	.2223	-3.0682	.7158
UtilHed	-1.4021	1.0506	-1.3346	.1829	-3.4684	.6641
Int_1	2.4126	1.3742	1.7556	.0800	-.2901	5.1154
mNovelSe	.1721	.1721	.9999	.3180	-.1664	.5107
Int_2	.2967	.2215	1.3390	.1814	-.1391	.7324
Int_3	.2922	.2380	1.2277	.2204	-.1759	.7604
Int_4	-.6364	.3056	-2.0827	.0380	-1.2374	-.0355

Product terms key:

Int_1 : Alreco x UtilHed
 Int_2 : Alreco x mNovelSe
 Int_3 : UtilHed x mNovelSe
 Int_4 : Alreco x UtilHed x mNovelSe

Test(s) of highest order unconditional interaction(s):

R2-chng	F	df1	df2	p	
X*W*Z	.0115	4.3377	1.0000	353.0000	.0380

Focal predict: Alreco (X)
 Mod var: UtilHed (W)
 Mod var: mNovelSe (Z)

Test of conditional X*W interaction at value(s) of Z:

mNovelSe	Effect	F	df1	df2	p
3.0000	.5033	.8183	1.0000	353.0000	.3663
4.5000	-.4514	1.3548	1.0000	353.0000	.2452
5.6700	-1.1960	4.4991	1.0000	353.0000	.0346

Model : 3

Y : FreeTria

X : Alreco

W : UtilHed

Z : mNovelSe

Sample

Size: 361

OUTCOME VARIABLE:

FreeTria

Model Summary

R	R-sq	MSE	F	df1	df2	p
.2174	.0472	3.5383	2.5007	7.0000	353.0000	.0161

Model

	coeff	se	t	p	LLCI	ULCI
constant	4.8169	.7913	6.0875	.0000	3.2607	6.3731
Alreco	-2.0361	1.0386	-1.9605	.0507	-4.0786	.0065
UtilHed	-2.0565	1.1342	-1.8131	.0707	-4.2872	.1742
Int_1	3.4942	1.4836	2.3552	.0191	.5763	6.4120
mNovelSe	-.0827	.1858	-.4453	.6564	-.4482	.2827
Int_2	.5422	.2392	2.2669	.0240	.0718	1.0126
Int_3	.5011	.2570	1.9499	.0520	-.0043	1.0065
Int_4	-.9236	.3299	-2.7996	.0054	-1.5724	-.2748

Product terms key:

Int_1 : Alreco x UtilHed
 Int_2 : Alreco x mNovelSe
 Int_3 : UtilHed x mNovelSe
 Int_4 : Alreco x UtilHed x mNovelSe

Test(s) of highest order unconditional interaction(s):

R2-chng	F	df1	df2	p	
X*W*Z	.0212	7.8376	1.0000	353.0000	.0054

 Focal predict: Alreco (X)
 Mod var: UtilHed (W)
 Mod var: mNovelSe (Z)

Test of conditional X*W interaction at value(s) of Z:

mNovelSe	Effect	F	df1	df2	p
3.0000	.7234	1.4504	1.0000	353.0000	.2293
4.5000	-.6620	2.5004	1.0000	353.0000	.1147
5.6700	-1.7426	8.1947	1.0000	353.0000	.0045

Conditional effects of the focal predictor at values of the moderator(s):

UtilHed	mNovelSe	Effect	se	t	p	LLCI	ULCI
.0000	3.0000	-.4095	.4009	-1.0215	.3077	-1.1979	.3789
.0000	4.5000	.4038	.2965	1.3618	.1741	-.1794	.9869
.0000	5.6700	1.0382	.4581	2.2663	.0240	.1372	1.9391
1.0000	3.0000	.3139	.4473	.7017	.4833	-.5658	1.1936
1.0000	4.5000	-.2582	.2956	-.8736	.3829	-.8395	.3231
1.0000	5.6700	-.7044	.4009	-1.7572	.0798	-1.4929	.0840

Data for visualizing the conditional effect of the focal predictor:

Paste text below into a SPSS syntax window and execute to produce plot.

DATA LIST FREE/

Alreco UtilHed mNovelSe FreeTria .

BEGIN DATA.

.0000	.0000	3.0000	4.5686
1.0000	.0000	3.0000	4.1591
.0000	.0000	4.5000	4.4445
1.0000	.0000	4.5000	4.8483
.0000	.0000	5.6700	4.3477
1.0000	.0000	5.6700	5.3859
.0000	1.0000	3.0000	4.0154
1.0000	1.0000	3.0000	4.3293
.0000	1.0000	4.5000	4.6429
1.0000	1.0000	4.5000	4.3847
.0000	1.0000	5.6700	5.1324
1.0000	1.0000	5.6700	4.4279

END DATA.

GRAPH/SCATTERPLOT=

UtilHed WITH FreeTria BY Alreco /PANEL ROWVAR= mNovelSe .

Graph

