



MASTERARBEIT | MASTER'S THESIS

Titel | Title

The Influence of Conversational Recommender Systems on
User Trust: An Experimental Study

verfasst von | submitted by

Laura Valentina Modre BA BSc

angestrebter akademischer Grad | in partial fulfilment of the requirements for the degree of
Master of Science (MSc)

Wien | Vienna, 2025

Studienkennzahl lt. Studienblatt | Degree
programme code as it appears on the
student record sheet:

UA 066 840

Studienrichtung lt. Studienblatt | Degree
programme as it appears on the student
record sheet:

Masterstudium Psychologie

Betreut von | Supervisor:

Univ.-Prof. Dipl.-Psych. Dr. Robert Böhm

To my family, for their enduring support;
and to my partner, for his unwavering encouragement.

Table of Contents

List of Abbreviations.....	1
Introduction.....	2
Overview of Recommender Systems.....	4
From Filters to Conversation: Conversational Recommender Systems	5
Key Design Factors Shaping Trust in Conversational Recommender Systems	5
Psychological Foundations of Trust.....	7
Defining Trust in AI Contexts.....	8
Comparing Human and AI Trust.....	9
A Socio-Technical Perspective on AI Trust	10
Heuristic vs. Deliberative Trust Processes.....	11
Psychological Antecedents of Trust	12
Explainability in AI.....	12
User Control in AI Interactions.....	15
Anthropomorphism in AI.....	17
Research Questions and Hypotheses	20
Direct Effect of Conversational Recommender Systems on Trust	20
Mediating Roles of Explainability, Control & Anthropomorphism.....	21
Exploratory Questions	23
Method	25
Experimental Design and Procedure.....	25
Power Analysis & Sample Size.....	26
Participants and Recruitment	27
Participant Characteristics and Demographics	28
Measures & Data Collection	28
Data Analysis	29
Assumption Checks	30
Conducting the Parallel Mediation	31

Sensitivity Analyses	32
Exploratory & Qualitative Analyses	32
Data Quality & Inclusion Criteria	33
Results	33
Assumption Checks	33
Appropriateness of Variables	33
Linearity	33
Multicollinearity	34
Homoscedasticity	34
Normality	34
Independence	34
Descriptives	35
Descriptive Results of Survey Responses	35
Model Summary	36
Hypothesis Tests	36
Robustness & Sensitivity Checks	38
Exploratory Outcomes	40
Qualitative Feedback	41
Discussion	42
Connecting Results to Current Literature	42
Explainability Without Enhancement	42
The Paradox of Conversational Control	43
The Modest Role of Anthropomorphism	43
Theoretical Contributions	43
Limitations	44
Evolving Research Design	44
Reactance and Individual Differences	45
Temporal and Contextual Limitations	46
Methodological Constraints	46

Generalizability	46
Implications.....	47
Practical Recommendations for Conversational Recommender Systems Design	47
Research Directions	48
Future Research Avenues	48
References.....	50
List of Figures	59
List of Tables.....	59
Appendix.....	60
Abstract Deutsch.....	60
Abstract English.....	61
Ethics and Open Science Statement.....	62
Limesurvey Study Introduction	63
Limesurvey Exit Survey	65
R Script and Results of Data Analysis	71

List of Abbreviations

AI	Artificial Intelligence
BRS	Baseline Recommender System
CRS	Conversational Recommender System
DV	Dependent Variable
EU	European Union
EU AI Acts	European Union Artificial Intelligence Act
GDPR	General Data Protection Regulation
IV	Independent Variable
LLM	Large Language Model
MV	Mediating Variable
RS	Recommender System
STS Theory	Socio-Technical-Systems Theory
XAI	Explainable Artificial Intelligence

Introduction

Interest and advancements in artificial intelligence (AI) have fluctuated since its “birth” in the mid-20th century (Haenlein & Kaplan, 2019). This fluctuation can be explained by diverging classifications and opinions of what makes a system or machine intelligent and at what point it reaches this intelligence (Haenlein & Kaplan, 2019; Kaplan et al., 2023). Despite inconsistent attention, significant technological strides were made, most prominently from the 2010s onwards, resulting in a global AI race among companies and an increase in public awareness of both AI’s opportunities and risks (Kiela et al., 2021).

In spite of its long history, there is no unified definition of what exactly AI entails (Samoili et al., 2020). However, combining various common understandings, the following working definition can be derived. AI refers to technical systems designed to act similarly to humans. Rather than simply following pre-defined rules through code, these systems are able to perceive their environment, interpret vast amounts of data, reason based on their derived knowledge, and decide which actions to take to best achieve a specific goal (Samoili et al., 2020; Haenlein & Kaplan, 2019; Lukyanenko et al., 2022). Moreover, AI systems are able to adapt their behavior by analyzing actions and outcomes in a process of “learning”. Therefore, AI emulates human cognitive abilities and decision-making processes to complete various tasks. While applications of AI range, it is frequently deployed as technology that “increase[s] human capability” (Kelly et al., 2023, p. 1) by saving time, increasing efficiency, and facilitating decision-making (Bach et al., 2024; Kelly et al., 2023; Solomonides et al., 2022).

Notably, AI was originally understood as a computational mirror of the human brain and behavior, rooted in cognitive science (Dubova et al., 2022; Gweon et al., 2023). Yet, recent advancements in AI have focused very strongly on technical components, giving less attention to human-centered theory and benchmarks (Atas et al., 2021). Because AI ultimately seeks to model humans, an approach incorporating cognitive, social, and ethical considerations is integral to building systems that are not only efficient but also aligned with human values (Atas et al., 2021; Tran et al., 2021). Embedding cognitive science and psychological theory in AI is therefore essential for trustworthy, human-centered systems – a gap addressed by this thesis by linking core psychological principles to emerging technologies.

One of the most widely distributed applications of AI is recommender systems (RS). Particularly in the digital world, the use of RSs is pervasive as they are utilised in diverse domains such as music and e-commerce (Jannach & Jugovac, 2019). A variety of technical metrics, such as accuracy, precision, and latency, are used to evaluate and improve the quality of these systems (Atas et al., 2021; Tran et al., 2021). More recently, “beyond-accuracy” indicators that capture novelty, diversity, and other psychological phenomena are also integrated in interdisciplinary research to assess and improve the value and performance of these systems (Jannach, 2023; Jannach & Jugovac, 2019; Tran et al., 2021). Comprehensive evaluation and development of RSs, therefore, requires a multidisciplinary lens that combines algorithmic precision with psychological insight.

Because every interaction is inherently reciprocal – a mutual give-and-take (American Psychological Association, 2025) – there are promising opportunities for developing more interactive RSs based on psychological research (Gao et al., 2021; Pramod & Bafna, 2022). Early studies show that interactive RSs that facilitate information exchange, meet users’ needs, and empower them to make informed decisions by “establish[ing] rapport and trust” (Jannach, 2023, p. 20) can improve both objective recommendation quality and subjective user attitudes such as satisfaction and trust (Christakopoulou et al., 2016, 2018; Jannach, 2023; Tu et al., 2024). This paradigm shift towards reciprocal, psychology-informed recommendation forms the conceptual foundation of the present study, which addresses how such interactions influence users’ trust.

As in any interaction, whether human-to-human or human-to-machine, numerous psychological processes are involved in guiding these interactions and their outcomes. Among these is trust (Tu et al., 2024). Trust is a central component in relationships and interactions as it reduces perceived risk, fosters smoother exchanges, and boosts behavioral intentions (Lukyanenko et al., 2022; Silva et al., 2023). Trust is multi-faceted and – in AI contexts – is shaped by cues such as explainability – explanations of system operations and outputs (Hoffman, Miller, et al., 2023; Hoffman, Mueller, et al., 2023; Pai et al., 2022), sense of control – users’ level of agency in AI interactions (Lockey et al., 2021; Vantrepotte et al., 2022), and anthropomorphism, the human-likeness of a system (Lee & See, 2004; Sidlauskiene et al., 2023). Because these factors can influence how users interact with AI and how they interpret its motives and output, they can act as antecedents to the formation of trust (Lee & Turban, 2001). Because trust is critical for engaging with another party, the formation and maintenance of trust in AI are ultimately essential for system adoption and reliance (Bach

et al., 2024; Lockey et al., 2021). Understanding this multi-factorial nature of trust in AI is therefore crucial for building systems that people are willing to use. Therefore, AI research is increasingly concerned with studying trust and how to increase the trustworthiness of AI (Bach et al., 2024; European Commission, 2019).

Despite growing interest in RSs and AI trust, little is known about how interactive RS, such as conversational recommender systems (CRS), systematically affect user trust (Jannach, 2023; Pramod & Bafna, 2022). Moreover, the interplay of system design elements, such as explainability, user control, and anthropomorphism, and their role as antecedents of trust, is still underexplored within CRSs. Because RSs and CRSs guide purchase decisions in e-commerce, clarifying how these design elements interact and impact buyers is critical for aligning system actions with user needs and commercial goals. This thesis addresses this gap by investigating how CRSs, and the elements explainability, control, and anthropomorphism, influence trust. Findings will inform the development of more user-centered RSs, ultimately fostering greater adoption and engagement.

Overview of Recommender Systems

RSs play an essential role in today's digital experiences (Wang & Benbasat, 2008). RSs collect, analyze, and integrate data about users' preferences, interests, and previous online behaviors to filter and customize the information presented to them, thereby facilitating navigation and decision-making in digital spaces (Alyari & Navimipour, 2018; Bobadilla et al., 2013; Pramod & Bafna, 2022). RSs are employed in different domains such as news, travel, music, and – most visibly – in e-commerce (Lin et al., 2019; Wang & Benbasat, 2008). In e-commerce, RSs narrow the item pool to a manageable shortlist, thereby influencing what customers ultimately purchase.

Early RS research primarily focused on accuracy: predicting the product, article, or song with the highest relevance to the user (Atas et al., 2021; Jesse & Jannach, 2021). Indeed, accuracy boosts click-through rates and conversion. Yet, an accuracy-first approach overlooks key psychological variables that influence how and if users adopt the system and its decisions (Atas et al., 2021; Tran et al., 2021). Increasing work, therefore, calls for metrics beyond accuracy to also capture factors underlying the user experience, thereby improving the overall efficiency and success of RSs.

From Filters to Conversation: Conversational Recommender Systems

With the increasing shift towards AI-driven interfaces, businesses are leveraging AI-enhanced recommendation systems to enhance user experiences and optimize the recommendation process (Jannach, 2023; Sidlauskiene et al., 2023). One of the most promising developments in this space includes the deployment of CRSs. CRSs are systems designed to provide personalized recommendations to users based on an interactive dialog (Chen et al., 2024; Christakopoulou et al., 2016; Jannach et al., 2021). Rather than solely operating based on more “static” or one-sided filtering techniques, CRSs combine a dialog module with the underlying recommendation engine to engage users in multi-turn conversations. In these, users can clarify preferences and goals, refine filters, and receive information and decision support from the conversational agent (Jannach, 2023; Sidlauskiene et al., 2023). This enables more dynamic elicitation and refinement of user preferences and more precise and personalized recommendations (Christakopoulou et al., 2018; Gao et al., 2021; Jannach et al., 2021). As such, CRSs – often classified as conversational (virtual) or chatbots – are particularly prominent in e-commerce, where they are designed to facilitate interactions with users and thereby generate value for both users and system deployers (Tan & Liew, 2020).

Evidence of the positive effects of CRSs is encouraging for practice. Various studies have assessed user reactions and satisfaction when interacting with conversational agents (chatbots) (Pizzi et al., 2021; Sidlauskiene et al., 2023; Siro et al., 2024; Tan & Liew, 2020). In e-commerce, for instance, CRSs have been linked to stronger purchase intent and increased user trust (Singh et al., 2024; Tan & Liew, 2020). Industry trends mirror this research, as retailers increasingly embed chatbots in their webshops, positioning them as digital shopping assistants designed to enhance engagement, provide customer support, and ultimately increase sales conversions. In conclusion, these converging insights suggest that CRSs are increasingly becoming supportive and trust-building engines that may provide a competitive advantage in modern e-commerce.

Key Design Factors Shaping Trust in Conversational Recommender Systems

The success of RSs in shaping user decisions not only depends on state-of-the-art algorithms but also on the psychological effects they have on their users. One of the core elements of this is trust. Previous research shows that users reporting higher trust in a RS are more likely to use the system and follow its recommendations (Ebrahimi et al., 2022;

Hoffman, Mueller, et al., 2023). Various design elements that can influence user trust recur across the RSs and CRSs literature (Lin et al., 2019; Saßmannshausen et al., 2023). Among these, the following three stand out as both theoretically grounded and practically relevant.

First, explainability provides users with an understanding of the system's reasoning and operations (Ferrario & Loi, 2022; Liu & Wang, 2025; Lockey et al., 2021). This reduces the black box phenomenon, which may lead users to mistrust system decisions. In non-conversational e-commerce RSs, users typically see recommendations and explanations such as "People like you also bought X". These can be described as static, one-shot rationales of how the RS generated its recommendations (Hernandez-Bocanegra & Ziegler, 2023; Weitz et al., 2019). In contrast, CRSs can generate interactive rationales through the chatbot disclosing its logic, asking follow-up questions, and refining recommendations based on the user-chatbot dialog (Christakopoulou et al., 2018; Jannach, 2023; Zhu et al., 2024). This results in two complementary layers of explainability: direct, or local explainability (e.g., "Based on your indication that two people live in your household, a washing machine with a load of 6L matches your needs."), and indirect, or global explainability, which arises throughout the conversation as users infer patterns and build mental models of the system's operations and architecture (Ferrario & Loi, 2022; Gao et al., 2021). Because both layers can be present in CRSs, CRSs have the potential to increase explainability and, thereby, trust (Hernandez-Bocanegra & Ziegler, 2023).

Second, while reducing choice overload, receiving only a static list of recommendations can feel prescriptive, making users contingent on what the system decides without being able to actively influence this outcome. This can lead to reactance or a sense that their freedom and autonomy are withdrawn (Liu & Wang, 2025; Pizzi et al., 2021). CRSs can mitigate this tension by allowing users to steer the recommendation process, or at least being active contributors. Because users can accept, reject, or edit suggestions, and clearly communicate their needs and preferences, the interactive, conversational approach in these more advanced systems places the user in a more controllable position. This also reframes the system from an "authority" to a "collaborative assistant", enhancing user autonomy and fostering the sense that recommendation outcomes are co-created (Pizzi et al., 2021). This heightened sense of control has been linked to stronger user trust and greater willingness to follow the CRSs' suggestions (Christakopoulou et al., 2016, 2018; Jannach et al., 2021).

Third, including a conversational interface adds specific human-like cues to an AI system. These include natural language, turn-taking, greetings, human names, avatars, and even empathy that boost social presence (Araujo, 2018; Qiu & Benbasat, 2009). Even when users know that their interaction partner is a machine – a chatbot – these cues tap into the human schema. This refers to the mental model people use to navigate social interactions: people interpret the chatbot as a social actor and therefore converse with it as they would with a human, and even attribute intentions, values, or empathy to it (Morana et al., 2020; Singh et al., 2024). Higher perceived anthropomorphism has been shown to increase engagement, trust, and acceptance of CRS recommendations, even in retail settings, where shoppers generally still prefer human advice (Pizzi et al., 2021; Singh et al., 2024; Tan & Liew, 2020). It therefore appears to be essential to incorporate anthropomorphic design to support the establishment of rapport and trust with the user.

To summarize, CRSs transform more classical approaches of user-profile-based filtering into a dialog-based process that can elicit richer information on preferences and needs, tailor recommendations interactively, and serve users like virtual assistants. Their value, especially when it comes to fostering trust, may, however, depend on how explainable the system is or appears, how much control the user retains or senses, and how human-like it is or is perceived. Having laid the technical foundation and definition of CRSs relevant to the present study, the subsequent section considers the psychological fundamentals of trust and the presented antecedents.

Psychological Foundations of Trust

Trust has been studied in various contexts, both in relation to technology and outside of it. Particularly, the study of trust in AI, automation, and other technological systems has increased in recent years, as there has been a growing awareness of its significance for system adoption and reliance, and, as an extension, for system development and deployment. Trust is particularly critical in AI-enhanced systems as it can be described as the psychological foundation by which users decide, both actively and passively if they will accept a system's decisions or follow its recommendations (Kaplan et al., 2023; Lockey et al., 2021). This section, therefore, clarifies what trust means in human-AI interactions, and discusses more deeply why the previously mentioned design levers – explainability, sense of control, and anthropomorphism – can shape trust, particularly in the context of CRSs. This lays the groundwork for the research question and hypotheses of this thesis.

Importantly, the use of AI in such systems not only adds complexity not only to the computing processes and functionalities, but also to the establishment and understanding of trust in these systems. Unlike traditional technology, AI systems operate with greater autonomy and unpredictability, often making decisions that are difficult to explain or interpret (Kaplan et al., 2023). Considering this, AI entails both opportunities and risks that are different from other technologies and therefore poses unique challenges in fostering user trust, requiring specifically dedicated research and development (Bach et al., 2024; Kaplan et al., 2023).

The advent of trustworthy AI, and how user trust can be promoted and sustained in human-AI interactions can be seen as part of this proposition. With increasingly complex systems, therefore, also comes an increasingly complex study of trust and its influencing factors, which has been identified as a research gap by Bach and colleagues (2024), particularly in collaborative AI environments, in which the user-AI relationship is central. This chapter explores key concepts and components of trust, with a focus on human-AI interactions. It also examines the psychological factors influencing trust, laying the foundation for the trust model tested in this study.

Defining Trust in AI Contexts

Trust has been conceptualized diversely, often depending on the discipline and application context. Yet, there is still no existing unified definition of the construct (Bach et al., 2024; Körber, 2018; Lockey et al., 2021). Nevertheless, in the context of AI and technology, some conceptualizations have persisted more than others. Among these is Mayer et al.'s (1995) model of organizational trust, which views trust as the willingness to “be vulnerable to the actions of another party based on the expectation that the other will perform a particular action important to the trustor” (p. 712), even when the trustor has no influence or control over the trustee. Being the most cited and accepted trust model, it has not only been applied in research on organizational development and management, but also in automation and AI (Körber, 2018; Lee & See, 2004). Although initially developed within research on organizational behavior, its applicability in human-technology interaction contexts is warranted by the following conditions (Mayer et al., 1995; McKnight et al., 2002).

1. Dyadic interaction – Trust, or its counterpart mistrust, arises in interactions between two parties: the trustor (party extending trust) and the trustee (party receiving trust). Moreover, these interactions tend to be characterized by interdependence – mutual

dependence on one another (Mayer et al., 1995). When people interact with technology, this can also be categorized as a dyadic interaction in which two entities are mutually dependent on their input and output. Hence, the extension of trust by a user is required for successful exchanges with technology.

2. Risk – Trust arises in situations of risk and uncertainty, requiring users to rely on the other party without having complete control over their actions. In technological, or AI-driven systems, this risk is heightened due to system opacity, unpredictability, and system autonomy (Kaplan et al., 2023). Therefore, trusting the system despite these inherent risks is critical for interacting with it.

Synthesizing these components, trust can also be defined as "the attitude that an agent will help achieve an individual's goals in a situation characterized by uncertainty and vulnerability" (Lee & See, 2004, p. 54). Moreover, psychologically, trust has been conceptualized as an affect, attitude, state, belief, and behavior (Lee & See, 2004; Lockey et al., 2021). Importantly, however, rather than reflecting conflicting perspectives, they can collectively be summarized in a way that defines trust as both a psychological state and a behavioral outcome. Integrating the multidimensional nature and strands of trust, this study accepts trust as a psychological state in which users accept vulnerability to a RS because they expect it to support their goals despite the presence of uncertainty and having limited control over the system. In other words, users interacting with RSs generally do not have direct control over the output (recommendations) or system operations, do not know what this output will be (uncertainty), are thereby vulnerable to the actions of the system, but nevertheless expect that the system's decisions will be beneficial for them.

Comparing Human and AI Trust

Although interpersonal and technology trust share similarities, such as the aspects of vulnerability, positive expectations, risk, uncertainty, and limited control, their foundations differ (Hoff & Bashir, 2015). This refers primarily to specific attributes of the trustee or sources through which trust is established (Hoff & Bashir, 2015). While interpersonal trust often depends on ability, goodness of character, benevolence, and integrity, technology trust is contingent on performance, process, and purpose (Hoff & Bashir, 2015; Lee & See, 2004). Looking more closely at technology trust, trust based on performance depends on how reliable, competent, and functional the technology and its performance (outcome) are. Second, process-based trust refers to a system's operations, and how appropriate they are for

the situation and the user's goals. Third, trust is influenced by the extent to which a system's design is positively oriented towards the operator, that is, the "intentions" of the system to provide value for the user (Lee & See, 2004). Recognizing the distinction between interpersonal and technology trust is essential for modelling how specific technologies and system features can enhance or reduce user trust.

Beyond differences in trust sources and foundations, the sequence in which trust develops frequently also differs for humans and technology, although this is also subject to interpersonal variances (Hoff & Bashir, 2015). In human interactions and relationships, trust tends to start cautiously – it initially depends on how predictable the other's behavior is, with little risk-taking on the trustor's part. Interaction partners "test" the predictability and reliability of one another's behavior, and only after repeated evidence of trustworthiness, trust is increasingly extended and maintained (Hoff & Bashir, 2015). However, this development of trust also differs depending on dispositional, motivational, and situational factors, where some personality traits, for instance, have been linked to higher general trusting tendencies (Weiss et al., 2022). In contrast, users often show a positive bias towards technology (Hoff & Bashir, 2015). That is, they often assume near-flawless performance, leading them to grant technology a level of faith that is not always granted when first interacting with another person. This early faith only decreases when errors or system failures occur. Thus, when system performance and accuracy are consistent, users tend to continuously exhibit trust (Hoff & Bashir, 2015). Of course, this pattern is also variable, often shaped by individual psychological traits, such as propensity to trust or technology affinity, the context, and specific design cues and functionalities presented by a system (Hoff & Bashir, 2015; Küper & Krämer, 2024). Thus, while contrasts can be identified in the trust development process, ultimately, subjective perception and evaluation rather than objective metrics alone often determine whether trust is built and maintained.

A Socio-Technical Perspective on AI Trust

Trust in technology and AI is dependent on the harmony between the user and the system. As such, trust in AI is co-produced by the technology and the user (Körber, 2018; Lin et al., 2019). Socio-technical systems (STS) theory provides a framework for understanding how both technological components (performance, design) and human factors (attitudes, perceptions, and experiences) contribute to trust formation (Bostrom & Heinen, 1977; Lin et al., 2019). It argues that algorithmic accuracy and reliable performance alone are insufficient.

Instead, how users perceive and interact with the systems is a key component to ensure trust and all its positive consequences are formed. Hence, designing and deploying AI systems should not simply be viewed as a technological undertaking, but also as a social and user-oriented one (Andras et al., 2018; Bostrom & Heinen, 1977).

In CRSs, this duality is particularly explicit as on the technical level, algorithmic logic and accuracy are essential to deliver high-quality recommendations, while on the social level, how the user perceives and interprets the system and its recommendations can strongly shape trust (Pramod & Bafna, 2022; Siro et al., 2024). Designing for trust in CRSs, therefore, means developing both layers in a way that they reinforce each other. For instance, relevant recommendations may lose value if the user distrusts the system or has a negative experience due to the interface design. STS theory, therefore, also serves as a reminder that trust is dynamic: to establish and maintain trust, there must be a balance between technical components and social responses. Accordingly, the present study examines how different psychological mediators in (C)RS may influence the socio-technical fit and translate it into trust.

Heuristic vs. Deliberative Trust Processes

Before moving to the discussion of the three key mediating variables under analysis, it is important to understand how humans process information and how this processing may subsequently lead to trust. Justifications – “good reasons” (Lockey et al., 2021, p. 5464) – play a critical role in determining whether users extend or withhold trust in a system (Lockey et al., 2021). These justifications can be recognized either consciously or subconsciously. As established previously, justifications can arise through system-based cues and system-directed assessments, such as reliability (technical features), social presence (design features), data-protection safeguards (socio-ethical considerations), or through user traits such as previous experience, personality, and baseline propensity to trust (Bach et al., 2024; Bao et al., 2021; Hoff & Bashir, 2015; Lockey et al., 2021). Notably, when users sense or have these “good reasons”, they will be more likely to trust and follow the system.

Cognitive science shows that judgments are formed via two tracks: automatic, instantaneous heuristics and slower, analytic reasoning (Hoff & Bashir, 2015; Jesse & Jannach, 2021; Tversky & Kahneman, 1974). While the former are often used when cognitive resources are scarce or few resources are expended, the latter often occur in higher-stakes situations. Moreover, while heuristics can be triggered by subtle interface or output cues,

more thorough analysis requires effortful and conscious considerations (Tversky & Kahneman, 1974). This collaboration and reciprocity facilitate information processing and decision-making, particularly in situations characterized by risk and uncertainty (Tversky & Kahneman, 1974). Since technology, and particularly interactions with (C)RSs in e-commerce, sit at the intersection of impulse and well-considered purchases, their design should satisfy both modes of thinking to enable instant trust and assurance in the system, while also being able to withstand reflective analysis (Hoff & Bashir, 2015). Perceived explainability, sense of control, and anthropomorphism can be defined as sources for such judgment, making it critical to thoroughly understand their underlying processes and effects. Hence, these factors and how they relate to trust are examined more closely in the subsequent sections.

Psychological Antecedents of Trust

As previously established, trust in AI arises dynamically and is typically shaped by various factors. Some of these factors can provide understandability, agency, or signal social presence, thereby enhancing users' trust in a system. In particular, explainability, sense of control, and anthropomorphism can serve as critical antecedents that foster trust in various AI technologies, including CRSs. Considering this, the psychological foundations of these antecedents are discussed in the following subsections.

Explainability in AI

AI is often regarded as a “black box” (Ferrario & Loi, 2022, p. 1459), meaning its operations are so complex that they cannot be fully understood or explained. This makes it difficult even for system engineers to thoroughly trace the steps from input to output (Ferrario & Loi, 2022; Lockey et al., 2021). This opacity is troublesome, particularly considering that transparency and explainability of AI operations and decisions is widely regarded as a central component of user trust (Hoffman, Mueller, et al., 2023; Lockey et al., 2021). Previous research shows that when users understand how a system works or why it produced a certain output, uncertainty reduces and trust increases, which in turn also enhances system adoption (Lockey et al., 2021). Explainability also appears to play a significant role within CRSs as users must decide, sometimes in seconds, whether to follow the system's recommendations (Govea et al., 2024; Weitz et al., 2019; Zhu et al., 2024). This study, therefore, considers perceived explainability as a key psychological factor in the relationship between CRS interactions and user trust.

Research in explainable AI (XAI) has grown rapidly in recent years (Hoffman, Mueller, et al., 2023; Weitz et al., 2019) with the aim to develop methods that enable both system designers and developers, and users to better understand AI systems, their operations, capabilities, and risks (Hoffman, Miller, et al., 2023). Good comprehension of AI is not only critical for improving the performance of these systems, but also has important implications for users and society at large. Particularly when AI is used in complex situations or critical domains such as healthcare, law enforcement, and the judiciary, understanding AI decisions and thereby understanding how to use these tools appropriately, is essential for successful human-AI collaboration (Hoffman, Miller, et al., 2023).

The increasing demand for explainability in AI is not only reflected in academic research and system development, but it has also gained traction at the policy level. Various regulatory frameworks developed across the globe, such as the European Union's (EU) General Data Protection Regulation (GDPR) or the EU AI Act, reinforce users' rights to explainability in a variety of contexts, such as AI-driven decision-making or profiling (Lobner et al., 2021; Pai et al., 2022; Pavlidis, 2024). Among other requirements, these regulations mandate that explanations are provided in a way that is understandable to users, ensuring fairness, accessibility, and transparent processing of AI logic and outcomes (Lobner et al., 2021). This highlights the growing importance of making AI understandable to non-experts, ensuring that users have the ability to evaluate the processes and results delivered by AI and to determine their trustworthiness.

XAI research offers multiple layers and methods for explainability. Designing and developing inherently interpretable models that include elements of explainability, such as decision trees, rule-based systems, or linear models that illustrate the AI's logic directly, is one method to address explainability (Ferrario & Loi, 2022). However, these representations tend to be highly mathematical, presenting explanations that are incomprehensible to the average user (Hoffman, Miller, et al., 2023; Weitz et al., 2019). Moreover, including extensive sets of variables and intricate decision pathways leads to increasing complexity, reducing the interpretability even more (Ferrario & Loi, 2022). To solve this issue and reduce this explainability gap, more recent XAI techniques approximate explanations of opaque models with simpler means, making understanding of AI-driven actions and outcomes more accessible to users (Ferrario & Loi, 2022; Pai et al., 2022). It is important, however, to acknowledge that this simplification comes at the expense of completeness – a tradeoff users may nevertheless be willing to accept if it improves comprehension of the AI (Ferrario & Loi,

2022; Zhu et al., 2024). Although scholars and engineers must understand the intricate inner workings of AI as best as possible, this logic must ultimately also be translated into clear, user-centered language to enable non-experts to also grasp and confidently engage with the system.

As established in the previous chapter, within the methods of explainability, there is a distinction between global (architectural) explainability, which may explain a system's overall structure, training data, and *inner workings*, and local (procedural) explainability, which aims to explain specific decisions or outcomes (Ferrario & Loi, 2022; Hoffman, Miller, et al., 2023; Hoffman, Mueller, et al., 2023). While theoretically distinct, the two forms often overlap in practice, as global insights can support local decision transparency, and vice versa (Hoffman, Miller, et al., 2023; Saßmannshausen et al., 2023). Importantly, both have been linked to greater trust and user satisfaction, which is imperative for real-world applications such as within the context of CRSs (Govea et al., 2024; Hoffman, Miller, et al., 2023; Mueller et al., 2020).

Psychologically, the importance of explainability in AI and its relationship with trust is manifold; however, the *need to understand* may be one of the most significant underpinnings. In general, people show a desire to understand the world around them and the forces that influence certain outcomes (Hoffman, Miller, et al., 2023; Tversky & Kahneman, 1974). This intensifies in high-stakes decisions or whenever they lack prior experience (Ferrario & Loi, 2022; Lockey et al., 2021). Considering this, users may approach technology and AI with an intrinsic need to understand how it works and what its purpose is. Explanations can satisfy this need by providing information, thereby improving the user's knowledge and understanding of a system (Wang & Benbasat, 2008).

The knowledge gap can be closed in two ways. First, users may learn and understand the system gradually through hands-on experience, repeated use, and interactions. Second, they may learn and understand instantly through direct explanations provided by the system (Wang & Benbasat, 2008). When people interact with specific technologies or AI for the first time, the latter route dominates if explicit explanations are given. By converting opaque computations into clear, understandable explanations, this supports structured reasoning processes users may need to validate a system and to have justifications for following and using it (Zhu et al., 2024).

Although explicit explanations may provide a clearer path to explainability, a more gradual, interactive explanation process can be identified in many CRSs. Here, knowledge sharing is not a static event, but rather a dynamic exchange (Hernandez-Bocanegra & Ziegler, 2023; Wang & Benbasat, 2008). First, when required, users can ask the system directly for explanations. Second, with every conversational turn, the user refines their mental model of the system's logic, resulting in an iterative, procedural process that deepens understanding (Christakopoulou et al., 2018; Wang & Benbasat, 2008). Importantly, the richer this dialog, the more opportunities the user has to learn, providing a stronger foundation for trust.

In both methods – direct and procedural explainability – added knowledge can reduce algorithmic opacity, reduce uncertainty, and increase perceptions of reliability, thereby enhancing trust (Bach et al., 2024). Effective XAI should therefore support both pathways to translate complex model logic into insights that equip users to understand and validate the systems.

User Control in AI Interactions

The level of interactivity between users and systems also appears to play a significant role in facilitating trust (Bach et al., 2024; Küper & Krämer, 2024; Wang & Benbasat, 2008). However, this may not only be the case because interactivity can enhance understanding in terms of explainability, but also because it can enhance users' sense of control over the interaction and system output (Ferrario & Loi, 2022; Wang & Benbasat, 2008). Previous research on trust in technology and AI has shown that trust is influenced by the extent to which users perceive themselves as having control over the system, its decisions and outcomes, and the interaction process (Küper & Krämer, 2024). Within RSs, this research has consistently shown that users who feel in control exhibit higher levels of trust towards that system and are therefore more likely to follow its recommendations (Hoff & Bashir, 2015; Lukyanenko et al., 2022).

Users' sense of control appears to be shaped by a system's level of interactivity with the user. This aligns with theories in human-computer interaction. This research treats interactivity as a multidimensional construct that incorporates perceived control – the ability to influence system behavior –, among other components, in its definition (Lee et al., 2015). High interactivity enables users to manipulate a system's operations and, as a result, its outputs. This seems to foster autonomy, control, and trust (Wang & Benbasat, 2008). The

amount of user involvement, and resultingly, control in the decision-making loop, can therefore positively influence perceived control and trust in a system (Hoff & Bashir, 2015).

Perceived control can reduce uncertainty, thereby playing a key role in early trust formation (Silva et al., 2023). This is because sense of control can make the system appear more predictable as it is not only the system's operations, but also the user's input and actions that lead to certain outcomes. This predictability can consequently decrease user uncertainty because there is a reduced risk of being "misguided" by the system when users also have direct influence (Silva et al., 2023). As predictability rises, user anxiety and perceived risk decrease, fostering confidence not only in the system but also in the act of using it (Küper & Krämer, 2024). In this way, sense of control can reduce uncertainty and increase trust.

Sense of control refers to an individual's perception that their own actions determine specific outcomes (Tapal et al., 2017; Van Elk et al., 2015). It therefore relates to the perception or belief of how much someone feels responsible for an outcome. In psychological research, sense of control is closely related to the concept of agency (Tapal et al., 2017). Although the two are frequently used interchangeably, attempts have also been made to distinguish them (Tapal et al., 2017; Van Elk et al., 2015). While sense of control refers primarily to the outcome of one's actions, sense of agency refers mostly to "authorship" of one's actions – how much a person feels in control of their own actions and decisions. While sense of agency emphasizes the behavioral process, sense of control extends beyond this, encompassing control over both the process and its consequences (Tapal et al., 2017; Van Elk et al., 2015).

Nevertheless, a clear distinction between agency and control has not been established. Indeed, some research finds these constructs practically indistinguishable (Bart et al., 2023; Vantrepotte et al., 2022). Moreover, agency has been defined as the control one experiences or feels about one's action and their effects in the environment, suggesting a continuum and close relationship rather than a sharp divide (Bart et al., 2023; Berberian et al., 2012; Vantrepotte et al., 2022). In the current study, sense of control is therefore understood as a unified construct that captures both the user's influence on the interactive process and its ultimate outcomes, combining elements of both common sense of control and sense of agency definitions.

The experience of control depends on a link between one's behavior and results in the external world (Vantrepotte et al., 2022). When a high degree of automation completes

operations, such as algorithmic filtering and recommending in RSs, it becomes challenging for users to clearly discern who performed which actions and the extent to which these actions respectively affected specific results (Berberian et al., 2012; Vantrepotte et al., 2022). Previous research has demonstrated that people's sense of control varies in grade. That is, different levels of control may be perceived and experienced by individuals (Berberian et al., 2012). In human-technology interactions, this level changes depending on the degree of automation and interactivity. As automation increases and interactivity or collaboration decreases, perceived control tends to decline (Berberian et al., 2012; Vantrepotte et al., 2022). Building on these findings, the present research investigates how varying levels of interactivity and "control-giving" in AI-based systems shape users' sense of control, and, consequently, their trust in the system.

Although users often actively and consciously operate specific AI systems, they do not necessarily process this control directly (Liu & Wang, 2025). Rather, much of control's trust-enhancing effect appears to work intuitively, outside of deliberate consideration or awareness (Liu & Wang, 2025). Hence, sense of control can also be classified as a heuristic signal, automatically cueing trust formation.

Given the strong association between sense of control and trust, it is crucial for developers and designers to consider ways to enhance perceived user control. In particular, well-designed interfaces that enable and promote customization, seek and provide feedback, and allow interactivity can contribute to a greater sense of control (Lee et al., 2015). Similar to the right to explainable and transparent AI introduced previously, the EU's Ethics Guidelines for Trustworthy AI also classify respect for human autonomy as a foundational principle, reinforcing the necessity of enabling and empowering human agency and control in interactions with AI (European Commission, 2019; Lukyanenko et al., 2022). By prioritizing user control in technology design, the groundwork for trust can therefore be laid.

Anthropomorphism in AI

Interactions and dialog are complex reciprocal actions that establish trust if successful (Tu et al., 2024). As previously determined, trust not only motivates users to engage and reengage with others, but also reduces uncertainty and anxiety. This is even more important when relying on information provided or delegating decisions to another party (Lukyanenko et al., 2022; Tu et al., 2024). The previous sections introduce explainability and sense of control as levers that can influence trust formation, particularly in human-AI contexts. The

two can be described as cognitive and autonomy cues, respectively. Additionally, social signals embedded in an AI's design can also have a substantial impact (Tu et al., 2024). For instance, when anthropomorphic cues are present in system design, trust in these systems may be enhanced (Araujo, 2018; Tu et al., 2024).

Anthropomorphism can be defined as the attribution of human traits, intentions, emotions, or actions to non-human objects (Guthrie, 2025; Lockey et al., 2021; Qiu & Benbasat, 2009). Recent advances in AI – from generative AI, large language models (LLMs), to humanoid robotics – have intensified interest in how social, relational, and affective design elements may shape user attitudes and behaviors (Bach et al., 2024; Lockey et al., 2021; Qiu & Benbasat, 2009). For instance, it has been explored how anthropomorphic elements, ranging from a robot's animacy (Bartneck et al., 2009) to a chatbot's empathetic language, turn-taking, and politeness (Morana et al., 2020; Singh et al., 2024) can boost trust, satisfaction, identification with the technology, and system adoption intent (Lee & See, 2004; Qiu & Benbasat, 2009). This is because anthropomorphism adds an element of humanity to a system that gives it “character”.

More precisely, anthropomorphic cues appear to appeal to users' cognitive of humanity, creating a strong foundation for trust (Sidlauskiene et al., 2023). Humans naturally rely on schemas when navigating the world and making judgments about unfamiliar entities or situations. Schemas are mental frameworks – representations and clusters of knowledge – that enable efficient processing of experiences (Sidlauskiene et al., 2023). When interacting with technology exhibiting anthropomorphic features, the human schema is automatically activated, leading them to perceive the technology as a social entity. This process can reduce uncertainty about the system's intentions and capabilities, as users can draw on their existing knowledge of human-like behavior to interpret the AI's actions and decisions. Consequently, the activation of this familiar schema through anthropomorphic design can serve as a psychological pathway to trust formation by transforming a more opaque computational interaction into a comprehensible “social” exchange (Sidlauskiene et al., 2023).

In addition to the human schema, anthropomorphism also operates via other heuristic mechanisms facilitating quick, intuitive trust formation through the perception of social presence and likability (Bartneck et al., 2009; Morana et al., 2020; Sidlauskiene et al., 2023). Social presence, which can be defined as the perception of warmth, sociability, and human contact during an interaction, emerges when users subconsciously process anthropomorphic

features and classify them as a social actor (Lockey et al., 2021; Morana et al., 2020; Singh et al., 2024). This makes human-technology interactions more engaging, personal, familiar, and therefore more comfortable. Previous research has shown that when AI systems successfully convey social presence through human-like avatars, natural language, and conversational patterns, users experience greater emotional engagement and enjoyment (Chong et al., 2021; Etemad-Sajadi, 2016; Pizzi et al., 2021; Sidlauskiene et al., 2023). Thus, anthropomorphism can function as an antecedent to trust based on intuitive positive impressions of social presence.

The trust-enhancing effects of anthropomorphic elements have been shown to be particularly valuable in e-commerce, where research has repeatedly shown that consumers are more likely to trust, engage with, and return to platforms that successfully emulate human-like social interactions (McKnight et al., 2002; Pizzi et al., 2021; Tu et al., 2024). This stems from anthropomorphism's ability to create a sense of interpersonal rapport, making users feel as if they are receiving assistance from a person (a "shopping assistant") rather than an impersonal algorithm (Singh et al., 2024; Tu et al., 2024). Therefore, the strategic implementation of anthropomorphic design elements can be a powerful approach to bridging the gap between technological capability and human psychological needs in commercial AI applications.

These findings underscore the importance of incorporating anthropomorphic design features in technology to foster trust. The more users perceive a system to be anthropomorphic, the more users tend to trust the system, paralleling trust formation in human-to-human interactions (Bach et al., 2024; Lockey et al., 2021; Morana et al., 2020). However, this effectiveness depends on contextual appropriateness and user expectations, as excessive anthropomorphism can also trigger discomfort and reduce trust (Bartneck et al., 2009; Morana et al., 2020). This paradox highlights the delicate calibration required in anthropomorphic design to enhance trust without triggering negative reactions.

In summary, anthropomorphism emerges as a key psychological mechanism that fosters trust through the activation of human schemas and the cultivation of social presence. Importantly, there may be a synergy, or at least parallels, between explainability, sense of control, and anthropomorphism. While explainability primarily addresses cognitive needs for understanding, control satisfies needs for autonomy, and anthropomorphism meets social and emotional needs for connection and familiarity. Given the described effectiveness of these

design elements in fostering trust across different tools and contexts where trust and satisfaction also directly impact business outcomes (Jannach & Jugovac, 2019; Pizzi et al., 2021), there is a continuous need to systematically examine the effects of these variables. Combining these insights with the trend towards CRSs, the subsequent section presents the research question and empirical hypotheses, derived to investigate the effect of perceived explainability, sense of control, and perceived anthropomorphism on user trust in CRSs. The ensuing results shall provide practical insights for designing more trustworthy and user-centered recommendation systems.

Research Questions and Hypotheses

Combining the insights gained in the literature review, the present study aims at studying the effects of a CRS on user trust within the domain of e-commerce. More specifically, it aims at assessing whether the variables perceived explainability, sense of control, and perceived anthropomorphism have a mediating effect on user trust, depending on the type of recommender system users interact with – a more novel system classified as a CRS or a baseline recommender system (BRS). The primary research question addressed in this paper is therefore:

RQ1: How do conversational recommender systems influence users' trust in the system? To answer this research question, the following hypotheses were formulated.

Direct Effect of Conversational Recommender Systems on Trust

Drawing on socio-technical systems and trust formation theory, CRSs can alter the user-system relationship by transforming the recommendation process from a passive-user, one-dimensional instance into an active, collaborative process, which has been hypothesized and shown to increase trust (Singh et al., 2024; Tu et al., 2024). Other forms of RSs operate primarily through one-sided algorithmic filtering that presents users with static recommendation lists, positioning users as passive recipients of system-generated recommendations (Baizal et al., 2020; Jannach, 2023). In contrast, conversational systems establish active bidirectional communication that enables users to directly articulate preferences, clarify needs, and receive feedback, creating a more interactive experience that mirrors human-to-human interactions (Jannach, 2023; Siro et al., 2024). This shift from system prescription to collaborative assistance appears to address fundamental psychological needs that, as a result, may enhance trust (Baizal et al., 2020; Jannach, 2023; Küper & Krämer, 2024).

Therefore, based on insights from research on trust formation and CRSs, users engaging with CRSs are expected to show higher levels of trust than those interacting with static, baseline systems. Thus, the following hypothesis was derived.

H1: Interacting with a conversational recommender system is associated with significantly higher user trust than interacting with a baseline recommender system.

Mediating Roles of Explainability, Control & Anthropomorphism

The theoretical foundation for expecting higher trust in interactions with CRSs rests on several psychological principles that distinguish interactive from static system interactions. Therefore, while the direct relationship between CRSs and trust provides important insight into direct system effectiveness and success, understanding the psychological mechanisms through which potentially enhanced trust can emerge is also crucial for both theoretical advancement and practical system design. Building on the research insights from the previously introduced variables, this study argues that the trust-enhancing capacity of CRSs can be explained through the mediation of perceived explainability, sense of control, and perceived anthropomorphism. This mediation model, the paths of which are depicted in Figure 1, was derived to provide a more nuanced understanding of how system design elements that align with each mediating variable influence user trust.

First, the iterative, dialog-based nature of conversational can enable dynamic information and explanation exchange, allowing users to ask for clarifications, challenge recommendations, and understand the system through experience (Christakopoulou et al., 2018; Hernandez-Bocanegra & Ziegler, 2023). Enhanced explainability can play an important role in fostering trust by addressing users' need to understand system operations, goals, and decision-making logic, thereby providing justifications for accepting and acting upon system recommendations (Govea et al., 2024; Hoffman, Miller, et al., 2023; Hoffman, Mueller, et al., 2023; Pai et al., 2022). Moreover, the shift from static to interactive explanation represents a great advancement in XAI, as it transforms explainability from a predetermined system explanation into a collaborative sense-making process between system and user (Hernandez-Bocanegra & Ziegler, 2023).

Empirical evidence demonstrates that explainability can enhance users' understanding, confidence, and trust in RSs, with interactive approaches showing particularly strong relationships with positive user evaluations (Bao et al., 2021; Govea et al., 2024; Pu &

Chen, 2007). Based on this, the assumption is that the increased level of interactivity provided by a CRS increases explainability and, therefore, trust in the system. In other words, it is assumed that perceived explainability is a mediator of trust in a CRS, resulting in the following hypotheses.

H2a: Interacting with a CRS is associated with significantly higher perceived explainability than interacting with a BRS.

H2b: Higher perceived explainability is associated with significantly higher user trust.

H2c: Perceived explainability mediates the association between interacting with a CRS and user trust.

Second, the interactive format fundamentally shifts users from passive recipients to active collaborators in the recommendation process. This shift is expected to enhance users' sense of control over both the interaction process and the recommendation outcome (Jannach, 2023; Vantrepotte et al., 2022). This enhanced sense of control when interacting with complex systems can facilitate confidence, acceptance, and trust of these systems (Vantrepotte et al., 2022). As CRSs increase user engagement and situational involvement (Ebrahimi et al., 2022) through direct input, preference clarification, and iterative recommendation refinement, this participatory approach can address fundamental needs for autonomy and control, which have been identified as important concepts in trust formation (Berberian et al., 2012; Pizzi et al., 2021; Van Elk et al., 2015).

With the enhanced interactivity of CRSs, users may gain a deeper understanding of the causal relationship between their input and the system's output, which may enhance predictability and trust in the current as well as future interactions (Tapal et al., 2017). As such, this can provide a feeling of empowerment, or control, that baseline systems cannot provide due to their unidirectional information flow (Berberian et al., 2012; Vantrepotte et al., 2022). Therefore, sense of control is hypothesized to serve as a mediator between CRSs and trust. Based on this, the following hypotheses were derived.

H3a: Interacting with a CRS is associated with significantly greater sense of control than interacting with a BRS.

H3b: Higher sense of control is associated with significantly higher user trust.

H3c: Sense of control mediates the association between interacting with a CRS and user trust.

Third, conversational interfaces introduce natural language patterns, turn-taking dynamics, response feedback, empathy, and additional human-like features (names, avatar images, and so on), which can activate users' human schema and establish perceptions of human interactions. Thus, anthropomorphism can be a key factor in boosting trust in AI systems (Morana et al., 2020; Singh et al., 2024). This anthropomorphism may serve as a critical social-affective pathway to trust formation by transforming purely technical interactions into seemingly interpersonal exchanges that users can more easily navigate with the help of established social cognitive frameworks and models (Araujo, 2018; Sidlauskiene et al., 2023).

Unlike non-anthropomorphized systems, conversational systems incorporate human-like features that signal social presence, enhance rapport and likability, and lead users to attribute human characteristics to technical systems (Bartneck et al., 2009; Sidlauskiene et al., 2023). When users perceive systems as human-like, they may intuitively ascribe interpersonal trust or use interpersonal trust heuristics, leading to faster and more enhanced trust formed on the basis of social familiarity rather than solely based on technical competence (Araujo, 2018; Qiu & Benbasat, 2009). CRSs incorporate anthropomorphic cues, thereby possibly creating the perception of human-like advisors (Jannach, 2023; Jannach et al., 2021; Singh et al., 2024). Therefore, anthropomorphism is expected to be a key mediator through which CRSs increase trust, leading to the following hypotheses.

H4a: Interacting with a CRS is associated with significantly higher perceived anthropomorphism than interacting with a BRS.

H4b: Higher perceived anthropomorphism is associated with significantly higher user trust.

H4c: Perceived anthropomorphism mediates the association between interacting with a CRS and user trust.

Exploratory Questions

Beyond examining the psychological processes underlying trust formation in CRSs, this study also seeks to explore the behavioral implications of enhanced trust in CRSs. This is done by investigating if increased trust also leads to significant differences in recommendation acceptance behaviors. Trust has been shown to serve as an antecedent to decision-making and behavioral outcomes because users who experience greater trust in systems are more likely to follow these recommendations (Kelly et al., 2023; Tu et al., 2024). This pathway follows established models of technology acceptance, where attitudes towards

systems are associated with behavioral intentions and subsequent usage behavior (Bao et al., 2021; Küper & Krämer, 2024; Mayer et al., 1995). Given that CRSs are hypothesized to enhance trust, it is therefore also assumed that the positive relationship between CRSs and trust also leads to greater recommendation acceptance.

This behavioral exploration is categorized as an exploratory hypothesis rather than a formal one due to the complexity of factors influencing decision-making beyond trust, such as inter-individual differences and contextual constraints (Arnott & Gao, 2019; Rastogi et al., 2022; Wang & Benbasat, 2008). While trust can play a role in decision-making, alone, it may not be sufficient to override other variables. Nevertheless, investigating whether outcomes of trust in CRSs translate to behavioral outcomes provides valuable insight into further practical implications and the real-world effectiveness of CRSs. Therefore, the following research question and hypothesis will be examined.

RQ2: Does interacting with a CRS vs. a BRS have a positive effect on recommendation acceptance?

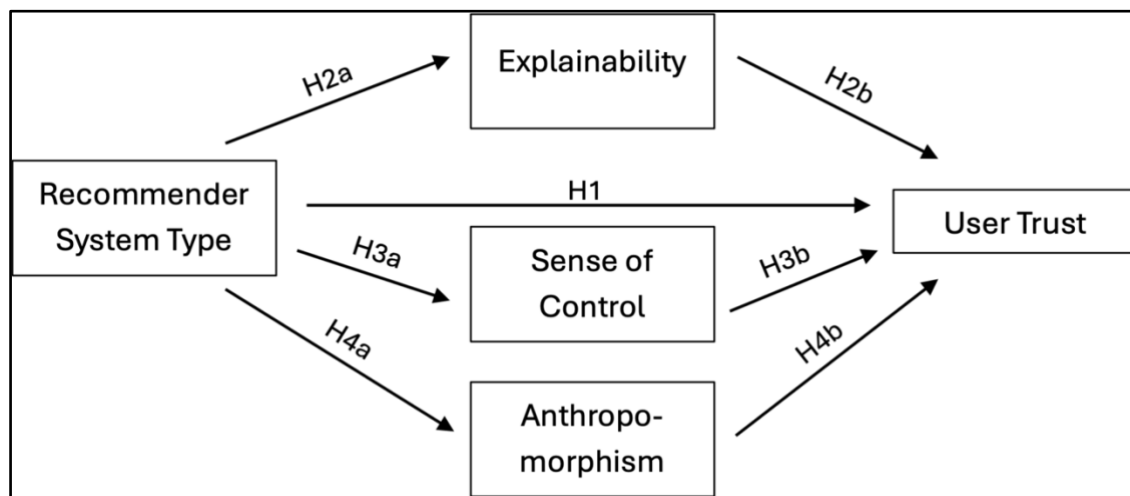


Figure 1. Illustration of the Conceptual Parallel Mediation Model

Method

To evaluate how CRSs influence user trust, and if perceived explainability, sense of control, and perceived anthropomorphism mediate these effects, an online experimental study was conducted. The following section describes the study's design and analyses.

Experimental Design and Procedure

A between-participants experimental design was implemented to test the research questions and hypotheses. The study was conducted as an online user study in Austria, using a simulated e-commerce interface modeled on a leading Austrian price-comparison site. One independent variable (IV) – the type of recommender system participants interacted with – was considered and varied between two groups, a non-conversational baseline recommender system and a conversational recommender system. Participants were randomly assigned to one of the two conditions to ensure that any observed differences in the outcome or mediating variables could be attributed more confidently to the mode of interaction rather than specific participant features.

Upon consenting to participate in the study, each participant was asked to read the standardized instructions to align focus and understanding of the scenario and the simulated RS environment. In the instructions, participants were tasked to imagine they needed to purchase a new washing machine – the only product category displayed by both systems. They were then asked to use the assigned e-commerce platform to find a suitable item. As can be seen in Figure 2, in the BRS condition, users refined results (recommendations) via manual filters, enabling them to set features such as load capacity, price, and energy efficiency according to their needs.

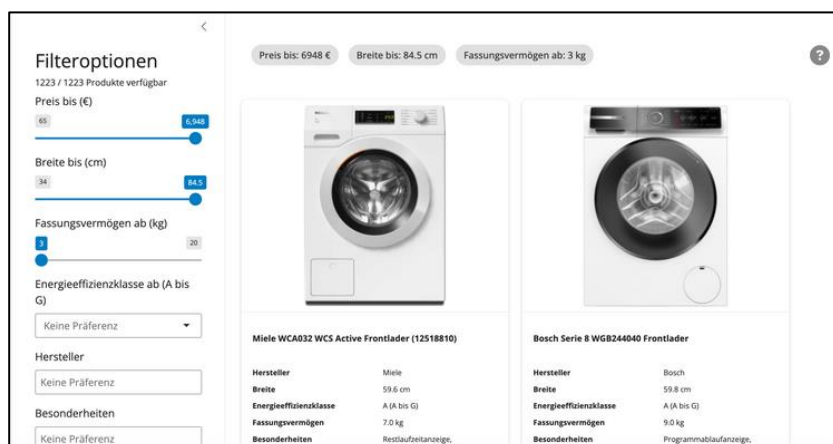


Figure 2. Illustration of the Baseline Recommender System Platform

In the CRS condition, depicted in Figure 3, participants communicated with a chatbot, which elicited their preferences through this dialog and set appropriate filters automatically. Both conditions shared identical layouts, product images and filter options, differing only in how recommendations were generated based on two different filtering techniques in order to isolate the effect of the RS type.

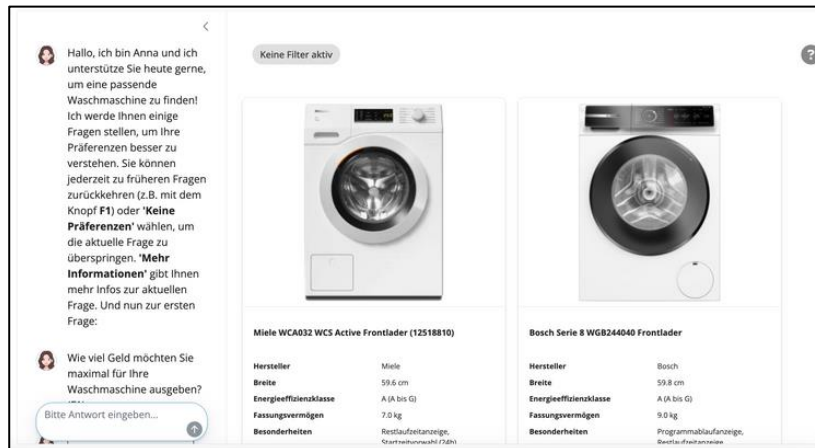


Figure 3. Illustration of the Conversational Recommender System Platform

After completing this interaction, ending either with or without a product selection depending on whether an appropriate item could be found, participants completed an exit questionnaire. This questionnaire captured perceptions of explainability, control, anthropomorphism, and trust. Additionally, demographic information (age, gender, occupation, education, German language skills), IT affinity, and online shopping preferences and behaviors were collected. An open-ended item at the end of the survey invited participants to share any additional comments or suggestions.

Power Analysis & Sample Size

A Monte Carlo power analysis for indirect effects was conducted to determine the required sample size for a parallel mediation model. The analysis was performed using a web application developed by Schoemann and colleagues (2017a, 2017b) using the “Set Power, Vary N” option, with the aim to achieve a power of .95 at $\alpha = .05$. Assuming a three-parallel-mediators model, the simulations carried out 10,000 replications with 20,000 draws per replication at a medium effect size ($r = .30$). Standard coefficients were used as the input method and the random seed was left at the default (1234), as proposed by the authors. The standard deviation was maintained at 1.00 for all input parameters.

The minimum and maximum sample size (n) was varied iteratively to refine power estimation and find the range that met the desired power. The final target sample size was obtained given a minimum n of 130, a maximum n of 150, and sample size steps set at 1. Given the parameters, the final analysis yielded a target sample size range of $n = 135$ to $n = 141$. To account for potential variability in data collection, the upper bound ($n = 141$) was defined as the final target sample size.

Participants and Recruitment

Participants were recruited through three channels between April 30 and May 22, 2025: snowball sampling through private networks and social media, a computer science course at an Austrian technical university, and Prolific (www.prolific.com) [22.05.2025]. Participation via snowballing was voluntary, with no financial compensation or other incentives provided. Participants who completed the study as part of their university course received bonus course credit. Twelve participants were recruited via Prolific. They received financial compensation (£1).

Eligibility criteria included age (<18 years) and German proficiency (B2+). Because the experiment was conducted as an online user study in a web-based application with the interface being limited to computer screens, participation required access to a desktop or laptop to ensure data quality. Mobile participants were prompted to switch devices.

To ensure data quality and meaningful interpretation of the results, data was included in the analysis only if the following criteria were met: For the baseline group (BRS), participants must have demonstrated appropriate interaction with the system by engaging with the filtering options (setting at least one filter). For the treatment group (CRS), participants must have interacted with the chatbot in a way that resulted in automated filtering (setting of at least one background filter). Additionally, participants must have completed the exit survey in full, providing responses to all items. Data that did not meet these criteria was excluded from the analysis. Those participants who completed the study in under 80s were also excluded.

Given these inclusion criteria, 77 participants were excluded from the study. That is, from 221 total responses, 77 did not meet the necessary criteria. This resulted in a final sample of 144 participants. While this sample size was larger than the calculated sample size, to avoid bias in the removal of three additional cases, the slightly larger sample was accepted, also considering that this would increase rather than decrease the study's power.

Recruitment efforts ensured that voluntary participation and informed consent were given. Thus, clear eligibility requirements, consent forms, and instructions were presented before any data collection to prevent usability issues.

Participant Characteristics and Demographics

Participants ranged in age from 21 to 65 years ($M = 28.78$, $SD = 8.12$). The sample was predominantly male ($n = 100$, 69.4%) with 44 female participants (30.6%). The majority of participants held bachelor's degrees ($n = 94$, 65.3%), followed by master's degrees ($n = 29$, 20.1%) and Matura (high school certificate) ($n = 14$, 9.7%). Four (2.8%) had completed apprenticeships, two (1.4%) held doctoral degrees, and one participant (0.7%) reported other educational qualifications. Most participants were students ($n = 108$, 75.0%). Employment status varied, with the largest group working part-time ($n = 64$, 44.4%), followed by full-time employment ($n = 29$, 20.1%), unemployment but seeking work ($n = 18$, 12.5%), self-employment ($n = 15$, 10.4%), unemployment and not seeking work ($n = 10$, 6.9%), and other ($n = 5$, 3.5%). Three participants (2.1%) did not specify employment status.

Measures & Data Collection

Data was collected through two complementary sources: behavioral logs (RS interactions) and self-reported measures (survey responses). Behavioral logs were primarily collected for the exploratory analysis. For this, participants' recommendation acceptance was captured by recording if they selected a washing machine they wished to "purchase". This was measured through click data at the end of each RS interaction.

To test the primary hypotheses, including the mediation model, the four key variables – trust, perceived explainability, sense of control, and perceived anthropomorphism – were measured in the exit survey, which was administered through LimeSurvey (<https://www.limesurvey.org/de>) [22.05.2025]. Items were measured on a 5-point Likert scale ranging from 1 ("stimme gar nicht zu" / "strongly disagree") to 5 ("stimmte voll zu" / "strongly agree") to capture user perceptions of each variable independently. The brevity of this part of the questionnaire was selected to balance depth with participant time constraints, ensuring high completion rates without compromising the precision of the psychological measures. This single-item approach has been warranted by other research as well to offer a less time-consuming and monotonous momentary assessment (Körber, 2018; Lee & See, 2004).

The dependent variable (DV) – trust – was measured with the item “Ich vertraue dem System” (“I trust the system”), extracted from the German Trust in Automation (German TiA) scale (Körber, 2018). The full scale has been validated against other established scales and was therefore used as the basis for the item used in the present study (Hoff & Bashir, 2015; Jian et al., 2000; Merritt et al., 2013).

The mediating variables were measured as follows. First, perceived explainability was measured using the item “Ich verstehe, wie das System funktioniert” (“I understand how the system works”), adapted and translated from the Explanation Satisfaction Scale (Hoffman, Mueller, et al., 2023). Sense of control was measured with the item “Ich habe Kontrolle darüber, wie das System handelt” (“I have control over the system’s actions”), which was adapted from the German Sense of Positive Agency Scale (SoPA), a subscale of the German Sense of Agency Scale (SoAS) (Bart et al., 2023), and the Sense of Agency Rating Scale (SOARS) (Polito et al., 2013). Last, perceived anthropomorphism was measured with the item “Das System wirkt menschlich” (“The system seems human-like”), which was extracted from the measurement scale for anthropomorphism presented by Sidlauskiene and colleagues (2023).

Finally, demographic variables were collected. These included age, gender, student status, employment status, and level of German skills. Additional information was collected about participants’ shopping preferences and IT affinity. IT affinity was measured on a 5-point Likert scale with the item “Ich beschäftige mich gern genauer mit technischen Systemen” (“I like to occupy myself in greater detail with technical systems”). This item was taken from the German version of the Affinity for Technology Interaction (ATI) Scale (Franke et al., 2019). Shopping preferences were gathered with multiple items specifically designed for this survey. More specifically, users were primarily asked about their use of price comparison sites, which reflected items of interest to an external research party. Importantly, all behavioral and survey data were linked via anonymous user access keys, preserving privacy while ensuring proper links between all data points.

Data Analysis

The following section outlines the data analysis procedure of the current study. All tests conducted to assess the hypotheses and thereby provide meaningful answers to the research questions are described.

Assumption Checks

Before conducting the analysis, the appropriate statistical assumptions for (parallel) mediation (regression-based analyses) were assessed.

Appropriateness of Variables. Ensuring that each model variable is suitable for the specific analysis is an important first step. This assessment involved examining the scale properties of each variable: the IV (type of RS), the three MV (perceived explainability, sense of control, and perceived anthropomorphism), and the DV (system trust). Attention was given to ensuring that the variable types were compatible with the regression-based procedures underlying the PROCESS macro and that the measurement scales provided sufficient variability for meaningful analysis.

Linearity. The assumption of linearity requires that the relationship between predictor and outcome variables is linear, ensuring that changes in the predictor variable correspond to proportional changes in the outcome variable (Hayes, 2022; Kane & Ashbaugh, 2017). Following established procedures for mediation models, separate linear regression models were conducted for each of the model paths. For each regression, residuals versus predicted values plots were generated and inspected for non-linear patterns, with random scatter around zero indicating satisfied linearity assumptions.

Multicollinearity. Multicollinearity was examined to ensure that the mediators were not excessively intercorrelated. If this was the case, the model's ability to distinguish the individual effects of each mediator on the DV could be compromised, leading to unstable parameter estimates (Fein et al., 2022; Shrestha, 2020). This assumption is particularly critical in parallel mediation models where multiple mediators are included simultaneously. To assess multicollinearity, Variance Inflation Factor (VIF) values were calculated. VIF values below 5.0 indicate acceptable levels of multicollinearity, with values close to 1.0 suggesting minimal correlation among predictors (Shrestha, 2020). VIF scores were calculated based on a correlation matrix among the mediators using Pearson product-moment correlations and a multiple linear regression model where all three mediators simultaneously predicted trust.

Homoscedasticity. To ensure that the variance of residuals remains constant across all levels of predicted values, homoscedasticity needed to be assessed. This assumption requires estimation errors to be relatively equal across all predicted Y (dependent variable) values, as heteroscedasticity can affect the standard errors of regression coefficients and

compromise the validity of significance tests (Hayes, 2022; Kane & Ashbaugh, 2017). To assess this, a comprehensive approach was employed using both visual inspection and statistical testing. Visually, a residual vs. fitted value plot (scatterplots) was examined, where homoscedasticity is indicated by residuals being randomly scattered around the horizontal line, rather than showing a specific pattern (Kane & Ashbaugh, 2017). Additionally, the Goldfeld-Quandt test, Breusch-Pagan test, and White test were conducted as a formal assessment (Hayes, 2022; Onifade & Olanrewaju, 2020).

Normality. Normality was assessed to evaluate if estimation errors were normally distributed (Kane & Ashbaugh, 2017). However, the PROCESS macro's use of bootstrapping procedures reduces reliance on this assumption, as these methods are robust to non-normal distributions (Hayes, 2022; Kane & Ashbaugh, 2017). Visual assessment of normality included examination of a Q-Q plot, where normally distributed data should fit with the diagonal line (Kane & Ashbaugh, 2017), as well as a histogram to assess distribution shape. Additionally, the Shapiro-Wilk test was conducted as a formal statistical assessment of normality, where a non-significant result ($p > .05$) would indicate satisfaction of the normality assumption (Ghasemi & Zahediasl, 2012).

Independence. Independence of observations is critical to ensure that errors associated with each participant are independent from errors from other participants (Hayes, 2022; Kane & Ashbaugh, 2017). This assumption requires that cases should not be systematically related to each other. The independence assumption was evaluated by examining the recruitment strategy, data collection context, participant interaction possibilities, and potential clustering effects within the sample.

Conducting the Parallel Mediation

Data analysis was conducted in R (Posit Software, 2025) and JASP (JASP Team, 2020), with key R packages visible in the R code in the appendix. Most specifically and importantly, to test the main hypotheses, a parallel mediation analysis was conducted using the PROCESS macro for R (Hayes, 2022; Kane & Ashbaugh, 2017). In this model, RS type served as the IV, while perceived explainability, sense of control, and perceived anthropomorphism served as mediating variables (MVs), and trust served as the DV. The mediation analysis assessed (a) the direct effect of the IV on the DV – the direct path of RS type to trust –, (b) the effect of each mediator on the DV, and (c) the indirect paths from the IV (system type) to the DV (trust) through each MV (perceived explainability, sense of

control, perceived anthropomorphism). Statistical significance of results was determined using the standard $p < .05$ criterion.

Sensitivity Analyses

Following examination of the key assumptions for parallel mediation analysis and after conducting the mediation analysis, a series of sensitivity analyses were conducted to assess the robustness and boundaries of results. The following analyses were completed.

Influential-Case / Outlier Diagnostics. To assess whether a small number of cases (participant observations) exerted disproportionate influence on the mediation results, an influential case analysis was conducted using Cook's distance (Cook, 1977).

Robust Estimators – Checking Heteroscedasticity. Considering mixed results of the assessment of homoscedasticity (presented in further detail in the Results), two complementary robust estimators were applied to evaluate if the regression results were sensitive to classical error assumptions for mediation (homoscedasticity) – heteroscedasticity-consistent 3 (HC3) standard errors, and M-estimation (Huber estimation) (Anandhi & Prabhu, 2023; Jochmans, 2022).

Moderated Mediation: Technology Affinity. To determine if the indirect or direct effects of the IV on the DV depended on participants' technology affinity, a moderated parallel mediation model was estimated. Technology affinity, measured on a 5-point Likert scale in the exit survey, was mean-centered and added as a moderator to all paths.

Exploratory & Qualitative Analyses

Finally, to explore whether the DV influenced participants' likelihood of selecting a recommended item, a binary measure of recommendation success (acceptance) was analyzed, coded as selection = TRUE and no selection = FALSE. First, an independent samples t-test compared acceptance rates between the BRS and CRS conditions. Due to the binary nature of the outcome variable and the small expected cell frequencies, Fisher's exact test was also employed to check robustness. Free-text responses voluntarily provided at the end of the survey were thematically reviewed to provide qualitative insights, primarily, to broaden points for discussion.

One last question was formulated and examined post hoc. It was examined if people's level of reported trust influenced recommendation acceptance. This was assessed via correlation analysis (linear regression). This analysis was included to illuminate how not only

interface design, but also underlying attitudes may affect user behavior and trust in different types of RSs.

Data Quality & Inclusion Criteria

To ensure data quality and meaningful interpretation of the results, data was included in the analysis only if the following criteria were met: For the baseline group (BRS), participants must have demonstrated appropriate interaction with the system by engaging with the filtering options (setting at least one filter). For the treatment group (CRS), participants must have interacted with the chatbot in a way that resulted in automated filtering (setting of at least one background filter). Additionally, participants must have completed the exit survey in full, providing responses to all items. Data that did not meet these criteria was excluded from the analysis. Participants who reported German proficiency below a B2 level and those participants who completed the study in under 80s were also excluded.

Results

The following section presents the outcomes of the mediation and exploratory analyses, beginning with checks of key statistical assumptions, and then moving to descriptives, hypothesis testing results, and supplemental exploratory tests.

Assumption Checks

The assumption checks conducted before the mediation analysis revealed the following about the variables and data.

Appropriateness of Variables

All variables in the mediation model demonstrated appropriate characteristics for this analytical procedure. The IV (RS type) was coded as a dichotomous variable, which is suitable for mediation analysis. The three mediators as well as the DV (trust) were measured on a 5-point Likert scale. As this provides sufficient variability in responses for regression analysis, their treatment as continuous variables is warranted (Kane & Ashbaugh, 2017).

Linearity

Linearity was satisfied across all mediator-outcome relationships. Residual vs. fitted plots for each regression model demonstrated random scatter with no systematic patterns, curves, or funnels, confirming linear relationships between variables. Additionally, with

binary measures of the independent variable (type of RS), the linearity assumption was automatically satisfied for this predictor.

Multicollinearity

The multi-method assessment revealed acceptable multicollinearity levels, with VIF scores being well below the threshold of 5.0: perceived explainability (VIF = 1.14), sense of control (VIF = 1.29), perceived anthropomorphism (VIF = 1.19). Being above a VIF of 1.0, the mediators show slight intercorrelation. Nevertheless, strong multicollinearity was not present.

Homoscedasticity

Informal, visual inspection of the residuals vs. fitted values plot revealed evidence of heteroscedasticity due to systematic scatter around the horizontal zero line, indicating violation of homoscedasticity. Heteroscedasticity can affect the standard errors of regression coefficients and reduce the validity of statistical tests (Hayes, 2022). In contrast, statistical tests yielded contradicting results with the Goldfeld-Quandt test ($GQ = 0.66, p = .95$), Breusch-Pagan Test ($BP = 4.24, p = .38$), and White test ($BP = 4.63, p = .09$) failing to reject the null hypothesis, indicating homoscedasticity. Given the convergent statistical evidence and the PROCESS macro's robust inference option using bootstrapping procedures, which can buffer against homoscedasticity violations, mediation analysis is still possible (Hayes, 2022). Nevertheless, to further address potential concerns regarding homoscedasticity identified through visual inspection, robust standard errors (HC3) were additionally computed for the subsequent analysis.

Normality

The Shapiro-Wilk Test indicated a statistically significant departure from normality ($W = 0.97, p = .01$). The Q-Q plot and histogram reveal moderate adherence to normality. Combined with PROCESS's use of bootstrapping, the fact that this assumption is rarely met in practice (Hayes, 2022) and the argument that departures from normality should not substantially affect results with adequate sample sizes and bootstrapping, these mixed results are nevertheless regarded as sufficient for mediation analysis.

Independence

The independence of observations assumption was considered satisfied based on the data collection methodology and sampling procedures employed. Participants were recruited

through multiple independent channels, thereby minimizing systematic clustering effects. The web-based study format ensured that participants completed the experiment individually at their own computers without direct interaction or influence from other participants.

Descriptives

Descriptive Results of Survey Responses

Participants were distributed relatively equally across both conditions (69 participants in the BRS, 75 in the CRS condition). Descriptive results of the main study variables are presented in Table 1.

Participants overall reported moderately high levels of trust in both, the baseline condition ($M = 3.88$, $SD = 1.04$) and conversational condition ($M = 3.89$, $SD = 0.89$), with negligible differences between groups ($d = 0.01$), on a 5-point Likert scale. Similarly, perceived explainability was rated highly in both conditions, with identical means for the baseline ($M = 4.41$, $SD = 1.03$) and conversational ($M = 4.41$, $SD = 0.77$) RSs ($d = 0.01$).

Perceived anthropomorphism was rated relatively low across both conditions, with minimal differences between baseline ($M = 2.58$, $SD = 1.31$) and conversational ($M = 2.51$, $SD = 1.13$) systems ($d = -0.06$). The most notable difference emerged for sense of control, where participants in the baseline condition reported higher perceived control ($M = 4.09$, $SD = 0.92$) compared to those in the conversational condition ($M = 3.60$, $SD = 1.13$), representing a moderate effect size ($d = -0.47$).

Table 1

Descriptive Results of Survey Responses

Variable	Baseline RS				Conversational RS				d
	n	M	SD	CI (95%)	n	M	SD	CI (95%)	
Trust	69	3.88	1.04	3.64–4.13	75	3.89	0.89	3.69–4.10	0.01
Explainability	69	4.41	1.03	4.16–4.65	75	4.41	0.77	4.24–4.59	0.01
Sense of Control	69	4.09	0.92	3.87–4.31	75	3.60	1.13	3.34–3.86	-0.47
Anthropomorphism	69	2.58	1.31	2.26–2.89	75	2.51	1.13	2.25–2.77	-0.06

Model Summary

The overall mediation model was statistically significant ($F(4, 139) = 9.84, p < .001, R^2 = .22$), explaining 22% of the variance in user trust. The model fit indices demonstrated adequate model performance with acceptable residual standard error ($SE = 0.74$).

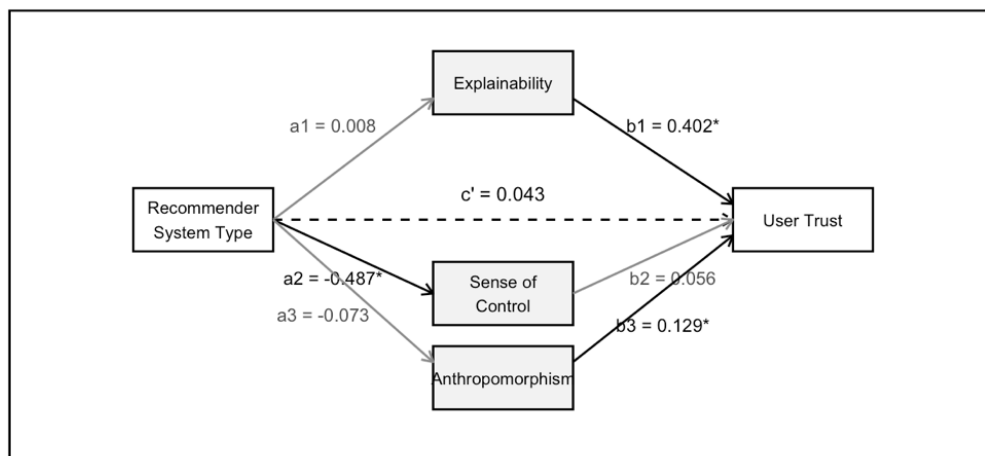
Figure 4 displays the parallel mediation model with unstandardized path coefficients. The diagram illustrates the direct path from recommender system type to user trust (c' path) and the indirect paths through the three mediators: perceived explainability, sense of control, and perceived anthropomorphism. Asterisks and black, bold arrows indicate statistically significant paths ($p < .05$).

Hypothesis Tests

Considering H1, the direct effect of RS type on user trust was not statistically significant ($\beta = 0.04, SE = 0.15, p = .77, 95\% CI [-0.25, 0.34]$). With a p -value above .05 and a CI including 0, H1 was not supported. Interacting with a CRS was not associated with significantly higher user trust than interacting with a BRS.

Examining the effects of the mediating variables, the analysis produced the following results.

Perceived Explainability. The effect of interacting with a CRS in comparison to a BRS was not statistically significant ($\beta = 0.01, SE = 0.15, p = .96, 95\% CI [-0.29, 0.31]$). H2a was therefore not supported. Interacting with a CRS was not associated with significantly higher perceived explainability than interacting with a BRS. H2b – higher perceived explainability being associated with significantly higher user trust, however was supported. Perceived explainability demonstrated a statistically significant positive effect on user trust ($\beta = 0.40, SE = 0.09, p < .00, 95\% CI [0.23, 0.57]$ with a medium to large effect size. The mediation, H2c however, was also not supported ($ab = 0.00, SE = 0.06, p = .95, 95\% CI [-0.11, 0.14]$), suggesting that perceived explainability did not mediate the association between interacting with a CRS (vs. a BRS) and user trust.



Note. Values are unstandardized path coefficients (B).
 Paths from recommender system type to mediators are labelled a1–a3.
 Paths from mediators to trust are labelled b1–b3.
 The dashed line is the direct effect (c').
 * significant paths with $p < .05$

Figure 4. Results of the Mediation Model Path Diagram

Sense of Control. The effect of RS type on people's sense of control was significant, however, in the opposite direction than hypothesized ($\beta = -0.49$, $SE = 0.17$, $p = .01$, 95% $CI [-0.83, -0.15]$). H3a was not supported due to the opposite direction of the effect. The effect of participants' sense of control on trust, H3b, was also not significant ($\beta = 0.06$, $SE = 0.08$, $p = .48$, 95% $CI [-0.10, 0.21]$). Higher sense of control was therefore not associated with significantly higher user trust. Finally, the indirect effect of RS type on trust through sense of control was not statistically significant either ($ab = -0.03$, $SE = 0.04$, $p = .17$, 95% $CI [-0.13, 0.05]$), with the bootstrap confidence interval including zero, indicating no mediation effect. H3c was therefore also not supported.

Perceived Anthropomorphism. The RS type did not have a significant effect on perceived anthropomorphism ($\beta = -0.07$, $SE = 0.20$, $p = .72$, 95% $CI [-0.48, 0.33]$). H4a, the hypothesis that interacting with a CRS would be associated with significantly higher perceived anthropomorphism than interacting with a BRS was not supported. However, perceived anthropomorphism did have a positive effect on user trust ($\beta = 0.13$, $SE = 0.07$, $p = .05$, 95% $CI [0.00, 0.26]$), however only with a small effect. H4b was therefore supported. Lastly, the indirect effect of RS type on trust via anthropomorphism (H4c) was not supported, with results being ($ab = -0.01$, $SE = 0.03$, $p = .30$, 95% $CI [-0.08, 0.05]$).

Table 2 presents a comprehensive overview of hypothesis testing results. Of the ten hypotheses tested, only two were supported (H2b and H4b). The primary hypothesis (H1) regarding the direct effect of CRS on trust was not supported. None of the proposed mediation effects (H2c, H3c, H4c) was statistically significant. Notably, H3a demonstrated a significant effect in the opposite direction than predicted, with CRS users reporting lower sense of control than BRS users.

Table 2

Overview of Hypotheses Test Results

Hypothesis	Prediction	Result	p	95 % CI	Support
H1 CRS → Trust (direct)	$\beta > 0$	$\beta = 0.043$	.772	[-0.251, 0.337]	Not supported
H2a CRS → Explainability	$\beta > 0$	$\beta = 0.008$	.960	[-0.292, 0.307]	Not supported
H2b Explainability → Trust	$\beta > 0$	$\beta = 0.402$	<.001	[0.233, 0.570]	Supported
H2c CRS → Trust via Explainability	Mediation	ab = 0.003	.948	[-0.106, 0.143]	Not supported
H3a CRS → Sense of Control	$\beta > 0$	$\beta = -0.487$	.005	[-0.828, -0.146]	Opposite direction
H3b Sense of Control → Trust	$\beta > 0$	$\beta = 0.056$	.482	[-0.102, 0.214]	Not supported
H3c CRS → Trust via Sense of Control	Mediation	ab = -0.027	.172	[-0.127, 0.051]	Not supported
H4a CRS → Anthropomorphism	$\beta > 0$	$\beta = -0.073$	.720	[-0.476, 0.329]	Not supported
H4b Anthropomorphism → Trust	$\beta > 0$	$\beta = 0.129$	.048	[0.001, 0.256]	Supported
H4c CRS → Trust via Anthropomorphism	Mediation	ab = -0.009	.300	[-0.075, 0.051]	Not supported

Note. CRS = Conversational Recommender System.

β = standardized coefficient for direct paths.

ab = indirect (mediated) effect.

CI = confidence interval.

Shaded rows indicate supported hypotheses.

Robustness & Sensitivity Checks

For assessing influential cases and outliers, the rule-of-thumb cutoff $4/n$ for Cook's distance was used, resulting in a cutoff value of 0.028, with $n = 144$, which identified 13

respondents who exceeded the threshold (cases in row 3, 44, 49, 54, 61, 63, 70, 74, 75, 88, 112, 118, 137 in the dataset). Rerunning the mediation model without these cases produced results that were largely unchanged, as seen in Table 3. Results show that after influential cases were removed (*B* Trimmed), the mediation path on trust via explainability remained stable, whereas the small positive association between anthropomorphism and trust was reduced, indicating that this relationship may be driven by a few influential cases in the initial analysis. No other results changed significantly after removal of outliers, indicating overall robustness of the model to influential observations.

The linear relation between covariates (IV and mediators) and trust was re-estimated with two alternative estimators to evaluate sensitivity to heteroskedasticity. With HC3 heteroskedasticity consistent standard errors, coefficient estimates were identical to ordinary least squares from the initial model, while standard errors were adjusted upward. The path from explainability to trust remained statistically significant ($B = 0.40$, $SE = 0.09$, $p < .00$), anthropomorphism showed a modest effect that reached the .05 significance level ($B = 0.13$, $SE = 0.07$, $p = .05$), and neither sense of control ($B = 0.06$, $SE = 0.08$, $p = .48$) nor the direct effect of RS type ($B = 0.04$, $SE = 0.15$, $p = .77$) was significant. With a Huber M-estimator, similar results could be observed for explainability ($B = 0.43$, $SE = 0.09$, $p < .00$), sense of control ($B = 0.07$, $SE = 0.08$, $p = .24$), and the direct effect of the type of RS ($B = 0.05$, $SE = 0.15$, $p = .76$). However, with the M-estimator, the path of anthropomorphism decreased slightly and was no longer significant ($B = 0.12$, $SE = 0.07$, $t = 1.75$, $p = .08$). Explainability was therefore stable across both estimators, anthropomorphism was sensitive to the weighting, and the other paths remained the same.

The moderated mediation analysis revealed no moderation by participants' reported level of technology affinity. For the direct effect of RS type on trust, the effect was non-significant for low ($B = -0.06$, $SE = 0.22$, $p = .78$, 95% $CI [-0.5, 0.37]$), medium ($B = 0.00$, $SE = 0.15$, $p = 1.00$, 95% $CI [-0.31, 0.30]$), and high ($B = 0.05$, $SE = 0.20$, $p = .80$, 95% $CI [-0.34, 0.44]$) technology affinity. Similarly, none of the indirect effects varied across participants with low, medium, or high technological affinity. For the path via explainability, the effect was non-significant for low ($B = -0.09$, $BootSE = 0.10$, 95% $CI [-0.30, 0.10]$), medium ($B = -0.00$, $BootSE = 0.06$, 95% $CI [-0.11, 0.14]$), and high ($B = 0.1$, $BootSE = 0.10$, 95% $CI [-0.06, 0.32]$) levels of technology affinity. The same results could be observed for through sense of control, where indirect effects remained non-significant at low ($B = -0.01$, $BootSE = 0.09$, 95% $CI [-0.19, 0.18]$), medium ($B = -0.02$, $BootSE = 0.05$, 95% $CI [-0.12,$

0.07]), and high ($B = -0.02$, $\text{BootSE} = 0.05$, 95% $CI [-0.13, 0.06]$) affinity levels. Finally, for the pathway of anthropomorphism, no significant indirect effects were found at low ($B = -0.02$, $\text{BootSE} = 0.05$, 95% $CI [-0.15, 0.04]$), medium ($B = -0.01$, $\text{BootSE} = 0.03$, 95% $CI [-0.08, 0.05]$), or high ($B = 0.04$, $\text{BootSE} = 0.06$, 95% $CI [-0.07, 0.16]$) levels of technology affinity. These findings indicate that neither the direct nor the indirect effects of RS type on user trust were dependent on participants' level of technology affinity.

Table 3

Influential-Cases Diagnostics: Path Coefficients With and Without High Influence Cases

Path	B (All Cases)	p	B (Trimmed)	p
RS Type → Trust (c')	0.043	.772	0.060	.652
Explainability → Trust (b_1)	0.402	< .001	0.454	< .001
Sense of Control → Trust (b_2)	0.056	.482	0.087	.244
Anthropomorphism → Trust (b_3)	0.129	.049	0.079	.168

Note. B = unstandardised coefficient.

Trimmed model excludes 13 cases with Cook's D > 4/n.

Exploratory Outcomes

Recommendation Acceptance. The exploratory analysis examined whether recommender system type influenced recommendation acceptance behavior. Of the original 144 participants, interaction data was only available for 100 participants. The remaining 44 data points may have been lost due to settings of tracking blockers on participants' devices. Nevertheless, missing interaction data was distributed relatively equally across the two experimental conditions, with 24 missing data points in the BRS and 20 missing in the CRS group. Descriptive analysis revealed high recommendation acceptance rates in both conditions, with 97.8% (44/45 participants) and 96.4% (53/55 participants) recommendation acceptance (item selections) in the BRS and CRS groups, respectively. This shows slightly higher recommendation acceptance within BRS interactions. For more thorough evaluation, Fisher's exact test also indicated no statistically significant difference in recommendation acceptance between system types ($p = 1.00$, $OR = 0.61$, 95% $CI [0.010, 11.99]$). Similarly, an independent samples t-test comparing mean recommendation acceptance rates between

groups was also not statistically significant ($t = 0.418, p = .68, 95\% CI [-0.05, 0.08]$). The analysis therefore revealed no positive effect of interacting with a CRS on recommendation acceptance in comparison to interacting with a BRS.

The Influence of Trust on Recommendation Acceptance. No statistically significant correlation could be observed between system trust and recommendation acceptance ($r = .05, p = .61, 95\% CI [-.15, .25]$). Hence, higher system trust did not show a significant relationship with recommendation acceptance.

Qualitative Feedback

Responses of the open-ended comment field included at the end of the exit survey were reviewed to gain understanding of participants' spontaneous reactions to the study and system. Responses were grouped into high-level themes through a basic thematic analysis. Four recurrent topics emerged.

1. **Anthropomorphism:** Several users noted that the CRS felt like a driven question-and-answer scheme rather than a free, open dialog, reducing the feeling of anthropomorphism. For instance, it was noted that they were not able to ask questions and received no natural responses by the chatbot when they diverged from the questions at hand.
2. **Interactivity and control:** CRS sessions were generally described as engaging, however, overly system-led, with participants requesting greater ability to take initiative and have more active power.
3. **Product information and navigation:** Some users enjoyed the clean presentation and ease of use, reporting that the systems felt very intuitive. In the CRS condition, one participant also noted the efficiency of the chatbot, appreciating that there were no waiting times for chatbot responses. However, other users also criticized the user interface and user experience, finding it unappealing. In both groups, participants desired more explicit information about product features or recommendation rationales.
4. **Session length:** Overall, sessions in the CRS condition tended to run longer. While some participants appreciated this as it allowed deeper exploration, others viewed it as excessive.

Discussion

This study aimed to investigate how conversational recommender systems influence user trust in e-commerce contexts, with a particular focus on the mediation effects of perceived explainability, sense of control, and perceived anthropomorphism. Drawing on the state-of-the-art literature, it was hypothesized that interactions with a CRS would increase user trust, in part, through these psychological mechanisms in comparison to interactions with a baseline system.

However, results revealed a more complex picture than assumed. Of the mediation hypotheses tested, only two were supported. Namely, perceived explainability (H2b) and perceived anthropomorphism (H4b) both showed positive associations with system trust. This is consistent with existing literature on trust antecedents. However, with both of these antecedents, only the indirect effect between mediator and trust as the outcome variable was statistically significant: the mediating effect of the antecedent on the relationship between system type and trust could not be identified. Critically, the primary hypothesis (H1) regarding the direct effect of CRSs on trust was not supported. Notably, participants in the CRS condition reported a significantly lower sense of control than those in the BRS condition, which directly contradicts H3a.

Connecting Results to Current Literature

Examining the study's findings more closely, it is also important to situate them within the broader literature. Thus, the sometimes unexpected results are discussed in the context of previous work in the following sub-section, which also ends with a discussion of how this research contributes to the existing literature.

Explainability Without Enhancement

Both systems achieved high explainability ratings, suggesting that both approaches provided sufficient transparency about system operations. This aligns with research indicating that architectural explainability – understanding how a system works – can be achieved through various interface designs (Ferrario & Loi, 2022). Importantly, the strong relationship between explainability and trust, regardless of the system type, confirms established findings in XAI research concerning the fundamental importance of system transparency for trust formation across different systems and technologies (Govea et al., 2024; Pai et al., 2022).

The Paradox of Conversational Control

The finding that CRS users experienced less sense of control may challenge assumptions about how conversational interfaces enhance user agency. While previous literature suggests that interactive dialog systems increase user control (Christakopoulou et al., 2018; Jannach, 2023; Pizzi et al., 2021), the results of this study indicate the opposite. Indeed, these results align more closely with automation research that suggests that increased system autonomy can diminish perceived human agency (Berberian et al., 2012; Vantrepotte et al., 2022). The conversational interface, despite its interactive nature, may have created a perception of reduced control because users delegated filtering decisions to the system (chatbot) rather than directly filtering themselves. In the present context, the CRS's automated filter-setting based on conversational input from the user may have positioned users as more passive recipients of system decisions rather than as truly active contributors or controllers of the recommendation process.

The Modest Role of Anthropomorphism

Minimal ratings of perceived anthropomorphism across both conditions indicate that neither system successfully triggered strong human-like perceptions, which is different from previous conversational AI research highlighting the anthropomorphic potential of chatbots (Araujo, 2018; Singh et al., 2024). Nevertheless, the significant positive relationship between anthropomorphism and trust supports the importance of human-like cues in technology trust and acceptance, even when such cues are subtle.

Theoretical Contributions

The results contribute several important insights to the literature of trust in AI. First, they show that increased interactivity does not automatically translate to enhanced trust. This challenges basic assumptions about conversational interfaces and highlights the need for a more nuanced understanding of how different interaction modalities affect user perceptions.

Second, the study reinforces the importance of explainability and anthropomorphism as predictors of trust but results question their roles as mediators of system design effects. This suggests that these factors may operate more as general trust antecedents rather than specific mechanisms through which interface design influences trust.

Third, the findings highlight the complexity of sense of control in automated systems. The paradoxical reduction in perceived control within the more interactive system

underscores the need to carefully consider how automation features affect user agency, even when designed to enhance user engagement.

However, all of these results may also be subject to the study design. That is, the design of this research may have been inadequate to fully address specific hypotheses. This is further discussed in the following section.

Limitations

There are several potential limitations to this study that may have affected the results. These limitations are presented here.

Evolving Research Design

A significant limitation of this study stems from the evolution of the study design during the study development process. Originally, the baseline condition was intended to represent a more “traditional” collaborative filtering RS that would allow minimal user control and include little explainability. However, practical constraints, including the cold start problem – the inability to generate meaningful recommendations without prior user interactions, preferences, or profiles – required the study group to switch to a filter-based system for the baseline condition. This design change fundamentally altered the nature of the comparison, as naturally, being able to set filters would enhance user control. Instead of contrasting a conversational system with a passive one, the study ultimately compared two “interactive” systems, though with different interaction modalities and different levels of “natural” interaction (filter vs. conversation). However, because the filter-based approach provided users with direct, granular control over recommendation parameters in comparison to the treatment group, in which the underlying system set these parameters based on chatbot interactions, sense of control was understandably higher in the control group than expected at the time of formulating the hypotheses.

This limitation also extends to perceived explainability, as filtering also makes the recommendation process highly transparent and comprehensible to users. When users set manual filters for various item parameters, they have increased insight into how and why recommendations are generated, i.e., the system logic becomes apparent as users can trace the causal relationship between their filter choices and resulting recommendations (architectural explainability) (Ferrario & Loi, 2022).

Moreover, due to development and implementation constraints, the CRS implemented in this study followed a system-driven, question-based paradigm rather than a truly interactive approach. Questions and question order were pre-defined, and user responses were mapped to specific, pre-defined filter values. This aligns with one of four paradigms for CRSs, namely the system-driven approach (Jannach et al., 2021). Here, the system leads the conversation and the user merely replies, thereby constraining the level of interactivity and control. Additionally, this approach prevents free user queries or topic shifts, making the system also appear less anthropomorphic. Some qualitative feedback from the study confirmed frustration with this one-way flow as users desired a more mixed-initiative design in which both system and user can dynamically initiate topics and shape the conversation. Different from initial expectations and therefore diverging from truly interactive, mixed-initiative CRSs (Jannach et al., 2021), the utilized CRSs may have inadvertently also felt more controlling and less anthropomorphic, reducing participants' sense of control and perceived anthropomorphism.

Reactance and Individual Differences

The exogenous assignment of recommender system types may have led to reactance effects. Some participants may inherently prefer filter over chatbot-based RSs, or vice versa, and may have therefore been less interested in and attentive to the system interactions and subsequent survey. Such a heterogeneity in system preference would suggest that the benefits of (conversational) RSs may be contingent on individual differences in interaction preferences and technology attitudes. However, as these are factors not captured in the current study design, no clear assumptions can be made here.

Another interindividual difference that may have influenced interest and outcomes may have been participants' domain knowledge about washing machines, the context of this study. Users with high product knowledge might have preferred the efficiency of filtering, while those with little knowledge may have benefited more from the conversational system, which was designed to provide more guidance and contextual support in a domain in which consumers typically have little experience or know-how. The study did not account for these knowledge differences, which may have masked differential effects of system type across user segments.

Temporal and Contextual Limitations

The study captured user perceptions at a single point in time following a brief interaction with an AI system. Trust in AI, however, develops dynamically through repeated interactions (Hoff & Bashir, 2015). The temporal limitations of the current design prevent understanding of how trust, explainability, and control perceptions might evolve with extended system use. This is particularly relevant for conversational systems, where users in different contexts, such as e-commerce, may need time to understand the chatbot's capabilities and develop effective interaction patterns. Additionally, the simulated e-commerce environment, while carefully designed to mimic authentic shopping experiences, lacked real purchase consequences. This may have reduced the stakes of the decision-making process and influenced how participants evaluated trust-related factors.

Methodological Constraints

Despite employing mediation analysis, the cross-sectional design limits causal inferences about the relationships between variables. Mediation analysis, while informative about indirect effects, does not establish true causal pathways (Hayes, 2022). The observed associations could therefore reflect common underlying factors or other causation patterns that were not captured by this analytical approach.

Furthermore, the assumption violations identified during analysis, particularly regarding homoscedasticity and normality, indicate that the data did not fully meet ideal conditions for mediation analysis, even though these violations were addressed through robust estimation checks.

The use of single-item measures for key constructs, while necessary for practical reasons, represents a further limitation concerning the methodology. Although the items were adapted from validated scales, single-item measures likely do not fully capture the complexity of multidimensional constructs such as trust and its antecedents and may therefore have led to sensitivity concerns.

Generalizability

The sample exhibited limited demographic diversity, with a majority of male participants (69.4%), and 75% of students, many of whom had a computer science background. This composition may limit generalizability to broader consumer populations,

particularly given that trust in technology can vary significantly across demographic groups and technical expertise levels.

Beyond concerns regarding demographics, as the study focused on a single product category (washing machines) within a specific cultural context (Austria) employing two specific AI applications (filter- and chatbot-based RSs), results may not extend to other product domains, cultural settings, or AI systems.

Implications

The divergent results carry meaningful insights for both practitioners and scholars. In this section, the findings are summarized into more concrete design, research implications, and guidelines for future CRS developers. Finally, avenues for future research are outlined.

Practical Recommendations for Conversational Recommender Systems Design

Despite the unexpected results, this study offers several practical insights for RS designers, and more specifically, e-commerce system and platform designers. First, the positive relationship between explainability and trust underscores the critical importance of enhancing transparency in recommendation systems, regardless of interface modality. Clear, accessible explanations of how recommendations are generated are essential and should therefore be prioritized by companies and systems designers, especially when consequences such as financial decisions (spending money on item purchases) are involved.

Second, the findings that filter-based interfaces can provide higher perceived control suggest that traditional interface elements retain significant value even as conversational systems gain traction. Fully removing established features and controls with chatbots risks alienating users who prefer direct manipulation of filters. A hybrid approach that preserves user agency – also in choosing between conventional filters and chat-driven recommendations – may therefore be optimal to preserve user agency, cater to diverse interaction preferences, and mitigate potential reactance from users who do not want to use chatbots.

Third, the results and constraints of this research underscore the importance of grounding new feature integrations in empirical user research rather than solely relying on state-of-the-art literature or other best practices. Integrating chatbots or other elements without validating user needs, preferences, and expectations may inadvertently drive

technology aversion. Thus, new systems should ideally be piloted to ensure that new paradigms and trends truly enhance the user experience.

Finally, considering the design constraints of the present study only enabling a system-driven conversational system, the observed results point to the necessity of leveraging the full capabilities of LLMs to develop flexible, mixed-initiative dialog-based systems that closely mirror natural human-to-human interactions. In e-commerce contexts, these systems may then indeed serve as virtual shopping assistants, possibly enhancing trust and satisfaction even more.

Research Directions

The findings highlight the need for more nuanced theoretical frameworks for understanding the effects of conversational AI and CRSs. The assumption that interactivity enhances positive user outcomes requires reconsideration in light of the control paradox observed in this study. Future development and research should therefore account for the trade-offs between system autonomy and user agency in conversational interfaces. Moreover, common models of CRSs should continue to evolve from system-driven dialogs to truly more mixed approaches (Jannach et al., 2021). Regarding this, the significance of cross-disciplinary research is also evident, as traditional system-based accuracy benchmarks often fail to capture user appreciation. Therefore, researchers should continue to consider various of evaluation criteria, ranging beyond traditional metrics to also include psychological constructs.

The study also underlines the importance of carefully matching research designs to theoretical propositions. The practical constraints of the experimental user study that necessitated design changes illustrate how real-world implementation challenges can affect theoretical testing, suggesting that researchers may need more flexible theoretical frameworks that can accommodate diverse system architectures.

Future Research Avenues

Future research could employ longitudinal designs that capture trust development over extended interaction periods. Such studies could reveal whether the control paradox observed in this study persists as users become more familiar with conversational interfaces or whether trust patterns shift with experience.

Research could also incorporate more factors of individual differences, particularly attitudes towards technology, domain knowledge, or interaction preferences. This could help identify which user groups may benefit most from different interaction modes. Such a user-centered approach can moreover inform the development of more personalized RSs that adapt interface elements based on user characteristics.

Experimental designs that manipulate specific system features, such as explainability modes or anthropomorphic cues, while holding other features constant, may also provide more precise insights into the individual and interactive effects of specific design elements on user trust.

Importantly, cross- and interdisciplinary research such as this, at the intersection of computer science and technology, should continue to investigate not only the effects and implications of technology on humans, but also of human needs, preferences, and tendencies on technology design, implementation, and acceptance. Moreover, multi-method and multi-context approaches that combine quantitative and qualitative measures, as well as those studies that are close to application and reality, may provide an even deeper understanding of the mechanisms at work in human-technology interactions.

References

- Alyari, F., & Navimipour, N. (2018). Recommender systems: A systematic review of the state of the art literature and suggestions for future research. *Kybernetes*, 47(5), 985–1017. <https://doi.org/10.1108/K-06-2017-0196>
- Anandhi, P., & Prabhu, D. (2023). The robust regression estimators: Performance & evaluation. *International Journal of Statistics and Applied Mathematics*, 8(6), 83–87. <https://doi.org/10.22271/math.2023.v8.i6a.1444>
- Andras, P., Esterle, L., Guckert, M., Han, T., Lewis, P., Milanovic, K., Payne, T., Perret, C., Pitt, J., Powers, S., Urquhart, N., & Wells, S. (2018). Trusting intelligent machines: Deepening trust within socio-technical systems. *IEEE Technology and Society Magazine*, 37(4), 76–83. <https://doi.org/10.1109/MTS.2018.2876107>
- American Psychological Association. (2025). *Dictionary of Psychology*. <https://dictionary.apa.org/>
- Araujo, T. (2018). Living up to the chatbot hype: The influence of anthropomorphic design cues and communicative agency framing on conversational agent and company perceptions. *Computers in Human Behavior*, 85, 183–189. <https://doi.org/10.1016/j.chb.2018.03.051>
- Arnott, D., & Gao, S. (2019). Behavioral economics for decision support systems researchers. *Decision Support Systems*, 122, 113063. <https://doi.org/10.1016/j.dss.2019.05.003>
- Atas, M., Felfernig, A., Polat-Erdeniz, S., Popescu, A., Tran, T., & Uta, M. (2021). Towards psychology-aware preference construction in recommender systems: Overview and research issues. *Journal of Intelligent Information Systems*, 57(3), 467–489. <https://doi.org/10.1007/s10844-021-00674-5>
- Bach, T., Khan, A., Hallock, H., Beltrão, G., & Sousa, S. (2024). A systematic literature review of user trust in AI-enabled systems: An HCI perspective. *International Journal of Human–Computer Interaction*, 40(5), 1251–1266. <https://doi.org/10.1080/10447318.2022.2138826>
- Baizal, Z., Widyanoro, D., & Maulidevi, N. (2020). Computational model for generating interactions in conversational recommender system based on product functional requirements. *Data & Knowledge Engineering*, 128, 101813. <https://doi.org/10.1016/j.datak.2020.101813>
- Bao, Y., Cheng, X., De Vreede, T., & De Vreede, G.-J. (2021). Investigating the relationship between AI and trust in human-AI collaboration. In T. Bui (Ed.), *Proceedings of the*

- 54th Hawaii International Conference on System Sciences (HICSS) (pp. 607616). University of Hawaii at Manoa. <https://doi.org/10.24251/HICSS.2021.074>
- Bart, V., Wenke, D., & Rieger, M. (2023). A German translation and validation of the sense of agency scale. *Frontiers in Psychology*, 14, 1199648. <https://doi.org/10.3389/fpsyg.2023.1199648>
- Bartneck, C., Kulić, D., Croft, E., & Zoghbi, S. (2009). Measurement instruments for the anthropomorphism, animacy, likeability, perceived intelligence, and perceived safety of robots. *International Journal of Social Robotics*, 1(1), 71–81. <https://doi.org/10.1007/s12369-008-0001-3>
- Berberian, B., Sarrazin, J.-C., Le Blaye, P., & Haggard, P. (2012). Automation technology and sense of control: A window on human agency. *PLoS ONE*, 7(3), e34075. <https://doi.org/10.1371/journal.pone.0034075>
- Bobadilla, J., Ortega, F., Hernando, A., & Gutiérrez, A. (2013). Recommender systems survey. *Knowledge-Based Systems*, 46, 109–132. <https://doi.org/10.1016/j.knosys.2013.03.012>
- Bostrom, R., & Heinen, J. (1977). MIS problems and failures: A socio-technical perspective, part II: The application of socio-technical theory. *MIS Quarterly*, 1(4), 11-28. <https://doi.org/10.2307/249019>
- Chen, X., Wang, Y., & Yang, J. (2024). UCRI: A unified conversational recommender system based on item-guided conditional generation. *IEEE Intelligent Systems*, 39(1), 46–55. <https://doi.org/10.1109/MIS.2023.3330367>
- Chong, T., Yu, T., Keeling, D., & De Ruyter, K. (2021). AI-chatbots on the services frontline addressing the challenges and opportunities of agency. *Journal of Retailing and Consumer Services*, 63, 102735. <https://doi.org/10.1016/j.jretconser.2021.102735>
- Christakopoulou, K., Beutel, A., Li, R., Jain, S., & Chi, E. (2018). Q&R: A two-stage approach toward interactive recommendation. In R. Rosales & J. Tiang Tang (Eds.), *Proceedings of the 24th ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD '18)* (pp. 139–148). Association for Computing Machinery. <https://doi.org/10.1145/3219819.3219894>
- Christakopoulou, K., Radlinski, F., & Hofmann, K. (2016). Towards conversational recommender systems. In A. Smola, C. Aggarwal, D. Shen & R. Rastogi (Eds.), *Proceedings of the ACM SIGKDD International Conference on Knowledge Discovery and Data Mining (KDD '16)* (pp. 815–824). Association for Computing Machinery. <https://doi.org/10.1145/2939672.2939746>

- Cook, R. (1977). Detection of influential observation in linear regression. *Technometrics*, 19(1), 15–18. <https://doi.org/10.1080/00401706.1977.10489493>
- Dubova, M., Galesic, M., & Goldstone, R. (2022). Cognitive science of augmented intelligence. *Cognitive Science*, 46(12), e13229. <https://doi.org/10.1111/cogs.13229>
- Ebrahimi, S., Ghasemaghaei, M., & Benbasat, I. (2022). The impact of trust and recommendation quality on adopting interactive and non-interactive recommendation agents: A meta-analysis. *Journal of Management Information Systems*, 39(3), 733–764. <https://doi.org/10.1080/07421222.2022.2096549>
- Etemad-Sajadi, R. (2016). The impact of online real-time interactivity on patronage intention: The use of avatars. *Computers in Human Behavior*, 61, 227–232. <https://doi.org/10.1016/j.chb.2016.03.045>
- European Commission. Directorate General for Communications Networks, Content and Technology. High Level Expert Group on Artificial Intelligence. (2019). *Ethics guidelines for trustworthy AI*. Publications Office. <https://data.europa.eu/doi/10.2759/346720>
- Fein, E., Gilmour, J., Machin, T., & Hendry, L. (2022). *Statistics for research students: An open access resource with self-Tests and illustrative examples*. University of Queensland. <https://doi.org/10.26192/Q7985>
- Ferrario, A., & Loi, M. (2022). How explainability contributes to trust in AI. In F. Zuiderveen Borgesius, M. Kearns, K. Lum & A. Wu (Eds.), *2022 Proceedings of the 2022 ACM Conference on Fairness, Accountability, and Transparency (FAccT '22)* (pp. 1457–1466). Association for Computing Machinery. <https://doi.org/10.1145/3531146.3533202>
- Franke, T., Attig, C., & Wessel, D. (2019). A personal resource for technology interaction: Development and validation of the affinity for technology interaction (ATI) scale. *International Journal of Human-Computer Interaction*, 35(6), 456–467. <https://doi.org/10.1080/10447318.2018.1456150>
- Gao, C., Lei, W., He, X., De Rijke, M., & Chua, T.-S. (2021). Advances and challenges in conversational recommender systems: A survey. *AI Open*, 2, 100–126. <https://doi.org/10.1016/j.aiopen.2021.06.002>
- Ghasemi, A., & Zahediasl, S. (2012). Normality tests for statistical analysis: A guide for non-statisticians. *International Journal of Endocrinology and Metabolism*, 10(2), 486–489. <https://doi.org/10.5812/ijem.3505>

- Govea, J., Gutierrez, R., & Villegas-Ch, W. (2024). Transparency and precision in the age of AI: Evaluation of explainability-enhanced recommendation systems. *Frontiers in Artificial Intelligence*, 7, 1410790. <https://doi.org/10.3389/frai.2024.1410790>
- Guthrie, S. (2025). Anthropomorphism. In G. Lotha, E. Rodriguez, M. Stefon & G. Young (Eds.), *Encyclopedia Britannica*. <https://www.britannica.com/topic/anthropomorphism>
- Gweon, H., Fan, J., & Kim, B. (2023). Socially intelligent machines that learn from humans and help humans learn. *Philosophical Transactions of the Royal Society A: Mathematical, Physical and Engineering Sciences*, 381(2251), 20220048. <https://doi.org/10.1098/rsta.2022.0048>
- Haenlein, M., & Kaplan, A. (2019). A brief history of artificial intelligence: On the past, present, and future of artificial intelligence. *California Management Review*, 61(4), 5–14. <https://doi.org/10.1177/0008125619864925>
- Hayes, A. (2022). *Introduction to mediation, moderation, and conditional process analysis: A regression-based approach* (3rd edition). Guilford Press.
- Hernandez-Bocanegra, D., & Ziegler, J. (2023). Explaining recommendations through conversations: Dialog model and the effects of interface type and degree of interactivity. *ACM Transactions on Interactive Intelligent Systems*, 13(2), 1–47. <https://doi.org/10.1145/3579541>
- Hoff, K., & Bashir, M. (2015). Trust in automation: Integrating empirical evidence on factors that influence trust. *Human Factors: The Journal of the Human Factors and Ergonomics Society*, 57(3), 407–434. <https://doi.org/10.1177/0018720814547570>
- Hoffman, R., Miller, T., Klein, G., Mueller, S., & Clancey, W. (2023). Increasing the value of XAI for users: A psychological perspective. *KI - Künstliche Intelligenz*, 37(2–4), 237–247. <https://doi.org/10.1007/s13218-023-00806-9>
- Hoffman, R., Mueller, S., Klein, G., & Litman, J. (2023). Measures for explainable AI: Explanation goodness, user satisfaction, mental models, curiosity, trust, and human-AI performance. *Frontiers in Computer Science*, 5, 1096257. <https://doi.org/10.3389/fcomp.2023.1096257>
- Jannach, D. (2023). Evaluating conversational recommender systems: A landscape of research. *Artificial Intelligence Review*, 56(3), 2365–2400. <https://doi.org/10.1007/s10462-022-10229-x>
- Jannach, D., & Jugovac, M. (2019). Measuring the business value of recommender systems. *ACM Transactions on Management Information Systems*, 10(4), 16, 1–23. <https://doi.org/10.1145/3370082>

- Jannach, D., Manzoor, A., Cai, W., & Chen, L. (2021). A survey on conversational recommender systems. *ACM Computing Surveys*, 54(5), 105, 1–36. <https://doi.org/10.1145/3453154>
- JASP Team. (2020). *JASP* (Version 0.14.1) [Computer software]. <https://jasp-stats.org/previous-versions/>
- Jesse, M., & Jannach, D. (2021). Digital nudging with recommender systems: Survey and future directions. *Computers in Human Behavior Reports*, 3, 1–14. <https://doi.org/10.1016/j.chbr.2020.100052>
- Jian, J.-Y., Bisantz, A., & Drury, C. (2000). Foundations for an empirically determined scale of trust in automated systems. *International Journal of Cognitive Ergonomics*, 4(1), 53–71. https://doi.org/10.1207/S15327566IJCE0401_04
- Jochmans, K. (2022). Heteroscedasticity-robust inference in linear regression models with many covariates. *Journal of the American Statistical Association*, 117(538), 887–896. <https://doi.org/10.1080/01621459.2020.1831924>
- Kane, L., & Ashbaugh, A. (2017). Simple and parallel mediation: A tutorial exploring anxiety sensitivity, sensation seeking, and gender. *The Quantitative Methods for Psychology*, 13(3), 148–165. <https://doi.org/10.20982/tqmp.13.3>.
- Kaplan, A., Kessler, T., Brill, J., & Hancock, P. (2023). Trust in artificial intelligence: Meta-analytic findings. *Human Factors: The Journal of the Human Factors and Ergonomics Society*, 65(2), 337–359. <https://doi.org/10.1177/00187208211013988>
- Kelly, S., Kaye, S.-A., & Oviedo-Trespalacios, O. (2023). What factors contribute to the acceptance of artificial intelligence? A systematic review. *Telematics and Informatics*, 77, 101925. <https://doi.org/10.1016/j.tele.2022.101925>
- Kiela, D., Bartolo, M., Nie, Y., Kaushik, D., Geiger, A., Wu, Z., Vidgen, B., Prasad, G., Singh, A., Ringshia, P., Ma, Z., Thrush, T., Riedel, S., Waseem, Z., Stenetorp, P., Jia, R., Bansal, M., Potts, C., & Williams, A. (2021). *Dynabench: Rethinking benchmarking in NLP*. <https://doi.org/10.48550/arXiv.2104.14337>
- Körber, M. (2018). Theoretical considerations and development of a questionnaire to measure trust in automation. In S. Bagnara, R. Tartaglia, S. Albolino, T. Alexander & Y. Fujita (Eds.), *Proceedings 20th Triennial Congress of the International Ergonomics Association (IEA 1028)* (pp. 1–20). Springer. <https://doi.org/10.31234/osf.io/nfc45>
- Küper, A., & Krämer, N. (2024). Psychological traits and appropriate reliance: Factors shaping trust in AI. *International Journal of Human–Computer Interaction*, 41(7), 4115–4131. <https://doi.org/10.1080/10447318.2024.2348216>

- Lee, D., Moon, J., Kim, Y., & Yi, M. (2015). Antecedents and consequences of mobile phone usability: Linking simplicity and interactivity to satisfaction, trust, and brand loyalty. *Information & Management*, 52(3), 295–304. <https://doi.org/10.1016/j.im.2014.12.001>
- Lee, J., & See, K. A. (2004). Trust in automation: Designing for appropriate reliance. *Human Factors: The Journal of the Human Factors and Ergonomics Society*, 46(1), 50–80. https://doi.org/10.1518/hfes.46.1.50_30392
- Lee, M., & Turban, E. (2001). A trust model for consumer internet shopping. *International Journal of Electronic Commerce*, 6(1), 75–91. <https://doi.org/10.1080/10864415.2001.11044227>
- Lin, X., Wang, X., & Hajli, N. (2019). Building e-commerce satisfaction and boosting sales: The role of social commerce trust and its antecedents. *International Journal of Electronic Commerce*, 23(3), 328–363. <https://doi.org/10.1080/10864415.2019.1619907>
- Liu, W., & Wang, Y. (2025). Evaluating trust in recommender systems: A user study on the impacts of explanations, agency attribution, and product types. *International Journal of Human-Computer Interaction*, 41(2), 1280–1292. <https://doi.org/10.1080/10447318.2024.2313921>
- Lobner, S., Tesfay, W., Nakamura, T., & Pape, S. (2021). Explainable machine learning for default privacy setting prediction. *IEEE Access*, 9, 63700–63717. <https://doi.org/10.1109/ACCESS.2021.3074676>
- Lockey, S., Gillespie, N., Holm, D., & Someh, I. (2021). A review of trust in artificial intelligence: Challenges, vulnerabilities and future directions. In T. Bui (Ed.), *Proceedings of the 54th Hawaii International Conference on System Sciences (HICSS)* (pp. 5463–5472). University of Hawaii at Manoa. <https://doi.org/10.24251/HICSS.2021.664>
- Lukyanenko, R., Maass, W., & Storey, V. (2022). Trust in artificial intelligence: From a foundational trust framework to emerging research opportunities. *Electronic Markets*, 32(4), 1993–2020. <https://doi.org/10.1007/s12525-022-00605-4>
- Mayer, R., Davis, J., & Schoorman, F. (1995). An integrative model of organizational trust. *The Academy of Management Review*, 20(3), 709. <https://doi.org/10.2307/258792>
- McKnight, D., Choudhury, V., & Kacmar, C. (2002). Developing and validating trust measures for e-commerce: An integrative typology. *Information Systems Research*, 13(3), 334–359. <https://doi.org/10.1287/isre.13.3.334.81>

- Merritt, S., Heimbaugh, H., LaChapell, J., & Lee, D. (2013). I trust it, but I don't know why: Effects of implicit attitudes toward automation on trust in an automated system. *Human Factors: The Journal of the Human Factors and Ergonomics Society*, 55(3), 520–534. <https://doi.org/10.1177/0018720812465081>
- Morana, S., Gnewuch, U., Jung, D., & Granig, C. (2020). The effect of anthropomorphism on investment decision-making with robo-advisor chatbots. In S. Newell, N. Pouloudi & E. van Heck (Eds.), *Proceedings of the Twenty-Eighth European Conference on Information Systems (ECIS2020)* (pp. 1–18). https://aisel.aisnet.org/ecis2020_rp/63
- Mueller, S., Veinott, E., Hoffman, R., Klein, G., Alam, L., Mamun, T., & Clancey, W. (2020). *Principles of Explanation in Human-AI Systems*.
- Onifade, O., & Olanrewaju, S. (2020). Investigating performances of some statistical tests for heteroscedasticity assumption in generalized linear model: A monte carlo simulations study. *Open Journal of Statistics*, 10(3), 453–493. <https://doi.org/10.4236/ojs.2020.103029>
- Pai, K.-C., Su, S.-A., Chan, M.-C., Wu, C.-L., & Chao, W.-C. (2022). Explainable machine learning approach to predict extubation in critically ill ventilated patients: A retrospective study in central Taiwan. *BMC Anesthesiology*, 22(1), 351. <https://doi.org/10.1186/s12871-022-01888-y>
- Pavlidis, G. (2024). Unlocking the black box: Analysing the EU artificial intelligence act's framework for explainability in AI. *Law, Innovation and Technology*, 16(1), 293–308. <https://doi.org/10.1080/17579961.2024.2313795>
- Pizzi, G., Scarpi, D., & Pantano, E. (2021). Artificial intelligence and the new forms of interaction: Who has the control when interacting with a chatbot?. *Journal of Business Research*, 129, 878–890. <https://doi.org/10.1016/j.jbusres.2020.11.006>
- Polito, V., Barnier, A., & Woody, E. (2013). Developing the sense of agency rating scale (SOARS): An empirical measure of agency disruption in hypnosis. *Consciousness and Cognition*, 22(3), 684–696. <https://doi.org/10.1016/j.concog.2013.04.003>
- Posit Software, PBC. (2025). RStudio (Version 2025.05.0+496). [Computer software]. <https://posit.co/products/open-source/rstudio/?sid=1>
- Pramod, D., & Bafna, P. (2022). Conversational recommender systems techniques, tools, acceptance, and adoption: A state of the art review. *Expert Systems with Applications*, 203, 117539. <https://doi.org/10.1016/j.eswa.2022.117539>

- Pu, P., & Chen, L. (2007). Trust-inspiring explanation interfaces for recommender systems. *Knowledge-Based Systems*, 20(6), 542–556. <https://doi.org/10.1016/j.knosys.2007.04.004>
- Qiu, L., & Benbasat, I. (2009). Evaluating anthropomorphic product recommendation agents: A social relationship perspective to designing information systems. *Journal of Management Information Systems*, 25(4), 145–182. <https://doi.org/10.2753/MIS0742-1222250405>
- Rastogi, C., Zhang, Y., Wei, D., Varshney, K., Dhurandhar, A., & Tomsett, R. (2022). Deciding Fast and Slow: The Role of Cognitive Biases in AI-assisted Decision-making. *Proceedings of the ACM on Human-Computer Interaction* 6 (CSCWI) (pp. 1–22). <https://doi.org/10.1145/3512930>
- Samoili, S., López Cobo M., Gómez, E., De Prato, G., Martínez-Plumed, F., & Delipetrev, B. (2020). *AI watch: Defining artificial intelligence: Towards an operational definition and taxonomy of artificial intelligence*. Publications Office. <https://data.europa.eu/doi/10.2760/382730>
- Saßmannshausen, T., Burggräf, P., Hassenzahl, M., & Wagner, J. (2023). Human trust in otherware – a systematic literature review bringing all antecedents together. *Ergonomics*, 66(7), 976–998. <https://doi.org/10.1080/00140139.2022.2120634>
- Schoemann, A., Boulton, A., & Short, S. (2017a). Determining power and sample size for simple and complex mediation models. *Social Psychological and Personality Science*, 8(4), 379–386. <https://doi.org/10.1177/1948550617715068>
- Schoemann, A., Boulton, A., & Short, S. (2017b). Monte carlo power analysis for indirect effects. [Shiny web application] https://schoemanna.shinyapps.io/mc_power_med/
- Shrestha, N. (2020). Detecting multicollinearity in regression analysis. *American Journal of Applied Mathematics and Statistics*, 8(2), 39–42. <https://doi.org/10.12691/ajams-8-2-1>
- Sidlauskiene, J., Joye, Y., & Auruskeviciene, V. (2023). AI-based chatbots in conversational commerce and their effects on product and price perceptions. *Electronic Markets*, 33(1), 24. <https://doi.org/10.1007/s12525-023-00633-8>
- Silva, S., De Ciccio, R., Vlačić, B., & Elmashhara, M. (2023). Using chatbots in e-retailing – how to mitigate perceived risk and enhance the flow experience. *International Journal of Retail & Distribution Management*, 51(3), 285–305. <https://doi.org/10.1108/IJRDM-05-2022-0163>

- Singh, C., Dash, M., Sahu, R., & Kumar, A. (2024). Investigating the acceptance intentions of online shopping assistants in E-commerce interactions: Mediating role of trust and effects of consumer demographics. *Heliyon*, 10(3), e25031. <https://doi.org/10.1016/j.heliyon.2024.e25031>
- Siro, C., Aliannejadi, M., & De Rijke, M. (2024). Understanding and predicting user satisfaction with conversational recommender systems. *ACM Transactions on Information Systems*, 42(2), 1–37. <https://doi.org/10.1145/3624989>
- Solomonides, A., Koski, E., Atabaki, S., Weinberg, S., McGreevey, J., Kannry, J., Petersen, C., & Lehmann, C. (2022). Defining AMIA’s artificial intelligence principles. *Journal of the American Medical Informatics Association*, 29(4), 585–591. <https://doi.org/10.1093/jamia/ocac006>
- Tan, S.-M., & Liew, T. (2020). Designing embodied virtual agents as product specialists in a multi-product category e-commerce: The roles of source credibility and social presence. *International Journal of Human–Computer Interaction*, 36(12), 1136–1149. <https://doi.org/10.1080/10447318.2020.1722399>
- Tapal, A., Oren, E., Dar, R., & Eitam, B. (2017). The sense of agency scale: A measure of consciously perceived control over one’s mind, body, and the immediate environment. *Frontiers in Psychology*, 8, 1552. <https://doi.org/10.3389/fpsyg.2017.01552>
- Tran, T., Felfernig, A., & Tintarev, N. (2021). Humanized recommender systems: State-of-the-art and research issues. *ACM Transactions on Interactive Intelligent Systems*, 11(2), 1–41. <https://doi.org/10.1145/3446906>
- Tu, T., Palepu, A., Schaekermann, M., Saab, K., Freyberg, J., Tanno, R., Wang, A., Li, B., Amin, M., Tomasev, N., Azizi, S., Singhal, K., Cheng, Y., Hou, L., Webson, A., Kulkarni, K., Mahdavi, S. S., Semturs, C., Gottweis, J., Barral, J., Chou, K., Corrado, G., Matias, Y., Karthikesalingam A., & Natarajan, V. (2024). *Towards Conversational Diagnostic AI*. <https://doi.org/10.48550/arXiv.2401.05654>
- Tversky, A., & Kahneman, D. (1974). Judgment under uncertainty: Heuristics and biases. *Science*, 185(4157), 1124–1131. <https://doi.org/10.1126/science.185.4157.1124>
- Van Elk, M., Rutjens, B., & Van Der Pligt, J. (2015). The development of the illusion of control and sense of agency in 7- to-12-year old children and adults. *Cognition*, 145, 1–12. <https://doi.org/10.1016/j.cognition.2015.08.004>
- Vantrepotte, Q., Berberian, B., Pagliari, M., & Chambon, V. (2022). Leveraging human agency to improve confidence and acceptability in human-machine interactions. *Cognition*, 222, 105020. <https://doi.org/10.1016/j.cognition.2022.105020>

- Wang, W., & Benbasat, I. (2008). Attributions of Trust in Decision Support Technologies: A Study of Recommendation Agents for E-Commerce. *Journal of Management Information Systems*, 24(4), 249–273. <https://doi.org/10.2753/MIS0742-1222240410>
- Weiss, A., Burgmer, P., & Hofmann, W. (2022). The experience of trust in everyday life. *Current Opinion in Psychology*, 44, 245–251. <https://doi.org/10.1016/j.copsyc.2021.09.016>
- Weitz, K., Schiller, D., Schlagowski, R., Huber, T., & André, E. (2019). “Do you trust me?”: Increasing user-trust by integrating virtual agents in explainable AI interaction design. In H. Buschmeier, G. Lucas & S. Kopp (Eds.), *Proceedings of the 19th ACM International Conference on Intelligent Virtual Agents (IVA '19)* (pp. 7–9). Association for Computing Machinery. <https://doi.org/10.1145/3308532.3329441>
- Zhu, X., Xia, X., Wu, Y., & Zhao, W. (2024). Enhancing explainable recommendations: Integrating reason generation and rating prediction through multi-task learning. *Applied Sciences*, 14(18), 8303. <https://doi.org/10.3390/app14188303>

List of Figures

Figure 1. Illustration of the Conceptual Parallel Mediation Model	24
Figure 2. Illustration of the Baseline Recommender System Platform.....	25
Figure 3. Illustration of the Conversational Recommender System Platform	26
Figure 4. Results of the Mediation Model Path Diagram	37

List of Tables

Table 1 Descriptive Results of Survey Responses	35
Table 2 Overview of Hypotheses Test Results	38
Table 3 Influential-Cases Diagnostics: Path Coefficients With and Without High Influence Cases	40

Appendix

Abstract Deutsch

Als eine Anwendung von Künstlicher Intelligenz (KI) werden Conversational Recommender Systems (CRSs, „Dialog-basierte Empfehlungssysteme“) zunehmend in Online-Shopping Kontexten eingesetzt, um Nutzererfahrungen zu verbessern und dadurch Verkäufe zu steigern. Traditionell liegt der Evaluationsfokus primär auf technischen und funktionalen Aspekten. Darüber hinaus untersucht diese Studie den Zusammenhang von CRS-Interaktionen und Nutzer:innen Vertrauen. Insbesondere wird untersucht ob wahrgenommene Erklärbarkeit, Kontrollempfinden, und wahrgenommener Anthropomorphismus das Nutzervertrauen beeinflussen und ob Interaktionen mit CRSs im Vergleich zu nicht-Dialog-basierten Systemen einen positiven Effekt auf diese Einflussfaktoren haben. Diese Fragen wurden in einem Zwischensubjekt Experiment untersucht, bei dem 144 Teilnehmer:innen mit einem CRS oder einem baseline recommender system (BRS) auf einer simulierten E-Commerce-Plattform interagierten. Die zentralen Variablen wurden mittels Fragebogen erhoben. Die Ergebnisse der parallelen Mediationsanalyse zeigen, dass wahrgenommene Erklärbarkeit und Anthropomorphismus positiv mit Vertrauen assoziiert sind. Dabei gab es jedoch keine Gruppenunterschiede. Auffällig ist, dass Teilnehmende in der CRS-Bedingung über ein geringeres Kontrollempfinden berichteten, was einerseits auf ein Kontrollparadoxon oder auf Limitationen des Studiendesigns hinweisen kann. Keine der Mediationshypothesen war statistisch signifikant. Die Ergebnisse der Studie deuten darauf hin, dass Entwickler von Empfehlungssystemen – insbesondere von CRSs – auf eine Balance zwischen technischen und psychologisch fundierten Designelementen achten sollen. Implikationen für zukünftige Forschung sowie die Gestaltung von RSs unterstreichen daher die Bedeutung von interdisziplinärer Forschung, um sowohl technische als auch nutzerorientierte Aspekte zu optimieren.

Keywords: Conversational Recommender Systems, Vertrauen, Künstliche Intelligenz, E-Commerce

Abstract English

Conversational recommender systems (CRSs) – as an application of artificial intelligence (AI) – are increasingly used in e-commerce to improve user experiences and thereby enhance sales conversions. Going beyond accuracy measures typically employed in CRS evaluations, this thesis investigates how CRS interactions can foster user trust in comparison to non-conversational systems. More specifically, this research investigates how perceived explainability, sense of control, and perceived anthropomorphism influence user trust and if CRSs have a positive effect on these antecedents. This was examined by employing a between-participants experimental design in which 144 participants interacted with either a CRS or a baseline recommender system (BRS) in a simulated e-commerce platform. The main variables were assessed using self-report data gathered through an online survey. Results from parallel mediation analysis show that perceived explainability and perceived anthropomorphism are positively associated with higher user trust, however, this was independent of the type of recommender system (RS). Notably, participants interacting with the CRS reported lower sense of control, suggesting either a control paradox or limitations in the study design. Additionally, none of the mediation hypotheses was significant. The results suggest that recommender system (RS) developers, particularly CRS developers, should carefully balance the integration of psychologically founded design features. Implications for future research and RS development therefore include exploring design strategies that optimize both technical and user-oriented aspects, highlighting the importance of interdisciplinary research.

Keywords: Conversational Recommender Systems, Trust, Artificial Intelligence, E-commerce

Ethics and Open Science Statement

This research project underwent ethical review and was subsequently approved by the Departmental Review Board of the Department of Occupational, Economic, and Social Psychology of the University of Vienna under the processing number 2025/S/004. Members of the Departmental Review Board include Univ.-Prof. Dr. Robert Böhm, Univ.-Prof. Dr. Arnd Florack, Univ.-Prof. Dr. Veronika Job, and Univ.-Prof. Dr. Christian Korunka.

The preregistration of this study, anonymized survey data and R analysis scripts, as well as the decision letter of the ethical review and declaration of the use of artificial intelligence are publicly available via osf: <https://osf.io/62m9v/files/osfstorage>.

Limesurvey Study Introduction

15/06/2025, 18:36

CDL RecSys - Studie zu digitalen Empfehlungssystemen

Studie zu digitalen Empfehlungssystemen

Herzlich Willkommen!

Liebe_r Studienteilnehmer_in,

vielen Dank für Ihr Interesse an unserer Studie zu digitalen Empfehlungssystemen!

Diese Studie wird vom **Christian Doppler Labor für Recommender Systems (CDL-RecSys)** an der TU Wien im Rahmen eines Forschungsprojekts durchgeführt. Im Zuge dieses Projekts wird eine Masterarbeit an der Fakultät für Psychologie der Universität Wien verfasst.

Bitte lesen Sie die nachfolgenden Informationen zum Datenschutz sorgfältig durch.

Nachdem Sie den Text gelesen haben, können Sie durch Ankreuzen der Checkbox Ihre Einwilligung zur Teilnahme an der Studie geben oder diese ablehnen.

In dieser Umfrage sind 4 Fragen enthalten.

Einverständniserklärung zur Teilnahme an der Studie

Bitte bestätigen Sie hiermit, dass Sie freiwillig an dieser Studie teilnehmen möchten.

Diese Studie besteht aus folgenden Teilen:

- Eingangsfragebogen (LimeSurvey)
- Interaktion mit einem webbasierten Forschungsprototyp
- Abschlussfragebogen (LimeSurvey)

Wichtige Hinweise:

- Die Teilnahme ist **freiwillig**. Sie können die Studie jederzeit und ohne Angabe von Gründen abbrechen.
- Ihre Teilnahme erfolgt **anonym**; Ihre Antworten und Nutzungsdaten sind nicht auf Ihre Person rückführbar.
- Für Studierende der Lehrveranstaltung „Recommender Systems“ gibt es die Möglichkeit, Bonuspunkte zu erhalten. Zu diesem Zweck wird im Abschlussfragebogen Ihre E-Mail-Adresse abgefragt und ausschließlich für die Zuordnung der Bonuspunkte verwendet. Nach Abschluss der LV-Bewertung wird diese Adresse vollständig gelöscht.
- Nach Abschluss der Studie werden **alle erhobenen Daten anonymisiert** und für wissenschaftliche Zwecke verwendet sowie unter einer offenen Lizenz (CC-BY) veröffentlicht. Ihre Anonymität bleibt hierbei vollständig gewahrt.

Kontakt für Rückfragen zur Studie:

- a11812688@unet.univie.ac.at
- recsys-contact@ec.tuwien.ac.at

Bitte bestätigen Sie nun durch Ankreuzen der Checkbox, dass Sie die oben stehenden Informationen gelesen haben und an der Studie teilnehmen möchten:

Möchten Sie an dieser Studie teilnehmen?

*

Bitte wählen Sie nur eine der folgenden Antworten aus:

☐ Ja

☐ Nein

Anleitung

Bitte lesen Sie die folgende Anleitung sorgfältig durch, um erfolgreich an der Studie teilzunehmen.

Sobald Sie auf „Weiter“ klicken, werden Sie **automatisch zur Website des Forschungsprototyps weitergeleitet**, auf der die eigentliche Studie beginnt.

<https://survey.ec.tuwien.ac.at/index.php?r=admin/printablesurvey/sa/index/surveyid/115986/lang/de>

1/2

Wichtig: Diese Studie ist nur auf einem Laptop oder Desktop-Computer durchführbar. **Eine Teilnahme über ein Smartphone oder Tablet ist leider nicht möglich.** Bitte stellen Sie sicher, dass Sie einen Computer verwenden. Um die uneingeschränkte Funktionalität des Systems zu gewährleisten, achten Sie bitte außerdem darauf, dass in Ihrem **Browser JavaScript aktiviert ist und Sie das Tracking für die Studie nicht automatisch blockieren.**

Aufgabe: Die Website, auf die Sie weitergeleitet werden, ist eine **Online-Shopping-Plattform für den Kauf von Waschmaschinen**. Stellen Sie sich vor, Ihre eigene Waschmaschine ist vor Kurzem defekt geworden und Sie sind nun auf der Suche nach einem Ersatz für Ihr Zuhause. Nutzen Sie die Plattform, um eine passende Waschmaschine zu finden. Verwenden Sie die Website so, wie Sie es auch bei einer echten Kaufentscheidung tun würden, damit Sie zu einer fundierten Auswahl kommen.

Am Ende der Nutzung der Plattform werden Sie gefragt, ob Sie sich für den Kauf einer Waschmaschine entscheiden würden. Dabei handelt es sich lediglich um eine Befragung – Sie lösen damit keinen Kauf aus und werden zu keinem tatsächlichen Onlineshop weitergeleitet.

Anschließend werden Sie automatisch zu einem kurzen Abschlussfragebogen weitergeleitet. Bitte beantworten Sie auch diesen sorgfältig und vollständig, da Ihre Antworten für unsere Forschung von großer Bedeutung sind.

Wenn Sie die Anleitung gelesen haben und bereit sind, klicken Sie jetzt auf „Weiter“, um direkt mit der Studie zu beginnen.

User ID *

Your assigned group is: {groupAssignment}

10.06.2025 – 13:45

Senden Sie Ihre Umfrage ein.

Vielen Dank für die Beantwortung des Fragebogens.

Limesurvey Exit Survey

15/06/2025, 18:28

CDL RecSys - Studie zu digitalen Empfehlungssystemen (cont.)

Studie zu digitalen Empfehlungssystemen (cont.)

In dieser Umfrage sind 19 Fragen enthalten.

System Wahrnehmung

Nun folgen ein paar Fragen zu Ihrer Wahrnehmung der Website, mit der Sie interagiert haben.

Die Website wird folgend als "System" bezeichnet.

Bitte geben Sie auf einer Skala von 1 bis 5 an, inwieweit Sie den folgenden Aussagen zustimmen (1 = stimme gar nicht zu, 5 = stimme voll zu).

*

Bitte wählen Sie die zutreffende Antwort für jeden Punkt aus:

	1 (stimme gar nicht zu)	2 (stimme wenig zu)	3 (weder noch)	4 (stimme etwas zu)	5 (stimme voll zu)
Ich vertraue dem System.	☐	☐	☐	☐	☐
Ich verstehe, wie das System funktioniert.	☐	☐	☐	☐	☐
Das System wirkt menschlich.	☐	☐	☐	☐	☐
Ich habe Kontrolle darüber, wie das System handelt.	☐	☐	☐	☐	☐
Ich habe den Eindruck, dass das System meine Antworten insgesamt korrekt übernommen hat.	☐	☐	☐	☐	☐

Demographische Daten

Im Folgenden bitten wir Sie noch um ein paar Angaben Ihrer persönlichen Daten.

Bitte geben Sie Ihr Alter an. *

Bitte geben Sie Ihre Antwort hier ein:

Welchem Geschlecht fühlen Sie sich zugehörig? *

Bitte wählen Sie nur eine der folgenden Antworten aus:

- ☐ Weiblich
☐ Männlich
☐ Divers
☐ Keine Angabe

Sind Sie aktuell Student*in? *

Bitte wählen Sie nur eine der folgenden Antworten aus:

- ☐ Ja
☐ Nein

Wenn Sie diese Studie im Rahmen der Lehrveranstaltung Recommender Systems der TU Wien durchführen, geben Sie bitte Ihre Matrikelnummer an, um die Zusatzpunkte für Ihre Bewertung zu erhalten.

Bitte geben Sie Ihre Antwort hier ein:

Abgesehen von einem Studium, welche der folgenden Kategorien beschreibt am besten Ihren Beschäftigungsstatus? *

Bitte wählen Sie nur eine der folgenden Antworten aus:

- ☐ Teilzeit
☐ Vollzeit
☐ Selbstständig
☐ Nicht angestellt, auf der Suche nach Arbeit
☐ Nicht angestellt, nicht auf der Suche nach Arbeit
☐ Pension
☐ Keine Angabe
☐ Sonstiges

Bitte geben Sie hier Ihre E-Mail Adresse an:

Bitte geben Sie Ihre Antwort hier ein:

Welchen höchsten Bildungsgrad haben Sie erreicht? *

Bitte wählen Sie nur eine der folgenden Antworten aus:

- ☐ Pflichtschulabschluss
- ☐ Lehre
- ☐ Matura
- ☐ Hochschulabschluss - Bachelor
- ☐ Hochschulabschluss - Master / Magister / Diplom
- ☐ Doktorat / PhD
- ☐ Keine Angabe
- ☐ Sonstiges

Wie würden Sie Ihre Deutschkenntnisse einschätzen? *

Bitte wählen Sie nur eine der folgenden Antworten aus:

- ☐ Keine Kenntnisse
- ☐ A1 (Anfänger)
- ☐ A2 (Grundlagenkenntnisse)
- ☐ B1 (Fortgeschritten)
- ☐ B2 (Selbstständige Sprachverwendung)
- ☐ C1 (Fachkundige Sprachkenntnisse)
- ☐ C2 (Muttersprachlich)

Ich beschäftige mich gern genauer mit technischen Systemen. *

Bitte wählen Sie nur eine der folgenden Antworten aus:

- ☐ Stimme gar nicht zu
- ☐ Stimme eher nicht zu
- ☐ Weder noch
- ☐ Stimme eher zu
- ☐ Stimme ganz zu

(Online) Shopping Verhalten

Zuletzt bitten wir Sie noch, folgende Fragen zu Ihrem (Online) Shoppingverhalten zu beantworten.

Wo kaufen Sie vorrangig Produkte außerhalb des alltäglichen Bedarfs (z.B. Elektrogeräte, Kleidung)? *

Bitte wählen Sie alle zutreffenden Antworten aus:

- ☐ In Läden / Stores
- ☐ Online, auf Seiten des Herstellers (z.B. Samsung, H&M)
- ☐ Online, auf Handelsplattformen (z.B. Amazon, Zalando)
- ☐ Online Second-Hand Plattformen
- ☐ Second-Hand Läden
- ☐ Sonstige
- ☐ Keine Angabe

Wo starten Sie den Online-Einkauf? *

Bitte wählen Sie alle zutreffenden Antworten aus:

- ☐ Google
- ☐ Amazon
- ☐ Preisvergleichsseite
- ☐ Sonstige
- ☐ Keine Angabe

Nutzen Sie Preisvergleichsplattformen? *

Bitte wählen Sie nur eine der folgenden Antworten aus:

- ☐ Ja
- ☐ Nein

Welche Preisvergleichsplattform nutzen Sie am häufigsten? *

Bitte wählen Sie nur eine der folgenden Antworten aus:

- ☐ Geizhals
- ☐ Idealo
- ☐ Google Shopping
- ☐ Sonstige
- ☐ keine Angabe

Wie oft nutzen Sie Preisvergleichsplattformen? *

Bitte wählen Sie nur eine der folgenden Antworten aus:

- ☐ mehrmals wöchentlich
☐ einmal im Monat
☐ einmal im Jahr
☐ seltener
☐ nie

Auf welchem Endgerät nutzen Sie Preisvergleichsplattformen? *

Bitte wählen Sie alle zutreffenden Antworten aus:

- ☐ Desktop
☐ Mobiltelefon
☐ Tablet
☐ Sonstige
☐ Keine Angabe

Auf welcher Plattform nutzen Sie Preisvergleichservices? *

Bitte wählen Sie alle zutreffenden Antworten aus:

- ☐ Desktop Browser
☐ Mobil Browser
☐ Mobil App
☐ keine Angabe

Mail Updates

Vielen Dank für Ihre Teilnahme an unserer Studie.

Falls Sie in Zukunft über Neuigkeiten zu den Ergebnissen der Studie informiert werden möchten, können Sie eine Referenz-E-Mail-Adresse angeben. Ebenfalls können Sie uns optional gerne allgemeines Feedback zu unserer Studie oder dem System hinterlassen. Andernfalls fahren Sie bitte mit „Absenden“ fort.

Bitte geben Sie hier Ihre E-Mail Adresse an:

Bitte geben Sie Ihre Antwort hier ein:

15/06/2025, 18:30

CDL RecSys - Studie zu digitalen Empfehlungssystemen (cont.)

Allgemeines Feedback:

Bitte geben Sie Ihre Antwort hier ein:

10.06.2025 – 13:45

Senden Sie Ihre Umfrage ein.

Vielen Dank für die Beantwortung des Fragebogens.

R Script and Results of Data Analysis

1. Load Data and Packages

```
# Packages needed for data import, cleaning, analysis, tables, and plots

library(readxl)

Study_Data <- read_excel("/Users/lauramodre/Documents/MA /Masterarbeit/Master_thesis_R/Study Results.xlsx")

View(Study_Data)

library(readr)
library(dplyr)
library(tidyr)
library(psych)
library(lavaan)
library(ggplot2)
library(apaTables)
library(effsize)
library(ggpubr)
library(car)
library(processR)
library(janitor)
library(knitr)
library(kableExtra)
library(stringr)
library(lmtest)
```

2. Data Inspection and Cleaning

```
# 1. Number of observations (study participants, rows) and variables (columns)
cat("Initial number of observations:", nrow(Study_Data), "\n")
## Initial number of observations: 145
cat("Number of variables:", ncol(Study_Data), "\n")
## Number of variables: 42

# 2. Summary
summary(Study_Data)

##   respondent_ID      rs_type      system_trust      system_understanding
##   Min.      : 81.0    Min.      :0.0000    Length:145    Length:145
##   1st Qu.:136.0    1st Qu.:0.0000    Class :character    Class :character
##   Median :187.0    Median :1.0000    Mode  :character    Mode  :character
```

```

## Mean :187.5 Mean :0.5172
## 3rd Qu.:241.0 3rd Qu.:1.0000
## Max. :289.0 Max. :1.0000
## system_anthropomorphism system_control system_accuracy age
## Length:145 Length:145 Length:145 Min. :21.00
## Class :character Class :character Class :character 1st Qu.:24.00
## Mode :character Mode :character Mode :character Median :26.00
## Mean :28.74
## 3rd Qu.:29.00
## Max. :65.00
## gender student_status work work_other
## Length:145 Length:145 Length:145 Length:145
## Class :character Class :character Class :character Class :character
## Mode :character Mode :character Mode :character Mode :character
##
##
##
## highest_education education_other german_level technology_affinity
## Length:145 Length:145 Length:145 Length:145
## Class :character Class :character Class :character Class :character
## Mode :character Mode :character Mode :character Mode :character
##
##
##
## buy_in_stores buy_online_seller buy_online_platform
## Length:145 Length:145 Length:145
## Class :character Class :character Class :character
## Mode :character Mode :character Mode :character
##
##
##
## buy_online_secondhand buy_secondhand_shop buy_other
## Length:145 Length:145 Length:145
## Class :character Class :character Class :character
## Mode :character Mode :character Mode :character
##
##
##
## buy_n/A start_Google start_Amazon

```

```

## Length:145          Length:145          Length:145
## Class :character    Class :character    Class :character
## Mode :character     Mode :character     Mode :character
##
##
##
## start_price_comparison start_other          start_N/A
## Length:145          Length:145          Length:145
## Class :character     Class :character    Class :character
## Mode :character      Mode :character     Mode :character
##
##
##
## price_comparison_use price_comparison_site comparison_usage_frequency
## Length:145          Length:145          Length:145
## Class :character     Class :character    Class :character
## Mode :character      Mode :character     Mode :character
##
##
##
## comparison_site_desktop comparison_site_phone price_comparison_tablet
## Length:145          Length:145          Length:145
## Class :character     Class :character    Class :character
## Mode :character      Mode :character     Mode :character
##
##
##
## price_comparison_tool_other price_comparison_tool_N/A
## Length:145          Length:145
## Class :character     Class :character
## Mode :character      Mode :character
##
##
##
## price_comparison_desktop_browser price_comparison_mobile_browser
## Length:145          Length:145
## Class :character     Class :character
## Mode :character      Mode :character
##

```

```
##
##
## price_comparison_mobile_app price_comparison_platform_N/A
## Length:145 Length:145
## Class :character Class :character
## Mode :character Mode :character
##
##
##
## Allgemeines Feedback: Gesamtzeit
## Length:145 Min. : 64.74
## Class :character 1st Qu.: 123.31
## Mode :character Median : 175.97
## Mean : 223.05
## 3rd Qu.: 272.16
## Max. :1804.80
# 3. Missing-value check
colSums(is.na(Study_Data)) #Counts missing per column
## respondent_ID rs_type
## 0 0
## system_trust system_understanding
## 0 0
## system_anthropomorphism system_control
## 0 0
## system_accuracy age
## 70 0
## gender student_status
## 0 0
## work work_other
## 0 140
## highest_education education_other
## 0 144
## german_level technology_affinity
## 0 0
## buy_in_stores buy_online_seller
## 0 0
## buy_online_platform buy_online_secondhand
## 0 0
## buy_secondhand_shop buy_other
```

```
##          0          0
##          buy_n/A          start_Google
##          0          0
##          start_Amazon          start_price_comparison
##          0          0
##          start_other          start_N/A
##          0          0
##          price_comparison_use          price_comparison_site
##          0          39
##          comparison_usage_frequency          comparison_site_desktop
##          39          0
##          comparison_site_phone          price_comparison_tablet
##          0          0
##          price_comparison_tool_other          price_comparison_tool_N/A
##          0          0
## price_comparison_desktop_browser price_comparison_mobile_browser
##          0          0
##          price_comparison_mobile_app          price_comparison_platform_N/A
##          0          0
##          Allgemeines Feedback:          Gesamtzeit
##          117          0
# 4. Remove time under 80s (inclusion criteria) (ID 198)
Study_Data_Clean <- Study_Data %>%
  filter(respondent_ID != 198)

cat("Observations removed:", nrow(Study_Data) - nrow(Study_Data_Clean), "\n")
## Observations removed: 1
cat("Final sample size:", nrow(Study_Data_Clean), "\n") #Show final sample size
## Final sample size: 144
# 5. Update the main dataset
Study_Data <- Study_Data_Clean
```

3. Variable Preparation

```
# Defining Key Variables

key_variables <- c("rs_type", "system_trust", "system_understanding", "system_anthr
opomorphism", "system_control", "system_accuracy")

# Extracting numerical values from Likert responses and saving them as numbers
```



```

extract_numeric <- function(x) {
  case_when(
    is.na(x) ~ NA_real_,
    is.numeric(x) ~ x,
    TRUE ~ as.numeric(str_extract(x, "^\\d+"))
  )
}

Study_Data <- Study_Data %>%
mutate(
  across(
    .cols = c("system_trust", "system_understanding", "system_anthropomorphism", "system_control", "system_accuracy"),
    .fns = ~ as.numeric(str_extract(.x, "^\\d+")),
    .names = "{.col}")
  )

# Converting Likert responses to technology affinity to numerical values
Study_Data <- Study_Data %>%
  mutate(
    across(
      .cols = contains("technology_affinity"),
      .fns = ~ case_when(
        str_detect(.x, regex("stimme gar nicht zu", ignore_case = TRUE)) ~ 1,
        str_detect(.x, regex("stimme eher nicht zu", ignore_case = TRUE)) ~ 2,
        str_detect(.x, regex("weder noch", ignore_case = TRUE)) ~ 3,
        str_detect(.x, regex("stimme eher zu", ignore_case = TRUE)) ~ 4,
        str_detect(.x, regex("stimme ganz zu", ignore_case = TRUE)) ~ 5,
        TRUE ~ NA_real_
      ),
      .names = "{.col}"
    ),

# Converting German responses to English
# Yes/No responses
    across(
      .cols = c("student_status", "buy_in_stores", "buy_online_seller", "buy_online_platform", "buy_online_secondhand", "buy_secondhand_shop", "buy_other", "buy_n/A", "start_Google", "start_Amazon", "start_price_comparison", "start_other", "start_N/A", "price_comparison_use", "comparison_site_desktop", "comparison_site_phone", "price_comparison_tablet", "price_comparison_tool_other", "price_comparison_tool_N/

```

```

A", "price_comparison_desktop_browser", "price_comparison_mobile_browser", "price_
_comparison_mobile_app", "price_comparison_platform_N/A"),

.fns = ~ case_when(
  grepl("^ja$", .x, ignore.case = TRUE) ~ "yes",
  grepl("^nein$", .x, ignore.case = TRUE) ~ "no",
  TRUE ~ .x
)
),
#Gender
across(
  .cols = c("gender"),
  .fns = ~ case_when(
    grepl("^Weiblich$", .x, ignore.case = TRUE) ~ "female",
    grepl("^Männlich$", .x, ignore.case = TRUE) ~ "male",
    TRUE ~ .x
  )
),
#Price Comparison Site Usage Frequency
across(
  .cols = c("comparison_usage_frequency"),
  .fns = ~ case_when(
    grepl("^einmal im Jahr$", .x, ignore.case = TRUE) ~ "once/year",
    grepl("^einmal im Monat$", .x, ignore.case = TRUE) ~ "once/month",
    grepl("^mehrmals wöchentlich$", .x, ignore.case = TRUE) ~ "multiple times/wee
k",
    grepl("^seltener$", .x, ignore.case = TRUE) ~ "rarely",
    grepl("^nie$", .x, ignore.case = TRUE) ~ "never",
    TRUE ~ .x
  )
)
)
)

```

4. Sample Description

```

#Number of Participants
nrow(Study_Data)
## [1] 144

# Participant Age
Study_Data$age <- as.numeric(Study_Data$age)
describe(Study_Data$age)

# Frequency Table for Demographics

```

```

table(Study_Data$gender)
##
## female    male
##      44     100

table(Study_Data$student_status) #Student = yes, Non-Student = no
##
## no yes
##  36 108

table(Study_Data$work)
##
##                               Keine Angabe
##                               3
##      Nicht angestellt, auf der Suche nach Arbeit
##                               18
## Nicht angestellt, nicht auf der Suche nach Arbeit
##                               10
##                               Selbstständig
##                               15
##                               Sonstiges
##                               5
##                               Teilzeit
##                               64
##                               Vollzeit
##                               29

table(Study_Data$highest_education)
##
##                               Doktorat / PhD
##                               2
##      Hochschulabschluss - Bachelor
##                               94
## Hochschulabschluss - Master / Magister / Diplom
##                               29
##                               Lehre
##                               4
##                               Matura
##                               14
##                               Sonstiges
##                               1

table(Study_Data$german_level)

```

```
##
## B2 (Selbstständige Sprachverwendung)      C1 (Fachkundige Sprachkenntnisse)
##                                           7                               10
##                C2 (Muttersprachlich)
##                                           127
```

5. Assumption Testing for Mediation Analysis

```
# Definition of key variables for mediation
X <- "rs_type"
Y <- "system_trust"
M1 <- "system_understanding"
M2 <- "system_anthropomorphism"
M3 <- "system_control"

# Assumption 1: Linearity
# Regression 1: X predicting Y (total effect c)

if(all(c(X, Y) %in% names(Study_Data))) {
  reg1 <- lm(Study_Data[[Y]] ~ Study_Data[[X]])

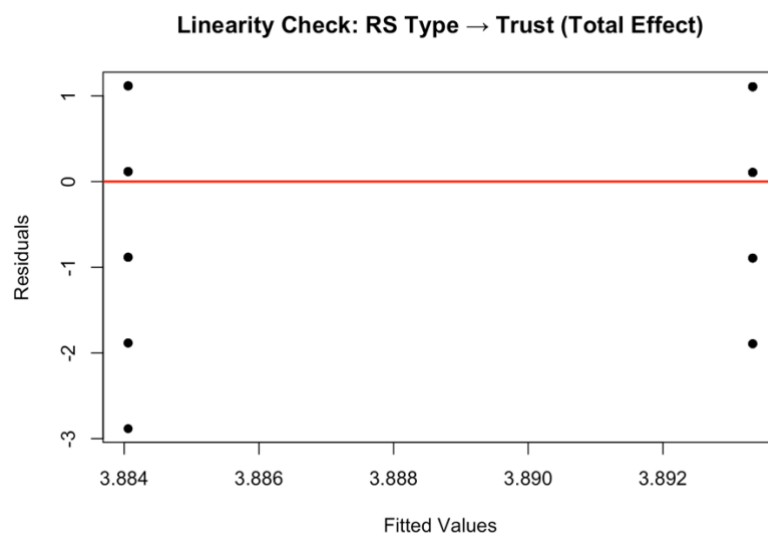
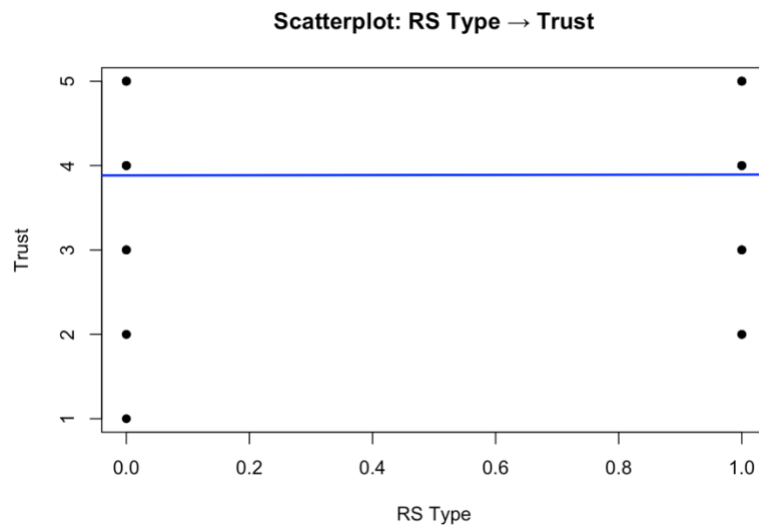
  print("Model Summary:")
  print(summary(reg1))

  # Scatterplot with regression line
  plot(Study_Data[[X]], Study_Data[[Y]],
       main = "Scatterplot: RS Type → Trust",
       xlab = "RS Type", ylab = "Trust",
       pch = 16)
  abline(reg1, col = "blue", lwd = 2)

  # Residuals vs Fitted plot
  plot(fitted(reg1), residuals(reg1),
       main = "Linearity Check: RS Type → Trust (Total Effect)",
       xlab = "Fitted Values",
       ylab = "Residuals",
       pch = 16)
  abline(h = 0, col = "red", lwd = 2)}

## [1] "Model Summary:"
##
```

```
## Call:
## lm(formula = Study_Data[[Y]] ~ Study_Data[[X]])
##
## Residuals:
##      Min       1Q   Median       3Q      Max
## -2.8841 -0.8841  0.1067  1.1067  1.1159
##
## Coefficients:
##              Estimate Std. Error t value Pr(>|t|)
## (Intercept)    3.884058   0.116166  33.435  <2e-16 ***
## Study_Data[[X]] 0.009275   0.160964   0.058   0.954
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Residual standard error: 0.9649 on 142 degrees of freedom
## Multiple R-squared:  2.338e-05, Adjusted R-squared: -0.007019
## F-statistic: 0.003321 on 1 and 142 DF, p-value: 0.9541
```



```
# Regression 2: X predicting M1 (a1 path)

if(all(c(X, M1) %in% names(Study_Data))) {
  reg2 <- lm(Study_Data[[M1]] ~ Study_Data[[X]])

  print("Model Summary:")
  print(summary(reg2))

  # Scatterplot with regression line
  plot(Study_Data[[X]], Study_Data[[M1]],
       main = "Scatterplot: RS Type → Understanding",
       xlab = "RS Type", ylab = "Understanding",
       pch = 16)
  abline(reg2, col = "blue", lwd = 2)
```

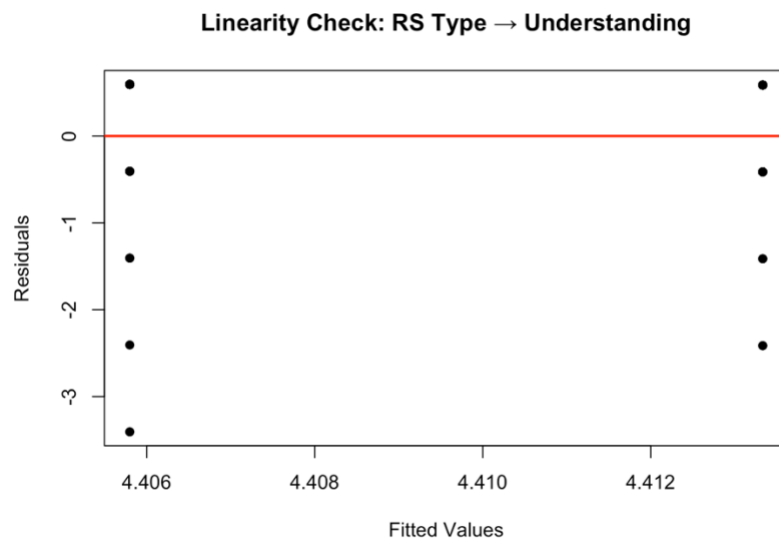
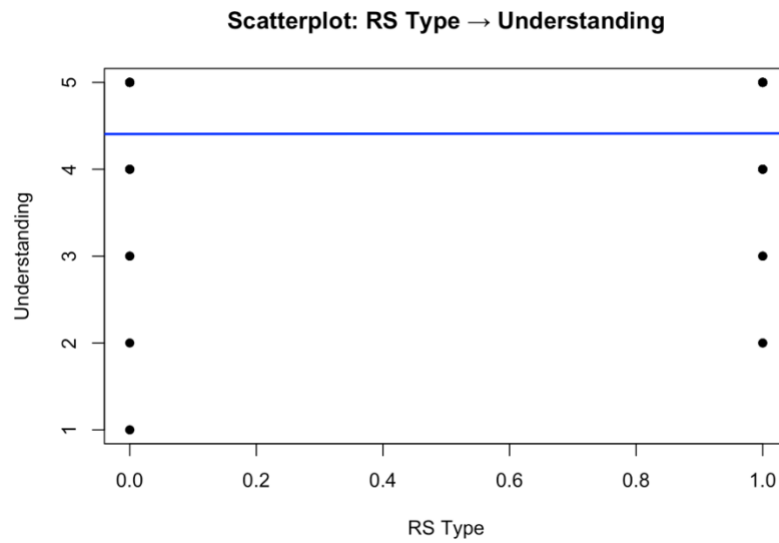
```

# Residuals vs Fitted plot
plot(fitted(reg2), residuals(reg2),
     main = "Linearity Check: RS Type → Understanding",
     xlab = "Fitted Values",
     ylab = "Residuals",
     pch = 16)

abline(h = 0, col = "red", lwd = 2)}

## [1] "Model Summary:"
##
## Call:
## lm(formula = Study_Data[[M1]] ~ Study_Data[[X]])
##
## Residuals:
##      Min       1Q   Median       3Q      Max
## -3.4058 -0.4133  0.5867  0.5942  0.5942
##
## Coefficients:
##              Estimate Std. Error t value Pr(>|t|)
## (Intercept)    4.405797   0.109194  40.35  <2e-16 ***
## Study_Data[[X]] 0.007536   0.151303   0.05   0.96
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Residual standard error: 0.907 on 142 degrees of freedom
## Multiple R-squared:  1.747e-05, Adjusted R-squared:  -0.007025
## F-statistic: 0.002481 on 1 and 142 DF,  p-value: 0.9603

```



```
# Regression 3: X predicting M2 (a2 path)

if(all(c(X, M2) %in% names(Study_Data))) {
  reg3 <- lm(Study_Data[[M2]] ~ Study_Data[[X]])

  print("Model Summary:")
  print(summary(reg3))

  # Scatterplot with regression line
  plot(Study_Data[[X]], Study_Data[[M2]],
       main = "Scatterplot: RS Type → Anthropomorphism",
       xlab = "RS Type", ylab = "Anthropomorphism",
       pch = 16)
```



```

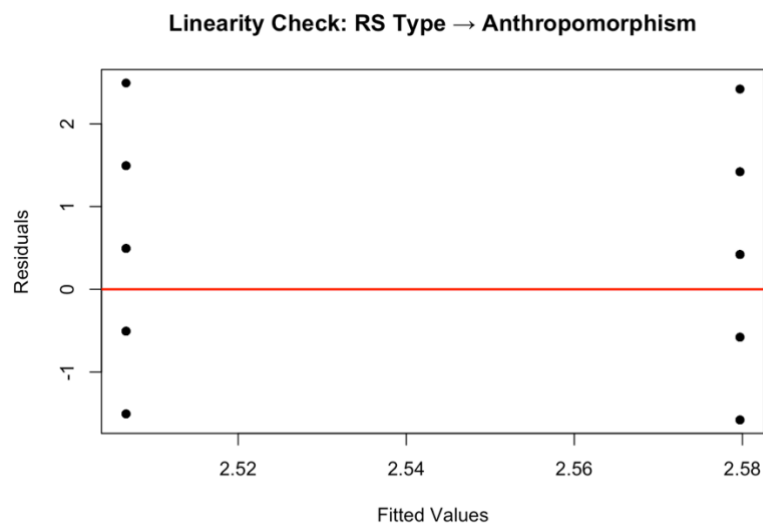
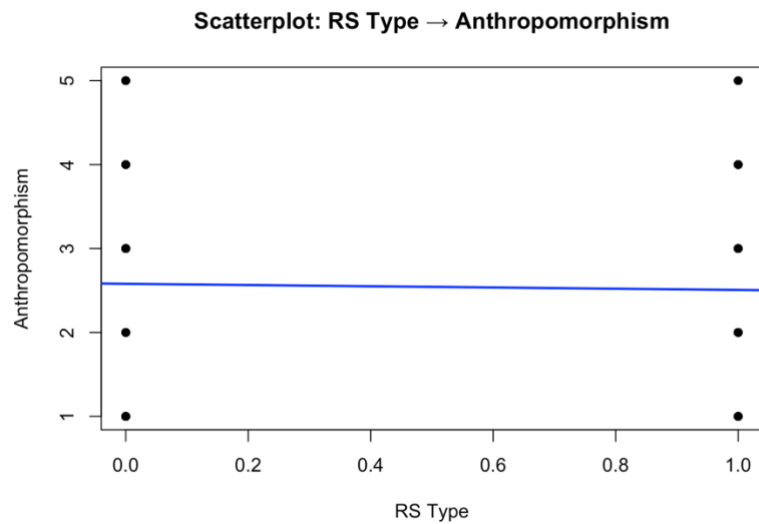
abline(reg3, col = "blue", lwd = 2)

# Residuals vs Fitted plot
plot(fitted(reg3), residuals(reg3),
     main = "Linearity Check: RS Type → Anthropomorphism",
     xlab = "Fitted Values",
     ylab = "Residuals",
     pch = 16)

abline(h = 0, col = "red", lwd = 2)}

## [1] "Model Summary:"
##
## Call:
## lm(formula = Study_Data[[M2]] ~ Study_Data[[X]])
##
## Residuals:
##      Min       1Q   Median       3Q      Max
## -1.57971 -0.81145 -0.04319  0.49333  2.49333
##
## Coefficients:
##              Estimate Std. Error t value Pr(>|t|)
## (Intercept)    2.57971    0.14694  17.556  <2e-16 ***
## Study_Data[[X]] -0.07304    0.20361  -0.359    0.72
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Residual standard error: 1.221 on 142 degrees of freedom
## Multiple R-squared:  0.0009055, Adjusted R-squared:  -0.00613
## F-statistic: 0.1287 on 1 and 142 DF, p-value: 0.7203

```



```
# Regression 4: X predicting M3 (a3 path)
print("4. REGRESSION: X → M3 (Control)")
## [1] "4. REGRESSION: X → M3 (Control)"
print("-----")
## [1] "-----"

if(all(c(X, M3) %in% names(Study_Data))) {
  reg4 <- lm(Study_Data[[M3]] ~ Study_Data[[X]])

  print("Model Summary:")
  print(summary(reg4))

  # Scatterplot with regression line
  plot(Study_Data[[X]], Study_Data[[M3]],
       main = "Scatterplot: RS Type → Control",
```

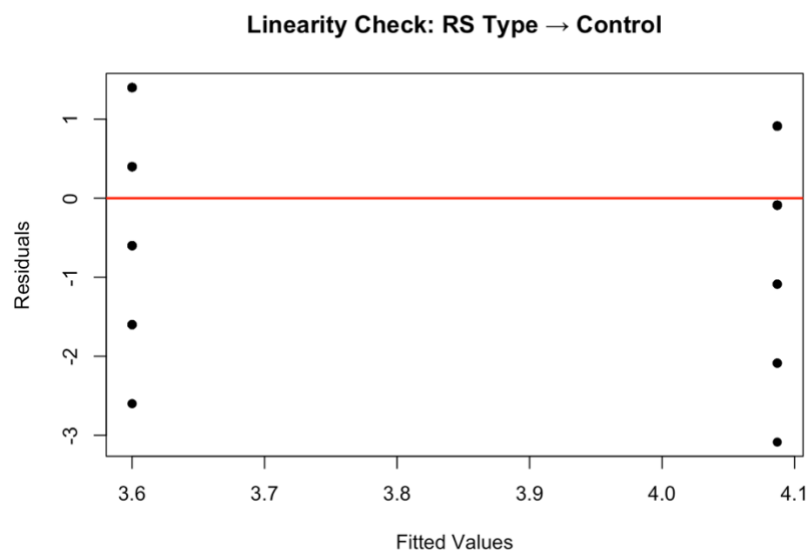
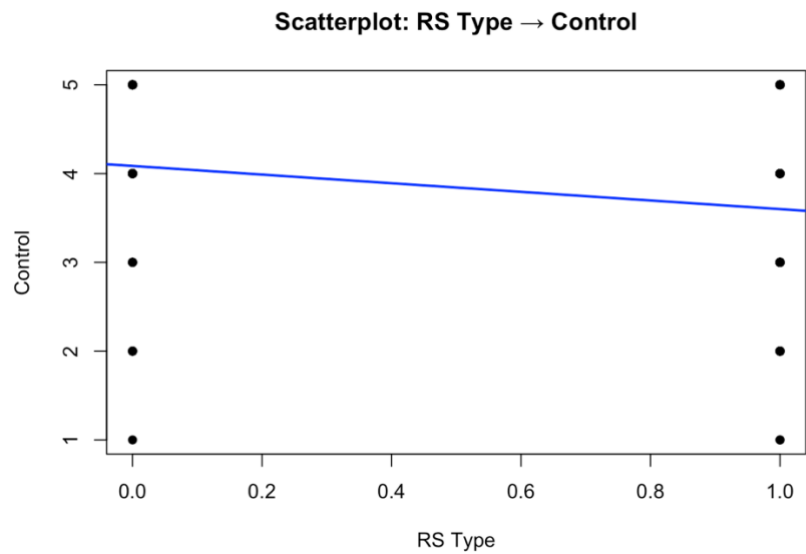
```

      xlab = "RS Type", ylab = "Control",
      pch = 16)
abline(reg4, col = "blue", lwd = 2)

# Residuals vs Fitted plot
plot(fitted(reg4), residuals(reg4),
     main = "Linearity Check: RS Type → Control",
     xlab = "Fitted Values",
     ylab = "Residuals",
     pch = 16)
abline(h = 0, col = "red", lwd = 2)}

## [1] "Model Summary:"
##
## Call:
## lm(formula = Study_Data[[M3]] ~ Study_Data[[X]])
##
## Residuals:
##      Min       1Q   Median       3Q      Max
## -3.08696 -0.60000 -0.08696  0.91304  1.40000
##
## Coefficients:
##              Estimate Std. Error t value Pr(>|t|)
## (Intercept)      4.0870     0.1243  32.870 < 2e-16 ***
## Study_Data[[X]] -0.4870     0.1723  -2.826  0.00539 **
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Residual standard error: 1.033 on 142 degrees of freedom
## Multiple R-squared:  0.05326,    Adjusted R-squared:  0.04659
## F-statistic: 7.989 on 1 and 142 DF,  p-value: 0.005387

```



```
# Regression 5: M1 predicting Y (b1 path)

if(all(c(M1, Y) %in% names(Study_Data))) {
  reg5 <- lm(Study_Data[[Y]] ~ Study_Data[[M1]])

  print("Model Summary:")
  print(summary(reg5))

  # Scatterplot with regression line
  plot(Study_Data[[M1]], Study_Data[[Y]],
       main = "Scatterplot: Understanding → Trust",
       xlab = "Understanding", ylab = "Trust",
       pch = 16)
```

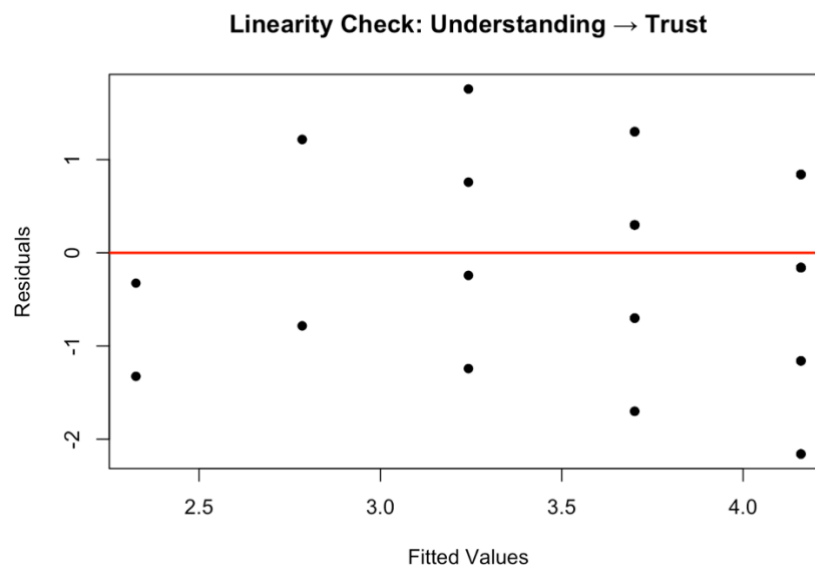
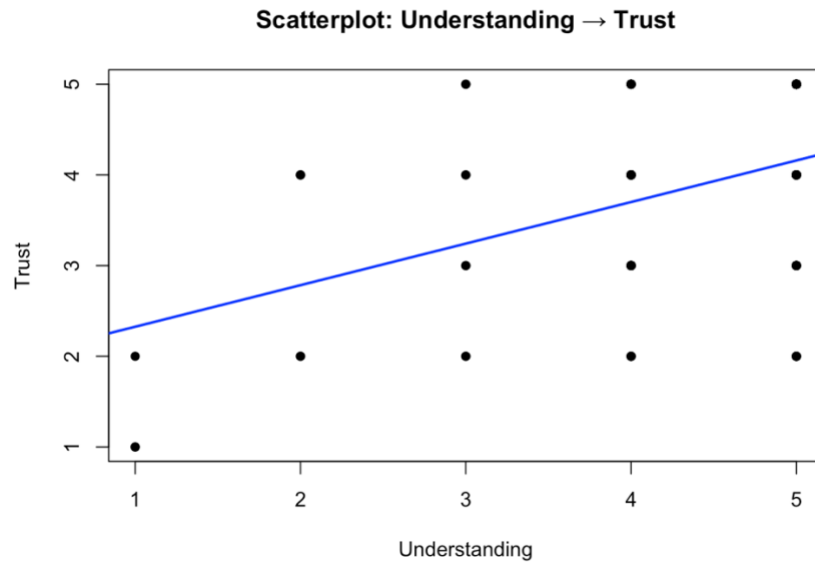
```

abline(reg5, col = "blue", lwd = 2)

# Residuals vs Fitted plot
plot(fitted(reg5), residuals(reg5),
     main = "Linearity Check: Understanding → Trust",
     xlab = "Fitted Values",
     ylab = "Residuals",
     pch = 16)

abline(h = 0, col = "red", lwd = 2)}
## [1] "Model Summary:"
##
## Call:
## lm(formula = Study_Data[[Y]] ~ Study_Data[[M1]])
##
## Residuals:
##      Min       1Q   Median       3Q      Max
## -2.1595 -0.2426 -0.1595  0.8405  1.7574
##
## Coefficients:
##              Estimate Std. Error t value Pr(>|t|)
## (Intercept)      1.86738    0.36260     5.15 8.54e-07 ***
## Study_Data[[M1]]  0.45842    0.08056     5.69 6.99e-08 ***
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Residual standard error: 0.8708 on 142 degrees of freedom
## Multiple R-squared:  0.1857, Adjusted R-squared:  0.1799
## F-statistic: 32.38 on 1 and 142 DF,  p-value: 6.995e-08

```



```
# Regression 6: M2 predicting Y (b2 path)

if(all(c(M2, Y) %in% names(Study_Data))) {
  reg6 <- lm(Study_Data[[Y]] ~ Study_Data[[M2]])

  print("Model Summary:")
  print(summary(reg6))

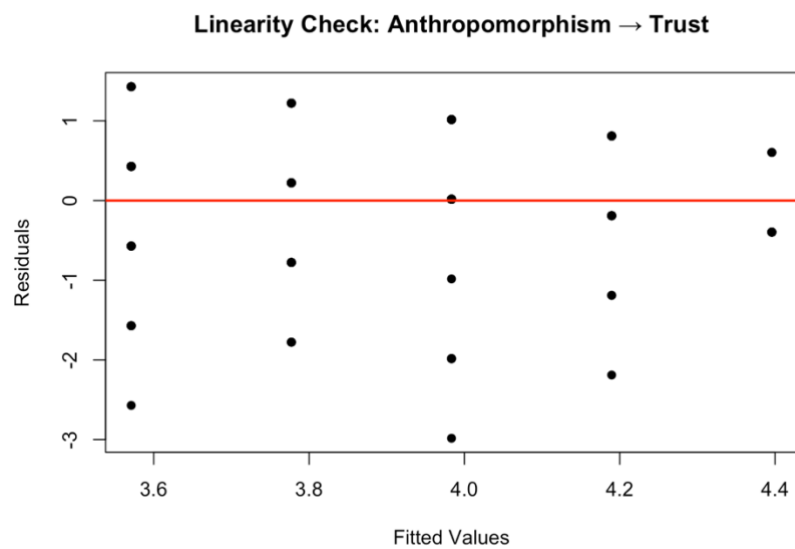
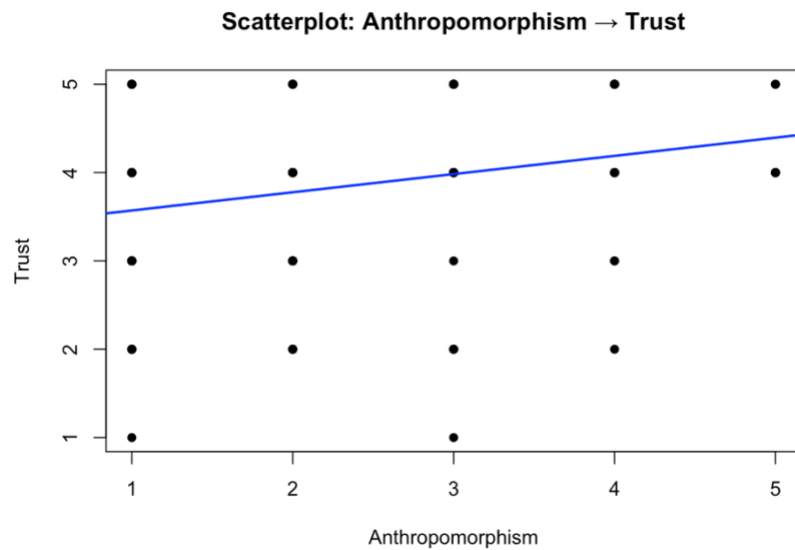
  # Scatterplot with regression line
  plot(Study_Data[[M2]], Study_Data[[Y]],
       main = "Scatterplot: Anthropomorphism → Trust",
       xlab = "Anthropomorphism", ylab = "Trust",
```

```

    pch = 16)
abline(reg6, col = "blue", lwd = 2)

# Residuals vs Fitted plot
plot(fitted(reg6), residuals(reg6),
     main = "Linearity Check: Anthropomorphism → Trust",
     xlab = "Fitted Values",
     ylab = "Residuals",
     pch = 16)
abline(h = 0, col = "red", lwd = 2)}
## [1] "Model Summary:"
##
## Call:
## lm(formula = Study_Data[[Y]] ~ Study_Data[[M2]])
##
## Residuals:
##      Min       1Q   Median       3Q      Max
## -2.98341 -0.57097  0.01659  0.65571  1.42903
##
## Coefficients:
##              Estimate Std. Error t value Pr(>|t|)
## (Intercept)    3.36475    0.18027  18.665 < 2e-16 ***
## Study_Data[[M2]] 0.20622    0.06401   3.221  0.00158 **
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Residual standard error: 0.9315 on 142 degrees of freedom
## Multiple R-squared:  0.0681, Adjusted R-squared:  0.06154
## F-statistic: 10.38 on 1 and 142 DF,  p-value: 0.001582

```



```
# Regression 7: M3 predicting Y (b3 path)

if(all(c(M3, Y) %in% names(Study_Data))) {
  reg7 <- lm(Study_Data[[Y]] ~ Study_Data[[M3]])

  print("Model Summary:")
  print(summary(reg7))

  # Scatterplot with regression line
  plot(Study_Data[[M3]], Study_Data[[Y]],
       main = "Scatterplot: Control → Trust",
       xlab = "Control", ylab = "Trust",
       pch = 16)
```



```

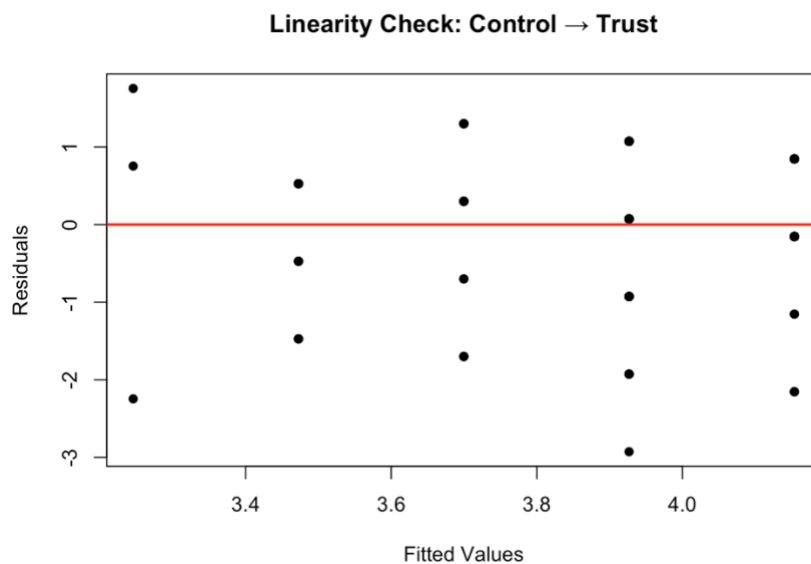
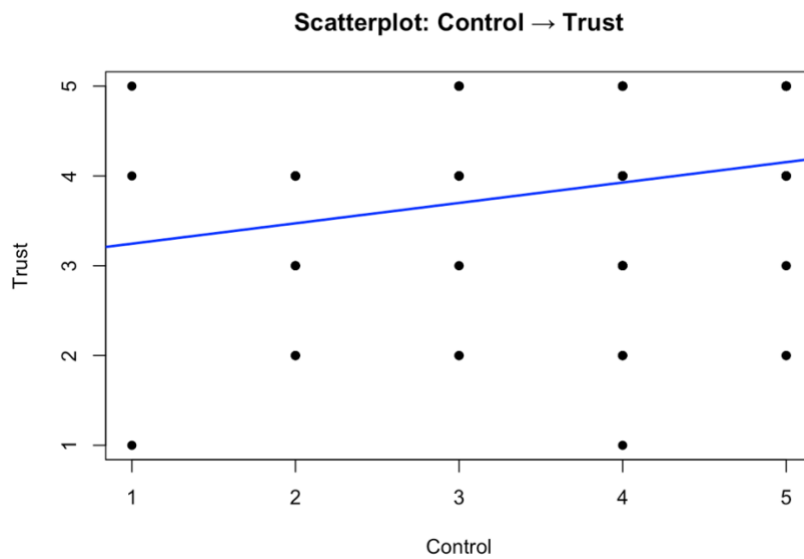
abline(reg7, col = "blue", lwd = 2)

# Residuals vs Fitted plot
plot(fitted(reg7), residuals(reg7),
     main = "Linearity Check: Control → Trust",
     xlab = "Fitted Values",
     ylab = "Residuals",
     pch = 16)

abline(h = 0, col = "red", lwd = 2)}

## [1] "Model Summary:"
##
## Call:
## lm(formula = Study_Data[[Y]] ~ Study_Data[[M3]])
##
## Residuals:
##      Min       1Q   Median       3Q      Max
## -2.92674 -0.47257  0.07326  0.84618  1.75451
##
## Coefficients:
##              Estimate Std. Error t value Pr(>|t|)
## (Intercept)      3.01840     0.29367   10.278 < 2e-16 ***
## Study_Data[[M3]]  0.22708     0.07387    3.074  0.00253 **
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Residual standard error: 0.9344 on 142 degrees of freedom
## Multiple R-squared:  0.0624, Adjusted R-squared:  0.0558
## F-statistic: 9.451 on 1 and 142 DF,  p-value: 0.002532

```



```
# Regression 8: X and all mediators predicting Y (combined model)

if(all(c(X, M1, M2, M3, Y) %in% names(Study_Data))) {
  reg8 <- lm(Study_Data[[Y]] ~ Study_Data[[X]] + Study_Data[[M1]] + Study_Data[[M2]]
    + Study_Data[[M3]])

  print("Model Summary:")
  print(summary(reg8))

  # Residuals vs. fitted plot for the combined model
  plot(fitted(reg8), residuals(reg8),
       main = "Linearity Check: Combined Model (RS Type + Mediators → Trust)",
       xlab = "Fitted Values",
```

```

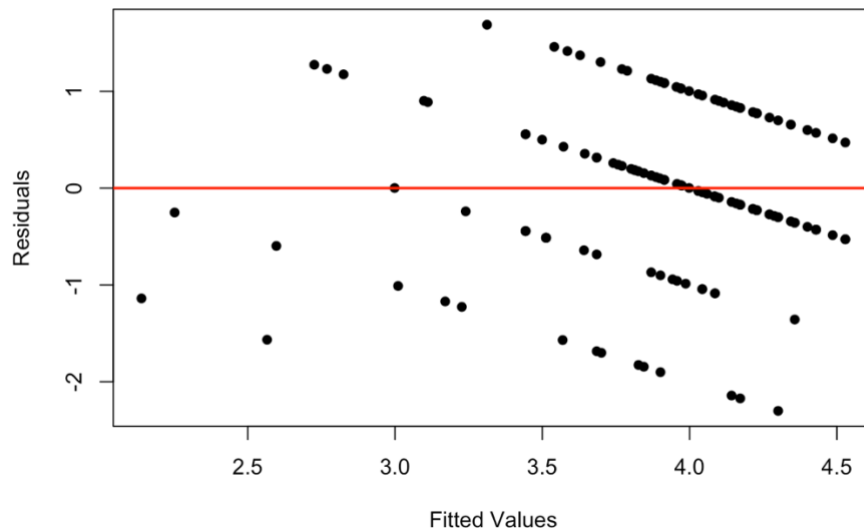
      ylab = "Residuals",
      pch = 16)

abline(h = 0, col = "red", lwd = 2)
}

## [1] "Model Summary:"
##
## Call:
## lm(formula = Study_Data[[Y]] ~ Study_Data[[X]] + Study_Data[[M1]] +
##     Study_Data[[M2]] + Study_Data[[M3]])
##
## Residuals:
##      Min       1Q   Median       3Q      Max
## -2.30059 -0.44327  0.00143  0.70674  1.68754
##
## Coefficients:
##              Estimate Std. Error t value Pr(>|t|)
## (Intercept)      1.55266    0.40562   3.828 0.000195 ***
## Study_Data[[X]]   0.04308    0.14850   0.290 0.772141
## Study_Data[[M1]]  0.40153    0.08533   4.705 6.05e-06 ***
## Study_Data[[M2]]  0.12873    0.06458   1.993 0.048181 *
## Study_Data[[M3]]  0.05633    0.07997   0.704 0.482328
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Residual standard error: 0.861 on 139 degrees of freedom
## Multiple R-squared:  0.2207, Adjusted R-squared:  0.1983
## F-statistic: 9.844 on 4 and 139 DF, p-value: 4.844e-07

```

Linearity Check: Combined Model (RS Type + Mediators → Trust)



```
# Assumption 2: Multicollinearity

# Method 1: Correlation matrix among mediators
cor_matrix <- Study_Data %>%
  dplyr::select(system_understanding, system_control, system_anthropomorphism) %>%
  cor(use = "complete.obs")

print("Correlation Matrix among Mediators:")
## [1] "Correlation Matrix among Mediators:"
print(round(cor_matrix, 3))
##
##               system_understanding system_control
## system_understanding             1.000           0.343
## system_control                   0.343           1.000
## system_anthropomorphism          0.197           0.391
##
##               system_anthropomorphism
## system_understanding             0.197
## system_control                   0.391
## system_anthropomorphism          1.000
# Method 2: VIF
vif_model <- lm(system_trust ~ system_understanding + system_control +
  system_anthropomorphism, data = Study_Data)

# Calculate VIF
print("Variance Inflation Factors:")
## [1] "Variance Inflation Factors:"
```

```

vif(vif_model)

##      system_understanding      system_control system_anthropomorphism
##              1.139058              1.292448              1.187028
# Method 3: Tolerance (1/VIF)
print("Tolerance values:")
## [1] "Tolerance values:"
1/vif(vif_model)

##      system_understanding      system_control system_anthropomorphism
##              0.8779183              0.7737253              0.8424401
# Assumption 3: Homoscedasticity

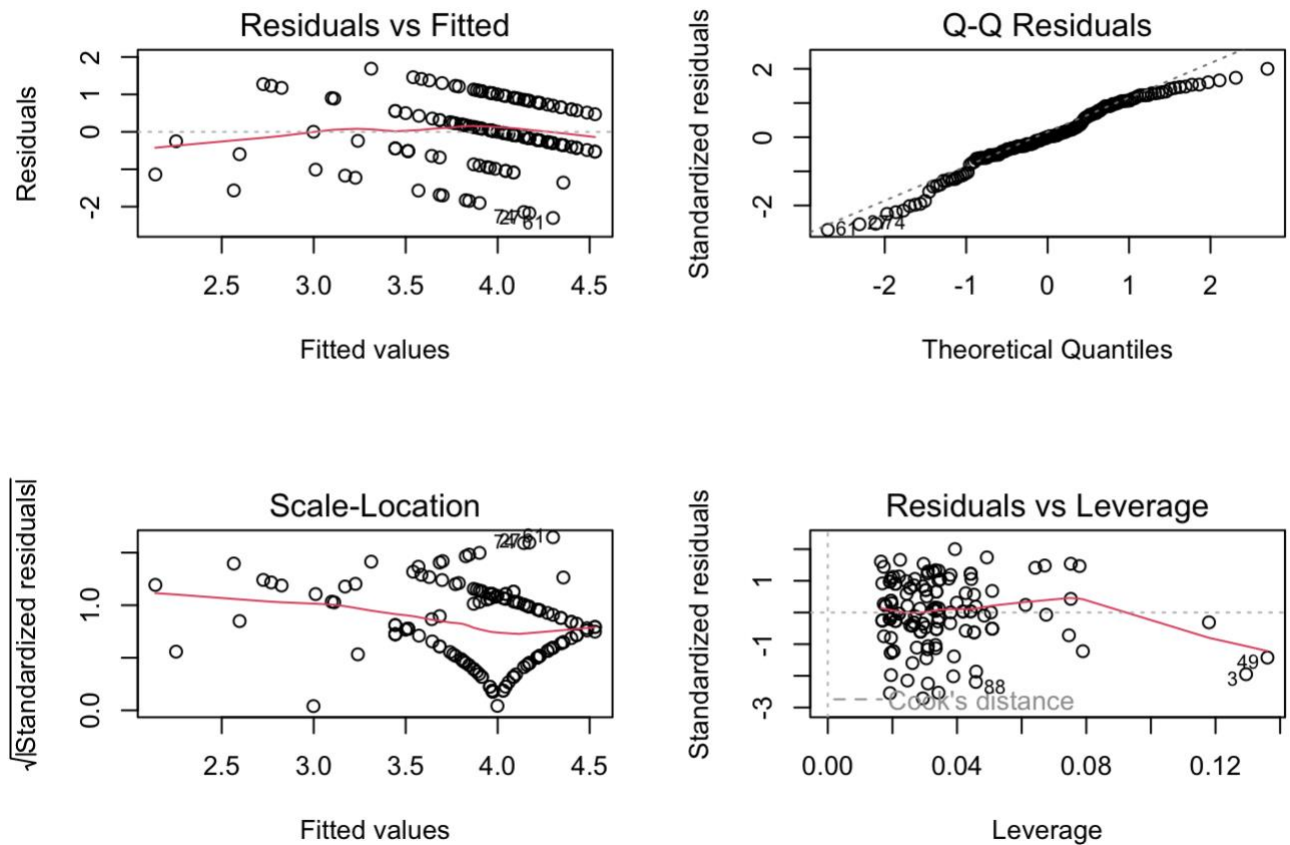
full_model <- lm(system_trust ~ rs_type + system_understanding +
                  system_control + system_anthropomorphism,
                  data = Study_Data)

# Summary of the model
summary(full_model)

##
## Call:
## lm(formula = system_trust ~ rs_type + system_understanding +
##      system_control + system_anthropomorphism, data = Study_Data)
##
## Residuals:
##      Min       1Q   Median       3Q      Max
## -2.30059 -0.44327  0.00143  0.70674  1.68754
##
## Coefficients:
##              Estimate Std. Error t value Pr(>|t|)
## (Intercept)      1.55266    0.40562   3.828 0.000195 ***
## rs_type           0.04308    0.14850   0.290 0.772141
## system_understanding  0.40153    0.08533   4.705 6.05e-06 ***
## system_control      0.05633    0.07997   0.704 0.482328
## system_anthropomorphism 0.12873    0.06458   1.993 0.048181 *
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1
##
## Residual standard error: 0.861 on 139 degrees of freedom
## Multiple R-squared:  0.2207, Adjusted R-squared:  0.1983
## F-statistic: 9.844 on 4 and 139 DF, p-value: 4.844e-07

```

```
# Method 1: Visual inspection - residuals vs fitted plot
par(mfrow = c(2, 2))
plot(full_model)
```

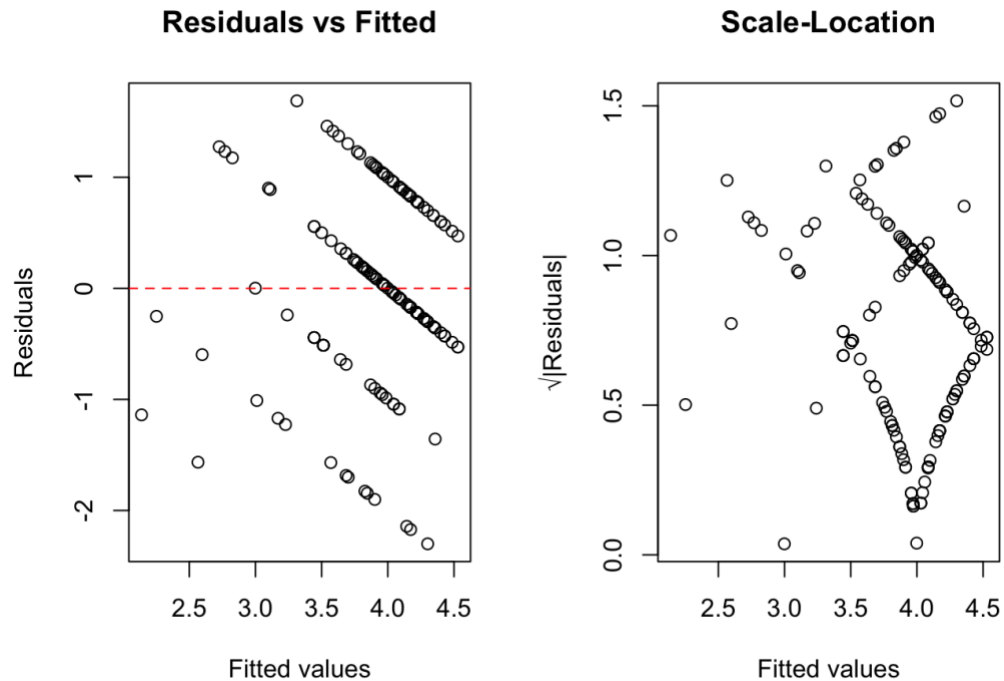


```
# Individual plots
par(mfrow = c(1, 2))

# Residuals vs Fitted
plot(fitted(full_model), residuals(full_model),
     xlab = "Fitted values",
     ylab = "Residuals",
     main = "Residuals vs Fitted")
abline(h = 0, col = "red", lty = 2)

# Scale-Location plot
plot(fitted(full_model), sqrt(abs(residuals(full_model))),
     xlab = "Fitted values",
     ylab = " $\sqrt{|\text{Residuals}|}$ ",
```

```
main = "Scale-Location")
```



```
# Method 2: Statistical test - Breusch-Pagan test
library(lmtest)
bp_test <- bptest(full_model)
print("Breusch-Pagan test for homoscedasticity:")
## [1] "Breusch-Pagan test for homoscedasticity:"
print(bp_test)
##
## studentized Breusch-Pagan test
##
## data: full_model
## BP = 4.2353, df = 4, p-value = 0.3751
# Assumption 4: Normality

# Reset plot parameters
par(mfrow = c(1, 2))

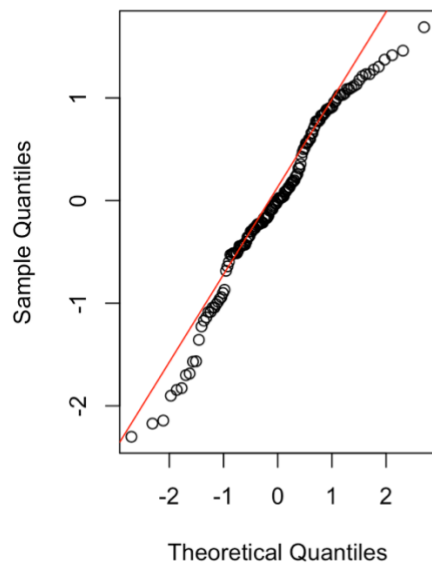
# Method 1: Q-Q plot
qqnorm(residuals(full_model), main = "Q-Q Plot of Residuals")
qqline(residuals(full_model), col = "red")

# Method 2: Histogram with normal curve
```

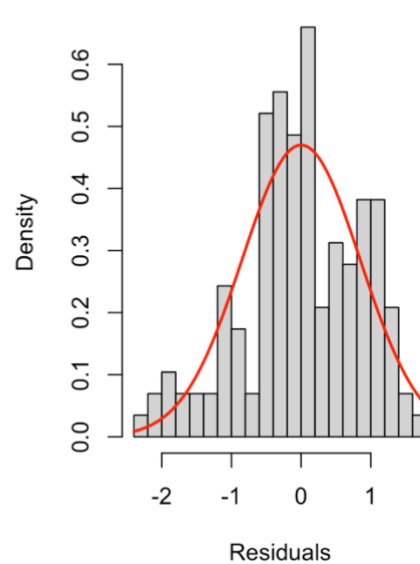
```
hist(residuals(full_model),
     breaks = 20,
     probability = TRUE,
     main = "Histogram of Residuals",
     xlab = "Residuals")

curve(dnorm(x, mean = mean(residuals(full_model)),
           sd = sd(residuals(full_model))),
      add = TRUE,
      col = "red",
      lwd = 2)
```

Q-Q Plot of Residuals



Histogram of Residuals



```
# Method 3: Statistical tests for normality
# Shapiro-Wilk test
shapiro_test <- shapiro.test(residuals(full_model))
print("Shapiro-Wilk test for normality:")
## [1] "Shapiro-Wilk test for normality:"
print(shapiro_test)
##
##  Shapiro-Wilk normality test
##
## data:  residuals(full_model)
## W = 0.97211, p-value = 0.004886
```

6. Descriptive Statistics


```

library(dplyr)
library(tidyr)
library(effsize)
library(knitr)
library(kableExtra)

#1. Keep only the two coded groups
clean_df <- Study_Data %>% filter(rs_type %in% c(0, 1))

#2. Variables
vars <- c("system_trust",
          "system_understanding",
          "system_control",
          "system_anthropomorphism")

pretty <- c(
  system_trust = "Trust",
  system_understanding = "Explainability",
  system_control = "Sense of Control",
  system_anthropomorphism = "Anthropomorphism"
)

# 3. Descriptives per group
desc <- clean_df %>%
  mutate(rs_type = factor(rs_type,
                          levels = c(0, 1),
                          labels = c("Baseline", "Conversational"))) %>%
  pivot_longer(all_of(vars),
               names_to = "Variable",
               values_to = "Value") %>%
  group_by(rs_type, Variable) %>%
  summarise(
    n = sum(!is.na(Value)),
    M = mean(Value, na.rm = TRUE),
    SD = sd(Value, na.rm = TRUE),
    SE = SD / sqrt(n),
    LL = M - qt(.975, df = n - 1) * SE,
    UL = M + qt(.975, df = n - 1) * SE,
    .groups = "drop"
  )

```

```

)

# 4. Make groups to columns
wide <- desc %>%
  pivot_wider(id_cols = Variable,
              names_from = rs_type,
              values_from = c(n, M, SD, LL, UL),
              names_sep = "_")

# 5. Add CI ranges & Cohen's d
table_df <- wide %>%
  mutate(
    CI_Baseline = sprintf("%.2f-%.2f", LL_Baseline, UL_Baseline),
    CI_Conversational = sprintf("%.2f-%.2f", LL_Conversational, UL_Conversational)
  ) %>%
  # Cohen's d (Conversational - Baseline)
  rowwise() %>%
  mutate(
    d = cohen.d(clean_df[[Variable]][clean_df$rs_type == 1],
                clean_df[[Variable]][clean_df$rs_type == 0],
                na.rm = TRUE)$estimate %>% round(2)
  ) %>%
  ungroup() %>%
  # Order & pretty labels
  mutate(Variable = pretty[Variable],
         Variable = factor(Variable,
                           levels = c("Trust", "Explainability",
                                       "Sense of Control", "Anthropomorphism"))) %>%
  arrange(Variable) %>%
  dplyr::select( # FIX: Use dplyr::select explicitly
    Variable,
    n_Baseline, M_Baseline, SD_Baseline, CI_Baseline,
    n_Conversational, M_Conversational, SD_Conversational, CI_Conversational,
    d
  ) %>%
  mutate(across(where(is.numeric), round, 2)) %>%
  as.data.frame()

row.names(table_df) <- NULL

```

```

# 6. Column labels
col_labels <- c("Variable",
               "n", "M", "SD", "CI (95%)",
               "n", "M", "SD", "CI (95%)",
               "d")

# 7. Table 1 : Descriptive Results of Survey Responses
table1 <- kable(table_df,
               col.names = col_labels,
               align      = c("l", rep("c", length(col_labels) - 1)),
               format     = "html") %>%
add_header_above(
  c(" " = 1,
    "Baseline RS"      = 4,
    "Conversational RS" = 4,
    " " = 1),
  bold = TRUE
) %>%
kable_styling(full_width = FALSE,
              position = "left",
              html_font = "Times New Roman") %>%
row_spec(0, bold = TRUE,
        extra_css = "border-top:2px solid black; border-bottom:1px solid black;"
) %>%
row_spec(nrow(table_df),
        extra_css = "border-bottom:1px solid black;") %>%
footnote(
  general_title = "Note.",
  general       = c(
    "n = number of participants within each recommender system condition."),
  footnote_as_chunk = TRUE
)

save_kable(table1, file = "Table1_Descriptive_Results.png", zoom = 3)

table1

```

Variable	Baseline RS				Conversational RS				d
	n	M	SD	CI (95%)	n	M	SD	CI (95%)	
Trust	693	3.88	1.04	3.64–4.13	753	3.89	0.89	3.69–4.10	0.01
Explainability	694	4.11	1.03	4.16–4.65	754	4.10	0.77	4.24–4.59	0.01
Sense of Control	694	0.90	0.92	3.87–4.31	753	0.60	1.13	3.34–3.86	-0.47
Anthropomorphism	692	5.81	1.31	2.26–2.89	752	5.11	1.13	2.25–2.77	-0.06

Note. n = number of participants within each recommender system condition.

7. Mediation Analysis

```
# Mediation Analysis with PROCESS Macro
source("~/Downloads/processv43 (1)/PROCESS v4.3 for R/process.R")

##
## ***** PROCESS for R Version 4.3.1 *****
##
##           Written by Andrew F. Hayes, Ph.D.  www.afhayes.com
##   Documentation available in Hayes (2022). www.guilford.com/p/hayes3
##
## *****
##
## PROCESS is now ready for use.
## Copyright 2020-2023 by Andrew F. Hayes ALL RIGHTS RESERVED
## Workshop schedule at http://haskayne.ucalgary.ca/CCRAM
##
set.seed(123)

mediation_results <- process(
```

```

data = Study_Data,
y = "system_trust",
x = "rs_type",
m = c("system_understanding", "system_control", "system_anthropomorphism"),
model = 4,
boot = 5000,
seed = 123
)
##
## ***** PROCESS for R Version 4.3.1 *****
##
##           Written by Andrew F. Hayes, Ph.D.  www.afhayes.com
##   Documentation available in Hayes (2022). www.guilford.com/p/hayes3
##
## *****
##
## Model : 4
##   Y : system_trust
##   X : rs_type
##   M1 : system_understanding
##   M2 : system_control
##   M3 : system_anthropomorphism
##
## Sample size: 144
##
## Custom seed: 123
##
## *****
## Outcome Variable: system_understanding
##
## Model Summary:
##           R      R-sq      MSE      F      df1      df2      p
##      0.0042   0.0000   0.8227   0.0025   1.0000  142.0000   0.9603
##
## Model:
##           coeff      se      t      p      LLCI      ULCI
## constant   4.4058   0.1092  40.3484  0.0000   4.1899   4.6217
## rs_type     0.0075   0.1513   0.0498  0.9603  -0.2916   0.3066

```

```
##
## *****
## Outcome Variable: system_control
##
## Model Summary:
##           R      R-sq      MSE      F      df1      df2      p
##      0.2308    0.0533    1.0667    7.9885    1.0000   142.0000    0.0054
##
## Model:
##           coeff      se      t      p      LLCI      ULCI
## constant    4.0870    0.1243   32.8695    0.0000    3.8412    4.3328
## rs_type    -0.4870    0.1723   -2.8264    0.0054   -0.8275   -0.1464
##
## *****
## Outcome Variable: system_anthropomorphism
##
## Model Summary:
##           R      R-sq      MSE      F      df1      df2      p
##      0.0301    0.0009    1.4898    0.1287    1.0000   142.0000    0.7203
##
## Model:
##           coeff      se      t      p      LLCI      ULCI
## constant    2.5797    0.1469   17.5560    0.0000    2.2892    2.8702
## rs_type    -0.0730    0.2036   -0.3587    0.7203   -0.4755    0.3295
##
## *****
## Outcome Variable: system_trust
##
## Model Summary:
##           R      R-sq      MSE      F      df1      df2      p
##      0.4698    0.2207    0.7413    9.8436    4.0000   139.0000    0.0000
##
## Model:
##           coeff      se      t      p      LLCI
## constant    1.5527    0.4056    3.8279    0.0002    0.7507
## rs_type      0.0431    0.1485    0.2901    0.7721   -0.2505
## system_understanding    0.4015    0.0853    4.7054    0.0000    0.2328
## system_control      0.0563    0.0800    0.7045    0.4823   -0.1018
## system_anthropomorphism    0.1287    0.0646    1.9933    0.0482    0.0010
```

```

##                                ULCI
## constant                      2.3546
## rs_type                      0.3367
## system_understanding          0.5703
## system_control                0.2144
## system_anthropomorphism       0.2564
##
## *****
## Bootstrapping progress:
##      |
| 0% |
|
|                                     | 1% |
|>                                     | 1% |
|>                                     | 2% |
|>>                                    | 2% |
|                                     |>>
| 3% |                                     |>>
| 4% |                                     |>>>
| 4% |                                     |>>>
| 5% |                                     |>>>
| 6% |                                     |>>>
>                                     | 6% |
|>>>>                                     | 7% |
|>>>>>                                    | 7% |
|>>>>>                                    | 8% |
|>>>>>                                    | 9% |
|>>>>>>                                   | 9% |
|>>>>>>                                   |10% |
|>>>>>>>                                  |10% |
|>>>>>>>                                  |11% |
|>>>>>>>                                  |12% |
|>>>>>>>>                                 |12% |
|>>>>>>>>                                 |13% |
|>>>>>>>>                                 |14% |
|>>>>>>>>>                                |14% |
|>>>>>>>>>                                |15% |
|>>>>>>>>>                                |15% |
|>>>>>>>>>                                |16% |
|>>>>>>>>>                                |17% |
|>>>>>>>>>                                |17% |
|>>>>>>>>>                                |18% |
|>>>>>>>>>                                |19% |
|>>>>>>>>>                                |19% |
|>>>>>>>>>>                               |20% |
|>>>>>>>>>>                               |20% |
|>>>>>>>>>>                               |21% |
|>>>>>>>>>>                               |22% |
|>>>>>>>>>>>                              |22% |
|>>>>>>>>>>>                              |23% |
|>>>>>>>>>>>                              |23% |
|>>>>>>>>>>>>                             |24% |
|>>>>>>>>>>>>                             |25% |
|>>>>>>>>>>>>                             |25% |
|>>>>>>>>>>>>>                            |26% |
|>>>>>>>>>>>>>                            |27% |
|>>>>>>>>>>>>>                            |27% |
|>>>>>>>>>>>>>>                           |28% |
|>>>>>>>>>>>>>>                           |28% |
|>>>>>>>>>>>>>>>                          |29% |
|>>>>>>>>>>>>>>>>                         |30% |
|>>>>>>>>>>>>>>>>>                       |30% |

```



```
##
## Indirect effect(s) of X on Y:
##           Effect      BootSE  BootLLCI  BootULCI
## TOTAL          -0.0338    0.0902   -0.1995    0.1594
## system_understanding    0.0030    0.0624   -0.1059    0.1434
## system_control         -0.0274    0.0433   -0.1271    0.0505
## system_anthropomorphism -0.0094    0.0297   -0.0745    0.0508
##
## ***** ANALYSIS NOTES AND ERRORS *****
##
## Level of confidence for all confidence intervals in output: 95
##
## Number of bootstraps for percentile bootstrap confidence intervals: 5000
# Standardize coefficients
Study_Data_std <- Study_Data %>%
  mutate(across(c(system_trust, system_understanding, system_control,
                    system_anthropomorphism), scale))

# Model 1: Effects on mediators
model_a1 <- lm(system_understanding ~ rs_type, data = Study_Data_std)
model_a2 <- lm(system_control ~ rs_type, data = Study_Data_std)
model_a3 <- lm(system_anthropomorphism ~ rs_type, data = Study_Data_std)

# Model 2: Full model for trust
model_full <- lm(system_trust ~ rs_type + system_understanding +
                  system_control + system_anthropomorphism,
                  data = Study_Data_std)

# Extract standardized coefficients
beta_a1 <- coef(model_a1)["rs_type"]
beta_a2 <- coef(model_a2)["rs_type"]
beta_a3 <- coef(model_a3)["rs_type"]
beta_full <- coef(model_full)
```

8. Mediation Analysis Results

```
# Table: Path Coefficients

path_coefficients_table <- function() {
  results <- data.frame(
```

```

Path = c("RS Type → Trust (Direct)",
        "RS Type → Explainability",
        "RS Type → Sense of Control",
        "RS Type → Anthropomorphism",
        "Explainability → Trust",
        "Sense of Control → Trust",
        "Anthropomorphism → Trust"),

B      = c(0.043,  0.008, -0.487, -0.073, 0.402, 0.056, 0.129),
Beta   = c(0.021,  0.004, -0.231, -0.030, 0.431, 0.062, 0.167),
SE      = c(0.149,  0.151,  0.172,  0.204, 0.085, 0.080, 0.065),
t       = c(0.290,  0.050, -2.826, -0.359, 4.705, 0.705, 1.993),
p       = c(".772", ".960", ".005", ".720", "< .001", ".482", ".048")
)

cat("\n\nTable \nPath Coefficients for the Parallel Mediation Model\n\n")

kable(
  results,
  caption = "<span style='font-family:\n\"Times New Roman\n\"; font-size:12pt; color:
black;'\n>
          <strong>Table&nbsp;1</strong><br>
          Path Coefficients of Parallel Mediation Model
          </span>",
  col.names = c("Path", "B", "β", "SE", "t", "p"),
  digits      = 3,
  align       = c("l", "r", "r", "r", "r", "r"),
  format      = "html",
  escape      = FALSE
) %>%
  kable_styling(full_width = FALSE, position = "left",
                html_font = "Times New Roman") %>%
  row_spec(0, bold = TRUE,
           extra_css = "border-top:2px solid black; border-bottom:1px solid black
;") %>%
  row_spec(nrow(results),
           extra_css = "border-bottom:1px solid black;") %>%
  footnote(
    general_title      = "Note.",
    general             = c("B = unstandardized coefficient.",
                           "β = standardized coefficient."),

```

```

        footnote_as_chunk = TRUE)
}

# Table : Indirect Effects

indirect_effects_table <- function() {

  indirect <- data.frame(
    Effect = c("Total indirect effect",
               "Via Explainability",
               "Via Sense of Control",
               "Via Anthropomorphism"),
    ab      = c(-0.034, 0.003, -0.027, -0.009),
    SE      = c(0.090, 0.062, 0.043, 0.030),
    CI      = c("[-0.200, 0.159]",
               "[-0.106, 0.143]",
               "[-0.127, 0.051]",
               "[-0.075, 0.051]"),
    p       = c(".613", ".948", ".172", ".300")
  )

  cat("\n\nTable \nIndirect Effects (Bootstrap, 5 000 samples)\n\n")

  kable(
    indirect,
    caption = "<span style='font-family:\"Times New Roman\"; font-size:12pt; color:black;'>
               <strong>Table&nbsp;2</strong><br>
               Indirect Effects of Parallel Mediation Model
               </span>",
    col.names = c("Effect", "ab", "SE", "CI&nbsp;(95&nbsp;%)", "p"),
    digits     = 3,
    align      = c("l", "r", "r", "c", "r"),
    format     = "html",
    escape     = FALSE
  ) %>%
  kable_styling(full_width = FALSE, position = "left",
                html_font = "Times New Roman") %>%
  row_spec(0, bold = TRUE,

```

```

        extra_css = "border-top:2px solid black; border-bottom:1px solid black
;") %>%
    row_spec(nrow(indirect),
        extra_css = "border-bottom:1px solid black;") %>%
    footnote(
        general_title      = "Note.",
        general             = "ab = indirect effect.",
        footnote_as_chunk = TRUE)
}

# Results Tables
path_coefficients_table()
##
##
## Table
## Path Coefficients for the Parallel Mediation Model

```

Table 1

Path Coefficients of Parallel Mediation Model

Path	B	β	SE	t	p
RS Type → Trust (Direct)	0.043	0.021	0.149	0.290	.772
RS Type → Explainability	0.008	0.004	0.151	0.050	.960
RS Type → Sense of Control	-0.487	-0.231	0.172	-2.826	.005
RS Type → Anthropomorphism	-0.073	-0.030	0.204	-0.359	.720
Explainability → Trust	0.402	0.431	0.085	4.705	
Sense of Control → Trust	0.056	0.062	0.080	0.705	.482
Anthropomorphism → Trust	0.129	0.167	0.065	1.993	.048

Note. B = unstandardized coefficient. β = standardized coefficient.

```

indirect_effects_table()
##
##

```

```
## Table
## Indirect Effects (Bootstrap, 5 000 samples)
```

Table 2
Indirect Effects of Parallel Mediation Model

Effect	ab	SE	CI (95 %)	p
Total indirect effect	-0.034	0.090	[-0.200, 0.159]	.613
Via Explainability	0.003	0.062	[-0.106, 0.143]	.948
Via Sense of Control	-0.027	0.043	[-0.127, 0.051]	.172
Via Anthropomorphism	-0.009	0.030	[-0.075, 0.051]	.300

Note. ab = indirect effect.

```
# Figure 4: Path Diagram

create_apa_path_diagram_integrated <- function() {
  nodes <- data.frame(
    name = c("Recommender\nSystem Type",
             "Explainability",
             "Sense of\nControl",
             "Anthropomorphism",
             "User Trust"),
    x = c(1.5, 4.5, 4.5, 4.5, 7.5),
    y = c(2.5, 3.8, 1.8, 1.0, 2.5)
  )

  lab_a1 <- "a1 = 0.008"
  lab_a2 <- "a2 = -0.487*"
  lab_a3 <- "a3 = -0.073"
  lab_b1 <- "b1 = 0.402*"
  lab_b2 <- "b2 = 0.056"
  lab_b3 <- "b3 = 0.129*"
  lab_c <- "c' = 0.043"

  ggplot() +
```

```

geom_rect(data = nodes,
          aes(xmin = x-.7, xmax = x+.7, ymin = y-.3, ymax = y+.3),
          fill = c("white", "grey95", "grey95", "grey95", "white"),
          color = "black") +
geom_text(data = nodes,
          aes(x = x, y = y, label = name), size = 2.8) +

# direct path (c')
geom_segment(aes(x = 2.2, y = 2.5, xend = 6.8, yend = 2.5),
             linetype = "dashed",
             arrow = arrow(length = unit(0.2, "cm"))) +
geom_text(aes(x = 4.5, y = 2.75, label = lab_c), size = 3.2) +

# a-paths
geom_segment(aes(x=2.2,y=2.5,xend=3.8,yend=3.65),
             arrow=arrow(length=unit(0.2,"cm")), colour="grey55") +
geom_text(aes(x=2.8,y=3.4,label=lab_a1), size=3, color="grey30") +

geom_segment(aes(x=2.2,y=2.5,xend=3.8,yend=1.8),
             arrow=arrow(length=unit(0.2,"cm"))) +
geom_text(aes(x=2.8,y=1.95,label=lab_a2), size=3) +

geom_segment(aes(x=2.2,y=2.5,xend=3.8,yend=1.15),
             arrow=arrow(length=unit(0.2,"cm")), colour="grey55") +
geom_text(aes(x=2.7,y=1.6,label=lab_a3), size=3, color="grey30") +

# b-paths
geom_segment(aes(x=5.2,y=3.65,xend=6.8,yend=2.65),
             arrow=arrow(length=unit(0.2,"cm"))) +
geom_text(aes(x=6.2,y=3.4,label=lab_b1), size=3) +

geom_segment(aes(x=5.2,y=1.8,xend=6.8,yend=2.5),
             arrow=arrow(length=unit(0.2,"cm")), colour="grey55") +
geom_text(aes(x=6.2,y=1.95,label=lab_b2), size=3, color="grey30") +

geom_segment(aes(x=5.2,y=1.15,xend=6.8,yend=2.35),
             arrow=arrow(length=unit(0.2,"cm"))) +
geom_text(aes(x=6.1,y=1.5,label=lab_b3), size=3) +

```

```

xlim(0,9) + ylim(0.4,4.3) +

# Add title and caption
labs(
  title = "Figure 4",
  subtitle = "Results of Mediation Model Path Diagram",
  caption = "Note. Values are unstandardized path coefficients (B).\nPaths from
recommender system type to mediators are labelled a1-a3.\nPaths from mediators to t
rust are labelled b1-b3.\nThe dashed line is the direct effect (c').\n * significan
t paths with p < .05"
) +

theme_void() +

# APA-style formatting
theme(
  # Title formatting
  plot.title = element_text(
    family = "Times New Roman",
    size = 12,
    face = "bold",
    hjust = 0,
    margin = margin(b = 5)
  ),
  # Subtitle formatting
  plot.subtitle = element_text(
    family = "Times New Roman",
    size = 12,
    face = "italic",
    hjust = 0,
    margin = margin(b = 15)
  ),
  # Caption/notes formatting
  plot.caption = element_text(
    family = "Times New Roman",
    size = 10,
    hjust = 0,
    margin = margin(t = 15)
  ),
  # Border and spacing
  panel.border = element_rect(color = "black", fill = NA, size = 1),

```



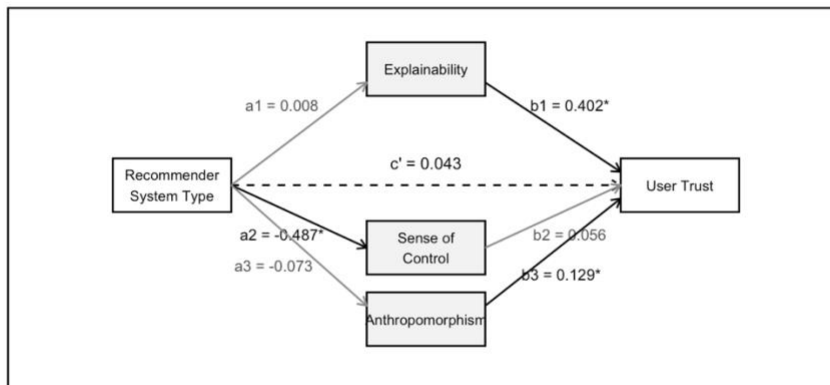
```

    plot.margin = margin(20, 20, 20, 20),
    plot.background = element_rect(color = "black", fill = "white", size = 1)
  )
}

# Create and display
figure4 <- create_apa_path_diagram_integrated()
print(figure4)

```

Figure 4
Results of Mediation Model Path Diagram



Note. Values are unstandardized path coefficients (B).
 Paths from recommender system type to mediators are labelled a1–a3.
 Paths from mediators to trust are labelled b1–b3.
 The dashed line is the direct effect (c').
 * significant paths with $p < .05$

```

# Save the complete figure as PNG
ggsave(
  filename = "Figure4_Mediation_Path_Diagram.png",
  plot = figure4,
  width = 12,
  height = 7,
  dpi = 600,
  bg = "white"
)

```

9. Hypothesis-Testing Summary

```

# Table 2: Hypotheses Results

create_hypothesis_summary_table <- function() {

  hypothesis_data <- data.frame(

```

```

Hypothesis = c("H1 CRS → Trust (direct)",
               "H2a CRS → Explainability",
               "H2b Explainability → Trust",
               "H2c CRS → Trust via Explainability",
               "H3a CRS → Sense of Control",
               "H3b Sense of Control → Trust",
               "H3c CRS → Trust via Sense of Control",
               "H4a CRS → Anthropomorphism",
               "H4b Anthropomorphism → Trust",
               "H4c CRS → Trust via Anthropomorphism"),
Prediction = c("β > 0", "β > 0", "β > 0", "Mediation",
               "β > 0", "β > 0", "Mediation",
               "β > 0", "β > 0", "Mediation"),
Result = c("β = 0.043", "β = 0.008", "β = 0.402", "ab = 0.003",
           "β = -0.487", "β = 0.056", "ab = -0.027",
           "β = -0.073", "β = 0.129", "ab = -0.009"),
p = c(".772", ".960", "<.001", ".948",
       ".005", ".482", ".172",
       ".720", ".048", ".300"),
CI = c("[-0.251, 0.337]", "[-0.292, 0.307]", "[0.233, 0.570]", "[-0.106, 0.143]",
       "[-0.828, -0.146]", "[-0.102, 0.214]", "[-0.127, 0.051]",
       "[-0.476, 0.329]", "[0.001, 0.256]", "[-0.075, 0.051]"),
Support = c("Not supported", "Not supported", "Supported", "Not supported",
            "Opposite direction", "Not supported", "Not supported",
            "Not supported", "Supported", "Not supported")
)

cat("\n\nTable 2\n")
cat("Overview of Hypotheses Test Results\n\n")

kable(
  hypothesis_data,
  caption = "<span style='font-family:\"Times New Roman\"; font-size:12pt; color:black;'><strong>Table&nbsp;2.</strong> <br> Overview of Hypotheses Test Results</span>",
  col.names = c("Hypothesis", "Prediction", "Result", "p", "95 % CI", "Support"),
  align = c("l", "c", "c", "r", "c", "c"),
  format = "html",
  escape = FALSE
) %>%

```

```

kable_styling(full_width = FALSE,
              position     = "left",
              html_font    = "Times New Roman") %>%
row_spec(0, bold = TRUE,
        extra_css = "border-top:2px solid black; border-bottom:1px solid black
;") %>%
row_spec(nrow(hypothesis_data),
        extra_css = "border-bottom:1px solid black;") %>%
# shade rows whose Support == "Supported"
row_spec(which(hypothesis_data$Support == "Supported"),
        background = "#e8e8e8") %>%
footnote(
  general_title    = "Note.",
  general          = c(
    "CRS = Conversational Recommender System. <br>",
    "β = standardized coefficient for direct paths. <br>",
    "ab = indirect (mediated) effect. <br>",
    "CI = confidence interval. <br>",
    "Shaded rows indicate supported hypotheses. <br>"
  ),
  footnote_as_chunk = TRUE,
  escape            = FALSE
)
}

create_hypothesis_summary_table()
##
##
## Table 2
## Overview of Hypotheses Test Results

```

Table 2.

Overview of Hypotheses Test Results

Hypothesis	PredictionResult	p	95 % CI	Support
H1 CRS → Trust (direct)	β > 0	β = 0.043 .772	[-0.251, 0.337]	Not supported

Table 2.

Overview of Hypotheses Test Results

Hypothesis	Prediction	Result	p	95 % CI	Support
H2a CRS → Explainability	$\beta > 0$	$\beta = 0.008$	.960	[-0.292, 0.307]	Not supported
H2b Explainability → Trust	$\beta > 0$	$\beta = 0.402$	<.001	[0.233, 0.570]	Supported
H2c CRS → Trust via Explainability	Mediation	ab = 0.003	.948	[-0.106, 0.143]	Not supported
H3a CRS → Sense of Control	$\beta > 0$	$\beta = -0.487$	.005	[-0.828, -0.146]	Opposite direction
H3b Sense of Control → Trust	$\beta > 0$	$\beta = 0.056$	.482	[-0.102, 0.214]	Not supported
H3c CRS → Trust via Sense of Control	Mediation	ab = -0.027	.172	[-0.127, 0.051]	Not supported
H4a CRS → Anthropomorphism	$\beta > 0$	$\beta = -0.073$	.720	[-0.476, 0.329]	Not supported
H4b Anthropomorphism → Trust	$\beta > 0$	$\beta = 0.129$	.048	[0.001, 0.256]	Supported
H4c CRS → Trust via Anthropomorphism	Mediation	ab = -0.009	.300	[-0.075, 0.051]	Not supported

Note. CRS = Conversational Recommender System.

β = standardized coefficient for direct paths.

ab = indirect (mediated) effect.

CI = confidence interval.

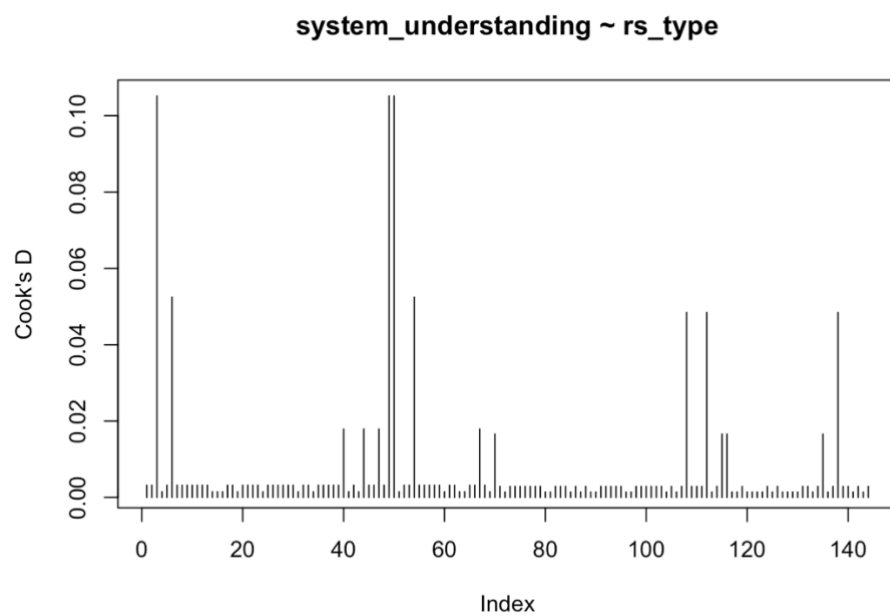
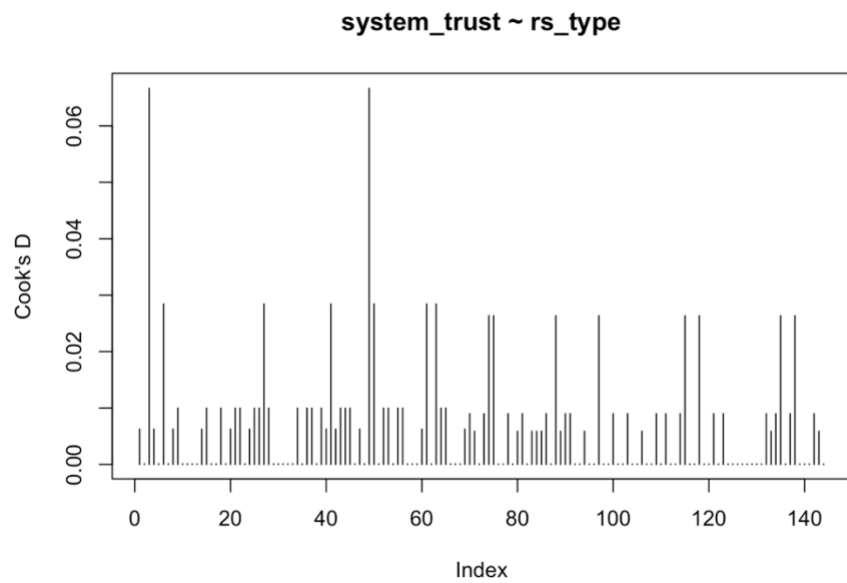
Shaded rows indicate supported hypotheses.

```
tbl4 <- create_hypothesis_summary_table()
##
##
```

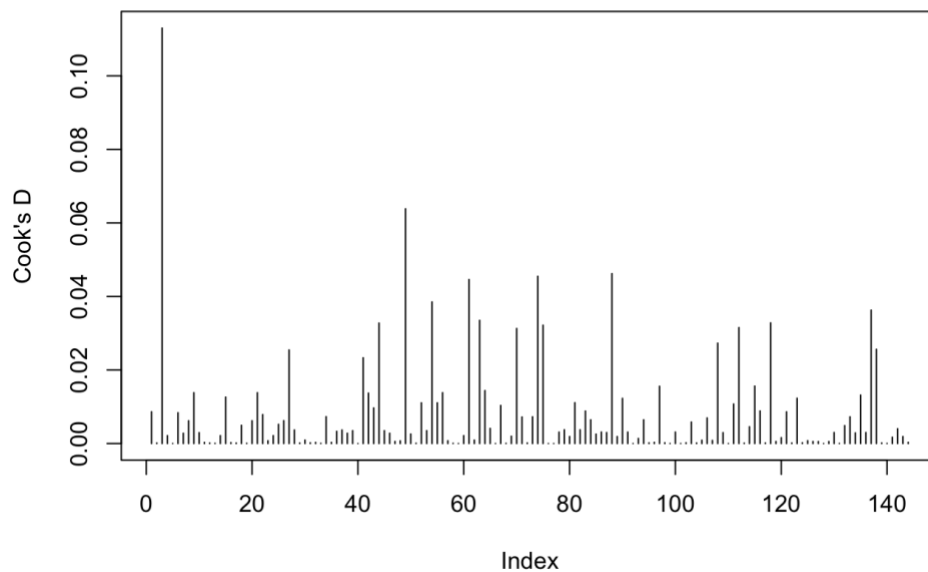
```
## Table 2
## Overview of Hypotheses Test Results
save_kable(tbl4, file = "Table2_Hypotheses.png", zoom = 3)
```

10. Sensitivity Analyses & Robustness Checks

```
# Influence diagnostics
ols_fun <- function(fml) {
  m <- lm(fml, Study_Data)
  plot(cooks.distance(m), type = "h",
       main = deparse(fml), ylab = "Cook's D")
  m
}
models <- list(
  ols_fun(system_trust ~ rs_type),
  ols_fun(system_understanding ~ rs_type),
  ols_fun(system_trust ~ rs_type + system_understanding +
          system_control + system_anthropomorphism)
)
```



**system_trust ~ rs_type + system_understanding + system_control +
system_anthropomorphism**



```
full_mod <- lm(system_trust ~ rs_type + system_understanding +
               system_control + system_anthropomorphism,
               data = Study_Data)

# Rule-of-thumb cutoff
cook_cut <- 4 / nrow(Study_Data)

infl_idx <- which(cooks.distance(full_mod) > cook_cut)
infl_idx                                     # rows to inspect
##    3  44  49  54  61  63  70  74  75  88 112 118 137
##    3  44  49  54  61  63  70  74  75  88 112 118 137
# Refit model without influential cases
full_mod_noinf <- update(full_mod, data = Study_Data[-infl_idx, ])

broom::tidy(full_mod)[, 1:5]
broom::tidy(full_mod_noinf)[, 1:5]
## Table 3: Influential Cases
library(dplyr)
library(kableExtra)
library(webshot2)

influence_data <- tibble::tribble(
  ~Path,                                     ~B_All, ~p_All,                                     ~B_Trim, ~p_Trim,
```

```

"RS Type → Trust (c')",          0.043, ".772",          0.060, ".652",
"Explainability → Trust (b1)",    0.402, "<&nbsp;&nbsp;&nbsp;.001",0.454,    "<&nbsp;&nbsp;&nbsp;p;.001",
"Sense of Control → Trust (b2)",    0.056, ".482",          0.087, ".244",
"Anthropomorphism → Trust (b3)",    0.129, ".049",          0.079, ".168"
)

influence_kable <- kable(
  influence_data,
  caption = "<span style='font-family:\"Times New Roman\"; font-size:12pt; color:black;'>
    <strong>Table&nbsp;&nbsp;&nbsp;3.</strong><br>
    Influential-Cases Diagnostics: Path Coefficients With and Without High Influence Cases
  </span>",
  col.names = c("Path",
    "B (All&nbsp;&nbsp;&nbsp;Cases)", "p",
    "B (Trimmed)", "p"),
  align      = c("l","r","r","r","r"),
  format     = "html",
  escape     = FALSE
) %>%
kable_styling(full_width = FALSE,
  position   = "left",
  html_font  = "Times New Roman") %>%
row_spec(0, bold = TRUE,
  extra_css = "border-top:2px solid black; border-bottom:1px solid black;"
) %>%
row_spec(nrow(influence_data),
  extra_css = "border-bottom:1px solid black;") %>%
row_spec(which(influence_data$Path == "Anthropomorphism → Trust (b3)"),
  background = "#e8e8e8") %>%
footnote(
  general_title = "Note.",
  general       = c(
    "B = unstandardised coefficient.<br>",
    "Trimmed model excludes 13 cases with Cook's&nbsp;&nbsp;&nbsp;D&nbsp;&nbsp;&nbsp;>&nbsp;&nbsp;&nbsp;4/n."
  ),
  footnote_as_chunk = TRUE,
  escape            = FALSE
)

```



```
writeLines(as.character(influence_kable), "influence_table.html")
webshot2::webshot("influence_table.html",
                  "influence_table.png",
                  vwidth = 820,
                  vheight = 380)
```

Table 3.
Influential-Cases Diagnostics: Path Coefficients With and Without High
Influence Cases

Path	B (All Cases)	p	B (Trimmed)	p
RS Type → Trust (c')	0.043	.772	0.060	.652
Explainability → Trust (b ₁)	0.402	< .001	0.454	< .001
Sense of Control → Trust (b ₂)	0.056	.482	0.087	.244
Anthropomorphism → Trust (b ₃)	0.129	.049	0.079	.168

Note. B = unstandardised coefficient.

Trimmed model excludes 13 cases with Cook's D > 4/n.

```
influence_kable
```

Table 3.

Influential-Cases Diagnostics: Path Coefficients With and Without High Influence Cases

Path	B (All Cases)	p	B (Trimmed)	p
RS Type → Trust (c')	0.043	.772	0.060	.652
Explainability → Trust (b ₁)	0.402	< .001	0.454	< .001
Sense of Control → Trust (b ₂)	0.056	.482	0.087	.244
Anthropomorphism → Trust (b ₃)	0.129	.049	0.079	.168

Note. B = unstandardised coefficient.

Trimmed model excludes 13 cases with Cook's D > 4/n.

```
# Robustness Checks
```

```

# Robust checks for heteroskedasticity (HC3) and heavy-tailed residuals (M-estimation)

library(sandwich)
library(lmtest)
library(MASS)

ols_full <- lm(system_trust ~ rs_type +
               system_understanding +
               system_control +
               system_anthropomorphism,
               data = Study_Data)

# 1. HC3 standard errors
hc3_out <- coeftest(ols_full, vcov = vcovHC(ols_full, type = "HC3"))
hc3_out      # print table of coefficients, HC3 SEs, t, p
##
## t test of coefficients:
##
##              Estimate Std. Error t value Pr(>|t|)
## (Intercept)      1.552661    0.499643   3.1075 0.0022881 **
## rs_type           0.043085    0.150008   0.2872 0.7743749
## system_understanding 0.401531    0.110929   3.6197 0.0004121 ***
## system_control     0.056334    0.084361   0.6678 0.5053910
## system_anthropomorphism 0.128735    0.062751   2.0515 0.0420946 *
## ---
## Signif. codes:  0 '***' 0.001 '**' 0.01 '*' 0.05 '.' 0.1 ' ' 1

# 2. Robust estimator (M-estimator)
robust_full <- rlm(system_trust ~ rs_type +
                  system_understanding +
                  system_control +
                  system_anthropomorphism,
                  data = Study_Data,
                  psi = psi.huber)

summary(robust_full)
##
## Call: rlm(formula = system_trust ~ rs_type + system_understanding +
##          system_control + system_anthropomorphism, data = Study_Data,
##          psi = psi.huber)

```

```
## Residuals:
##      Min       1Q   Median       3Q      Max
## -2.34334 -0.48218 -0.05874  0.65754  1.69797
##
## Coefficients:
##              Value Std. Error t value
## (Intercept)      1.4680 0.4185      3.5078
## rs_type           0.0481 0.1532      0.3137
## system_understanding 0.4300 0.0880      4.8836
## system_control      0.0649 0.0825      0.7865
## system_anthropomorphism 0.1165 0.0666      1.7484
##
## Residual standard error: 0.8211 on 139 degrees of freedom
```

11. Exploratory Analyses

```
# 1. Moderated Mediation: Technology Affinity

# Center moderator
Study_Data <- Study_Data %>%
  mutate(tech_aff_c = scale(technology_affinity, scale = FALSE))

# Moderation of technology affinity of all paths
mod59 <- process(
  data   = Study_Data,
  y      = "system_trust",
  x      = "rs_type",
  m      = c("system_understanding",
             "system_control",
             "system_anthropomorphism"),
  w      = "tech_aff_c",
  model  = 59,
  boot   = 5000,
  seed   = 123,
  center = 1
)
##
## ***** PROCESS for R Version 4.3.1 *****
##
##              Written by Andrew F. Hayes, Ph.D.  www.afhayes.com
```

```

## Documentation available in Hayes (2022). www.guilford.com/p/hayes3
##
## *****
##
## Model : 59
## Y : system_trust
## X : rs_type
## M1 : system_understanding
## M2 : system_control
## M3 : system_anthropomorphism
## W : tech_aff_c
##
## Sample size: 144
##
## Custom seed: 123
##
## *****
## Outcome Variable: system_understanding
##
## Model Summary:
##


|  | R      | R-sq   | MSE    | F      | df1    | df2      | p      |
|--|--------|--------|--------|--------|--------|----------|--------|
|  | 0.1705 | 0.0291 | 0.8102 | 1.3978 | 3.0000 | 140.0000 | 0.2461 |


##
## Model:
##


|            | coeff  | se     | t      | p      | LLCI    | ULCI   |
|------------|--------|--------|--------|--------|---------|--------|
| constant   | 0.0026 | 0.0750 | 0.0350 | 0.9722 | -0.1457 | 0.1510 |
| rs_type    | 0.0108 | 0.1502 | 0.0720 | 0.9427 | -0.2861 | 0.3077 |
| tech_aff_c | 0.0810 | 0.0693 | 1.1678 | 0.2449 | -0.0561 | 0.2181 |
| Int_1      | 0.2294 | 0.1390 | 1.6511 | 0.1010 | -0.0453 | 0.5042 |


##
## Product terms key:
## Int_1 : rs_type x tech_aff_c
##
## Test(s) of highest order unconditional interaction(s):
##


|     | R2-chng | F      | df1    | df2      | p      |
|-----|---------|--------|--------|----------|--------|
| X*W | 0.0189  | 2.7262 | 1.0000 | 140.0000 | 0.1010 |


##
## *****

```

```
## Outcome Variable: system_control
##
## Model Summary:
##           R      R-sq      MSE      F      df1      df2      p
##      0.2552    0.0651    1.0684    3.2510    3.0000   140.0000    0.0237
##
## Model:
##           coeff      se      t      p      LLCI      ULCI
## constant      0.0024    0.0862    0.0276    0.9780   -0.1680    0.1727
## rs_type      -0.4885    0.1725   -2.8323    0.0053   -0.8294   -0.1475
## tech_aff_c   -0.0242    0.0796   -0.3040    0.7616   -0.1817    0.1332
## Int_1         0.2083    0.1596    1.3053    0.1939   -0.1072    0.5238
##
## Product terms key:
## Int_1 : rs_type x tech_aff_c
##
## Test(s) of highest order unconditional interaction(s):
##           R2-chng      F      df1      df2      p
## X*W      0.0114      1.7038      1.0000   140.0000    0.1939
##
## *****
## Outcome Variable: system_anthropomorphism
##
## Model Summary:
##           R      R-sq      MSE      F      df1      df2      p
##      0.1654    0.0273    1.4711    1.3118    3.0000   140.0000    0.2730
##
## Model:
##           coeff      se      t      p      LLCI      ULCI
## constant      0.0033    0.1011    0.0325    0.9741   -0.1966    0.2032
## rs_type      -0.0789    0.2024   -0.3899    0.6972   -0.4790    0.3212
## tech_aff_c   -0.1158    0.0934   -1.2396    0.2172   -0.3006    0.0689
## Int_1         0.2878    0.1873    1.5372    0.1265   -0.0824    0.6581
##
## Product terms key:
## Int_1 : rs_type x tech_aff_c
##
## Test(s) of highest order unconditional interaction(s):
##           R2-chng      F      df1      df2      p
```

```

## X*W      0.0164      2.3630      1.0000    140.0000      0.1265
##
## *****
## Outcome Variable: system_trust
##
## Model Summary:
##           R      R-sq      MSE      F      df1      df2      p
##      0.4809    0.2312    0.7586    4.4785    9.0000   134.0000    0.0000
##
## Model:
##           coeff      se      t      p      LLCI
## constant      3.8957    0.0737   52.8886    0.0000    3.7500
## rs_type      -0.0016    0.1540   -0.0102    0.9918   -0.3063
## system_understanding    0.4102    0.0894    4.5870    0.0000    0.2333
## system_control      0.0384    0.0822    0.4675    0.6409   -0.1241
## system_anthropomorphism    0.1257    0.0661    1.9024    0.0593   -0.0050
## tech_aff_c     -0.0209    0.0688   -0.3029    0.7625   -0.1570
## Int_1      0.0566    0.1401    0.4040    0.6869   -0.2205
## Int_2      0.0266    0.0776    0.3421    0.7328   -0.1270
## Int_3      0.0278    0.0772    0.3607    0.7189   -0.1248
## Int_4      0.0627    0.0680    0.9215    0.3585   -0.0718
##           ULCI
## constant      4.0413
## rs_type      0.3031
## system_understanding    0.5870
## system_control      0.2009
## system_anthropomorphism    0.2564
## tech_aff_c      0.1153
## Int_1      0.3337
## Int_2      0.1801
## Int_3      0.1805
## Int_4      0.1971
##
## Product terms key:
## Int_1 : rs_type x tech_aff_c
## Int_2 : system_understanding x tech_aff_c
## Int_3 : system_control x tech_aff_c
## Int_4 : system_anthropomorphism x tech_aff_c
##

```

```
## Test(s) of highest order unconditional interaction(s):
##          R2-chng          F          df1          df2          p
## X*W      0.0009      0.1632      1.0000     134.0000      0.6869
## M1*W     0.0007      0.1170      1.0000     134.0000      0.7328
## M2*W     0.0007      0.1301      1.0000     134.0000      0.7189
## M3*W     0.0049      0.8491      1.0000     134.0000      0.3585
##
## *****
## Bootstrapping progress:
## |
| 0% |
|
|                                     | 1% |
|>                                     | 1% |
|>                                     | 2% |
|>>                                    | 2% |
|                                     |>>
| 3% |                                     |>>
| 4% |                                     |>>>
| 4% |                                     |>>>
| 5% |                                     |>>>
| 6% |                                     |>>>
>                                     | 6% |
|>>>>                                     | 7% |
|>>>>>                                    | 7% |
|>>>>>                                    | 8% |
|>>>>>                                    | 9% |
|>>>>>>                                   | 9% |
|>>>>>>                                   |10% |
|>>>>>>>                                  |10% |
|>>>>>>>                                  |11% |
|>>>>>>>                                  |12% |
|>>>>>>>>                                 |12% |
|>>>>>>>>                                 |13% |
|>>>>>>>>                                 |14% |
|>>>>>>>>>                                |14% |
|>>>>>>>>>                                |15% |
|>>>>>>>>>                                |15% |
|>>>>>>>>>                                |16% |
|>>>>>>>>>                                |17% |
|>>>>>>>>>                                |17% |
|>>>>>>>>>                                |18% |
|>>>>>>>>>                                |19% |
|>>>>>>>>>                                |19% |
|>>>>>>>>>>                               |20% |
|>>>>>>>>>>                               |20% |
|>>>>>>>>>>                               |21% |
|>>>>>>>>>>                               |22% |
|>>>>>>>>>>>                              |22% |
|>>>>>>>>>>>                              |23% |
|>>>>>>>>>>>                              |23% |
|>>>>>>>>>>>                              |24% |
|>>>>>>>>>>>>                             |25% |
|>>>>>>>>>>>>                             |25% |
|>>>>>>>>>>>>                             |26% |
|>>>>>>>>>>>>                             |27% |
|>>>>>>>>>>>>                             |27% |
|>>>>>>>>>>>>                             |28% |
|>>>>>>>>>>>>>                            |28% |
|>>>>>>>>>>>>>                            |29% |
|>>>>>>>>>>>>>>                           |30% |
|>>>>>>>>>>>>>>>                          |30% |
```



```

##      -0.0486   -0.0043    0.1546   -0.0280    0.9777   -0.3100    0.3014
##      0.9514    0.0523    0.1980    0.2639    0.7923   -0.3394    0.4440
##
## Conditional indirect effects of X on Y:
##
## INDIRECT EFFECT:
##
## rs_type      ->    system_understanding      ->    system_trust
##
## tech_aff_c    Effect      BootSE  BootLLCI  BootULCI
##      -1.0486   -0.0879    0.0952   -0.2826    0.0940
##      -0.0486   -0.0001    0.0635   -0.1092    0.1431
##      0.9514    0.0998    0.0975   -0.0561    0.3221
##
## INDIRECT EFFECT:
##
## rs_type      ->    system_control      ->    system_trust
##
## tech_aff_c    Effect      BootSE  BootLLCI  BootULCI
##      -1.0486   -0.0065    0.0882   -0.1885    0.1764
##      -0.0486   -0.0185    0.0453   -0.1174    0.0660
##      0.9514    -0.0188    0.0450   -0.1335    0.0559
##
## INDIRECT EFFECT:
##
## rs_type      ->    system_anthropomorphism      ->    system_trust
##
## tech_aff_c    Effect      BootSE  BootLLCI  BootULCI
##      -1.0486   -0.0228    0.0495   -0.1509    0.0425
##      -0.0486   -0.0114    0.0306   -0.0819    0.0459
##      0.9514    0.0361    0.0552   -0.0701    0.1599
##
## ***** ANALYSIS NOTES AND ERRORS *****
##
## Level of confidence for all confidence intervals in output: 95
##
## Number of bootstraps for percentile bootstrap confidence intervals: 5000
##
## W values in conditional tables are the 16th, 50th, and 84th percentiles.

```

```
##
## NOTE: The following variables were mean centered prior to analysis:
##          tech_aff_c rs_type system_understanding system_control system_anthropom
orphism
# Moderation Results
mod59$direct
## NULL
mod59$int_effects
## NULL
mod59$index_modmed
## NULL
# Recommendation Acceptance

# 1: Load libraries
library(readxl)
library(dplyr)
library(stringr)

# 2: Load Excel file
uncleaned_survey_data <- read_excel("Study Results_R_base.xlsx")

# 3: Rename the columns to match Study_Data

rename_mapping <- c("Bitte geben Sie auf einer Skala von 1 bis 5 an, inwieweit Sie
den folgenden Aussagen zustimmen (1 = stimme gar nicht zu, 5 = stimme voll zu). [I
ch vertraue dem System.]" = "system_trust",

                    "Bitte geben Sie auf einer Skala von 1 bis 5 an, inwieweit Sie d
en folgenden Aussagen zustimmen (1 = stimme gar nicht zu, 5 = stimme voll zu). [Ic
h verstehe, wie das System funktioniert.]" = "system_understanding",

                    "Bitte geben Sie auf einer Skala von 1 bis 5 an, inwieweit Sie d
en folgenden Aussagen zustimmen (1 = stimme gar nicht zu, 5 = stimme voll zu). [Ic
h habe Kontrolle darüber, wie das System handelt.]" = "system_control",

                    "Bitte geben Sie auf einer Skala von 1 bis 5 an, inwieweit Sie d
en folgenden Aussagen zustimmen (1 = stimme gar nicht zu, 5 = stimme voll zu). [Da
s System wirkt menschlich.]" = "system_anthropomorphism",

                    "Automaticalle pre-filled with the data from the URL parameters
(website/prototype). [Group]" = "rs_type",

                    "Automaticalle pre-filled with the data from the URL parameters
(website/prototype). [UID]" = "UID")

names(uncleaned_survey_data)[match(names(rename_mapping), names(uncleaned_survey_da
ta))] <- rename_mapping

# 4: Extract numeric values
extract_number <- function(x) {
```

```

    as.numeric(str_extract(x, "^\\d+"))
  }

uncleaned_survey_data$system_trust_num <- extract_number(uncleaned_survey_data$system_trust)

uncleaned_survey_data$system_understanding_num <- extract_number(uncleaned_survey_data$system_understanding)

uncleaned_survey_data$system_control_num <- extract_number(uncleaned_survey_data$system_control)

uncleaned_survey_data$system_anthropomorphism_num <- extract_number(uncleaned_survey_data$system_anthropomorphism)

# 5: Merge Study_Data with uncleaned data to get UIDs back
Study_Data_with_UID <- merge(Study_Data,
                             uncleaned_survey_data[, c("Antwort ID", "UID", "system_trust_num",
                                                         "system_understanding_num", "system_control_num",
                                                         "system_anthropomorphism_num", "rs_type")],
                             by.x = c("respondent_ID", "system_trust", "system_understanding",
                                       "system_control", "system_anthropomorphism", "rs_type"),
                             by.y = c("Antwort ID", "system_trust_num", "system_understanding_num",
                                       "system_control_num", "system_anthropomorphism_num", "rs_type"),
                             all.x = TRUE)

dim(Study_Data_with_UID)
## [1] 144 44
sum(is.na(Study_Data_with_UID$UID))
## [1] 0
sum(duplicated(Study_Data_with_UID$UID))
## [1] 0
# 6: Merge with interaction analytics
interaction_analytics <- read.csv("interaction_analytics.csv")

cat("interaction_analytics dimensions:", dim(interaction_analytics), "\n")
## interaction_analytics dimensions: 183 6
head(interaction_analytics)
final_dataset <- merge(Study_Data_with_UID, interaction_analytics,
                       by = "UID", all.x = TRUE)

```

```

dim(final_dataset)
## [1] 144  49
sum(is.na(final_dataset$success_occurred))
## [1] 44
# 7: Create analysis dataset for recommendation acceptance
analysis_dataset <- final_dataset[!is.na(final_dataset$success_occurred), ]

# Check final dataset
dim(analysis_dataset)
## [1] 100  49
table(analysis_dataset$rs_type, analysis_dataset$success_occurred)
##
##      False True
##    0      1   44
##    1      2   53
# 8: Run exploratory analyses
# Recommendation acceptance analysis
analysis_dataset$success_occurred <- as.logical(analysis_dataset$success_occurred)

success_rates <- analysis_dataset %>%
  group_by(rs_type) %>%
  summarise(
    n = n(),
    successes = sum(success_occurred),
    success_rate = mean(success_occurred),
    .groups = 'drop'
  )
print(success_rates)
## # A tibble: 2 × 4
##   rs_type      n successes success_rate
##   <dbl> <int>    <int>      <dbl>
## 1      0    45      44      0.978
## 2      1    55      53      0.964
# Statistical tests for recommendation acceptance
fisher.test(table(analysis_dataset$rs_type, analysis_dataset$success_occurred))
##
## Fisher's Exact Test for Count Data
##

```

```

## data:  table(analysis_dataset$rs_type, analysis_dataset$success_occurred)
## p-value = 1
## alternative hypothesis: true odds ratio is not equal to 1
## 95 percent confidence interval:
##    0.009990857 11.987297429
## sample estimates:
## odds ratio
##    0.6052111
t.test(success_occurred ~ rs_type, data = analysis_dataset)
##
##  Welch Two Sample t-test
##
## data:  success_occurred by rs_type
## t = 0.41833, df = 97.888, p-value = 0.6766
## alternative hypothesis: true difference in means between group 0 and group 1 is
not equal to 0
## 95 percent confidence interval:
##   -0.05294337  0.08122620
## sample estimates:
## mean in group 0 mean in group 1
##      0.9777778      0.9636364
# Gender effect on trust
t.test(system_trust ~ gender, data = Study_Data_with_UID)
##
##  Welch Two Sample t-test
##
## data:  system_trust by gender
## t = -0.94838, df = 79.719, p-value = 0.3458
## alternative hypothesis: true difference in means between group female and group
male is not equal to 0
## 95 percent confidence interval:
##   -0.5182947  0.1837493
## sample estimates:
## mean in group female  mean in group male
##      3.772727      3.940000
# Age effect on trust
cor.test(Study_Data_with_UID$age, Study_Data_with_UID$system_trust)
##
##  Pearson's product-moment correlation
##

```

```

## data: Study_Data_with_UID$age and Study_Data_with_UID$system_trust
## t = -0.03675, df = 142, p-value = 0.9707
## alternative hypothesis: true correlation is not equal to 0
## 95 percent confidence interval:
## -0.1665759 0.1605729
## sample estimates:
## cor
## -0.00308398
# Trust effect on recommendation acceptance
cor.test(analysis_dataset$system_trust, as.numeric(analysis_dataset$success_occurre
d))
##
## Pearson's product-moment correlation
##
## data: analysis_dataset$system_trust and as.numeric(analysis_dataset$success_occ
urred)
## t = 0.51851, df = 98, p-value = 0.6053
## alternative hypothesis: true correlation is not equal to 0
## 95 percent confidence interval:
## -0.1456084 0.2461945
## sample estimates:
## cor
## 0.05230571
# Exploration: Effect of Trust on Acceptance

# Point-biserial correlation
cor.test(analysis_dataset$system_trust, as.numeric(analysis_dataset$success_occurre
d))
##
## Pearson's product-moment correlation
##
## data: analysis_dataset$system_trust and as.numeric(analysis_dataset$success_occ
urred)
## t = 0.51851, df = 98, p-value = 0.6053
## alternative hypothesis: true correlation is not equal to 0
## 95 percent confidence interval:
## -0.1456084 0.2461945
## sample estimates:
## cor
## 0.05230571
# Logistic regression - does trust predict acceptance?

```

```

trust_acceptance_model <- glm(success_occurred ~ system_trust,
                              data = analysis_dataset,
                              family = binomial)

summary(trust_acceptance_model)

##
## Call:
## glm(formula = success_occurred ~ system_trust, family = binomial,
##      data = analysis_dataset)
##
## Coefficients:
##              Estimate Std. Error z value Pr(>|z|)
## (Intercept)    2.3694     2.1359   1.109   0.267
## system_trust    0.2899     0.5583   0.519   0.604
##
## (Dispersion parameter for binomial family taken to be 1)
##
##      Null deviance: 26.948  on 99  degrees of freedom
## Residual deviance: 26.694  on 98  degrees of freedom
## AIC: 30.694
##
## Number of Fisher Scoring iterations: 6
# Descriptive statistics: trust levels by acceptance status
analysis_dataset %>%
  group_by(success_occurred) %>%
  summarise(
    n = n(),
    mean_trust = mean(system_trust, na.rm = TRUE),
    sd_trust = sd(system_trust, na.rm = TRUE),
    min_trust = min(system_trust, na.rm = TRUE),
    max_trust = max(system_trust, na.rm = TRUE),
    .groups = 'drop'
  )
t.test(system_trust ~ success_occurred, data = analysis_dataset)

##
## Welch Two Sample t-test
##
## data:  system_trust by success_occurred
## t = -0.84056, df = 2.362, p-value = 0.4771
## alternative hypothesis: true difference in means between group FALSE and group T
RUE is not equal to 0

```



```
## 95 percent confidence interval:
##  -1.587793  1.003600
## sample estimates:
## mean in group FALSE  mean in group TRUE
##           3.666667           3.958763
```