



The *Who* in XAI: How AI Background Shapes Perceptions of AI Explanations

Upol Ehsan
ehsanu@gatech.edu
Georgia Institute of Technology
USA

Samir Passi
v-samirpassi@microsoft.com
Microsoft
USA

Q. Vera Liao
veraliao@microsoft.com
Microsoft Research
Canada

Larry Chan
larry.chan@illumio.com
Illumio
USA

I-Hsiang Lee
ethan@gatech.edu
Georgia Institute of Technology
USA

Michael Muller
michael_muller@us.ibm.com
IBM Research
USA

Mark O. Riedl
riedl@cc.gatech.edu
Georgia Institute of Technology
USA

ABSTRACT

Explainability of AI systems is critical for users to take informed actions. Understanding *who* opens the black-box of AI is just as important as opening it. We conduct a mixed-methods study of how two different groups—people with and without AI background—perceive different types of AI explanations. Quantitatively, we share user perceptions along five dimensions. Qualitatively, we describe how AI background can influence interpretations, elucidating the differences through lenses of appropriation and cognitive heuristics. We find that (1) both groups showed unwarranted faith in numbers for different reasons and (2) each group found value in different explanations beyond their intended design. Carrying critical implications for the field of XAI, our findings showcase how AI generated explanations can have negative consequences despite best intentions and how that could lead to harmful manipulation of trust. We propose design interventions to mitigate them.

CCS CONCEPTS

• **Human-centered computing** → **Empirical studies in HCI; User studies; Empirical studies in collaborative and social computing**; • **Computing methodologies** → **Artificial intelligence**.

KEYWORDS

Explainable AI, Human-Centered Explainable AI, User Characteristics

ACM Reference Format:

Upol Ehsan, Samir Passi, Q. Vera Liao, Larry Chan, I-Hsiang Lee, Michael Muller, and Mark O. Riedl. 2024. The *Who* in XAI: How AI Background Shapes Perceptions of AI Explanations. In *Proceedings of the CHI Conference*

on Human Factors in Computing Systems (CHI '24), May 11–16, 2024, Honolulu, HI, USA. ACM, New York, NY, USA, 32 pages. <https://doi.org/10.1145/3613904.3642474>

1 INTRODUCTION

As AI-driven systems increasingly power high-stakes decision-making in public domains such as healthcare [25, 67, 76, 98], finance [107, 118], law [15, 151, 158], and criminal justice [62, 80, 136], their explainability is critical for end-users to take informed and accountable actions [143]. Issues concerning explainability lie at the heart of Explainable AI (XAI), a research area that aims to provide human-understandable justifications for the system's behavior [2, 47, 61]. Explainability is not a new issue within AI [74, 115, 142], but the proliferation of Deep Learning and Reinforcement Learning based approaches—models of which are considered hard to interpret, even by experts—has led to remarkable growth in techniques that aim to “open” the AI opaque box [61].

While opening the opaque box is important, *who* interacts with the box also matters. Implicit in XAI is the question: “explainable to whom?” [45]. The *who* governs the most effective way of describing the *why* behind the decisions. Getting a situated understanding of how different *who*'s with different user characteristics matter in XAI is thus important. To give an illustrative example: riders (end-users) of a self-driving car have different user characteristics than its engineers (developers). Riders, many of whom are not AI experts, might not have the AI background that the engineers have and thus have different explainability needs and goals.

One's AI background is an impactful user characteristic in XAI because there is often a disparity in this characteristic between creators/developers and end-users, which can lead to usability failures and irresponsible design [123], even inequities [29]. Many end-users are unlikely to have AI backgrounds comparable to the creators of the technology [73]. Nonetheless, XAI developers tend to design explanations *as if* people like them are going to use their systems [112]. In fact, a majority of current deployments of XAI technologies serve AI engineers instead of end-users [7, 91]. This creates a consumer-creator gap, one between design intention and

Permission to make digital or hard copies of part or all of this work for personal or classroom use is granted without fee provided that copies are not made or distributed for profit or commercial advantage and that copies bear this notice and the full citation on the first page. Copyrights for third-party components of this work must be honored. For all other uses, contact the owner/author(s).

CHI '24, May 11–16, 2024, Honolulu, HI, USA

© 2024 Copyright held by the owner/author(s).

ACM ISBN 979-8-4007-0330-0/24/05

<https://doi.org/10.1145/3613904.3642474>

reality—how developers envision the AI explanations to get interpreted and how users actually perceive them. If we want to bridge this gap, we need to understand how user characteristics, such as AI background, impact it.¹

In this paper, we share *how* and *why* one’s AI background (or lack thereof) shapes their perceptions of AI explanations. Focusing on two groups, one with and one without an AI background, we found that (1) both groups had unwarranted faith in numbers, but exhibited it for different reasons and to differing degrees, with the AI group showing a higher propensity to over-trust numerical representations and potentially be misled by the presence of it. (2) Each group found explanatory values in different explanations that went beyond the usage we designed them for. These insights have potential negative implications like susceptibility to harmful manipulation of user trust.

We found these insights through a mixed-methods study where we probed for user perceptions of *three* types of AI-generated explanations: (1) natural language with justification (explaining the “why” behind the action), (2) natural language without justification (describing “what” the action was), and (3) numbers that determine the agent’s actions (akin to “transparent” AI). We measure perceptions along five dimensions: *confidence*, *intelligence*, *understandability*, *second chance*, and *friendliness*, which are grounded in related work around HCI, HRI, and XAI [17, 34, 36, 47, 160] and quantitatively share within- and between-group differences. Through qualitative analysis, we examined *how* AI background shaped each group’s interpretation of explanations and highlight *why* their perceptual differences might exist.

We elucidate the *why* behind the group differences using the conceptual lenses of *heuristics* (mental shortcuts) [75, 140] and *appropriation* (users’ repurposing of a design) [40, 117, 138]. In light of the findings, we share concrete design implications around mitigating the risks of overreliance on numbers which can potentially lead to negative consequences such as over-trust on XAI systems. We share broader lessons around how our insights can help re-imagine AI education and mitigate potential harmful manipulation with explanations. By bringing conscious awareness of the group differences to the human-centered design of XAI systems, we address the AI creator-consumer gap by making the following contributions:

- We quantify user preferences (*what*) of three types of AI explanations along five dimensions of user perceptions.
- We qualitatively situate *how* one’s AI background (or lack thereof) influences the perception of the explanations.
- We elucidate *why* the group differences or similarities might exist and interpret them through the conceptual lenses of heuristics and appropriation.
- Using our findings, we identify potentially negative consequences (like harmful manipulation of user perceptions and over-trust in XAI systems) and propose mitigation strategies.

2 BACKGROUND

In this section, we review related work in the field of XAI salient to the paper, highlight the need to attend to XAI’s sociotechnical

¹By highlighting the creator-consumer gap in XAI, we do not mean to undermine the diversity of stakeholders in the ecosystem. By calling attention to extreme ends, we are highlighting the severity of the gap while fully acknowledging the ecosystem’s diversity. See, e.g.: [60, 111, 116, 130].

dimensions and human-centered perspectives, and discuss HCI work studying how user background shapes users’ perception of and needs for technology that motivated our work.

2.1 Explainable AI

While the origin of “explainable AI” can be traced back to expert systems in the 1980s [150], the field of XAI has been undergoing a resurgence due to the proliferation of complex Deep Learning models. Although there is a current lack of consensus on the meaning of explainability and related terms such as interpretability [10, 135], XAI work shares a common goal of making the AI systems’ decisions or behaviors understandable by people [2, 47]. Among other dimensions to map the landscape of technical XAI approaches, the field differentiates between methods to build directly interpretable models and methods to generate explanations for opaque-box models [56, 95, 134, 169] (for a detailed overview see recent survey papers [2, 10, 61]). While simpler models such as linear regression and decision-tree are typically considered directly interpretable but low-performing, recent work (e.g. [35, 165]) focuses on developing new algorithms that produce “clear-box” models allowing “under the hood” inspection without sacrificing performance.

In contrast, *explanation generation* methods—used in this paper—aim to explain models that are not directly human-understandable (e.g., deep neural networks). They are often post-hoc techniques [47, 96, 112, 134, 169] that could be applied after model building. Typically, these methods rely on distilling a simpler model from the input and output [105, 134] or meta-knowledge about the model [47] to generate explanations that approximate the model’s behavior. While there is, by design, a loss of scrutability, these methods allow the flexibility to make any model explainable, and thus have become popular and been applied to transforming simulation logs to explanations [157], intelligent tutoring systems [31], transforming AI plans into natural language [156], and translating multi-agent communication policies into natural language [9]. By privileging accessible understanding over revealing “under the hood” model mechanisms, explanation generation methods can be geared toward non-AI experts. In this paper, we focus on a specific explanation generation technique called *rationale generation* [47]—a process of producing a natural language explanation for agent behavior as if a human had performed the behavior and verbalized their inner monologue. While explanations can be in any modality, rationales are natural language-based, making it especially accessible for non-AI experts [47].

2.2 Towards Human-Centered XAI

There has been a growing recognition that XAI systems are often developed without an understanding of the recipients’ needs and characteristics [45, 112]. For instance, many of the XAI techniques created to support explainability needs during model development [134, 139] may break down when it comes to serving end-users with different needs [91]. It is imperative to follow human-centered approaches to understand the “personal, social, and cultural aspects” [71] of the recipients of AI explanations, especially since a monolithic view of the *who* may inadvertently risk dehumanization [23, 81]. Given deployment in high-stakes settings, AI systems designed without attending to the needs and values of

different stakeholders may also risk marginalizing certain groups or exacerbating existing inequities [11, 127].

The need for human-centered approaches in XAI has inspired increasing efforts among HCI and CSCW researchers, following the community’s long-standing tradition of designing and studying explainable computing systems [48, 52, 84, 93, 133]. Studies empirically evaluating XAI techniques in specific use contexts reveal the divergent needs and preferences of users [7, 22, 27, 41, 65, 77]. For example, while data scientists might need multiple XAI tools for a comprehensive understanding [65], simple explanations are often sufficient for AI-novices [27]. User characteristics such as cognitive load disposition [55] and general trust in AI [41] could moderate how users perceive AI explanations. Some studies further reveal potential drawbacks of AI explanations—how explanations could impose undesired cognitive burden [1], create a false sense of security and over-trust [55, 77], and how even placebo explanations (devoid of justificatory content) can engender trust in AI systems [49].

Researchers have begun to examine people’s cognitive process of interpreting AI explanations, which could help us understand atypical and misaligned user receptions of XAI. Recent work highlights the dual-process of cognition when people process AI explanations [20]. The dual-process theory [75, 132, 163] posits that people’s cognitive processes follow two systems: System 1 processes stimuli in a fast and automatic manner, whereas System 2 engages in deliberative and analytical thinking. System 1 often relies on heuristics (rules-of-thumb or mental shortcuts) that can be developed through past experiences. These heuristics, if applied inappropriately, should be considered cognitive biases [75]. In XAI, there is often an assumption that people mostly engage in analytic System 2 thinking whereas there is growing evidence that people mostly engage in System 1 thinking [20]. Negative consequences of cognitive biases such as over-trust in XAI could be attributed to a System 1 heuristic of associating explanations with AI competence [29, 124]. One way to mitigate these biases would be to use *Cognitive Forcing Functions* (CFFs)—interventions that disrupt heuristic reasoning and promote System 2 analytical thinking [32].

Human-centered XAI calls for pluralistic explanation design—for the *who*, not just the *what*, of XAI [45]. Recent XAI work has begun to differentiate major categories of XAI consumers, from builders to regulatory bodies [10, 38, 114, 153]. While this work is informative, there is still a dearth of empirical insights and understanding of the differences between these XAI users, and actionable design guidelines to support them.

A generative approach towards pluralistic design of XAI, we believe, is to develop a systematic understanding of *how* and *why* users with different characteristics (e.g., with or without AI backgrounds) form different perceptions (*what*) of XAI systems. Such insights can refine our understanding of *who* the humans are in XAI. To our knowledge, this has not yet been systematically explored in the specific context of XAI. Thus, our work adds to the discourse by adding empirical insights through a systematic exploration of two different *who*’s in XAI. Our work is also motivated by broader research studying individual differences that impact human-AI interaction, as reviewed below.

2.3 Individual differences in human-AI interaction

There is a rich history of transforming insights into how individual differences impact technology perception into the design of personalized, accessible, and inclusive technologies. For example, there is a large body of work on how various types of epistemic background, including computer literacy [39], digital literacy [12], and numeracy [72], impact user competency to use computing systems and ways to mitigate the competency gaps.

Recent work has paid attention to how user characteristics impact human-AI interaction. For example, studies on human-robot [33, 88, 146] and human-agent interaction [90, 92] found that user’s schema, whether an agent is seen as a utilitarian tool or a social entity, leads to noticeable differences in interactions with and evaluations of the agent. Research around AI fairness has identified various mediating factors such as users’ education level, fairness criteria, and general trust in ML [41, 162]. In short, individual differences and user characteristics shape perceptions of AI systems and should be appropriately understood and carefully accommodated when designing them.

A commonly studied user characteristic in Human-AI interaction is users’ knowledge about AI, often operationalized as AI programming experience [119], building ML models [55], or type of profession (e.g., data scientist) [65]. Recently Long and Magerko provided a concrete definition of AI literacy [104] with a set of core competencies to understand and use AI, and proposed design guidelines to mitigate the competency gaps. Recent work also examines the implications of background in AI as a determining factor of one’s *role* in an AI ecosystem. Motivated to “problematize the asymmetric relationship between technical experts and users”, Cheon and Su [28] highlighted the misalignment between the creators’ (roboticists’) vision of how consumers (end-users) interpret and use the product—a salient theme in this paper as well. McDonald and Pan [110] examined how CS students (likely to become AI developers) viewed ethical problems in AI, found substantial limitations, and called for a closer integration of ethics education with technical training.

Our understanding of how people’s AI backgrounds might impact their perceptions of AI explanations also draws from a long line of social science research on the relations between sensemaking [166] and professional knowledge [59]. Passi and Jackson [128], for instance, analyzed academic learning practices in the field of data science to show how students start “seeing” the world differently once they learn to work with algorithms and numbers. Through ethnographic fieldwork, they highlight how having a computational background enables students to gain actionable forms of “data vision”—ways of seeing that allow them to approach and analyze the world as data.

A person’s AI background (or a lack thereof)—a focal point of this paper—in fact, directly impacts user perceptions. In a recent ethnographic study [129], researchers found that data scientists and business analysts perceived an AI system’s accuracy score differently: business analysts saw the score as a measure of overall performance (good vs. bad), while data scientists perceived more granular insights into types of errors (false positives vs. false negatives). As people learn specific ways of *doing*, it also changes their

own ways of *knowing*—in fact, as we argue in this paper, people’s AI background impacts their perception of *what* it means to explain something and *how*.

Thus, the *who* questions are important, because people tend to interpret technologies differently, leading to different usages of those technologies [8, 18, 40, 117, 138]. The process of differential interpretation, followed by different usage, has been called *appropriation*—e.g., usage patterns that go beyond the original designers’ expectations [83, 154]. Dix listed six principles of user appropriation of technologies, including support for interpretive flexibility [40]. In this paper, we go beyond the insight that different people need different explanations. Building on the discourse of user characteristics in XAI [125, 149], we extend the literature through an in-depth exploration of how and why interpretation and use of explanations may differ between people with or without AI background.

3 STUDY DESIGN AND METHODS

We begin by sharing the research questions (RQs) followed by how we operationalize key aspects of the research design.

- RQ1: “Quantitatively, how do people with different AI backgrounds perceive different explanations given by AI agents?”
- RQ2: “Qualitatively, how and why do differences in AI background result in different or similar perceptions of explanations?”

We address these RQs by conducting a within-subjects experiment in which two groups of participants, with or without an AI background, see three versions of AI agents (depicted as robots in our study, Fig. 1) with different types of explanation. Using this 2x3 factorial design, we quantitatively measure the perceptions of the AI explanations and compare the differences between the two participant groups to address RQ1 (Section 4). Participants rank their preferences and justify their choices through open-ended text responses, which we qualitatively analyze to understand the underlying differences between the two groups to address RQ2 (Section 5). Below, we unpack how we operationalize three things: (1) explanation types, (2) user perceptions, and (3) backgrounds in AI.

3.1 Explanation Generation Method and Types

We begin with the *task environment* to situate the design of the AI agents. In the user study (task details in Section 3.4), participants watched 3 robots (AI agents) carry out an identical sequence of actions which differed only in the way that the AI agent “thinks out loud” about its actions. The robots need to navigate through a sequential decision-making environment—a field of rolling boulders and a river of flowing lava—to retrieve essential food supplies for trapped space explorers (Fig. 1). The robots thus need to observe a dynamic environment and think ahead in order to complete an objective. We chose a sequential environment because the explainability of sequential tasks is under-explored, while prior XAI work has explored non-sequential tasks (e.g., classification, captioning, etc. [161, 168, 170]). To solve the navigation problem, the agent used a Reinforcement Learning (RL) algorithm called *tabular Q-learning* [164]. Reinforcement Learning is both a promising technology for autonomous AI systems but also challenging from an explanation perspective. Tabular Q-learning agents attempt to learn the utility (called a Q-value for “quality” of the action) of different

actions in different situations. Once learning is complete, the agent makes decisions by picking the action with the highest Q-value. In our case, the fully trained agent solved the environment, generating an action trace that contains the sequence of steps needed to reach the goal (food supplies) without failure. Recall that our study has 3 robots (AI agents) with different types of explanation. In order to standardize their actions across experimental conditions, *all robots use the same trace*. This means that the decision-making mechanism underlying each robot is the same RL-based one. They *only differ* in how they explain their actions.

Explanation generation: The three different types of explanation are the within-subject variable in our study comparing perception differences between the AI and non-AI groups. Based on a review of related work, we chose three types of explanation that vary in their justification quality and representation modality (e.g., textual, numerical). With these considerations in mind, we set up the robots to express themselves in three ways: (1) natural language with justification (explaining the “why” behind the action; Fig 1, left), (2) natural language without justification (describing “what” the action was, Fig 1, center), and (3) numbers that determine the agent’s actions (akin to “transparent” AI, in this case showing Q-values, Fig 1, right). We share the explanation mechanisms and the attributes of the three robots below.

- The *Rationale-Generating (RG)* robot (Robot A in the study): this robot “thinks out loud” in natural language rationales explaining the “why” behind the action (#1 above). Our generation approach is similar to prior work in XAI and HRI [34, 43] where they use a neural machine translation (NMT) [106] approach to produce plausible rationales to explain sequential behavior. We extend and adapt it to fit our sequential environment depicting a space mission. The RG robot’s expressions aim to provide a *functional* understanding [102, 103] of the actions by appealing to functions or goals of the agent (e.g., Fig 1, left). The RG robot has language and justification.
- The *Action-Declaring (AD)* robot (Robot B in the study): this robot “thinks out loud” by stating its action in natural language without any justification. It’s a neutral option that simply states “what” the action is (#2 above). For instance, it states “I will move right” as it moves right. These outputs are generated from pre-fixed expressions that are triggered based on the agent’s action.
- The *Numerical-Reasoning (NR)* robot (Robot C in the study): this robot “thinks out loud” by simply outputting the numerical Q-values for the current state with no language component (#3 above). It is akin to a “transparent AI” where we directly look inside the opaque RL-box by observing its Q-values. Q-values can provide some transparency into the agent’s beliefs about each action’s relative utility (“quality”). However, Q-values themselves do not contain information on “why” one action has a higher utility than another. Note that we do not indicate that the numbers are Q-values in our study nor indicate which value is associated with which action.

While all robots “think out loud”, only the RG robot is designed to have any justificatory quality—it is designed to provide a functional understanding [102] of the “why” behind a robot’s action. The

Figure 1: The three robots navigating the task environment and explaining their actions. From left to right, the robots and their colors are: Rationale-Generation (in blue), Action-Declaring (in purple), and Numerical-Reasoning (in green). In the screenshot, each robot is taking the same action, but they are explaining it differently. The explanation text accompanying each robot is taken verbatim from the videos participants watched. To improve legibility, the text has been remastered to a higher resolution.



other two conditions provide baselines that should be considered as lacking justificatory qualities by design. The AD robot merely states “what” the action was (a neutral option). While NR could, in theory, provide a mechanistic understanding [103] of the “how”, the unlabelled numerical format should make its meaning difficult to access, if not impossible. While AD and NR do not have explanatory qualities by design, we do not know if or how participants will interpret them. Through the experiment, we are interested in whether and how these explanations invoke different perceptions in the two groups.

3.2 Grounding the Dimensions of User Perceptions

We now motivate & ground the dimensions of perceptions. To scope dimensions (measures) of perceptions appropriate for our use case, we engaged in an iterative filtering process—this process included (1) a systematic review of related work around trust, acceptance, and engagement of autonomous or AI systems followed by (2) informal interviews with six experts spanning HCI, AI, and HRI. The aforementioned process informed the adaptation of the following dimensions from classical technology acceptance models and emerging work in HRI and XAI literature [17, 34, 36, 47, 160].

One of the core goals of XAI is to make AI systems more understandable. The improved understandability (or lack thereof) can impact one’s trust or confidence in the system. In line with these facets, we adapted *understandability* and *confidence* from TAM & UTAUT [36, 160]. Prior work in HRI and AI shows that tolerance to failure [37] and perceived capability of the AI system [82] are impacted by how one perceives how intelligent the agent is; thus we added *intelligence* to our list of user perceptions. Recent work in autonomous system acceptance [144] shows that sociability factors are core markers of user adoption [94]. We adopted the dimension of *friendliness* or how friendly an agent appears because of its impact on relationship development [68, 152] and partnership [16, 21], which is essential for human-AI collaboration [120–122]. Emerging work in XAI and HRI [85, 89, 113, 141] suggests that how an AI agent communicates failure governs future collaboration relationships with humans. Thus, we add the notion of a *second chance* to understand how past failures impact future collaboration chances for XAI agents. In our study, participants ranked the robots along these *five* perception dimensions and justified their choices using open-ended text responses:

- (1) *Understandability*: Based on their explanations, I found each robot’s explanation of its actions understandable in the following order
- (2) *Confidence*: Based on their explanations, I would rank my confidence in each robot’s ability to do its task in the following order
- (3) *Intelligence*: Based on their explanations, I would rank each robot’s intelligence in the following order
- (4) *Friendliness*: Based on their explanations, I would rank the friendliness of each robot in the following order
- (5) *Second chance*: Each robot failed. Based on their explanations, I’d rank my willingness to give another chance to each robot in the following order

3.3 Participants: Operationalizing AI vs. non-AI Backgrounds

We now address how we operationalized the user background in our study. We acknowledge that AI background can be operationalized in different ways. As a *formative first step*, we use a “high contrast” approach where there is a stark difference between the groups. That is why, we focus on people *with* and *without* an AI background. This high-contrast approach provides a baseline understanding and creates the foundation for future granular operationalizations (e.g., years of AI experience). Below we share the motivation and operationalization of each group formation.

3.3.1 The AI Background Group.

For the AI group, we recruited participants who are students enrolled in CS programs and taking AI courses. Granted there are other ways to operationalize this group (e.g., recruiting AI practitioners), as a first step, we chose students because it allows us to explore how their AI coursework impacts the way they make sense of explanations from AI systems. Since professionalization starts with academic training [128], investigating the roots of one’s AI background and the impact on how students make sense of AI systems is important. In fact, as we show in this paper, *especially* during their learning phase, AI students adopt reasoning artifacts that impact their perceptions of explanations from AI systems in very specific ways, which has core implications for the future of XAI design (a point we elaborate in Section 6.3). Granted that not every AI student will go on to build AI systems but with the proliferation of AI systems in the workplace, a majority of these students could become stakeholders residing on the creation or development

end of the technology spectrum—as potential developers, designers, and managers of AI-based systems. As potential creators of AI systems, their perceptions matter in bridging the creator-consumer gap in XAI.

3.3.2 The Non-AI Background Group.

For the non-AI group, we recruited participants from Amazon Mechanical Turk (AMT). Carefully screened AMT participants have been shown to be representative of consumer research [14, 58, 79], which facilitates our goals of comparing potential creators (the AI group) with potential end-users (the non-AI group). We acknowledge that consumers too can and do have significant AI backgrounds, which is why we systematically screen out people in the non-AI group with any level of AI knowledge (using the screening process outlined in 3.3.3). Non-AI students, albeit an intuitive comparison group, would be a subset of the larger consumer base making them a good candidate for future granular investigations. Multi-disciplinary research has also shown how AMT participants can be reasonable alternatives to a university participant pool in terms of data integrity [13, 63, 145]. Weighing the affordances and limitations, AMT participants are thus a reliable and accessible comparison group, one that reasonably satisfies our initial desire to create a high-contrast comparison between potential creators and end-users of AI systems.

Our operationalization of the two groups' AI backgrounds aligns closely with the human-grounded evaluation proposed in [42], in which participants conduct controlled tasks to get a formative sense of the affordances of the explanations. We acknowledge that there are limitations to this experimental setup and that our insights should be scoped accordingly (more in Section 6.3). However, as we will see later, even with a carefully controlled task, we discover surprising, non-intuitive insights about how different groups interpret explanations.²

3.3.3 Recruitment and Screening Methods.

All participants received US \$10 for their time. The **AI group** had 96 US-based participants (39% self-identifying females; rest as males) taking 31.1 minutes on average for task completion. We recruited undergraduate students enrolled in an AI course at a large public research university located in the US. Typically taken in the 3rd year, this is a keystone course in an AI degree specialization track, implying that a significant number of students have expressed longitudinal interest in AI. The course is taught using the most widely adopted textbook (by Russell & Norvig [137] taught in over 1500 schools worldwide [6]) and using a widely adopted set of assignments. While there is no guarantee that all students will be future AI creators, the faculty believes that we can reasonably assume that many students aspire to have careers in the development of AI technology. In the course, students learn and implement many foundational AI concepts; for instance, Markov Decision Processes and Reinforcement Learning. Our study was deployed after students had taken exams on these concepts. For the **Non-AI group**, we recruited participants from Amazon Mechanical Turk (MTurk) using TurkPrime [97]. The non-AI group had 83 US-based participants

(46% self-identifying females; rest as males) taking 29.8 minutes on average for task completion.

Establishing group differences while controlling for shared characteristics: We systematically screened *both* groups and ensured they are *measurably* different. Every participant had to pass three parts of our screener: (1) A knowledge test on CS & AI (collaboratively developed with the course's teaching staff to ensure relevancy); (2) Self-reported knowledge levels in computer programming and AI using two 5-point Likert-scales; (3) confirmation of whether they have ever taken an AI class. The *cut-off points* were decided through iterative piloting and consultation with the teaching staff of the AI course. The **non-AI group**, by design, should have "No knowledge" [= 1] in both programming and AI along with no prior AI classes. Out of 5, the **AI-group** should score [≥ 4] ("moderate" knowledge or more) for (1) [knowledge test]; for (2) [self-reported AI knowledge], it should be [≥ 3] ("some knowledge" or more).

To establish that these *two groups are measurably different*, we performed statistical tests. Mann Whitney U-test with Bonferroni correction showed that the groups are different—for the AI knowledge test, $p < 2.2 \times 10^{-16}$; for the self-reported programming knowledge, $p < 2.2 \times 10^{-16}$. Thus, we were able to establish two "high contrast" groups. For further details, including screener questions, please refer to ?? in the Appendix, uploaded with Supplementary Materials. Beyond the differences, we controlled for factors such as age and education levels. All participants were adults up to 25 years old, using smartphones and laptops every day. The non-AI group self-reported an average education level of 4.95 (4= "Vocational Training", 5= "Some college/Associate's degree") while the AI group's average was 5. Mann Whitney U-tests showed that the groups did *not* differ along *age* ($p = 0.505$) and *education levels* ($p = 0.146$).

3.4 Procedure: Task Details

We now discuss study mechanics—after providing informed consent, participants watched an orientation video outlining the scenario. Participants were asked to imagine themselves as space explorers faced with a search-and-rescue mission involving robots. They face a life-and-death situation where they are stuck on a different planet and must remain inside a protective dome. Their *only* source of survival is a remote supply depot, which they cannot reach. They must rely on autonomous robots, *ones they cannot control*, to navigate through a field of boulders and a river of flowing lava to retrieve the essential food supplies (See Fig 1). Participants could only see a non-interactive video stream of their activities through their "space visors". This non-interactiveness aimed to heighten their sense of lack of control (and thereby reliance on the robots). Since the robots took identical actions during the task, participants were asked to pay special attention to the only differentiating factor—the way each robot explained its actions.

After orientation, participants watched 6 counterbalanced and randomized videos showing the three robots succeeding and failing to retrieve the essential supplies using identical sequences of actions. To mitigate the effects of preconceived notions, we did not use any descriptive names for the robots; instead, we introduced the robots as "Robot A," "Robot B," and "Robot C" for the RG, AD, and NR robots respectively. Participants, by design, had no idea about how

²*Practices for better data:* We share some strategies that helped us gather high-quality data and maintain rapport with our participants throughout the study (e.g., fair payment structure, active engagement participants, etc.) in the Appendix (A.1), uploaded with Supplementary Material.

each robot generated its expressions; for instance, participants were not informed that NR's numbers are Q-values. To them, it was just another different way of explaining the actions. As mentioned in 3.1, the goal was to see how and if how people can make sense of them even if they are unlabelled. After watching all the videos, participants ranked the robots (1st, 2nd, and 3rd– no ties allowed) along the five (5) dimensions of user perceptions highlighted above (in 3.2). After ranking on a dimension, participants justified and contextualized their ranking using a mandatory free-text response.

4 QUANTITATIVE RESULTS

We conducted within-group and between-group analyses. The *within*-group analyses give a sense of which robot's explanation was preferred along each dimension of user perception. The *between*-group analyses tell us how the two groups are the same or different when considering the robots across multiple dimensions and set the stage for the qualitative analysis in Section 5. Taken together, the quantitative results address RQ1 around quantifying the relative perceptions and preferences of the two groups along the five dimensions.

4.1 Within-group Comparisons

Figure 2 shows how each robot was ranked in each dimension. The “violin chart” visualization makes similarities and differences easier to distinguish. For example, large differences are exemplified in the *Friendliness* rankings given by the non-AI group participants (bottom, second left). A wide area at the top of the blue plot shows that the Rational Generating (RG) robot was ranked first by most. In contrast, the plot for the *Intelligence* dimension by the non-AI group (bottom, middle) shows that all robots received the first, second, or third rank comparably similar number of times.

For each group, we conducted five Friedman tests [171] of differences among repeated measures (one for each of the five dimensions). We used the maximum-type (max T) implementation of the Friedman test, which controls for the family-wise error rate [171] to determine *whether* any preferences between robots were detectable. To determine *which* robot was preferred, we used the Wilcoxon-Nemenyi-McDonald-Thompson test [66] to make pairwise comparisons between the robots, for each participant group, and within each dimension. The results are summarized in Table 1.

From the within-group analysis (Table 1 and Figure 2), we have some notable insights: across each dimension, the AI-background group unambiguously preferred RG to the other robots, particularly over AD. However, the RG robot is not the unanimous winner across each dimension for the non-AI group– here, participants show no preference between RG and AD in 4 out of the 5 dimensions (RG wins over AD in *Friendliness*). For the non-AI group, AD wins over NR across all dimensions except for *Intelligence*, which is a noteworthy dimension–the non-AI group showed no preferences between the robots, but the AI group felt NR was more intelligent than AD. This is the only time where NR wins over AD, highlighting an important point that we explore in our qualitative findings (Section 5.1 and 5.2) around the AI group's preference for numerical representations.

4.2 Between-group Comparisons

To detect group differences in the preference for the robots, we used Ordinal Logistic Regression (OLR) [5, 109, 155, 159], an extension of Logistic Regression when the response variable is ordinal. To model a 3-level categorical variable (Robot Type), OLR requires us to analyze two variables holding one as a reference (constant): in our case, we analyzed AD and NR, holding RG constant. We investigate the interaction effects as well as changes in reference levels in the OLR analysis that reveals the relative impact in ranking the robots between the groups. To investigate the effect of the dimensions, we explore the interaction effects between Robot Type (RG, AD, and NR) and the Participant Group (AI vs. non-AI) into the OLR model. Changing the reference levels of Robot Type and Participant Group allows us to isolate the interaction effects, which we interpret similarly to [70]. A full list of OLR tables is provided in the Appendix (Table 5– 10).

If we group everyone regardless of their AI backgrounds, we find that $Rank_{RG} > Rank_{AD} > Rank_{NR}$. The odds of receiving a higher ranking for the RG robot is 5.5 times that of the AD robot, whose odds are 2.66 times that of the NR robot (see Tables 3 and 4 in the Appendix (A.3)).

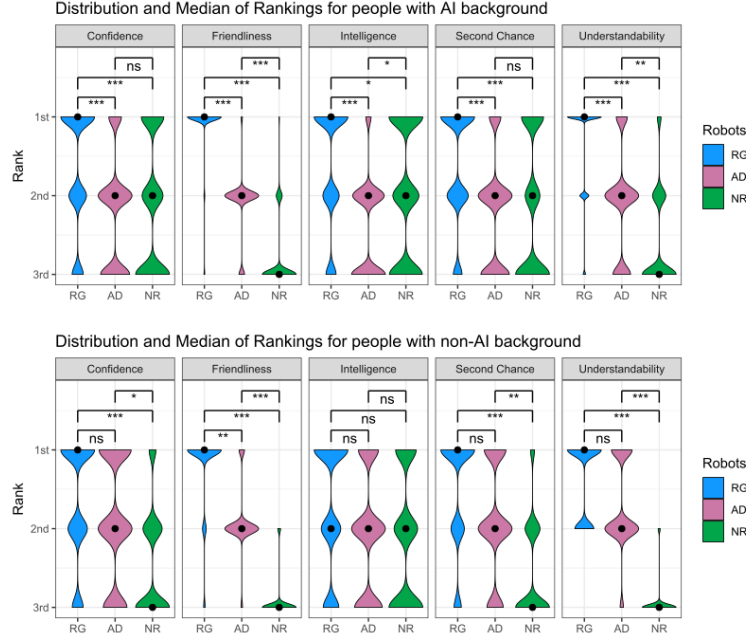
Between the groups, there is no significant difference in ranking RG Robot—it is always the top choice across all dimensions (Table 2, top row). However, the groups exhibit differences in their preferences when it comes to AD and NR. The Non-AI group shows more preference for the AD Robot (odds ratio= 1.986) while the AI group shows more preference for the NR Robot (odds ratio= $1.0/0.465 = 2.15$).

Table 2 also provides the odds ratio and the p -value for each Robot Type per dimension. For the RG Robot, all of its p -values are greater than 0.05 for each of the five dimensions resulting in no significant pattern of preference between the groups. Conversely, the p -values of AD Robot are all smaller than 0.05 meaning that, when it comes to ranking the AD Robot, there is a significant difference between the Non-AI group and the AI group. Moreover, all odds ratios related to AD Robot in Table 2 are greater than 1.0, indicating the Non-AI group shows a stronger preference for the AD Robot than the AI group for each of the five dimensions.

Last, for the NR robot, the p -values related to *Confidence*, *Second Chance*, and *Understandability* are also significant. The odds ratios for these dimensions are less than 1.0 indicating that the AI-group is more likely to rank the NR robot higher than the non-AI-group on these dimensions. There is no significant difference between the Non-AI group and the AI group when it comes to ranking the dimensions of *Friendliness* and *Intelligence*.

4.3 Conclusions of Quantitative Analyses

Overall, while we find that $Rank_{RG} > Rank_{AD} > Rank_{NR}$, there are significant differences in the AI and the Non-AI ranking behaviors. In other words: AI background does change perceptions of explanations along the dimensions we identified. In particular, having a background in AI correlates with having a greater preference for the NR Robot over the AD Robot. While AD wins over NR in most head-to-head comparisons, NR wins it on *Intelligence* for the AI

Figure 2: Distributions of rankings for each robot, in each dimension, separated by participant group

Note: The black horizontal bars indicate whether a pair of distributions is significantly different. ns = not significant, * $p < .05$, ** $p < .01$, *** $p < .001$. The width of each violin plot at each ranking level indicates the proportion of people who assigned that rank to that robot. The black bullet (•) refers to the median rank.

Table 1: Summary of p -values for pairwise comparisons, showing which robots were preferred.

Dimension	AI background		Non- AI background	
	Robot	p -value	Robot	p -value
<i>Confidence</i>	RG vs AD	< 0.001	RG vs AD	0.362
	RG vs NR	< 0.001	RG vs NR	< 0.001
	AD vs NR	0.869	AD vs NR	0.014
<i>Friendliness</i>	RG vs AD	< 0.001	RG vs AD	< 0.001
	RG vs NR	< 0.001	RG vs NR	< 0.001
	AD vs NR	< 0.001	AD vs NR	< 0.001
<i>Intelligence</i>	RG vs AD	< 0.001	RG vs AD	-
	RG vs NR	0.021	RG vs NR	-
	AD vs NR	0.011	AD vs NR	-
<i>Second Chance</i>	RG vs AD	< 0.001	RG vs AD	0.083
	RG vs NR	< 0.001	RG vs NR	< 0.001
	AD vs NR	0.902	AD vs NR	0.003
<i>Understandability</i>	RG vs AD	< 0.001	RG vs AD	0.187
	RG vs NR	< 0.001	RG vs NR	< 0.001
	AD vs NR	0.002	AD vs NR	< 0.001

Note: Favored robot and significant p -values are in bold for each pairwise comparison.

group. The two groups are indistinguishable concerning *Friendliness*. The next section helps us understand why these trends are present in the quantitative results and their implications.

5 QUALITATIVE ANALYSIS & FINDINGS

While quantitative analysis tells us *what* is different between groups, qualitative analysis sheds light on *how* and *why* those differences

manifest, addressing RQ2. Quantitative analysis puts RG as the clear winner; however, as we show in this section, the story is more nuanced and interesting, especially the underlying reasons behind the group differences between AD and NR. Below we describe the process of our qualitative analysis before moving on to the findings.

The qualitative data was coded according to principles of *grounded theory* analysis [24, 147]. Coding and analysis of the qualitative data were done by the first and second authors in two stages, each

Table 2: Robot Preference Summary across Dimensions (Non-AI [baseline] vs. AI)

Dimension	RG Robot		AD Robot		NR Robot	
	Odds ratio	<i>p</i> -value	Odds ratio	<i>p</i> -value	Odds ratio	<i>p</i> -value
All dimensions	0.855	0.323	1.986	< 0.001	0.465	< 0.001
Confidence	0.971	0.930	2.003	0.026	0.487	0.029
Friendliness	0.280	0.051	2.827	0.017	0.436	0.160
Intelligence	0.717	0.315	1.959	0.031	0.646	0.178
Second Chance	1.369	0.356	1.965	0.028	0.375	0.005
Understandability	0.659	0.255	3.088	0.005	0.114	0.006

Note: Odds ratios > 1.0 indicate the non-AI group prefers the robot more than the AI group.

consisting of multiple rounds of iteration. In the first round, the authors separately performed coding using *in-vivo* codes, which involves generating codes from the data (e.g., using participant phrases such as ‘the robot knew what it was doing’ and ‘easier to understand’ as codes). Through discussion, the authors generated a merged open coding scheme. In the second round, the authors analyzed the data using *axial codes*. Axial coding involves finding connections between open codes and classifying them into different categories (e.g., potential actionability of numbers). In the last step, the authors analyzed the different sets of axial codes, unified them into high-level themes, and consolidated them into selective codes.

After this, we grouped the data along the two groups (AI vs. non-AI) and our five (5) analytic perception dimensions that we motivated in Section 3.2 and also used for quantitative analysis (*confidence*, *friendliness*, *intelligence*, *second chance*, and *understandability*), allowing us to compare our findings across dimensions and groups. Finally, based on the grounded theory heuristic of “constant comparison” [57, 147], the authors frequently compared and contrasted axial codes across perception dimensions to tease out the similarities and differences between the different reasonings used by the AI and non-AI groups to make sense of the explanations provided by the three robots.

Below, we report on the most salient themes (selective codes) from our analysis. In particular, we showcase (1) how irrespective of their AI background both groups exhibited unwarranted faith in numbers for different reasons and (2) how each group saw explanatory value in explanations that were not designed with justificatory qualities. For each theme, we highlight distinct categories of reasoning (axial codes) used by the groups to justify their choice of robots across the different perception dimensions.

5.1 Unwarranted Faith in Numbers

Participants in both groups had unwarranted faith in numbers. However, their *extent* and *reasons* for doing so were *different*. On the one hand, AI group participants often ascribed more value to numbers than was justified. On the other hand, some non-AI group participants believed that numbers signaled intelligence even if participants could not capture the numbers’ meaning. Below we highlight two major ways in which participants misplaced faith in numbers.

The mere presence of numbers was associated with an algorithmic thinking process in the robot even when the meaning

of numbers was unclear. Both groups exhibit this perception of algorithmic thinking. We will begin with the AI group—this group ascribed higher-order cognitive abilities to the robot with numerical representations. Between AD and NR, the AI group found AD more understandable (Table 1) but deemed the NR robot as the more intelligent of the two (the only time NR wins over AD across all dimensions— (Table 1)). It seems contradictory that the AI group found the *less understandable* robot to be *more intelligent!*—the main reason for this is that the presence of numbers fostered the “assumption that [the NR robot] uses some sort of [an] algorithm” (A50) in the AI group. The perception of “under-the-hood math, boosted [NR’s] trustworthiness” (A49). Some explicitly compared AD and NR robots and concluded that the mathematical representation demonstrated a method for the NR robot’s behavior:

“With [the NR robot], while I did not understand its methodology, I could see that it was using some mathematical calculations to determine which way to move. [...] With [the AD robot], I could not see any methodology or signs of decision-making” (A23).

A small minority in the AI group indeed figured out that “the numbers 0-4 represented different actions, and [NR] would choose the action with the highest numerical value” (A23). An even smaller minority even guessed the numbers’ correct meaning—a “utility function and reward systems” (A64). *despite* the fact that the numbers were *unlabelled*. These participants projected meaning on the numbers, lending further credence to the role of data vision [59, 128].

What is surprising is their faith in numbers arose even when AI group participants did not “fully understand the logic behind [...NR robot’s] decision making.” (A43). **The AI group seems to have followed heuristic reasoning that associates mathematical representations with logic and intelligence**, e.g.: “logic must have been derived from a formula, [...which is] intelligent” (A54), or “Math [...had] an aura of intelligence” and “exact values” made the NR robot “feel smarter” (A16, A77, A75). This perception of logical thinking engendered unwarranted trust and made the NR robot seem like it “should theoretically succeed more than the others, making him more intelligent” (A37). The linkage between perceived logical thinking and higher cognitive abilities can elucidate why the AI group prefers NR over AD when it comes to *Intelligence* (Table 1). A few participants even claimed that they could “actually see the math that [the NR robot] was making decisions off of, [...making it

feel] more real” (A9). In fact, to them, it appeared as if the NR robot “*clearly had an algorithm* that worked [...] and] *seemed to know* what it was doing” even if they “did not know what it was going to do in the future” (A91, emphasis added).

Participants with AI background also viewed numbers as potentially actionable even when their meaning was unclear. *Actionability* refers to what one might do with the information to make sense of the robot’s behavior—“debug its faulty behavior, or predict its future behavior” (A76). While many highlighted that they could not “make sense of numbers right now, [they] believed that] in principle, [they] should be able to act on them in the future” (A39). The potential explanatory value in numbers was better than the vacuous statements of AD: “[The NR] robot gave mathematical results as explanation while [the AD robot] gave no explanation” (A19).

Many participants connected their AI background to their ability to work with numbers. They mentioned how their current “AI course [helped them] to understand what [NR robot] is doing and what the numbers might mean” (A52). Some even highlighted that if they “had a pen and some paper, writing down information, [they] could gleam [sic] some information based off the [numerical] patterns” (A28). But how *actionable* were NR’s numbers in actuality? The numbers were Q-values, which only indicate the relative strength of one action versus the others. Specifically, Q-values indicate the agent’s belief that certain actions lead to greater or lesser future reward, affording some amount of explanatory power. However, they cannot indicate *why* the agent has come to believe that one action is better than another—this information is not retained by the agent nor conveyed through the numbers. This did not deter the AI group from deferring to the authority of numbers. After all, while all the robots used the same AI algorithm to make decisions, the NR robot’s expressions seemed most ‘AI-like’.

These insights can help us understand why, even though both groups have misplaced faith in numbers, the AI group shows a higher inclination towards numbers (as evidenced in the quantitative results (Section 4.1, Table 1), the only time NR wins over AD happens when the AI group judges *Intelligence*).

Unable to access their meaning, the non-AI group associated numbers with the presence of a higher, more intelligent expression that, they argued, could only come from an intelligent agent. Since the NR robot was “communicating in a numerical language that’s too hard to understand”, the numbers had a “mystery and aura of higher intelligence” (NA22, NA33). The “language of numbers”, *because* of its “cryptic incomprehensibility”, signalled higher-order thinking (NA6, NA1): “*Because* I could not understand [NR’s] output, I deemed it to be intelligent” (NA30, emphasis added).

This can explain why the non-AI group showed no preferences between AD and NR when judging *Intelligence*, despite favoring AD over NR across all other dimensions (Table 1). To them, numbers signaled “precision”—an important quality of intelligence (NA8, 16, 21, 30, 34). Numbers gave the impression that the NR robot “was more technical” than others—its “precise numerical explanations” resulted from it “calculating everything” (NA53, NA21, NA12). To these participants, “anything that uses pure numbers is going to be more intelligent” because “numerical outputs are likely to be more precise [...] whereas textual representations involve a degree of uncertainty and subjectivity” (NA41, NA30). Such perceptions point

to how the modality of expression—numeric vs. textual—impacts perceptions of explanations from AI agents, where we see projections of normative notions (e.g., objective vs. subjective) in judging intelligence.

5.2 Unanticipated Explanatory Value

As designers, we had specific goals behind each robot’s mode of expression. As discussed in Section 3.1, RG was the only robot designed to have the justificatory quality to explain the *why* behind the robot’s actions, while AD and NR (as baselines) should be considered as lacking justificatory qualities by-design. RG won the overall competition (as discussed in Section 4.2), but what was surprising was that **both AI and non-AI groups found unanticipated explanatory value in AD’s declarative statements and NR’s numerical representations**. As we will show below, qualitative analysis revealed that the two groups had different *explanatory intent* (what they wanted to do with the explanation), which was closely associated with their AI background. On the one hand, the non-AI group found *affirmatory* value (confirmation of stable performance) in AD’s statements. On the other hand, the AI group overly ascribed *diagnostic* value (debugging in case of failure) to NR’s numbers even when they could not make sense of them.

For the non-AI group, their **desire for affirmation**—confirming the action without necessarily explaining it—played a key part in finding value in AD’s explanations. The affirmatory value manifests most clearly in their comparison between AD and NR in the dimensions of *Confidence* and *Understandability*. For both dimensions, the non-AI group preferred AD over NR (Table 1). Recall that AD merely declared its action, stating the *what*, not the *why*. Despite this, the non-AI group found value in the confirmatory information because there was alignment between what AD was doing and saying. It showed that “[AD] is consistent and nothing crazy is going on where it says it went right but in actuality, it went down” (NA34). For both dimensions, the non-AI group attributed greater value to AD’s declarative statements (compared to NR’s numbers) because it “at least said what its movements were going to be” (NA7). Its “brief,” “un-embellished,” and “easier to understand” language that got “straight to the point” boosted its *understandability* (NA14, NA23, NA28, NA7). In fact, AD’s “just the facts” (NA38) declarative and succinct nature were signs of *confidence* itself. It did not need to say much because “it knew what it was doing,”—evident in its lack of “any hesitation” and “business-like [style] focused on performing the task at hand” (NA17, NA41, NA11). In contrast, NR’s numbers were inaccessible to the non-AI group, thereby inactionable and valueless. In short, when we designed AD, there was no intention of offering value through confirmation; yet, **through their explanatory intent of affirmation, the non-AI group found value in AD’s explanations by interpreting them as signals for stable system performance**.

The AI group overly ascribed diagnostic value in NR’s numbers even when their meaning was unclear—they felt NR’s numbers had “diagnostic information that can be used to debug [the robot] in case of failure” (A39). The explanatory intent of *diagnosis* is a consequence of the AI group’s faith in numbers and related to their perceived actionability—what one might do with the information (as discussed above in 5.1). Analysis of the open-ended

text responses for the perception dimensions of *Intelligence*, *Second Chance*, and *Confidence* exhibits this explanatory intent most clearly. Recall that the AI group preferred NR over AD for Intelligence and felt there were no differences between them for Second Chance and Confidence (Table 1). This is in contrast to the non-AI group that preferred AD over NR for both of these dimensions. AI group members perceived NR's numbers to have more explanatory value simply because they felt they could do "more with numbers" (A30) in "cases of failure and troubleshooting" (A78).

For Intelligence, NR's numerical representation and "exact values" made it appear more "valuable" than AD's "inert" statements (A23, A51, A42). Even when it came to giving NR a second chance, numbers helped because the NR robot appeared to be "trying very hard since it provided [...] mathematical evidence of its moves" (A79). "Math-based decisions" (A2) made the NR robot appear "more reliable" (A8, A42) and worthy of another chance. Numbers inspired confidence because they signaled the existence of a "concrete methodology"—"the higher the number, the more optimal the move" (A66, A67), which added to the perceived explanatory value in them. Moreover, if there were a method, an algorithm, or a formula behind the NR's actions, the AI group participants believed that they could deduce it using the numbers:

"[NR] returned the most amount of data, and *if I were to understand what that numbers mean*, it'd be the most useful one to debug on repeating runs to analyze what it is doing and why" (A91, emphasis added).

The AI group participants connected their background to a desire to troubleshoot and repair things. Numbers, they felt, were more manipulable than language, catalyzing over-ascription of diagnostic value to them:

"As an engineer, I want to fix things. Numbers are concrete and objective, language is not. But you can manipulate numbers [...] With the AI stuff I'm learning, I can use it to diagnose [NR]." (A49)

Thus, we see, what Eriksson called the "seductive allure" [51, 101] exhibited through the perception of numbers in the AI group. This diagnostic intent showcases how the AI group over-ascribed explanatory value in numbers, highlighting how its unwarranted faith in numbers can manifest in potentially negative ways—a point we discuss in Sections 6.

In **summary**, *first*, both groups had undue faith in numbers for *different reasons* and to *different extents*. Both groups associated the presence of numbers to signal the presence of algorithmic thinking even when the meaning of numbers was unclear. The AI group seems to have employed heuristic reasoning, connecting mathematical representations with the notions of logic and intelligence. The non-AI group, lacking access to their meaning, correlated numbers with the existence of a higher, more intelligent expression. *Second*, both AI and non-AI groups found unanticipated explanatory value by appropriating AD's declarative statements and NR's numerical representations. The non-AI group had an explanatory intent of *affirmation* (confirming the action without explaining it). This enabled them to find value in AD's explanations by interpreting them as signals for stable system performance. Notably, the AD robot was not designed to offer value through confirmation. The AI group

overly ascribed *diagnostic* value in NR's numbers for potential debugging even when their meaning was unclear. The explanatory intent of *diagnosis* is a consequence of the AI group's unwarranted faith in numbers and related to their perceived actionability—what one might do with the information (as discussed above in 5.1).

6 DISCUSSION & IMPLICATIONS

In this section, we discuss possible causes of the observed differences between the groups, then discuss implications for designing XAI systems by accounting for users' background differences to bridge the AI creator-consumer gaps, and then the broader implications for explainable and responsible AI.

6.1 Discussion: User background impacting perception of XAI

We highlight two ways to interpret the potential causes behind group differences—**cognitive heuristics and appropriation**. First, we use the notion of *heuristics* based on the dual-process theory (reviewed in Section 2.2) to understand how unwarranted faith in numbers (Section 5.1) can emerge from different lines of thinking affected by one's AI background. Second, we incorporate the lens of *appropriation* to the unanticipated ways in which people find explanatory value (Section 5.2).

Heuristics and Faith in Numbers: Recent work (reviewed in 2.2 like [20, 55]) has highlighted that while XAI is often developed with an implicit assumption that the recipient will process each explanation through analytical System 2 thinking (from dual-process theory [75]), in reality people are more likely to rely on System 1 thinking by invoking *heuristics*—rules-of-thumb or mental shortcuts, which leads to biases and errors if applied inappropriately [75, 132, 163].

The notion of *heuristics* can help us understand the potential reasons behind the two groups' different faith in numbers. On the one hand, the AI group seemed to have an instinctual response to numerical values; they assumed that the numbers possessed all the information needed to manipulate, diagnose, and reverse engineer. There appear to be heuristics that associate numbers with logical intelligence, which could potentially be acted upon (e.g., through diagnosis). Such heuristics are likely formed and validated from their past experience working with numbers and algorithms. This heuristic is risky because, as we noted in Section 3.1, the numbers are Q-values and do not allow much actionability beyond an assessment of the quality of the actions available. We return to this risk and highlight design implications to address this in the next subsection.

On the other hand, lacking the AI background, some in the non-AI group seemed to have different heuristics where their very inability to understand complex numbers was associated with the presence of a higher-order intelligence. The non-AI group also lacks the requisite AI background to potentially access and deliberately think through NR's numbers (System 2 thinking). Thus, we see how different heuristics, tied to one's AI background, can lead people to the same outcome—faith in numbers. Note that these heuristics are not an exhaustive list, but a starting point to understand how group differences connect to their AI backgrounds.

Appropriation and Unexpected Explanatory Value: A second lens for understanding why participants found explanatory value in unexpected places is that of *appropriation*. As we shared in Sections 2.3 and 5.2, appropriation happens when end-users interpret and use technologies in ways not envisioned by designers [40, 117, 138, 154]. People do not passively absorb information—they interpret it, often processing it in unanticipated ways. This is what happened to AD’s declarative statements and NR’s numbers. Driven by different explanatory intents, each group appropriated the explanations in unanticipated ways. These different appropriations were, in part, a function of each group’s AI background that led participants to develop their own sense of how they can and cannot use explanations. What is striking is that the appropriation took place even in a controlled experiment like ours where the participants were not explicitly asked to take actions based on explanations. Even in passive interactions (robots engaged in one-way communication), participants envisioned themselves *using* the explanations in hypothetical scenarios. On the one hand, the AI group’s intent of *diagnosis* led many to envision scenarios where they would troubleshoot the robot. As a result, they might have misplaced or over-placed diagnostic value in NR’s numbers—even when they could not fully understand them. On the other hand, for the non-AI group, NR’s numbers are inaccessible, thus in-actionable. Their *affirmatory* intent drives them to find value in the confirmatory statements of AD, appropriating it as a signal for stable system behavior.

Our discussion extends the literature around individual differences and heuristics-based processing of XAI [20, 55, 124] in two ways: first, we explicitly highlight *what* heuristics people might use, *how* their AI background influences their thinking, and posit the *why* (underlying reasons) behind them. Second, we add the lenses of appropriation to the conversation, which has implications on user agency in design (points we touch on below). For both, we connect it to an important user characteristic—one’s AI background.

Both points around heuristics and appropriation highlight an important yet overlooked point around the *duality* of explanations. Explanations are *both products and processes* [99]. The product-centered view, common in psychology and XAI, often ignores the essential processes through which people make sense of explanations. However, the sensemaking process is as important as the explanation itself [100, 167]. Recent work calls for guiding the process of understanding explanations, especially for non-AI experts [29]. By investigating both AI and non-AI groups, our work extends the current XAI discourse— it directly speaks to the duality of explanations by focusing on both products (types of explanations from 3 robots) and processes (how one’s AI background influences the interpretation of explanations).

Taken together, these findings reveal that the design and use of AI explanations is as much in the eye of the beholder as it is in the minds of the designer— the user’s explanatory intent and common heuristics matter just as much as the designer’s intended goal. Users might find explanatory value where designers never intended to be and use them based on their explanatory intent. Contextually understanding the misalignment between designer goals and user intent is key to fostering effective human-AI collaboration, especially in XAI systems [54, 114]. We demonstrated how different groups find meaning in different places—even when the meaning is

misplaced. The ‘ability’ in explain-*ability* depends on *who* is looking at it and emerges from the meaning-making process between humans and explanations.

While the groups differed in other aspects, both preferred the RG robot because of its *command of language* exhibited through explanatory power (depth of reasoning) and variety (style and length). RG appeared to have a “good command of English deep enough to get at the why” (A29). Moreover, they found RG’s communication to be *relatable* and exhibited a *personality* even if the interaction was passive. It appeared to “include you in its thought process,” making you feel “as if you could talk to it” (NA11, A46). With RG robot, many felt “like you are having a conversation with a friend, analyzing some problem, and making the best decision for this problem.” (NA36) As RG navigated the terrain, its expressions included phrases like ‘I’m not just winning at life, I’m biwinning!’, which participants found humorous and engaging. Since participants felt RG was “the most humanlike in its explanations” (A35), they attributed *emotional intelligence* to it, which garnered empathy and support. Most people “rooted for [the RG robot because they] liked its personality, felt connected to it, and wanted it to succeed.” (NA27).

Given our findings, it might be tempting to conclude that RG-style explanations should be used at all times. This is where the *who* in XAI comes in— the *who* scopes *how* to show the *why*. If the explainability needs of the user are to achieve *functional understanding* [102], then RG-style explanations are strong contenders. RG-style explanations are also suitable if the goal is to make AI explanations accessible to non-AI experts [47]. However, if the goal is to provide a *mechanistic understanding* [103, 126], then RG-style explanations will not be actionable. Mechanistic understanding can be useful for debugging purposes. Here, a hybrid style could work where the functional understanding of RG is contextualized with the mechanistic affordances from Q-values in NR. Imagine an explanation format where the user can “unfold” the associated Q-values and ground the natural language explanation. There are engineering and resource costs for RG-style explanations. They need to be trained with human explanations, which can be a challenging data collection exercise for more complex environments. The full treatise of explanation design is beyond the scope of this paper; however, focusing on the needs and characteristics of the “who” is a good starting point.

6.2 Implications: Designing XAI for background differences

Here, we discuss the implications for designing XAI systems that accommodate users’ AI background differences. We focus on mitigating the risks of over-reliance on numbers which can potentially lead to negative consequences such as over-trust on AI systems [131]. Moreover, as those in the AI group are likely to be on the creation end while those in the non-AI group are likely to be towards the consumption and interaction end, our results have implications for bridging the AI creator-consumer gaps by bringing conscious awareness of the group differences to the design of XAI systems.

Our discussion around **heuristics** carries design implications for both groups. We noticed how the AI and non-AI groups utilize different heuristics (mental short-cuts) in System 1 (fast, automatic)

thinking to ascribe misplaced faith in them. Shifting people’s thinking from System 1 to System 2 (slow, deliberative) is not only an active area of research but is also a challenging one [26, 86]. There can be two actionable ways to tackle biases resulting from cognitive heuristics— first, locally *at the time* of decision-making, we can use Cognitive Forcing Functions (CFFs) [32] (prompts, delays, etc.) that can interject the heuristic reasoning, potentially allowing the person to engage in deliberative analytical thinking. Second, globally for future decision-making, we can utilize metacognitive strategies, often called cognitive forcing strategies, that include simulation training, increasing awareness of potential pitfalls of heuristics, etc.

To potentially prevent over-trust in numbers, we can do design interventions at the local and global levels. At a local design level, we can introduce CFFs that can break instinctive thinking patterns and promote mindful ones. What if the AI group members were prompted to reflect on their instinctive thinking through a combination of prompts and multi-modal explanations (blend between RG and NR)? Situating numbers in the context of language and vice versa can act as a CFF that could prompt the AI group members to reflect deliberatively (using System 2) and realize the limited nature of the numbers. At a global level, we can introduce simulation training that provides counterfactual (what-if) scenarios highlighting cases where numbers from the robots are erroneous or faulty (vs. correct Q-values). For the non-AI group members who associate the opacity of numbers with higher intelligence, we can introduce scenario-based examples that explicitly highlight how indecipherable numbers can also be gibberish and useless. The goal here is *not* to eliminate heuristic reasoning but to mitigate blind faith in a certain modality of explanation. Exposure to these scenarios can facilitate long-term calibration of trust in numbers.

Given the negative impact of cognitive biases, it might be tempting to exclusively design explanations that only promote System 2 thinking. This is also risky because it forces users to constantly engage in deliberative thinking, their satisfaction suffers due to higher cognitive friction [20]. We need to strike a balance between System 1 and 2 thinking to appropriately calibrate trust. To do so, we can bridge existing work in non-XAI settings (e.g., balancing System 1 and 2 thinking in clinical decision-making) and translate them to XAI use cases in a contextually relevant manner.

The **appropriation** of explanatory intent we saw from both groups also has important design implications. Our findings highlight that users will appropriate explanations regardless of careful design. This is not a bad thing. Given dynamic user goals and needs, it is impossible to preemptively exhaust all the interpretations. Our goal then should be to *support, not control* end-users by providing *resources* that mitigate improper or harmful appropriations.

To aid *appropriation-aware* design in XAI, an operationally viable path is to leverage emerging work on Seamful XAI [44] that translates the principles of Seamful Design from Ubiquitous Computing to XAI. Instead of hiding the inherent AI imperfections through seamless design, Seamful XAI turns imperfections into opportunities for explanations while enhancing user agency through appropriation. Simply put, “seams” are mismatches or gaps between what we ideally design the AI to do and how it functions when used in reality. Through a concrete decision process, Seamful XAI uses adversarial thinking like red-teaming to proactively find seams

and harness them to improve XAI. For instance, what if the lending officer of an AI-based lending system knew that the AI was trained on North American data but deployed in South Asia (the seam here is a data drift and a context shift)? The knowledge of the mismatch (seam) could help calibrate how they appropriate the AI’s output. For instance, they might override the AI’s output if it rejects the loan application of a wealthy farmer whose income is non-traditional compared to the income sources in the dataset. In our case, if the AI group were made aware of how little explanatory power Q-values have, it might have prevented them from ascribing diagnostic value to NR’s numbers. Last, a key part of designing for appropriation is learning *from* appropriation. Most XAI systems, including ours, are one-way systems that do not allow for user feedback. Feedback loops from users can highlight how, when, and why users are needing to appropriate XAI outputs and ensure the appropriation is appropriate.

6.3 Broader lessons: Explainable and Responsible AI

Our work has broader implications beyond the immediate design of XAI systems, especially around the discourse of responsible and explainable AI. Below, we share three main takeaways—how, despite best intentions, unanticipated negative consequences can emerge from AI-generated explanations & what we can do about it, how our insights can help re-imagine AI education, and how our insights can reframe how we operationalize AI adoption.

First, our findings illustrate *how despite best intentions, unsuspecting negative effects of AI-generated explanations can emerge* (e.g., unwarranted faith in numbers in 5.1). This is an instance of an **Explainability Pitfall** (EP)—“*unanticipated* negative downstream effects from adding AI explanations that emerge even when there is *no intention* to manipulate anyone” [46]. EPs are different from *dark patterns*, where there is an *intent* to deceive the user [19]. Recall that in our study design, we had no intention of deceiving anyone. While dark patterns have been explored in XAI (e.g., [30, 50]), EPs are *severely under-explored*. This paper sheds important light on this under-explored space.

For mitigation strategies, given EPs are often a consequence of *uncritical acceptance* [46], we can *shift* our explanation design philosophy to emphasize *critical reflection* (as opposed to acceptance) during interpretation of explanations. This aligns with Human-centered XAI work that advocates engendering trust via *reflection* [45]. Langer et al. [87] point out that people are likely to accept explanations without conscious attention if no effortful thinking is required from them. In Kahneman’s dual-process theory [75] terms, this means that if we do not invoke mindful and deliberative (system 2) thinking with explanations, we increase the likelihood of uncritical consumption. To trigger mindfulness, Langer et al. [87] recommend designing for “effortful responses” or “thoughtful responding.” We can utilize design interventions highlighted in Sec. 6.2 such as Cognitive Forcing Functions and using Seamful XAI strategies.

Second, our insights call attention to wider challenges with *academic practices of AI learning*. AI students exhibited unwarranted faith in and preference for numbers even when they were not readily comprehensible. In light of the importance of quantitative

practices in AI, this bias is not surprising. One effective way to address this would be to critically reflect on the way we educate students in AI. In particular, how do we ensure that students have a more critical eye towards the working and outputs of AI systems? This is where we see the impact of our focus on AI students (vs. fully formed AI practitioners)—by investigating how AI background can lead to the formation of certain heuristics, we have crucial insights on how we might address the issues from an academic training perspective. For instance, there is a need to introduce courses like critical data studies and human-centered data science that can provide much-needed reflective lenses to students to understand their own cognitive biases and the importance of thinking about the user during system design and development. By addressing issues during the formation of their "Data Vision" [128], and "Professional Vision" [59] more generally, we can revolutionize how we train the next generations of AI creators and mitigate the creator-consumer gap. If we are asking today's AI students to become future creators for consumers who think drastically differently from themselves, we need to empower our students with the lenses such that the gap between them and the stakeholders they will serve is less than what it is today. When extrapolating findings for the AI group, we should note that the AI-background population was drawn from students in a very typical introductory AI course, taught from the most widely adopted textbook (by Russell & Norvig [137] taught in over 1500 schools worldwide [6]) and using a widely adopted set of instructional assignments.

Third, our work carries broader implications on *reframing AI adoption* through re-examining the relationship between user trust and AI adoption. This re-examination has upstream (e.g., research) and downstream (e.g., industry practices) implications. There is an oft-unspoken yet dominant assumption that connects *adoption with acceptance*—where we view AI adoption emerging from user trust, which, in turn, emerges from user acceptance [45]. Acceptance is seen as a core tenet of trust-building (thereby adoption). There is often an uncritical push towards trust building without asking a fundamental question: is AI worthy of our trust? Moreover, *is acceptance the only way to build trust and get adoption?* How would we feel in a human relationship if trust hinged solely on accepting everything the other party said? There are many ways to build trust. The principles of HCXAI [45] suggest diverse foundations for trust building such as healthy skepticism from informed users, awareness of variability of around system decisions, and knowledge of what the AI *cannot* do (as opposed to the typical what the AI *can* do).

Opting to promote reflection in the user can begin the process of defamiliarization from acceptance-first approaches, reframing the discourse around AI adoption. Such mindset shifts in AI adoption have cascading societal ripple effects at both upstream (e.g., research) and downstream (e.g., industry practices) levels. For instance, if an organization embodies a *critically reflective* AI adoption mindset (as opposed to an *acceptance-driven* one), it could mitigate *perverse organizational incentives* such as prioritizing growth and adoption at all costs. This can reflexively catalyze a change in organizational culture towards responsible and accountable AI governance. Note that we are not advocating to eliminate building trust via acceptance; rather, we are encouraging reforms that go beyond

acceptance. More importantly, reflection and acceptance are not mutually exclusive and can work in tandem with each other. Reforming the dialogue around AI adoption and trust building promotes our resilience against detrimental pitfalls like over-estimating (or hyping) AI capabilities, which has cascading societal implications around holding these systems accountable.

Last, when transferring our findings to other XAI scenarios, instead of shooting for blanket generalizability, we should aim for transferability. Inherent in transferability is context-sensitivity [53, 64, 148]. Real-world XAI settings have different contexts, domains, and application areas—as such, the transferability of our findings to new domains is subject to context-applicability. Context permitting, especially in terms of user characteristics, the findings are likely to transfer. This is because the heuristics and appropriation strategies identified (in Sec. 6.2) have broad applicability. For example, our findings can shed light on why data scientists over-trust XAI methods like SHAP that have numerical-based outputs [78]. There are, however, limits to our study design and exhaustiveness of the setup, elaborated in Sec. 6.3

7 LIMITATIONS & FUTURE WORK

With this study, we have taken a formative step towards understanding how interpretations of AI-generated explanations differ between groups with or without AI background (a consequential user characteristic). Given this essential yet first step, the insights from our work should be scoped accordingly. A focus on AI students has helped understand how enculturation into AI alters people's perceptions of AI explanations and makes visible the presence of other kinds of – often similar – interpretations within a heterogeneous set of non-AI group participants. [C1, C2] However, our study has areas of improvement. First, as shared in Secs. 3.3.3 & 6.3 while the AI course experienced by the AI group is largely representative of curriculum across other institutions, future research could explore how differences in AI curriculum may impact people's AI background and perception of explanations. Second, in this study, we focused on using an RL-based AI agent to conduct a sequential decision-making task in a controlled environment. Future research could extend to other types of AI agents on other tasks (e.g., classification), especially in a situated sociotechnical context. Third, the quasi-experimental setup (common in behavioral research and neuroscience [108]) does not include random assignment of participants; future work could conduct Randomized Controlled Trials using the same setup to establish causal relationships. Fourth, we used systematic screening to ensure the two groups are measurably different while controlling for factors such as age and education levels (Sec. 3.3.3). Despite these efforts, confounds can still exist such as technological familiarity and AI use. Future research can conduct studies of user characteristics in isolation and their impact on the interpretation of AI explanations. Fifth, in this initial effort, we scoped our study in the context of three types of AI explanations—while they are informative, they are not exhaustive of the different types of explanations. Future work can explore how users with two or more different user characteristics (e.g., comparison with multiple facets of one's background) or more homogeneous and stratified AI backgrounds (e.g., years of AI programming experience) perceive explanations in different ways. For future iterations, we

are inspired by Agre’s design philosophy of Critical Technical Practice [3] where “*at least for the foreseeable future, [we] will require a split identity – one foot planted in the craft work of design and the other foot planted in the reflexive work of critique.*” [4]. Through this work, we have “planted one foot” in the work of design. Now, we seek to learn from and with the broader HCI and XAI communities as we “plant the other foot” in the self-reflective realm of critique.

8 CONCLUSIONS

In this paper, we focus on the *who* of XAI by investigating how two different groups of *whos*—people with and without a background in AI—perceive different types of AI explanations. Through a mixed-methods user study, we demonstrate how interpretations and perceptions of AI explanations differ between groups with or without AI background. Our mixed-methods analysis provides different levels of insight. *What* people prefer is relatively clear—they prefer natural language-based justificatory rationales. While the *what* is somewhat straightforward, the *why* and *how* behind their preferences are nuanced. Different perceptions (e.g., unanticipated explanatory value) sometimes arise from different forms of appropriation (e.g., diagnosis vs. affirmatory intent). The same perception (unwarranted faith in numbers) sometimes arises from different types of heuristics (associating numerical vs. incomprehensible reasoning with intelligence). Finally, both groups were very similar in their desire to engage with natural-language-based explanations.

Explainability of AI systems is crucial to instill appropriate user trust and facilitate recourse. Disparities in AI backgrounds have the potential to exacerbate the challenges arising from the differences between how designers imagine users will appropriate explanations vs. how users actually interpret and use them. We provided concrete design implications to mitigate the risk of over-reliance on numbers and broader lessons about the need for re-imagining AI education. By focusing on the *who* and not just the *what* of XAI, our work takes a formative step in advancing a pluralistic human-centered XAI discourse to help bridge the creator-consumer gap.

REFERENCES

- [1] Ashraf Abdul, Christian von der Weth, Mohan Kankanalli, and Brian Y Lim. 2020. COGAM: Measuring and Moderating Cognitive Load in Machine Learning Model Explanations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*. 1–14.
- [2] Amina Adadi and Mohammed Berrada. 2018. Peeking inside the black-box: A survey on Explainable Artificial Intelligence (XAI). *IEEE Access* 6 (2018), 52138–52160.
- [3] P Agre. 1997. Toward a critical technical practice: Lessons learned in trying to reform AI in Bowker. G., Star, S., Turner, W., and Gasser, L., eds, *Social Science, Technical Systems and Cooperative Work: Beyond the Great Divide*, Erlbaum (1997).
- [4] Philip Agre and Philip E Agre. 1997. *Computation and human experience*. Cambridge University Press.
- [5] Alan Agresti. 2003. *Categorical data analysis*. Vol. 482. John Wiley & Sons.
- [6] AIMA. [n.d.]. 1549 Schools Worldwide That Have Adopted AIMA. <https://aima.cs.berkeley.edu/adoptions.html>
- [7] Ahmed Alqaraawi, Martin Schuessler, Philipp Weiß, Enrico Costanza, and Nadia Berthouze. 2020. Evaluating saliency map explanations for convolutional neural networks: a user study. In *Proceedings of the 25th International Conference on Intelligent User Interfaces*. 275–285.
- [8] Junia Anacleto and Sidney Fels. 2013. Adoption and appropriation: A design process from HCI research at a Brazilian neurological hospital. In *IFIP Conference on Human-Computer Interaction*. Springer, 356–363.
- [9] Jacob Andreas, Anca Dragan, and Dan Klein. 2017. Translating neuralese. *arXiv preprint arXiv:1704.06960* (2017).
- [10] Alejandro Barredo Arrieta, Natalia Diaz-Rodriguez, Javier Del Ser, Adrien Ben-netot, Siham Tabik, Alberto Barbado, Salvador García, Sergio Gil-López, Daniel Molina, Richard Benjamins, et al. 2020. Explainable Artificial Intelligence (XAI): Concepts, taxonomies, opportunities and challenges toward responsible AI. *Information Fusion* 58 (2020), 82–115.
- [11] Solon Barocas and Andrew D Selbst. 2016. Big data’s disparate impact. *Cal. L. Rev.* 104 (2016), 671.
- [12] David Bawden et al. 2008. Origins and concepts of digital literacy. *Digital literacies: Concepts, policies and practices* 30, 2008 (2008), 17–32.
- [13] Tara S Behrend, David J Sharek, Adam W Meade, and Eric N Wiebe. 2011. The viability of crowdsourcing for survey research. *Behavior research methods* 43, 3 (2011), 800.
- [14] Frank R Bentley, Nediya Daskalova, and Brooke White. 2017. Comparing the reliability of Amazon Mechanical Turk and Survey Monkey to traditional market research surveys. In *Proceedings of the 2017 CHI Conference Extended Abstracts on Human Factors in Computing Systems*. 1092–1099.
- [15] Richard A Berk and Justin Bleich. 2013. Statistical procedures for forecasting criminal behavior: A comparative assessment. *Criminology & Pub. Pol’y* 12 (2013), 513.
- [16] Timothy W Bickmore and Rosalind W Picard. 2004. Towards caring machines. In *CHI’04 extended abstracts on Human factors in computing systems*. 1489–1492.
- [17] Reuben Binns, Max Van Kleek, Michael Veale, Ulrik Lyngs, Jun Zhao, and Nigel Shadbolt. 2018. ‘It’s Reducing a Human Being to a Percentage’: Perceptions of Justice in Algorithmic Decisions. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*. ACM, 377.
- [18] John Bowers. 2012. The logic of annotated portfolios: communicating the value of ‘research through design’. In *Proceedings of the Designing Interactive Systems Conference*. 68–77.
- [19] Harry Brignull, Marc Miquel, Jeremy Rosenberg, and James Offer. 2015. Dark Patterns-User Interfaces Designed to Trick People.
- [20] Zana Bućinca, Phoebe Lin, Krzysztof Z Gajos, and Elena L Glassman. 2020. Proxy tasks and subjective measures can be misleading in evaluating explainable ai systems. In *Proceedings of the 25th International Conference on Intelligent User Interfaces*. 454–464.
- [21] Angelo Cafaro, Hannes Högni Vilhjálmsón, Timothy Bickmore, Dirk Heylen, Kamilla Rún Jóhannsdóttir, and Gunnar Steinn Valgarðsson. 2012. First impressions: Users’ judgments of virtual agents’ personality and interpersonal attitude in first encounters. In *International conference on intelligent virtual agents*. Springer, 67–80.
- [22] Carrie J Cai, Jonas Jongejan, and Jess Holbrook. 2019. The effects of example-based explanations in a machine learning interface. In *Proceedings of the 24th International Conference on Intelligent User Interfaces*. 258–262.
- [23] Stevie Chancellor, Eric PS Baumer, and Munmun De Choudhury. 2019. Who is the “Human” in Human-Centered Machine Learning: The Case of Predicting Mental Health from Social Media. *Proceedings of the ACM on Human-Computer Interaction* 3, CSCW (2019), 1–32.
- [24] Kathy Charmaz. 2014. *Constructing Grounded Theory (Introducing Qualitative Methods series) 2nd Edition*. Sage, London.
- [25] Zhengping Che, Sanjay Purushotham, Robinder Khemani, and Yan Liu. 2016. Interpretable deep models for ICU outcome prediction. In *AMIA Annual Symposium Proceedings*, Vol. 2016. American Medical Informatics Association, 371.
- [26] Jim Q Chen and Sang M Lee. 2003. An exploratory cognitive DSS for strategic decision making. *Decision support systems* 36, 2 (2003), 147–160.
- [27] Hao-Fei Cheng, Ruotong Wang, Zheng Zhang, Fiona O’Connell, Terrance Gray, F Maxwell Harper, and Haiyi Zhu. 2019. Explaining decision-making algorithms through UI: Strategies to help non-expert stakeholders. In *Proceedings of the 2019 chi conference on human factors in computing systems*. 1–12.
- [28] EunJeong Cheon and Norman Makoto Su. 2017. Configuring the User: “Robots have Needs Too”. In *Proceedings of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing*. 191–206.
- [29] Michael Chromik, Malin Eiband, Felicitas Buchner, Adrian Krüger, and Andreas Butz. 2021. I Think I Get Your Point, AI! The Illusion of Explanatory Depth in Explainable AI. In *26th International Conference on Intelligent User Interfaces*. 307–317.
- [30] Michael Chromik, Malin Eiband, Sarah Theres Völkel, and Daniel Buschek. 2019. Dark Patterns of Explainability, Transparency, and User Control for Intelligent Systems.. In *IUI workshops*, Vol. 2327.
- [31] Mark Core, H. Chad Lane, Michael van Lent, Dave Gomboc, Steve Solomon, and Milton Rosenberg. 2006. Building Explainable Artificial Intelligence Systems. In *Proceedings of the 18th Innovative Applications of Artificial Intelligence Conference*. Boston, MA.
- [32] Pat Croskerry. 2003. Cognitive forcing strategies in clinical decisionmaking. *Annals of emergency medicine* 41, 1 (2003), 110–120.
- [33] Devleena Das, Siddhartha Banerjee, and Sonia Chernova. 2021. Explainable ai for robot failures: Generating explanations that improve user assistance in fault recovery. In *Proceedings of the 2021 ACM/IEEE International Conference on Human-Robot Interaction*. 351–360.
- [34] Devleena Das and Sonia Chernova. 2020. Leveraging rationales to improve human task performance. In *Proceedings of the 25th International Conference on Intelligent User Interfaces*. 510–518.

- [35] Sanjeeb Dash, Oktay Gunluk, and Dennis Wei. 2018. Boolean Decision Rules via Column Generation. *Advances in Neural Information Processing Systems* 31 (2018), 4655–4665.
- [36] Fred D Davis, Richard P Bagozzi, and Paul R Warshaw. 1989. User acceptance of computer technology: a comparison of two theoretical models. *Management science* 35, 8 (1989), 982–1003.
- [37] Munjal Desai, Poornima Kaniarasu, Mikhail Medvedev, Aaron Steinfeld, and Holly Yanco. 2013. Impact of robot failures and feedback on real-time trust. In *Proceedings of the 8th ACM/IEEE international conference on Human-robot interaction*. IEEE Press, 251–258.
- [38] Shipi Dhanorkar, Christine T Wolf, Kun Qian, Anbang Xu, Lucian Popa, and Yunyao Li. 2021. Who needs to know what, when?: Broadening the Explainable AI (XAI) Design Space by Looking at Explanations Across the AI Lifecycle. In *Designing Interactive Systems Conference 2021*. 1591–1602.
- [39] Andrea A DiSessa. 2001. *Changing minds: Computers, learning, and literacy*. MIT Press.
- [40] Alan Dix. 2007. Designing for appropriation. In *Proceedings of HCI 2007 The 21st British HCI Group Annual Conference University of Lancaster, UK 21*. 1–4.
- [41] Jonathan Dodge, Q Vera Liao, Yunfeng Zhang, Rachel KE Bellamy, and Casey Dugan. 2019. Explaining models: an empirical study of how explanations impact fairness judgment. In *Proceedings of the 24th International Conference on Intelligent User Interfaces*. 275–285.
- [42] Finale Doshi-Velez and Been Kim. 2017. Towards a rigorous science of interpretable machine learning. *arXiv preprint arXiv:1702.08608* (2017).
- [43] Upol Ehsan, Brent Harrison, Larry Chan, and Mark O. Riedl. 2018. Rationalization: A Neural Machine Translation Approach to Generating Natural Language Explanations. In *Proceedings of the AAAI Conference on Artificial Intelligence, Ethics, and Society*.
- [44] Upol Ehsan, Q Vera Liao, Samir Passi, Mark O Riedl, and Hal Daume III. 2022. Seamful XAI: Operationalizing Seamful Design in Explainable AI. *arXiv preprint arXiv:2211.06753* (2022).
- [45] Upol Ehsan and Mark O Riedl. 2020. Human-centered Explainable AI: Towards a Reflective Sociotechnical Approach. *arXiv preprint arXiv:2002.01092* (2020).
- [46] Upol Ehsan and Mark O Riedl. 2021. Explainability pitfalls: Beyond dark patterns in explainable AI. *arXiv preprint arXiv:2109.12480* (2021).
- [47] Upol Ehsan, Pradyumna Tambwekar, Larry Chan, Brent Harrison, and Mark Riedl. 2019. Automated Rationale Generation: A Technique for Explainable AI and its Effects on Human Perceptions. In *Proceedings of the International Conference on Intelligence User Interfaces*.
- [48] Malin Eiband, Daniel Buschek, and Heinrich Hussmann. 2021. How to support users in understanding intelligent systems? Structuring the discussion. In *26th International Conference on Intelligent User Interfaces*. 120–132.
- [49] Malin Eiband, Daniel Buschek, Alexander Kremer, and Heinrich Hussmann. 2019. The Impact of Placebic Explanations on Trust in Intelligent Systems. In *Extended Abstracts of the 2019 CHI Conference on Human Factors in Computing Systems (CHI EA '19)*. Association for Computing Machinery, New York, NY, USA, 1–6. <https://doi.org/10.1145/3290607.3312787>
- [50] Malin Eiband, Hanna Schneider, Mark Bilandzic, Julian Fazekas-Con, Mareike Haug, and Heinrich Hussmann. 2018. Bringing transparency design into practice. In *23rd international conference on intelligent user interfaces*. 211–223.
- [51] Kimmo Eriksson. 2012. The nonsense math effect. *Judgment and Decision Making* 7, 6 (2012), 746–749. <http://decisionsciencenews.com/sjdm/journal.sjdm.org/12/12810/jdm12810.pdf>
- [52] Motahhare Eslami, Aimee Rickman, Kristen Vaccaro, Amirhossein Aleyasen, Andy Vuong, Karrie Karahalios, Kevin Hamilton, and Christian Sandvig. 2015. “I Always Assumed That I Wasn’t Really That Close to [Her]”: Reasoning about Invisible Algorithms in News Feeds. In *Proceedings of the 33rd Annual ACM Conference on Human Factors in Computing Systems (CHI '15)*. Association for Computing Machinery, New York, NY, USA, 153–162. <https://doi.org/10.1145/2702123.2702556>
- [53] Deborah Finfgeld-Connett. 2010. Generalizability and transferability of meta-synthesis research findings. *Journal of advanced nursing* 66, 2 (2010), 246–254.
- [54] Iason Gabriel. 2020. Artificial Intelligence, Values, and Alignment. *Minds and Machines* 30, 3 (Sep 2020), 411–437. <https://doi.org/10.1007/s11023-020-09539-2>
- [55] Bhavya Ghai, Q Vera Liao, Yunfeng Zhang, Rachel Bellamy, and Klaus Mueller. 2021. Explainable Active Learning (XAL) Toward AI Explanations as Interfaces for Machine Teachers. *Proceedings of the ACM on Human-Computer Interaction* 4, CSCW3 (2021), 1–28.
- [56] Leilani H Gilpin, David Bau, Ben Z Yuan, Ayesha Bajwa, Michael Specter, and Lalana Kagal. 2018. Explaining explanations: An approach to evaluating interpretability of machine learning. *arXiv preprint arXiv:1806.00069* (2018).
- [57] Barney Glaser and Anselm Strauss. 1967. *The Discovery of Grounded Theory: Strategies for Qualitative Research*. Aldine Transactions, Chicago.
- [58] Joseph K Goodman and Gabriele Paolacci. 2017. Crowdsourcing consumer research. *Journal of Consumer Research* 44, 1 (2017), 196–210.
- [59] Charles Goodwin. 1994. Professional Vision. *American Anthropologist* 96, 3 (1994), 606–633.
- [60] Mary L Gray and Siddharth Suri. 2019. *Ghost work: How to stop Silicon Valley from building a new global underclass*. Eamon Dolan Books.
- [61] Riccardo Guidotti, Anna Monreale, Salvatore Ruggieri, Franco Turini, Fosca Giannotti, and Dino Pedreschi. 2018. A survey of methods for explaining black box models. *ACM computing surveys (CSUR)* 51, 5 (2018), 1–42.
- [62] Karen Hao. 2019. AI is sending people to jail – and getting it wrong. *MIT Technology Review* (21 January 2019). Retrieved 26-August-2019 from <https://www.technologyreview.com/s/612775/algorithms-criminal-justice-ai/>
- [63] David J Hauser and Norbert Schwarz. 2016. Attentive Turkers: MTurk participants perform better on online attention checks than do subject pool participants. *Behavior research methods* 48, 1 (2016), 400–407.
- [64] Gillian R Hayes. 2011. The relationship of action research to human-computer interaction. *ACM Transactions on Computer-Human Interaction (TOCHI)* 18, 3 (2011), 1–20.
- [65] Fred Hohman, Andrew Head, Rich Caruana, Robert DeLine, and Steven M Drucker. 2019. Gamut: A design probe to understand how data scientists understand machine learning models. In *Proceedings of the 2019 CHI conference on human factors in computing systems*. 1–13.
- [66] Myles Hollander, Douglas A Wolfe, and Eric Chicken. 2013. *Nonparametric statistical methods*. Vol. 751. John Wiley & Sons.
- [67] Andreas Holzinger, Chris Biemann, Constantinos S Pattichis, and Douglas B Kell. 2017. What do we need to build explainable AI systems for the medical domain? *arXiv preprint arXiv:1712.09923* (2017).
- [68] Chien-Ming Huang, Takamasa Iio, Satoru Satake, and Takayuki Kanda. 2014. Modeling and Controlling Friendliness for An Interactive Museum Robot.. In *Robotics: Science and Systems*. 12–16.
- [69] Lilly Irani and M Six Silberman. 2014. From critical design to critical infrastructure: lessons from turkopticon. *interactions* 21, 4 (2014), 32–35.
- [70] James Jaccard. 2001. *Interaction effects in logistic regression*. Vol. 135. SAGE Publications, Incorporated.
- [71] Alejandro Jaimes, Daniel Gatica-Perez, Nicu Sebe, and Thomas S Huang. 2007. Guest Editors’ Introduction: Human-Centered Computing—Toward a Human Revolution. *Computer* 40, 5 (2007), 30–34.
- [72] Jakob D Jensen, Andy J King, LaShara A Davis, and Lisa M Guntzville. 2010. Utilization of internet technology by low-income adults: the role of health literacy, health numeracy, and computer assistance. *Journal of aging and health* 22, 6 (2010), 804–826.
- [73] Weina Jin, Sheelagh Carpendale, Ghassan Hamarneh, and Diane Gromala. 2019. Bridging ai developers and end users: an end-user-centred explainable ai taxonomy and visual vocabularies. *Proceedings of the IEEE Visualization, Vancouver, BC, Canada* (2019), 20–25.
- [74] W Lewis Johnson. 1994. Agents that Learn to Explain Themselves.. In *AAAI*. 1257–1263.
- [75] Daniel Kahneman. 2011. *Thinking, fast and slow*. Macmillan.
- [76] Gajendra Jung Katuwal and Robert Chen. 2016. Machine learning model interpretability for precision medicine. *arXiv preprint arXiv:1610.09045* (2016).
- [77] Harmanpreet Kaur, Harsha Nori, Samuel Jenkins, Rich Caruana, Hanna Wallach, and Jennifer Wortman Vaughan. 2020. Interpreting Interpretability: Understanding Data Scientists’ Use of Interpretability Tools for Machine Learning. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems (CHI '20)*. Association for Computing Machinery, New York, NY, USA, 1–14. <https://doi.org/10.1145/3313831.3376219>
- [78] Harmanpreet Kaur, Harsha Nori, Samuel Jenkins, Rich Caruana, Hanna Wallach, and Jennifer Wortman Vaughan. 2020. Interpreting Interpretability: Understanding Data Scientists’ Use of Interpretability Tools for Machine Learning. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*. 1–14.
- [79] Aniket Kittur, Ed H Chi, and Bongwon Suh. 2008. Crowdsourcing user studies with Mechanical Turk. In *Proceedings of the SIGCHI conference on human factors in computing systems*. 453–456.
- [80] Jon Kleinberg, Himabindu Lakkaraju, Jure Leskovec, Jens Ludwig, and Sendhil Mullainathan. 2017. Human Decisions and Machine Predictions. *The Quarterly Journal of Economics* 133, 1 (2017), 237–293. <https://doi.org/10.1093/qje/qjx032>
- [81] Rob Kling and Susan Leigh Star. 1998. Human centered systems in the perspective of organizational and social informatics. *ACM SIGCAS Computers and Society* 28, 1 (1998), 22–29.
- [82] Tomoko Koda and Pattie Maes. 1996. Agents with faces: The effect of personification. In *Proceedings 5th IEEE International Workshop on Robot and Human Communication. RO-MAN'96 TSUKUBA*. IEEE, 189–194.
- [83] Alina Krischkowsky. 2018. *Technology in Uses: Negotiating the in-between of designer | user and design | use*. Ph.D. Dissertation. University of Salzburg, Salzburg, Austria.
- [84] Todd Kulesza, Simone Stumpf, Margaret Burnett, Sherry Yang, Irwin Kwan, and Weng-Keen Wong. 2013. Too much, too little, or just right? Ways explanations impact end users’ mental models. In *2013 IEEE Symposium on visual languages and human centric computing*. IEEE, 3–10.
- [85] Minae Kwon, Sandy H Huang, and Anca D Dragan. 2018. Expressing Robot Incapability. In *Proceedings of the 2018 ACM/IEEE International Conference on*

- Human-Robot Interaction*. ACM, 87–95.
- [86] Kathryn Ann Lambe, Gary O'Reilly, Brendan D Kelly, and Sarah Curristan. 2016. Dual-process cognitive interventions to enhance diagnostic reasoning: a systematic review. *BMJ quality & safety* 25, 10 (2016), 808–820.
 - [87] Ellen J Langer, Arthur Blank, and Ben Zion Chanowitz. 1978. The mindlessness of ostensibly thoughtful action: The role of "placebic" information in interpersonal interaction. *Journal of personality and social psychology* 36, 6 (1978), 635.
 - [88] Min Kyung Lee, Sara Kiesler, and Jodi Forlizzi. 2010. Receptionist or information kiosk: how do people talk with a robot?. In *Proceedings of the 2010 ACM conference on Computer supported cooperative work*. 31–40.
 - [89] Min Kyung Lee, Sara Kiesler, Jodi Forlizzi, Siddhartha Srinivasa, and Paul Rybski. 2010. Gracefully mitigating breakdowns in robotic services. In *Human-Robot Interaction (HRI), 2010 5th ACM/IEEE International Conference on*. IEEE, 203–210.
 - [90] Q Vera Liao, Matthew Davis, Werner Geyer, Michael Muller, and N Sadat Shami. 2016. What can you do? Studying social-agent orientation and agent proactive interactions with an agent for employees. In *Proceedings of the 2016 acm conference on designing interactive systems*. 264–275.
 - [91] Q Vera Liao, Daniel Gruen, and Sarah Miller. 2020. Questioning the AI: informing design practices for explainable AI user experiences. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*. 1–15.
 - [92] Q Vera Liao, Muhammed Mas-ud Hussain, Praveen Chandar, Matthew Davis, Yasaman Khazaeni, Marco Patricio Crasso, Dakuo Wang, Michael Muller, N Sadat Shami, and Werner Geyer. 2018. All work and no play? Conversations with a Question-and-Answer Chatbot in the Wild. In *Proceedings of the 2018 CHI Conference on Human Factors in Computing Systems*. 1–13.
 - [93] Brian Y Lim, Anind K Dey, and Daniel Avrahami. 2009. Why and Why Not Explanations Improve the Intelligibility of Context-aware Intelligent Systems. In *Proceedings of the SIGCHI Conference on Human Factors in Computing Systems (CHI '09)*. ACM, New York, NY, USA, 2119–2128. <https://doi.org/10.1145/1518701.1519023>
 - [94] Tung-Ching Lin, Sheng Wu, Jack Shih-Chieh Hsu, and Yi-Ching Chou. 2012. The integration of value-based adoption and expectation–confirmation models: An example of IPTV continuance intention. *Decision Support Systems* 54, 1 (2012), 63–75.
 - [95] Peter Lipton. 2001. What good is an explanation? In *Explanation*. Springer, 43–59.
 - [96] Zachary C Lipton. 2018. The mythos of model interpretability. *Queue* 16, 3 (2018), 31–57.
 - [97] Leib Litman, Jonathan Robinson, and Tzvi Abberbock. 2017. TurkPrime.com: A versatile crowdsourcing data acquisition platform for the behavioral sciences. *Behavior research methods* 49, 2 (2017), 433–442.
 - [98] Tyler J. Loftus, Patrick J. Tighe, Amanda C. Filiberto, Philip A. Efron, Scott C. Brakenridge, Alicia M. Mohr, Parisa Rashidi, Jr Upchurch, Gilbert R., and Azra Bihorac. 2020. Artificial Intelligence and Surgical Decision-making. *JAMA Surgery* 155, 2 (02 2020), 148–158. <https://doi.org/10.1001/jamasurg.2019.4917> arXiv:<https://jamanetwork.com/journals/jamasurgery/articlepdf/2756311/jamasurgery.16fab20a119a002041f>
 - [99] Tania Lombrozo. 2011. The instrumental value of explanations. *Philosophy Compass* 6, 8 (2011), 539–551.
 - [100] Tania Lombrozo. 2012. Explanation and abductive inference. *Oxford handbook of thinking and reasoning* (2012), 260–276.
 - [101] Tania Lombrozo. 2016. Explanatory Preferences Shape Learning and Inference. *Trends in Cognitive Sciences* 20, 10 (2016), 748 – 759. <https://doi.org/10.1016/j.tics.2016.08.001>
 - [102] Tania Lombrozo and Nicholas Z Gwynne. 2014. Explanation and inference: Mechanistic and functional explanations guide property generalization. *Frontiers in Human Neuroscience* 8 (2014), 700.
 - [103] Tania Lombrozo and Daniel Wilkenfeld. 2019. Mechanistic versus functional understanding. *Varieties of understanding: New perspectives from philosophy, psychology, and theology* (2019), 209.
 - [104] Duri Long and Brian Magerko. 2020. What is AI Literacy? Competencies and Design Considerations. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems* (Honolulu, HI, USA) (CHI '20). Association for Computing Machinery, New York, NY, USA, 1–16. <https://doi.org/10.1145/3313831.3376727>
 - [105] Scott M Lundberg and Su-In Lee. 2017. A unified approach to interpreting model predictions. In *Advances in neural information processing systems*. 4765–4774.
 - [106] Minh-Thang Luong, Hieu Pham, and Christopher D Manning. 2015. Effective approaches to attention-based neural machine translation. *arXiv preprint arXiv:1508.04025* (2015).
 - [107] Donald MacKenzie. 2018. Material Signals: A Historical Sociology of High-Frequency Trading. *Amer. J. Sociology* 123, 6 (2018), 1635–1683. <https://doi.org/10.1086/697318>
 - [108] Ioana E Marinescu, Patrick N Lawlor, and Konrad P Kording. 2018. Quasi-experimental causality in neuroscience and behavioural research. *Nature human behaviour* 2, 12 (2018), 891–898.
 - [109] Peter McCullagh. 1980. Regression models for ordinal data. *Journal of the Royal Statistical Society: Series B (Methodological)* 42, 2 (1980), 109–127.
 - [110] Nora McDonald and Shimei Pan. 2020. Intersectional AI: A Study of How Information Science Students Think about Ethics and Their Impact. *Proceedings of the ACM on Human-Computer Interaction* 4, CSCW2 (2020), 1–19.
 - [111] Milagros Miceli, Martin Schuessler, and Tianling Yang. 2020. Between Subjectivity and Imposition: Power Dynamics in Data Annotation for Computer Vision. *Proc. ACM Hum.-Comput. Interact.* 4, CSCW2, Article 115 (Oct. 2020), 25 pages. <https://doi.org/10.1145/3415186>
 - [112] Tim Miller. 2019. Explanation in artificial intelligence: Insights from the social sciences. *Artificial Intelligence* 267 (2019), 1–38.
 - [113] Nicole Mirnig, Gerald Stollnberger, Markus Miksch, Susanne Stadler, Manuel Giuliani, and Manfred Tscheligi. 2017. To err is robot: How humans assess and act toward an erroneous social robot. *Frontiers in Robotics and AI* 4 (2017), 21.
 - [114] Sina Mohseni, Niloofar Zarei, and Eric D. Ragan. 2020. A Multidisciplinary Survey and Framework for Design and Evaluation of Explainable AI Systems. arXiv:1811.11839 [cs.HC]
 - [115] Johanna D Moore and William R Swartout. 1988. *Explanation in expert systems: A survey*. Technical Report. UNIVERSITY OF SOUTHERN CALIFORNIA MARINA DEL REY INFORMATION SCIENCES INST.
 - [116] Michael Muller, Melanie Feinberg, Timothy George, Steven J. Jackson, Bonnie E. John, Mary Beth Kery, and Samir Passi. 2019. Human-Centered Study of Data Science Work Practices. In *Extended Abstracts of the 2019 CHI Conference on Human Factors in Computing Systems* (Glasgow, Scotland Uk) (CHI EA '19). Association for Computing Machinery, New York, NY, USA, 1–8. <https://doi.org/10.1145/3290607.3299018>
 - [117] Michael Muller, Katja Neureiter, Nervo Verdezoto, Alina Krischkowsky, Anna Maria Al Zubaidi-Polli, and Manfred Tscheligi. 2016. Collaborative appropriation: How couples, teams, groups and communities adapt and adopt technologies. In *Proceedings of the 19th ACM Conference on Computer Supported Cooperative Work and Social Computing Companion*. 473–480.
 - [118] John Murawski. 2019. Mortgage Providers Look to AI to Process Home Loans Faster. *Wall Street Journal* (18 March 2019). Retrieved 16-September-2020 from <https://www.wsj.com/articles/mortgage-providers-look-to-ai-to-process-home-loans-faster-11552899212>
 - [119] Chelsea M Myers, Anushay Furqan, and Jichen Zhu. 2019. The impact of user characteristics and preferences on performance with an unfamiliar voice user interface. In *Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems*. 1–9.
 - [120] Clifford Nass, BJ Fogg, and Youngme Moon. 1996. Can computers be teammates? *International Journal of Human-Computer Studies* 45, 6 (1996), 669–678.
 - [121] Clifford Nass and Youngme Moon. 2000. Machines and mindlessness: Social responses to computers. *Journal of social issues* 56, 1 (2000), 81–103.
 - [122] Clifford Nass, Jonathan Steuer, Lisa Henriksen, and D Christopher Dryer. 1994. Machines, social attributions, and ethopoeia: Performance assessments of computers subsequent to "self-" or "other-" evaluations. *International Journal of Human-Computer Studies* 40, 3 (1994), 543–559.
 - [123] Donald Norman. 2004. *The design of everyday things: Revised and expanded edition*. Basic books.
 - [124] Mahsan Nourani, Chiradeep Roy, Jeremy E Block, Donald R Honeycutt, Tahrima Rahman, Eric Ragan, and Vibhav Gogate. 2021. Anchoring Bias Affects Mental Model Formation and User Reliance in Explainable AI Systems. In *26th International Conference on Intelligent User Interfaces*. 340–350.
 - [125] Mahsan Nourani, Chiradeep Roy, Jeremy E. Block, Donald R. Honeycutt, Tahrima Rahman, Eric D. Ragan, and Vibhav Gogate. 2022. On the Importance of User Backgrounds and Impressions: Lessons Learned from Interactive AI Applications. *ACM Transactions on Interactive Intelligent Systems* 12, 4 (Dec. 2022), 1–29. <https://doi.org/10.1145/3531066>
 - [126] Andrés Páez. 2019. The pragmatic turn in explainable artificial intelligence (XAI). *Minds and Machines* 29, 3 (2019), 441–459.
 - [127] Samir Passi and Solon Barocas. 2019. Problem Formulation and Fairness. In *Proceedings of the ACM Conference on Fairness, Accountability, and Transparency (FAT* '19)*. ACM, New York, NY, 39–48.
 - [128] Samir Passi and Steven J. Jackson. 2017. Data Vision: Learning to See Through Algorithmic Abstraction. In *Proceedings of the 2017 ACM Conference on Computer Supported Cooperative Work and Social Computing (CSCW '17)*. ACM Press, 2436–2447.
 - [129] Samir Passi and Steven J. Jackson. 2018. Trust in Data Science: Collaboration, Translation, and Accountability in Corporate Data Science Projects. *Proceedings of the ACM on Human-Computer Interaction* 2, CSCW (nov 2018), 1–28.
 - [130] Samir Passi and Phoebe Sengers. 2020. Making data science systems work. *Big Data & Society* 7, 2 (2020), 2053951720939605. <https://doi.org/10.1177/2053951720939605>
 - [131] Samir Passi and Mihaela Vorvoreanu. 2022. *Overreliance on AI: Literature Review*. Technical Report MSR-TR-2022-12. Microsoft. <https://www.microsoft.com/en-us/research/publication/overreliance-on-ai-literature-review/>
 - [132] Richard E Petty and John T Cacioppo. 1986. The elaboration likelihood model of persuasion. In *Communication and persuasion*. Springer, 1–24.
 - [133] Emilee Rader and Rebecca Gray. 2015. Understanding User Beliefs About Algorithmic Curation in the Facebook News Feed. In *Proceedings of the 33rd*

- Annual ACM Conference on Human Factors in Computing Systems (CHI '15)*. Association for Computing Machinery, New York, NY, USA, 173–182. <https://doi.org/10.1145/2702123.2702174>
- [134] Marco Tulio Ribeiro, Sameer Singh, and Carlos Guestrin. 2016. Why should I trust you?: Explaining the predictions of any classifier. In *Proceedings of the 22nd ACM SIGKDD international conference on knowledge discovery and data mining*. ACM, 1135–1144.
- [135] Avi Rosenfeld and Ariella Richardson. 2019. Explainability in human-agent systems. *Autonomous Agents and Multi-Agent Systems* 33, 6 (2019), 673–705.
- [136] Cynthia Rudin, Caroline Wang, and Beau Coker. 2020. The Age of Secrecy and Unfairness in Recidivism Prediction. *Harvard Data Science Review* 2, 1 (31 3 2020). <https://doi.org/10.1162/99608f92.6ed64b30> <https://hdsr.mitpress.mit.edu/pub/7z10o269>.
- [137] Stuart J Russell and Peter Norvig. 2010. *Artificial intelligence a modern approach*. London.
- [138] Antti Salovaara. 2008. Inventing new uses for tools: a cognitive foundation for studies on appropriation. *Human Technology: An Interdisciplinary Journal on Humans in ICT Environments* (2008).
- [139] Ramprasaath R Selvaraju, Michael Cogswell, Abhishek Das, Ramakrishna Vedantam, Devi Parikh, and Dhruv Batra. 2017. Grad-cam: Visual explanations from deep networks via gradient-based localization. In *Proceedings of the IEEE international conference on computer vision*. 618–626.
- [140] Steven J Sherman and Eric Corty. 1984. Cognitive heuristics. (1984).
- [141] Elaine Short, Justin Hart, Michelle Vu, and Brian Scassellati. 2010. No fair!! an interaction with a cheating robot. In *2010 5th ACM/IEEE International Conference on Human-Robot Interaction (HRI)*. IEEE, 219–226.
- [142] Edward Shortliffe. 2012. *Computer-based medical consultations: MYCIN*. Vol. 2. Elsevier.
- [143] Alison Smith-Renner, Ron Fan, Melissa Birchfield, Tongshuang Wu, Jordan Boyd-Graber, Daniel S Weld, and Leah Findlater. 2020. No explainability without accountability: An empirical study of explanations and feedback in interactive ml. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*. 1–13.
- [144] Kwonsang Sohn and Ohbyung Kwon. 2020. Technology acceptance theories and factors influencing artificial intelligence-based intelligent products. *Telematics and Informatics* 47 (2020), 101324.
- [145] Jon Sprouse. 2011. A validation of Amazon Mechanical Turk for the collection of acceptability judgments in linguistic theory. *Behavior research methods* 43, 1 (2011), 155–167.
- [146] Sonja Stange and Stefan Kopp. 2020. Effects of a Social Robot's Self-Explanations on How Humans Understand and Evaluate Its Behavior. In *Proceedings of the 2020 ACM/IEEE international conference on human-robot interaction*. 619–627.
- [147] Anselm Strauss and Juliet M. Corbin. 1990. *Basics of Qualitative Research: Grounded Theory Techniques and Procedures*. Sage, New York.
- [148] Ernest T Stringer. 2007. *Action research third edition*. (2007).
- [149] Harini Suresh and John V Guttag. 2019. A framework for understanding unintended consequences of machine learning. *arXiv preprint arXiv:1901.10002* (2019).
- [150] William R Swartout. 1983. XPLAIN: A system for creating and explaining expert consulting programs. *Artificial intelligence* 21, 3 (1983), 285–325.
- [151] Sarah Tan, Rich Caruana, Giles Hooker, and Yin Lou. 2018. Distill-and-compare: Auditing black-box models using transparent model distillation. In *Proceedings of the 2018 AAAI/ACM Conference on AI, Ethics, and Society*. 303–310.
- [152] Linda Tickle-Degnen and Robert Rosenthal. 1990. The nature of rapport and its nonverbal correlates. *Psychological inquiry* 1, 4 (1990), 285–293.
- [153] Richard Tomsett, Dave Braines, Dan Harborne, Alun Preece, and Supriyo Chakraborty. 2018. Interpretable to whom? A role-based model for analyzing interpretable machine learning systems. *arXiv preprint arXiv:1806.07552* (2018).
- [154] Manfred Tscheligi, Alina Krishkowsky, Katja Neureiter, Kori Inkpen, Michael Muller, and Gunnar Stevens. 2014. Potentials of the "Unexpected" Technology Appropriation Practices and Communication Needs. In *Proceedings of the 18th international conference on supporting group work*. 313–316.
- [155] UCLA. 2014. Ordinal Logistic Regression | R Data Analysis Examples. <https://stats.idre.ucla.edu/r/dae/ordinal-logistic-regression/>
- [156] Michael van Lent, , Paul Carpenter, Ryan McAlinden, and Paul Brobst. 2005. Increasing Replayability with Deliberative and Reactive Planning. In *1st Conference on Artificial Intelligence and Interactive Digital Entertainment (AIIDE)*. Maria Del Rey, California, 135–140.
- [157] Michael Van Lent, William Fisher, and Michael Mancuso. 2004. An explainable artificial intelligence system for small-unit tactical behavior. In *Proceedings of the national conference on artificial intelligence*. Menlo Park, CA; Cambridge, MA; London; AAAI Press; MIT Press; 1999, 900–907.
- [158] Michael Veale and Reuben Binns. 2017. Fairer machine learning in the real world: Mitigating discrimination without collecting sensitive data. *Big Data & Society* 4, 2 (2017), 2053951717743530.
- [159] W. N. Venables and B. D. Ripley. 2002. *Modern Applied Statistics with S* (fourth ed.). Springer, New York. <http://www.stats.ox.ac.uk/pub/MASS4> ISBN 0-387-95457-0.
- [160] Viswanath Venkatesh, Michael G Morris, Gordon B Davis, and Fred D Davis. 2003. User acceptance of information technology: Toward a unified view. *MIS quarterly* (2003), 425–478.
- [161] Fei Wang, Mengqing Jiang, Chen Qian, Shuo Yang, Cheng Li, Honggang Zhang, Xiaogang Wang, and Xiaoou Tang. 2017. Residual Attention Network for Image Classification. In *Proceedings of the IEEE Conference on Computer Vision and Pattern Recognition*. 3156–3164.
- [162] Ruotong Wang, F Maxwell Harper, and Haiyi Zhu. 2020. Factors influencing perceived fairness in algorithmic decision-making: Algorithm outcomes, development procedures, and individual differences. In *Proceedings of the 2020 CHI Conference on Human Factors in Computing Systems*. 1–14.
- [163] Peter C Wason and J St BT Evans. 1974. Dual processes in reasoning? *Cognition* 3, 2 (1974), 141–154.
- [164] Christopher JCH Watkins and Peter Dayan. 1992. Q-learning. *Machine learning* 8, 3-4 (1992), 279–292.
- [165] Dennis Wei, Sanjeeb Dash, Tian Gao, and Oktay Gunluk. 2019. Generalized linear rule models. In *International Conference on Machine Learning*. PMLR, 6687–6696.
- [166] Karl E. Weick. 1995. *Sensemaking in Organizations*. SAGE Publications, Thousand Oaks, California.
- [167] Daniel A Wilkenfeld and Tania Lombrozo. 2015. Inference to the best explanation (IBE) versus explaining for the best inference (EBI). *Science & Education* 24, 9-10 (2015), 1059–1077.
- [168] Kelvin Xu, Jimmy Ba, Ryan Kiros, Kyunghyun Cho, Aaron Courville, Ruslan Salakhudinov, Rich Zemel, and Yoshua Bengio. 2015. Show, attend and tell: Neural image caption generation with visual attention. In *International conference on machine learning*. 2048–2057.
- [169] Jason Yosinski, Jeff Clune, Anh Nguyen, Thomas Fuchs, and Hod Lipson. 2015. Understanding neural networks through deep visualization. *arXiv preprint arXiv:1506.06579* (2015).
- [170] Quanzeng You, Hailin Jin, Zhaowen Wang, Chen Fang, and Jiebo Luo. 2016. Image captioning with semantic attention. In *Proceedings of the IEEE conference on computer vision and pattern recognition*. 4651–4659.
- [171] Achim Zeileis, Mark A Wiel, Kurt Hornik, and Torsten Hothorn. 2008. Implementing a class of permutation tests: the coin package. *Journal of statistical software* 28, 8 (2008), 1–23.

A APPENDIX

A.1 Best practices for data integrity and participant engagement

Data quality and engagement with user study participants are integral to research. Below we share how previous guidelines for fair and equitable work treatment for MTurk workers [69] helped us decide the payment structure, tips on engagement with participants, design motivations behind the task environment, and how we deployed and reviewed the task catering for a global audience. While these insights are transferable to other contexts, most of these practices are geared towards MTurk participants, given one has less control over the platform.

- (1) *Payment*: our study was not a micro-task that is traditionally deployed in MTurk. As a result, we calibrated our payment to reflect the task duration and expected effort. We tried our best to structure our study based on previous guidelines for fair and equitable work treatment for MTurk workers [69]. We strived to pay equal to or more than a minimum wage [at the time of deployment, the local minimum wage was \$8.5/hour]. We paid \$10 for a task budgeted for 45 minutes, making the hourly pay \$13.3. However, almost all participants took 30 mins to complete the task on average, making the effective hourly rate around \$20/hour.

As a policy, we disbursed payments within 48 hours of task completion. This robust turnaround time helped our reputation on Turkopticon, a forum for MTurk workers to engage in peer-to-peer assistance on job information and hold employers accountable for fair treatment [69]. Moreover, one researcher regularly engaged with workers on Turkopticon, answering questions and returning compliments. This engagement built rapport throughout the study. As a platform, TurkPrime allows internal messages between MTurk workers and employers by using a proxy userID and protecting the privacy of the worker. This feature allowed participants to communicate if they had internet issues or were running out of time. In such cases, we sent them a one-time link for completion. The same applied for participants in the AI background group who were not bound by AMT rules.

Every payment of a HIT had a thank you message attached to it. Every rejection had custom justifications backed by evidence. The research team created message templates based on major issues, which allowed for a quick turnaround time even with a custom message. For participants who failed to do the task despite best efforts, we paid them for their time even if we could not use their data. This equitable policy also made our HITs one of the most sought-after in the marketplace.

- (2) *Task environment and setup*:
 - (a) *Task orientation*: Pilot testing showed that participants preferred a multi-modal (e.g., video) orientation compared to a textual description of it. Therefore, we provided both modalities.
 - (b) *Task engagement*: Participant had to successfully pass attention checks and/or explicitly acknowledge they understood the instructions in a given module. On the backend,

we had timers to evaluate if a participant spent a reasonable amount of time on a particular section. For instance, if the video was 2 minutes long and a participant clicked through that section in 30 seconds, it triggered a review of their responses. These steps augment the quality of the experimental data. Moreover, for tasks that had to be rejected, these metrics also served as justificatory evidence.

- (c) *Design of the robots*: To mitigate effects of preconceived notions, we did not use any descriptive names for the robots; instead, we introduced the robots as “Robot A” [= the Rationale-Generation robot], “Robot B” [= the Action-Declaring robot], and “Robot C” [= the Numerical robot] (details on each of their attributes are below/above). To reduce any preferential treatment of robots based on their appearances, we need to standardize their appearances without sacrificing their distinctness. That is, the robots needed to look similar, but not to the point of indistinguishability. This insight came from pilot testing which indicated that if we made the robots identical in appearance but different in color, the cognitive load for recall was too high. Therefore, we iterated and struck a balance where all robots had wheels with a “car-like” structure (see Fig. 1) while they differed in color and shape.

- (3) *Deployment and Review*: Across both groups, we manually reviewed every response in the survey, especially the qualitative justifications provided by the participants. We deployed 10-15 tasks per day to allow for manual reviews. To facilitate outreach of our task to all time zones, using an automated scheduling system, we released 3 tasks every 3 hours over a 24-hour cycle. This improved the potential for global participation in our task.

Spamming is a serious issue when it comes to survey data. Here, the qualitative responses served a secondary screening purpose. Participants with good-faith efforts always had reasonable qualitative justifications. Those who had spamming intentions shared non-sensical and even comical qualitative responses; e.g., Movie titles and plots, snippets from Wikipedia, etc.

While all of these steps required considerable time, and effort, they paid off in the high data quality we received for a task lasting 30 minutes on average.

A.2 Participant Screening

Screening criteria, group makeup, and establishing group differences:

To ensure that the two groups were *measurably different* along the dimension of our investigation—AI background—we performed additional *screening* using a questionnaire with three components—(1) a knowledge test to get a baseline understanding of programming and AI competency. This test was collaboratively developed with the course’s teaching staff to calibrate the content relevancy and question difficulty. (2) Self-reported knowledge levels in (a) computer programming and (b) AI using two 5-point Likert-scales. (3) confirmation of whether they have ever taken an AI class.

To get empirically ground the cut-off points, we piloted the screener with 10 participants from each group to get a baseline understanding of the scores and completion time. For the *AI* group, a score 4 or more on the knowledge test along with self-reports of having “Moderate knowledge” or more [≥ 4] in programming and some knowledge or more [≥ 3] in AI were required. For the *non-AI* group, self-report of having “No knowledge” [= 1] in both programming and AI, along with no prior AI classes, were required. Using these criteria we formed the two groups.

The *AI background group* consisted of 96 adult students taking an AI class. On average, the task duration was 31.1 minutes. Participants received US \$10 for their time (\$20/hr rate). 39% of the participants self-identified as females while the rest identified as males. Participants reported an average education level of 5.25 (5= “Associate’s degree”, 6= “Bachelor’s Degree”). By design, all of them currently reside in the US. On the *screening criteria*, the AI students scored an average of 4.73 (out of 5) [$SD = 0.45$] on the *knowledge test*, self-reported “moderate knowledge” on *programming* ($M = 4.53$, $SD = 0.52$) [4= “Moderate Knowledge”, 5= “A lot of knowledge”] and “some knowledge” on *AI* ($M = 3.74$, $SD = 0.49$) [3= “Some knowledge”, 4= “Moderate Knowledge”].

The *Non-AI background group with MTurk participants* consisted of 83 adults, who were recruited from Amazon Mechanical Turk (AMT) through a management service called TurkPrime [97]. On average, the task duration was 29.8 minutes. Participants received US \$10 for their time (\$20/hr). 46% of the participants self-identified as females while the rest identified as males. Participants reported an average education level of 4.8 (4= “Vocational Training”, 5= Associate’s degree). We screened for participants who reside in the US. On the *screening criteria*, MTurk participants scored an average of 0.91 (out of 5) [$SD = 0.32$] on the *knowledge test*. By design, for *programming* and *AI*, we screened for people who reported as having “No knowledge” [=1] as well as *never taking* an AI class.

To establish that these two groups are *measurably different*, we performed statistical tests. For the knowledge test, the two groups were significantly different based on a two-sample Mann-Whitney U-test (even after Bonferroni correction, $p < 2.2 \times 10^{-16}$). Since the non-AI group systematically had only people with “No knowledge” [=1] in programming or AI, we performed one-sample Mann Whitney U-tests on the AI-group to compare its means against 1 [= “No knowledge”]. Even after Bonferroni correction, we found strong evidence that the two groups are different ($p < 2.2 \times 10^{-16}$, for both programming and AI scores). These results indicate that our

screening criteria have successfully established two groups that are measurably different in terms of their AI background.

The following aspects of the selection criteria facilitated formation of the two user groups:

The AI knowledge questionnaire: We developed the knowledge questionnaire iteratively using a participatory process involving Teaching Assistants for the class along with Graduate students familiar with the area. There are five (5) multiple-choice questions in total equalling 5 points. The first two questions are programming questions: the first is a question asking for the output of a simple print statement, and the second is one that asks for the output of a for-loop. The remaining three covered concepts in AI such as Markov Decision Processes, Reinforcement Learning, and Unsupervised Learning. By the time of deployment, the AI students had already gone through lectures covering the AI topics in the questionnaire. All the questions are inspired by or directly taken from past exam questions on various topics. For further details, please refer to section A.2.1 for the AI knowledge questionnaire.

To calibrate the relative difficulty of the knowledge test, we used a collaborative and iterative process until a consensus between the researchers and the teaching staff was reached. We expected that most students with satisfactory prerequisites (that contain fundamentals of programming) and current knowledge from the class should at least get 4 out of the 5 correct. This calibration appears to have been a reasonable one since all students naturally passed these thresholds. On the other hand, in order to be assigned to the non-AI background group, participants had to score less than or equal to 1 (out of 5). We expected that some participants, without any AI background, might be able to guess the output of the “print” statement question. However, it was unlikely that someone without a basic programming understanding would be able to answer the “for-loop” output question. Therefore, if someone correctly answers more than one, their AI knowledge background is not the type we need for members of the non-AI background group.

AI background measurement: The knowledge test was followed by two 5-point Likert-scale questions measuring the AI background for computer programming and AI concepts. The range of self-reported knowledge goes from “No knowledge” [= 1] to “A lot of knowledge” [= 5]. Each level of knowledge has a sentence clarifying the meaning behind the label. For illustrative purposes, here is an example from the AI scale: “No knowledge: I might be aware of AI, but have *no knowledge* about it.” Both scales had similar construction and wording. For further details, please refer to section A.2.1 for the AI background knowledge Likert-scales.

AI class: Finally, participants answer if they have ever taken any classes on Artificial Intelligence.

A.2.1 Screening questionnaire.

Here we share the survey instruments used to screen participants.

Knowledge test questionnaire

- (1) What would be the output of the following python program?

```
name = "Peter"
print("Hello " + name)
```

- (a) Peter
 - (b) Hello Peter
 - (c) Hello + Peter
 - (d) "Hello" + name
- (2) What would be the output of the following python program?
- ```

numbers = [2, 4]
for i in range(len(numbers)):
 print(numbers[i] + i)

```
- (a) 2  
5
  - (b) 2  
5  
8
  - (c) 2  
4
  - (d) 2  
4  
10
- (3) Which of the following is an unsupervised learning task?
- (a) Distinguishing pictures containing cats from pictures not containing cats
  - (b) Flagging text messages as appropriate or inappropriate
  - (c) Divide data points into different clusters without any labels available
  - (d) Predict the value of a house after training on a dataset with house features and values
- (4) What is the general goal of reinforcement learning?
- (a) Maximize potential or expected punishment
  - (b) Maximize potential or expected reward
  - (c) Get to the goal as soon as possible
  - (d) Avoid the most obstacles in any given state
- (5) In MDPs, the Markov assumption is that:
- (a) The current state is independent of all other states
  - (b) The current state depends only on the history of previous states and actions
  - (c) The current state depends on the full sequence of states and actions (past and future)
  - (d) The current state only depends on the immediate previous state and action

#### Computer Programming Background Knowledge.

When it comes to computer programming or coding, I believe I have

- (1) No knowledge: I might be *aware* of computer programs, but have *never* coded before
- (2) A little knowledge: I know *basic* concepts in programming, but have *never applied* it
- (3) Some knowledge: I have *applied* programming concepts by coding at least *once* before
- (4) Moderate knowledge: I *apply* programming concepts *somewhat frequently* for my work, class, or leisure
- (5) A lot of knowledge: I *apply* programming concepts *very frequently* or create cutting edge software

#### AI Background Knowledge.

When it comes to Artificial Intelligence (AI), I believe I have

- (1) No knowledge: I might be *aware* of AI, but have *no knowledge* about it
- (2) A little knowledge: I know *basic* concepts in AI, but have *never applied* it
- (3) Some knowledge: I have *applied* AI concepts by coding at least *once* before
- (4) Moderate knowledge: I *apply* AI concepts *somewhat frequently* for my work, class, or leisure
- (5) A lot of knowledge: I *apply* AI concepts *very frequently* or create cutting edge software

#### AI class.

Have you ever taken or are currently taking any classes on Artificial Intelligence?

- Yes
- No

### **A.3 OLR Summary Tables**

**Table 3: Summary of OLR with Ranking as Response and Robot Type as Predictor**

|                 | Value  | Std. Error | t- value | p- value          | Odds Ratio   |
|-----------------|--------|------------|----------|-------------------|--------------|
| <b>AD Robot</b> | -1.711 | 0.104      | -16.473  | <b>&lt; 0.001</b> | <b>0.181</b> |
| <b>NR_Robot</b> | -2.691 | 0.115      | -23.393  | <b>&lt; 0.001</b> | <b>0.068</b> |
| 1 2             | -2.366 | 0.093      | -25.567  | < 0.001           | 0.094        |
| 2 3             | -0.598 | 0.077      | -7.799   | < 0.001           | 0.5501       |

Note: RG\_Robot is the reference level.

**Table 4: OLR Summary with Robot Type**

|                            | Value   | Std. Error | t value | p value | Odds Ratio |
|----------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot | 1.7105  | 0.1038     | 16.4728 | 0       | 5.5319     |
| Numerical-Reasoning Robot  | -0.9807 | 0.1001     | -9.7980 | 0       | 0.3751     |
| 1 2                        | -0.6550 | 0.0695     | -9.4282 | 0       | 0.5195     |
| 2 3                        | 1.1129  | 0.0734     | 15.1592 | 0       | 3.0433     |

Note: Action-Declaring Robot is the reference level.

**Table 5: OLR Summary - Ref. Levels: Rationale-Generation Robot and AI Group**

|                                        | Value   | Std. Error | t value  | p value | Odds Ratio |
|----------------------------------------|---------|------------|----------|---------|------------|
| Action-Declaring Robot                 | -2.0246 | 0.1302     | -15.5462 | 0.0000  | 0.1320     |
| Numerical-Reasoning Robot              | -2.5072 | 0.1391     | -18.0221 | 0.0000  | 0.0815     |
| Non-AI Group                           | -0.1563 | 0.1582     | -0.9879  | 0.3232  | 0.8553     |
| Action-Declaring Robot:Non-AI Group    | 0.8422  | 0.2093     | 4.0243   | 0.0001  | 2.3214     |
| Numerical-Reasoning Robot:Non-AI Group | -0.6088 | 0.2264     | -2.6892  | 0.0072  | 0.5440     |
| 1 2                                    | -2.4515 | 0.1104     | -22.2009 | 0.0000  | 0.0862     |
| 2 3                                    | -0.6507 | 0.0965     | -6.7413  | 0.0000  | 0.5217     |

**Table 6: OLR Summary - Ref. Levels: Rationale-Generation Robot and Non-AI Group**

|                                    | Value   | Std. Error | t value  | p value | Odds Ratio |
|------------------------------------|---------|------------|----------|---------|------------|
| Action-Declaring Robot             | -1.1824 | 0.1674     | -7.0614  | 0.0000  | 0.3065     |
| Numerical-Reasoning Robot          | -3.1160 | 0.1893     | -16.4636 | 0.0000  | 0.0443     |
| AI Group                           | 0.1564  | 0.1582     | 0.9881   | 0.3231  | 1.1692     |
| Action-Declaring Robot:AI Group    | -0.8422 | 0.2093     | -4.0245  | 0.0001  | 0.4308     |
| Numerical-Reasoning Robot:AI Group | 0.6087  | 0.2264     | 2.6889   | 0.0072  | 1.8381     |
| 1 2                                | -2.2952 | 0.1362     | -16.8540 | 0.0000  | 0.1007     |
| 2 3                                | -0.4944 | 0.1258     | -3.9305  | 0.0001  | 0.6099     |

**Table 7: OLR Summary - Ref. Levels: Action-Declaring Robot and AI Group**

|                                         | Value   | Std. Error | t value | p value | Odds Ratio |
|-----------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot              | 2.0246  | 0.1302     | 15.5462 | 0e+00   | 7.5732     |
| Numerical-Reasoning Robot               | -0.4826 | 0.1224     | -3.9421 | 1e-04   | 0.6172     |
| Non-AI Group                            | 0.6859  | 0.1368     | 5.0127  | 0e+00   | 1.9855     |
| Rationale-Generation Robot:Non-AI Group | -0.8422 | 0.2093     | -4.0243 | 1e-04   | 0.4308     |
| Numerical-Reasoning Robot:Non-AI Group  | -1.4510 | 0.2125     | -6.8278 | 0e+00   | 0.2343     |
| 1 2                                     | -0.4269 | 0.0836     | -5.1091 | 0e+00   | 0.6526     |
| 2 3                                     | 1.3739  | 0.0904     | 15.2045 | 0e+00   | 3.9507     |

**Table 8: OLR Summary - Ref. Levels: Action-Declaring Robot and Non-AI Group**

|                                     | Value   | Std. Error | t value  | p value | Odds Ratio |
|-------------------------------------|---------|------------|----------|---------|------------|
| Rationale-Generation Robot          | 1.1824  | 0.1674     | 7.0613   | 0e+00   | 3.2622     |
| Numerical-Reasoning Robot           | -1.9336 | 0.1750     | -11.0468 | 0e+00   | 0.1446     |
| AI Group                            | -0.6859 | 0.1368     | -5.0127  | 0e+00   | 0.5037     |
| Rationale-Generation Robot:AI Group | 0.8422  | 0.2093     | 4.0245   | 1e-04   | 2.3215     |
| Numerical-Reasoning Robot:AI Group  | 1.4510  | 0.2125     | 6.8278   | 0e+00   | 4.2672     |
| 1 2                                 | -1.1127 | 0.1147     | -9.6973  | 0e+00   | 0.3287     |
| 2 3                                 | 0.6880  | 0.1123     | 6.1249   | 0e+00   | 1.9898     |

**Table 9: OLR Summary - Ref. Levels: Numerical-Reasoning Robot and AI Group**

|                                         | Value   | Std. Error | t value | p value | Odds Ratio |
|-----------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot              | 2.5072  | 0.1391     | 18.0222 | 0.0000  | 12.2708    |
| Action-Declaring Robot                  | 0.4826  | 0.1224     | 3.9422  | 0.0001  | 1.6203     |
| Non-AI Group                            | -0.7651 | 0.1620     | -4.7240 | 0.0000  | 0.4653     |
| Rationale-Generation Robot:Non-AI Group | 0.6088  | 0.2264     | 2.6891  | 0.0072  | 1.8382     |
| Action-Declaring Robot:Non-AI Group     | 1.4510  | 0.2125     | 6.8278  | 0.0000  | 4.2672     |
| 1 2                                     | 0.0557  | 0.0923     | 0.6037  | 0.5460  | 1.0573     |
| 2 3                                     | 1.8565  | 0.1033     | 17.9798 | 0.0000  | 6.4012     |

**Table 10: OLR Summary - Ref. Levels: Numerical-Reasoning Robot and Non-AI Group**

|                                     | Value   | Std. Error | t value | p value | Odds Ratio |
|-------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot          | 3.1160  | 0.1893     | 16.4637 | 0.0000  | 22.5559    |
| Action-Declaring Robot              | 1.9336  | 0.1750     | 11.0468 | 0.0000  | 6.9141     |
| AI Group                            | 0.7651  | 0.1620     | 4.7240  | 0.0000  | 2.1492     |
| Rationale-Generation Robot:AI Group | -0.6088 | 0.2264     | -2.6891 | 0.0072  | 0.5440     |
| Action-Declaring Robot:AI Group     | -1.4510 | 0.2125     | -6.8279 | 0.0000  | 0.2343     |
| 1 2                                 | 0.8208  | 0.1336     | 6.1447  | 0.0000  | 2.2724     |
| 2 3                                 | 2.6216  | 0.1444     | 18.1523 | 0.0000  | 13.7575    |

**Table 11: OLR Summary - Confidence - Ref. Levels: Type Rationale-Generation Robot and Group AI Group**

|                                        | Value   | Std. Error | t value | p value | Odds Ratio |
|----------------------------------------|---------|------------|---------|---------|------------|
| Action-Declaring Robot                 | -1.3753 | 0.2740     | -5.0185 | 0.0000  | 0.2528     |
| Numerical-Reasoning Robot              | -1.2607 | 0.2815     | -4.4783 | 0.0000  | 0.2835     |
| Non-AI Group                           | -0.0290 | 0.3297     | -0.0879 | 0.9299  | 0.9714     |
| Action-Declaring Robot:Non-AI Group    | 0.7232  | 0.4545     | 1.5913  | 0.1115  | 2.0610     |
| Numerical-Reasoning Robot:Non-AI Group | -0.6906 | 0.4668     | -1.4795 | 0.1390  | 0.5013     |
| 1 2                                    | -1.6705 | 0.2164     | -7.7194 | 0.0000  | 0.1882     |
| 2 3                                    | -0.1251 | 0.1996     | -0.6268 | 0.5308  | 0.8824     |

**Table 12: OLR Summary - Confidence - Ref. Levels: Type Rationale-Generation Robot and Group Non-AI Group**

|                                    | Value   | Std. Error | t value | p value | Odds Ratio |
|------------------------------------|---------|------------|---------|---------|------------|
| Action-Declaring Robot             | -0.6521 | 0.3654     | -1.7845 | 0.0743  | 0.5210     |
| Numerical-Reasoning Robot          | -1.9513 | 0.3800     | -5.1343 | 0.0000  | 0.1421     |
| AI Group                           | 0.0290  | 0.3297     | 0.0879  | 0.9299  | 1.0294     |
| Action-Declaring Robot:AI Group    | -0.7232 | 0.4545     | -1.5913 | 0.1115  | 0.4852     |
| Numerical-Reasoning Robot:AI Group | 0.6906  | 0.4668     | 1.4795  | 0.1390  | 1.9949     |
| 1 2                                | -1.6415 | 0.2777     | -5.9112 | 0.0000  | 0.1937     |
| 2 3                                | -0.0961 | 0.2651     | -0.3626 | 0.7169  | 0.9083     |

**Table 13: OLR Summary - Confidence - Ref. Levels: Type Action-Declaring Robot and Group AI Group**

|                                         | Value   | Std. Error | t value | p value | Odds Ratio |
|-----------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot              | 1.3753  | 0.2740     | 5.0185  | 0.0000  | 3.9561     |
| Numerical-Reasoning Robot               | 0.1146  | 0.2678     | 0.4279  | 0.6687  | 1.1214     |
| Non-AI Group                            | 0.6942  | 0.3127     | 2.2203  | 0.0264  | 2.0021     |
| Rationale-Generation Robot:Non-AI Group | -0.7232 | 0.4545     | -1.5913 | 0.1115  | 0.4852     |
| Numerical-Reasoning Robot:Non-AI Group  | -1.4138 | 0.4561     | -3.0995 | 0.0019  | 0.2432     |
| 1 2                                     | -0.2952 | 0.1873     | -1.5759 | 0.1150  | 0.7444     |
| 2 3                                     | 1.2501  | 0.1979     | 6.3175  | 0.0000  | 3.4908     |

**Table 14: OLR Summary - Confidence - Ref. Levels: Type Action-Declaring Robot and Group Non-AI Group**

|                                     | Value   | Std. Error | t value | p value | Odds Ratio |
|-------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot          | 0.6521  | 0.3654     | 1.7845  | 0.0743  | 1.9196     |
| Numerical-Reasoning Robot           | -1.2992 | 0.3688     | -3.5230 | 0.0004  | 0.2727     |
| AI Group                            | -0.6942 | 0.3127     | -2.2203 | 0.0264  | 0.4995     |
| Rationale-Generation Robot:AI Group | 0.7232  | 0.4545     | 1.5913  | 0.1115  | 2.0609     |
| Numerical-Reasoning Robot:AI Group  | 1.4138  | 0.4561     | 3.0995  | 0.0019  | 4.1115     |
| 1 2                                 | -0.9894 | 0.2603     | -3.8015 | 0.0001  | 0.3718     |
| 2 3                                 | 0.5560  | 0.2566     | 2.1668  | 0.0303  | 1.7436     |

**Table 15: OLR Summary - Confidence - Ref. Levels: Type Numerical-Reasoning Robot and Group AI Group**

|                                         | Value   | Std. Error | t value | p value | Odds Ratio |
|-----------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot              | 1.2607  | 0.2815     | 4.4783  | 0.0000  | 3.5278     |
| Action-Declaring Robot                  | -0.1146 | 0.2678     | -0.4279 | 0.6687  | 0.8918     |
| Non-AI Group                            | -0.7196 | 0.3304     | -2.1778 | 0.0294  | 0.4870     |
| Rationale-Generation Robot:Non-AI Group | 0.6906  | 0.4668     | 1.4795  | 0.1390  | 1.9949     |
| Action-Declaring Robot:Non-AI Group     | 1.4138  | 0.4561     | 3.0995  | 0.0019  | 4.1115     |
| 1 2                                     | -0.4098 | 0.1999     | -2.0501 | 0.0404  | 0.6638     |
| 2 3                                     | 1.1356  | 0.2076     | 5.4701  | 0.0000  | 3.1130     |

**Table 16: OLR Summary - Confidence - Ref. Levels: Type Numerical-Reasoning Robot and Group Non-AI Group**

|                                     | Value   | Std. Error | t value | p value | Odds Ratio |
|-------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot          | 1.9513  | 0.3800     | 5.1343  | 0.0000  | 7.0377     |
| Action-Declaring Robot              | 1.2992  | 0.3688     | 3.5230  | 0.0004  | 3.6664     |
| AI Group                            | 0.7196  | 0.3304     | 2.1778  | 0.0294  | 2.0536     |
| Rationale-Generation Robot:AI Group | -0.6906 | 0.4668     | -1.4795 | 0.1390  | 0.5013     |
| Action-Declaring Robot:AI Group     | -1.4138 | 0.4561     | -3.0995 | 0.0019  | 0.2432     |
| 1 2                                 | 0.3098  | 0.2666     | 1.1620  | 0.2452  | 1.3631     |
| 2 3                                 | 1.8552  | 0.2812     | 6.5966  | 0.0000  | 6.3927     |

**Table 17: OLR Summary - Friendliness - Ref. Levels: Type Rationale-Generation Robot and Group AI Group**

|                                        | Value   | Std. Error | t value  | p value | Odds Ratio |
|----------------------------------------|---------|------------|----------|---------|------------|
| Action-Declaring Robot                 | -5.8425 | 0.6241     | -9.3621  | 0.0000  | 0.0029     |
| Numerical-Reasoning Robot              | -9.1613 | 0.6886     | -13.3037 | 0.0000  | 0.0001     |
| Non-AI Group                           | -1.2732 | 0.6531     | -1.9494  | 0.0512  | 0.2799     |
| Action-Declaring Robot:Non-AI Group    | 2.3121  | 0.7844     | 2.9477   | 0.0032  | 10.0953    |
| Numerical-Reasoning Robot:Non-AI Group | 0.4423  | 0.8811     | 0.5021   | 0.6156  | 1.5564     |
| 1 2                                    | -7.4844 | 0.6292     | -11.8956 | 0.0000  | 0.0006     |
| 2 3                                    | -3.1140 | 0.5110     | -6.0945  | 0.0000  | 0.0444     |

**Table 18: OLR Summary - Friendliness - Ref. Levels: Type Rationale-Generation Robot and Group Non-AI Group**

|                                    | Value   | Std. Error | t value  | p value | Odds Ratio |
|------------------------------------|---------|------------|----------|---------|------------|
| Action-Declaring Robot             | -3.5301 | 0.5314     | -6.6431  | 0.0000  | 0.0293     |
| Numerical-Reasoning Robot          | -8.7179 | 0.7549     | -11.5482 | 0.0000  | 0.0002     |
| AI Group                           | 1.2734  | 0.6531     | 1.9497   | 0.0512  | 3.5730     |
| Action-Declaring Robot:AI Group    | -2.3128 | 0.7844     | -2.9485  | 0.0032  | 0.0990     |
| Numerical-Reasoning Robot:AI Group | -0.4435 | 0.8809     | -0.5035  | 0.6146  | 0.6418     |
| 1 2                                | -6.2112 | 0.5476     | -11.3434 | 0.0000  | 0.0020     |
| 2 3                                | -1.8407 | 0.4068     | -4.5247  | 0.0000  | 0.1587     |

**Table 19: OLR Summary - Friendliness - Ref. Levels: Type Action-Declaring Robot and Group AI Group**

|                                         | Value   | Std. Error | t value | p value | Odds Ratio |
|-----------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot              | 5.8429  | 0.6241     | 9.3622  | 0.0000  | 344.7825   |
| Numerical-Reasoning Robot               | -3.3185 | 0.3827     | -8.6712 | 0.0000  | 0.0362     |
| Non-AI Group                            | 1.0394  | 0.4342     | 2.3940  | 0.0167  | 2.8274     |
| Rationale-Generation Robot:Non-AI Group | -2.3129 | 0.7844     | -2.9486 | 0.0032  | 0.0990     |
| Numerical-Reasoning Robot:Non-AI Group  | -1.8692 | 0.7336     | -2.5481 | 0.0108  | 0.1542     |
| 1 2                                     | -1.6417 | 0.2599     | -6.3158 | 0.0000  | 0.1936     |
| 2 3                                     | 2.7288  | 0.3585     | 7.6120  | 0.0000  | 15.3144    |

**Table 20: OLR Summary - Friendliness - Ref. Levels: Type Action-Declaring Robot and Group Non-AI Group**

|                                     | Value   | Std. Error | t value | p value | Odds Ratio |
|-------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot          | 3.5301  | 0.5314     | 6.6431  | 0.0000  | 34.1267    |
| Numerical-Reasoning Robot           | -5.1877 | 0.6660     | -7.7898 | 0.0000  | 0.0056     |
| AI Group                            | -1.0394 | 0.4341     | -2.3941 | 0.0167  | 0.3537     |
| Rationale-Generation Robot:AI Group | 2.3125  | 0.7844     | 2.9483  | 0.0032  | 10.1000    |
| Numerical-Reasoning Robot:AI Group  | 1.8692  | 0.7336     | 2.5481  | 0.0108  | 6.4832     |
| 1 2                                 | -2.6811 | 0.4164     | -6.4392 | 0.0000  | 0.0685     |
| 2 3                                 | 1.6894  | 0.3421     | 4.9375  | 0.0000  | 5.4160     |

**Table 21: OLR Summary - Friendliness - Ref. Levels: Type Numerical-Reasoning Robot and Group AI Group**

|                                         | Value   | Std. Error | t value | p value | Odds Ratio |
|-----------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot              | 9.1614  | 0.6887     | 13.3033 | 0.0000  | 9522.2302  |
| Action-Declaring Robot                  | 3.3185  | 0.3827     | 8.6712  | 0.0000  | 27.6181    |
| Non-AI Group                            | -0.8299 | 0.5912     | -1.4038 | 0.1604  | 0.4361     |
| Rationale-Generation Robot:Non-AI Group | -0.4436 | 0.8809     | -0.5036 | 0.6146  | 0.6417     |
| Action-Declaring Robot:Non-AI Group     | 1.8693  | 0.7336     | 2.5481  | 0.0108  | 6.4835     |
| 1 2                                     | 1.6768  | 0.2812     | 5.9620  | 0.0000  | 5.3482     |
| 2 3                                     | 6.0473  | 0.4618     | 13.0936 | 0.0000  | 422.9598   |

**Table 22: OLR Summary - Friendliness - Ref. Levels: Type Numerical-Reasoning Robot and Group Non-AI Group**

|                                     | Value   | Std. Error | t value | p value | Odds Ratio |
|-------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot          | 8.7179  | 0.7549     | 11.5480 | 0.0000  | 6111.4635  |
| Action-Declaring Robot              | 5.1879  | 0.6660     | 7.7897  | 0.0000  | 179.0894   |
| AI Group                            | 0.8301  | 0.5912     | 1.4041  | 0.1603  | 2.2934     |
| Rationale-Generation Robot:AI Group | 0.4434  | 0.8809     | 0.5033  | 0.6148  | 1.5579     |
| Action-Declaring Robot:AI Group     | -1.8695 | 0.7336     | -2.5483 | 0.0108  | 0.1542     |
| 1 2                                 | 2.5068  | 0.5200     | 4.8208  | 0.0000  | 12.2655    |
| 2 3                                 | 6.8773  | 0.6364     | 10.8059 | 0.0000  | 970.0040   |

**Table 23: OLR Summary - Intelligence - Ref. Levels: Type Rationale-Generation Robot and Group AI Group**

|                                        | Value   | Std. Error | t value | p value | Odds Ratio |
|----------------------------------------|---------|------------|---------|---------|------------|
| Action-Declaring Robot                 | -1.8632 | 0.2799     | -6.6555 | 0.0000  | 0.1552     |
| Numerical-Reasoning Robot              | -0.9546 | 0.2801     | -3.4088 | 0.0007  | 0.3850     |
| Non-AI Group                           | -0.3321 | 0.3304     | -1.0053 | 0.3147  | 0.7174     |
| Action-Declaring Robot:Non-AI Group    | 1.0047  | 0.4542     | 2.2120  | 0.0270  | 2.7310     |
| Numerical-Reasoning Robot:Non-AI Group | -0.1047 | 0.4626     | -0.2262 | 0.8210  | 0.9006     |
| 1 2                                    | -1.7485 | 0.2186     | -7.9976 | 0.0000  | 0.1740     |
| 2 3                                    | -0.2121 | 0.2011     | -1.0547 | 0.2916  | 0.8089     |

**Table 24: OLR Summary - Intelligence - Ref. Levels: Type Rationale-Generation Robot and Group Non-AI Group**

|                                    | Value   | Std. Error | t value | p value | Odds Ratio |
|------------------------------------|---------|------------|---------|---------|------------|
| Action-Declaring Robot             | -0.8585 | 0.3631     | -2.3643 | 0.0181  | 0.4238     |
| Numerical-Reasoning Robot          | -1.0593 | 0.3716     | -2.8510 | 0.0044  | 0.3467     |
| AI Group                           | 0.3321  | 0.3304     | 1.0053  | 0.3148  | 1.3939     |
| Action-Declaring Robot:AI Group    | -1.0046 | 0.4542     | -2.2119 | 0.0270  | 0.3662     |
| Numerical-Reasoning Robot:AI Group | 0.1047  | 0.4626     | 0.2263  | 0.8210  | 1.1104     |
| 1 2                                | -1.4164 | 0.2748     | -5.1533 | 0.0000  | 0.2426     |
| 2 3                                | 0.1200  | 0.2651     | 0.4528  | 0.6507  | 1.1275     |

**Table 25: OLR Summary - Intelligence - Ref. Levels: Type Action-Declaring Robot and Group AI Group**

|                                         | Value   | Std. Error | t value | p value | Odds Ratio |
|-----------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot              | 1.8631  | 0.2799     | 6.6555  | 0.0000  | 6.4440     |
| Numerical-Reasoning Robot               | 0.9085  | 0.2704     | 3.3604  | 0.0008  | 2.4806     |
| Non-AI Group                            | 0.6725  | 0.3110     | 2.1623  | 0.0306  | 1.9592     |
| Rationale-Generation Robot:Non-AI Group | -1.0047 | 0.4542     | -2.2120 | 0.0270  | 0.3662     |
| Numerical-Reasoning Robot:Non-AI Group  | -1.1093 | 0.4502     | -2.4641 | 0.0137  | 0.3298     |
| 1 2                                     | 0.1147  | 0.1878     | 0.6106  | 0.5415  | 1.1215     |
| 2 3                                     | 1.6511  | 0.2051     | 8.0509  | 0.0000  | 5.2125     |

**Table 26: OLR Summary - Intelligence - Ref. Levels: Type Action-Declaring Robot and Group Non-AI Group**

|                                     | Value   | Std. Error | t value | p value | Odds Ratio |
|-------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot          | 0.8585  | 0.3631     | 2.3642  | 0.0181  | 2.3596     |
| Numerical-Reasoning Robot           | -0.2008 | 0.3593     | -0.5589 | 0.5762  | 0.8181     |
| AI Group                            | -0.6725 | 0.3110     | -2.1623 | 0.0306  | 0.5104     |
| Rationale-Generation Robot:AI Group | 1.0047  | 0.4542     | 2.2120  | 0.0270  | 2.7310     |
| Numerical-Reasoning Robot:AI Group  | 1.1093  | 0.4502     | 2.4641  | 0.0137  | 3.0323     |
| 1 2                                 | -0.5579 | 0.2530     | -2.2047 | 0.0275  | 0.5724     |
| 2 3                                 | 0.9785  | 0.2566     | 3.8132  | 0.0001  | 2.6605     |

**Table 27: OLR Summary - Intelligence - Ref. Levels: Type Numerical-Reasoning Robot and Group AI Group**

|                                         | Value   | Std. Error | t value | p value | Odds Ratio |
|-----------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot              | 0.9546  | 0.2801     | 3.4088  | 0.0007  | 2.5977     |
| Action-Declaring Robot                  | -0.9085 | 0.2704     | -3.3604 | 0.0008  | 0.4031     |
| Non-AI Group                            | -0.4368 | 0.3244     | -1.3467 | 0.1781  | 0.6461     |
| Rationale-Generation Robot:Non-AI Group | 0.1047  | 0.4626     | 0.2263  | 0.8210  | 1.1104     |
| Action-Declaring Robot:Non-AI Group     | 1.1093  | 0.4502     | 2.4641  | 0.0137  | 3.0323     |
| 1 2                                     | -0.7939 | 0.2021     | -3.9274 | 0.0001  | 0.4521     |
| 2 3                                     | 0.7425  | 0.2016     | 3.6836  | 0.0002  | 2.1013     |

**Table 28: OLR Summary - Intelligence - Ref. Levels: Type Numerical-Reasoning Robot and Group Non-AI Group**

|                                     | Value   | Std. Error | t value | p value | Odds Ratio |
|-------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot          | 1.0593  | 0.3716     | 2.8510  | 0.0044  | 2.8844     |
| Action-Declaring Robot              | 0.2008  | 0.3593     | 0.5589  | 0.5763  | 1.2224     |
| AI Group                            | 0.4368  | 0.3244     | 1.3467  | 0.1781  | 1.5477     |
| Rationale-Generation Robot:AI Group | -0.1047 | 0.4626     | -0.2263 | 0.8210  | 0.9006     |
| Action-Declaring Robot:AI Group     | -1.1093 | 0.4502     | -2.4641 | 0.0137  | 0.3298     |
| 1 2                                 | -0.3571 | 0.2628     | -1.3587 | 0.1742  | 0.6997     |
| 2 3                                 | 1.1793  | 0.2693     | 4.3787  | 0.0000  | 3.2522     |

**Table 29: OLR Summary - Potential - Ref. Levels: Type Rationale-Generation Robot and Group AI Group**

|                                        | Value   | Std. Error | t value | p value | Odds Ratio |
|----------------------------------------|---------|------------|---------|---------|------------|
| Action-Declaring Robot                 | -1.4669 | 0.2736     | -5.3619 | 0.0000  | 0.2306     |
| Numerical-Reasoning Robot              | -1.4049 | 0.2855     | -4.9210 | 0.0000  | 0.2454     |
| Non-AI Group                           | 0.3139  | 0.3403     | 0.9226  | 0.3562  | 1.3688     |
| Action-Declaring Robot:Non-AI Group    | 0.3615  | 0.4584     | 0.7888  | 0.4303  | 1.4355     |
| Numerical-Reasoning Robot:Non-AI Group | -1.2945 | 0.4854     | -2.6667 | 0.0077  | 0.2740     |
| 1 2                                    | -1.7680 | 0.2168     | -8.1558 | 0.0000  | 0.1707     |
| 2 3                                    | -0.1478 | 0.1969     | -0.7505 | 0.4530  | 0.8626     |

**Table 30: OLR Summary - Potential - Ref. Levels: Type Rationale-Generation Robot and Group Non-AI Group**

|                                    | Value   | Std. Error | t value | p value | Odds Ratio |
|------------------------------------|---------|------------|---------|---------|------------|
| Action-Declaring Robot             | -1.1053 | 0.3729     | -2.9642 | 0.0030  | 0.3311     |
| Numerical-Reasoning Robot          | -2.6994 | 0.4039     | -6.6836 | 0.0000  | 0.0672     |
| AI Group                           | -0.3139 | 0.3403     | -0.9224 | 0.3563  | 0.7306     |
| Action-Declaring Robot:AI Group    | -0.3616 | 0.4584     | -0.7889 | 0.4302  | 0.6966     |
| Numerical-Reasoning Robot:AI Group | 1.2944  | 0.4854     | 2.6666  | 0.0077  | 3.6489     |
| 1 2                                | -2.0820 | 0.2966     | -7.0203 | 0.0000  | 0.1247     |
| 2 3                                | -0.4617 | 0.2792     | -1.6534 | 0.0983  | 0.6302     |

**Table 31: OLR Summary - Potential - Ref. Levels: Type Action-Declaring Robot and Group AI Group**

|                                         | Value   | Std. Error | t value | p value | Odds Ratio |
|-----------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot              | 1.4668  | 0.2736     | 5.3616  | 0.0000  | 4.3354     |
| Numerical-Reasoning Robot               | 0.0620  | 0.2728     | 0.2271  | 0.8203  | 1.0639     |
| Non-AI Group                            | 0.6755  | 0.3077     | 2.1950  | 0.0282  | 1.9649     |
| Rationale-Generation Robot:Non-AI Group | -0.3614 | 0.4584     | -0.7885 | 0.4304  | 0.6967     |
| Numerical-Reasoning Robot:Non-AI Group  | -1.6561 | 0.4643     | -3.5668 | 0.0004  | 0.1909     |
| 1 2                                     | -0.3012 | 0.1880     | -1.6022 | 0.1091  | 0.7399     |
| 2 3                                     | 1.3191  | 0.2005     | 6.5806  | 0.0000  | 3.7401     |

**Table 32: OLR Summary - Potential - Ref. Levels: Type Action-Declaring Robot and Group Non-AI Group**

|                                     | Value   | Std. Error | t value | p value | Odds Ratio |
|-------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot          | 1.1053  | 0.3729     | 2.9641  | 0.0030  | 3.0201     |
| Numerical-Reasoning Robot           | -1.5941 | 0.3754     | -4.2465 | 0.0000  | 0.2031     |
| AI Group                            | -0.6755 | 0.3077     | -2.1951 | 0.0282  | 0.5089     |
| Rationale-Generation Robot:AI Group | 0.3616  | 0.4584     | 0.7889  | 0.4302  | 1.4356     |
| Numerical-Reasoning Robot:AI Group  | 1.6560  | 0.4643     | 3.5667  | 0.0004  | 5.2385     |
| 1 2                                 | -0.9767 | 0.2543     | -3.8407 | 0.0001  | 0.3766     |
| 2 3                                 | 0.6436  | 0.2510     | 2.5639  | 0.0104  | 1.9034     |

**Table 33: OLR Summary - Potential - Ref. Levels: Type Numerical-Reasoning Robot and Group AI Group**

|                                         | Value   | Std. Error | t value | p value | Odds Ratio |
|-----------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot              | 1.4049  | 0.2855     | 4.9210  | 0.0000  | 4.0752     |
| Action-Declaring Robot                  | -0.0620 | 0.2728     | -0.2271 | 0.8203  | 0.9399     |
| Non-AI Group                            | -0.9806 | 0.3456     | -2.8371 | 0.0046  | 0.3751     |
| Rationale-Generation Robot:Non-AI Group | 1.2945  | 0.4854     | 2.6667  | 0.0077  | 3.6491     |
| Action-Declaring Robot:Non-AI Group     | 1.6560  | 0.4643     | 3.5667  | 0.0004  | 5.2386     |
| 1 2                                     | -0.3631 | 0.2060     | -1.7628 | 0.0779  | 0.6955     |
| 2 3                                     | 1.2571  | 0.2162     | 5.8156  | 0.0000  | 3.5154     |

**Table 34: OLR Summary - Potential - Ref. Levels: Type Numerical-Reasoning Robot and Group Non-AI Group**

|                                     | Value   | Std. Error | t value | p value | Odds Ratio |
|-------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot          | 2.6994  | 0.4039     | 6.6837  | 0.0000  | 14.8709    |
| Action-Declaring Robot              | 1.5941  | 0.3754     | 4.2466  | 0.0000  | 4.9238     |
| AI Group                            | 0.9806  | 0.3456     | 2.8371  | 0.0046  | 2.6659     |
| Rationale-Generation Robot:AI Group | -1.2945 | 0.4854     | -2.6667 | 0.0077  | 0.2740     |
| Action-Declaring Robot:AI Group     | -1.6560 | 0.4643     | -3.5667 | 0.0004  | 0.1909     |
| 1 2                                 | 0.6174  | 0.2798     | 2.2063  | 0.0274  | 1.8541     |
| 2 3                                 | 2.2377  | 0.2983     | 7.5020  | 0.0000  | 9.3718     |

**Table 35: OLR Summary - Understandability - Ref. Levels: Type Rationale-Generation Robot and Group AI Group**

|                                        | Value   | Std. Error | t value  | p value | Odds Ratio |
|----------------------------------------|---------|------------|----------|---------|------------|
| Action-Declaring Robot                 | -2.6153 | 0.3274     | -7.9873  | 0.0000  | 0.0731     |
| Numerical-Reasoning Robot              | -4.2054 | 0.3661     | -11.4866 | 0.0000  | 0.0149     |
| Non-AI Group                           | -0.4175 | 0.3664     | -1.1393  | 0.2546  | 0.6587     |
| Action-Declaring Robot:Non-AI Group    | 1.5455  | 0.4949     | 3.1227   | 0.0018  | 4.6903     |
| Numerical-Reasoning Robot:Non-AI Group | -1.7585 | 0.7309     | -2.4059  | 0.0161  | 0.1723     |
| 1 2                                    | -3.5627 | 0.3026     | -11.7746 | 0.0000  | 0.0284     |
| 2 3                                    | -1.0644 | 0.2358     | -4.5135  | 0.0000  | 0.3449     |

**Table 36: OLR Summary - Understandability - Ref. Levels: Type Rationale-Generation Robot and Group Non-AI Group**

|                                    | Value   | Std. Error | t value | p value | Odds Ratio |
|------------------------------------|---------|------------|---------|---------|------------|
| Action-Declaring Robot             | -1.0698 | 0.3807     | -2.8100 | 0.0050  | 0.3431     |
| Numerical-Reasoning Robot          | -5.9638 | 0.6824     | -8.7391 | 0.0000  | 0.0026     |
| AI Group                           | 0.4175  | 0.3664     | 1.1393  | 0.2546  | 1.5182     |
| Action-Declaring Robot:AI Group    | -1.5455 | 0.4949     | -3.1227 | 0.0018  | 0.2132     |
| Numerical-Reasoning Robot:AI Group | 1.7584  | 0.7309     | 2.4058  | 0.0161  | 5.8032     |
| 1 2                                | -3.1452 | 0.3361     | -9.3583 | 0.0000  | 0.0431     |
| 2 3                                | -0.6469 | 0.2808     | -2.3041 | 0.0212  | 0.5237     |

**Table 37: OLR Summary - Understandability - Ref. Levels: Type Action-Declaring Robot and Group AI Group**

|                                         | Value   | Std. Error | t value | p value | Odds Ratio |
|-----------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot              | 2.6151  | 0.3274     | 7.9869  | 0.0000  | 13.6684    |
| Numerical-Reasoning Robot               | -1.5903 | 0.2993     | -5.3142 | 0.0000  | 0.2039     |
| Non-AI Group                            | 1.1277  | 0.3316     | 3.4009  | 0.0007  | 3.0884     |
| Rationale-Generation Robot:Non-AI Group | -1.5454 | 0.4949     | -3.1226 | 0.0018  | 0.2132     |
| Numerical-Reasoning Robot:Non-AI Group  | -3.3039 | 0.7154     | -4.6183 | 0.0000  | 0.0367     |
| 1 2                                     | -0.9474 | 0.2107     | -4.4964 | 0.0000  | 0.3877     |
| 2 3                                     | 1.5506  | 0.2308     | 6.7170  | 0.0000  | 4.7143     |

**Table 38: OLR Summary - Understandability - Ref. Levels: Type Action-Declaring Robot and Group Non-AI Group**

|                                     | Value   | Std. Error | t value | p value | Odds Ratio |
|-------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot          | 1.0698  | 0.3807     | 2.8099  | 0.0050  | 2.9147     |
| Numerical-Reasoning Robot           | -4.8941 | 0.6646     | -7.3645 | 0.0000  | 0.0075     |
| AI Group                            | -1.1281 | 0.3316     | -3.4020 | 0.0007  | 0.3237     |
| Rationale-Generation Robot:AI Group | 1.5456  | 0.4949     | 3.1229  | 0.0018  | 4.6909     |
| Numerical-Reasoning Robot:AI Group  | 3.3039  | 0.7153     | 4.6190  | 0.0000  | 27.2187    |
| 1 2                                 | -2.0755 | 0.2980     | -6.9645 | 0.0000  | 0.1255     |
| 2 3                                 | 0.4228  | 0.2591     | 1.6321  | 0.1027  | 1.5263     |

**Table 39: OLR Summary - Understandability - Ref. Levels: Type Numerical-Reasoning Robot and Group AI Group**

|                                         | Value   | Std. Error | t value | p value | Odds Ratio |
|-----------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot              | 4.2054  | 0.3661     | 11.4867 | 0.0000  | 67.0472    |
| Action-Declaring Robot                  | 1.5901  | 0.2993     | 5.3136  | 0.0000  | 4.9043     |
| Non-AI Group                            | -2.1759 | 0.6325     | -3.4403 | 0.0006  | 0.1135     |
| Rationale-Generation Robot:Non-AI Group | 1.7584  | 0.7309     | 2.4058  | 0.0161  | 5.8031     |
| Action-Declaring Robot:Non-AI Group     | 3.3039  | 0.7153     | 4.6189  | 0.0000  | 27.2188    |
| 1 2                                     | 0.6427  | 0.2166     | 2.9667  | 0.0030  | 1.9015     |
| 2 3                                     | 3.1410  | 0.2846     | 11.0345 | 0.0000  | 23.1261    |

**Table 40: OLR Summary - Understandability - Ref. Levels: Type Numerical-Reasoning Robot and Group Non-AI Group**

|                                     | Value   | Std. Error | t value | p value | Odds Ratio |
|-------------------------------------|---------|------------|---------|---------|------------|
| Rationale-Generation Robot          | 5.9638  | 0.6824     | 8.7390  | 0.0000  | 389.0826   |
| Action-Declaring Robot              | 4.8940  | 0.6646     | 7.3643  | 0.0000  | 133.4890   |
| AI Group                            | 2.1759  | 0.6325     | 3.4403  | 0.0006  | 8.8105     |
| Rationale-Generation Robot:AI Group | -1.7584 | 0.7309     | -2.4058 | 0.0161  | 0.1723     |
| Action-Declaring Robot:AI Group     | -3.3039 | 0.7153     | -4.6189 | 0.0000  | 0.0367     |
| 1 2                                 | 2.8186  | 0.5943     | 4.7431  | 0.0000  | 16.7534    |
| 2 3                                 | 5.3169  | 0.6257     | 8.4981  | 0.0000  | 203.7525   |