

Acceptance of Task Allocation to Intelligent Robots in Monetary and Moral Contexts: A Study of Independent and Collaborative Modes

Authors: Jiang Duo, Luo Zhenwang, Huang Weiqi, Luo Nanbao, Chen Yawen, Luo Zhenwang, Huang Weiqi

Date: 2025-06-26T15:23:44+00:00

Abstract

With the advancement of artificial intelligence technology, intelligent robots have begun to act as working agents in human work contexts. Do humans accept the allocation of work to intelligent robots? Intelligent robots possess moderate agency and low sentience, resulting in differences between the work tasks they can undertake and those performed by humans. Experiment 1 and Experiment 2 respectively explored whether there exist human-robot differences in the acceptance of work allocation within the contexts of monetary gain/loss and moral gain/loss work tasks. Results indicated that, in both monetary and moral tasks, people are more accepting of allocating loss-related work to robots and gain-related work to humans. Mind perception and responsibility perception mediated the effect of agent type on work allocation acceptance. Experiment 3 and Experiment 4 further investigated whether multi-agent collaborative work influences work allocation acceptance. Experimental results demonstrated that different collaborative work groups exhibited differentiated group mind, and group mind could also influence people's acceptance of work allocation through responsibility perception. The findings offer reference value for clarifying the status and role of intelligent robots in the social division of labor.

Full Text

Acceptance of Work Allocation to Intelligent Robots in Monetary and Moral Contexts: A Study Based on Independent and Collaborative Work Models

JIANG Duo, LUO Zhenwang, HUANG Weiqi, LUO Nanbao, CHEN Yawen

(School of Psychology, Shenzhen University, Shenzhen 518060, China)

Abstract

As artificial intelligence technology advances, intelligent robots have begun participating in human work as active agents. This raises a fundamental question: To what extent are humans willing to accept work allocation to intelligent robots? Intelligent robots possess moderate agency and low experience, enabling them to undertake tasks that differ from those suitable for humans. Experiments 1 and 2 explored whether acceptance of work allocation differs between humans and robots in monetary and moral contexts. Results showed that in both contexts, people more readily accept allocating loss-related work to robots while preferring to assign gain-related work to humans. Mind perception and responsibility perception mediated the relationship between agent type and work allocation acceptance. Experiments 3 and 4 further examined whether collaborative work modes influence acceptance of work allocation. Findings revealed that different collaborative teams exhibit distinct collective minds, which influence work allocation acceptance through responsibility perception. These results provide valuable insights for clarifying the role and position of intelligent robots in the social division of labor.

Keywords: work allocation, mind perception, responsibility perception, human-robot collaboration

Classification: B849

1. Introduction

Traditional industrial robots primarily execute repetitive tasks in structured environments through pre-programmed instructions, focusing on process standardization and efficiency improvement. With advances in sensor technology, path planning algorithms, and adaptive control systems, modern automated robots have gained greater task flexibility (enabling task switching through parameter adjustment) and environmental adaptability (allowing operation correction based on real-time sensor data), enabling them to perform more complex tasks such as high-precision assembly and dynamic obstacle avoidance (Siciliano & Khatib, 2016). However, their behavioral logic remains highly dependent on programming rules, limiting autonomous decision-making capabilities in unstructured scenarios. Artificial Intelligence (AI) technology simulates human cognitive processes to process information through algorithms (Longoni et al., 2019). This technology can equip automated robots with an “AI brain,” transforming them into intelligent robots capable of performing various thinking and decision-making tasks like humans (Gibney, 2024). Intelligent robots represent a new automation system combining AI technology with automated robotics, retaining the functionality of automated robots while integrating AI algorithms to facilitate autonomous learning and thinking, and possessing certain independent decision-making capabilities (Bowen & Morosan, 2018; Huang & Rust, 2018).

Currently, intelligent robots are increasingly penetrating various aspects of hu-

man work. When humans and intelligent robots simultaneously serve as independent work agents, work allocation between them becomes necessary. Addressing this allocation problem first requires clarifying which types of work humans are willing to accept being performed by intelligent robots within the social division of labor. Mind perception theory suggests that differences between intelligent robots and humans are reflected in mind capabilities (Huebner, 2010; Pazhoohi et al., 2023). These differences may influence perceptions and judgments about the responsibilities that humans and intelligent robots can assume in work, thereby shaping considerations of work allocation. Therefore, this study explores how to allocate work between humans and intelligent robots in ways that are more acceptable, based on mind perception theory.

Although intelligent robots' mind capabilities are lower than humans', their powerful computational and storage abilities create possibilities for complementarity between humans and robots, leading to new models of human-robot collaboration (Awad et al., 2020). Meanwhile, workplace collaboration also exists in human-human and robot-robot forms. When work agents shift from individual to group forms, how do the mind capabilities of various collaborative teams (including human-human, human-robot, and robot-robot teams) change compared to individual agents? How do these collaborative teams' mind capabilities affect the responsibilities they can assume and the acceptance of work allocation? This study also explores the collective mind capabilities possessed by collaborative teams and reveals how collective mind influences work allocation acceptance.

1.1 Mind Perception Theory

Mind refers to the capacity for thinking, feeling, and conscious behavior (Tharp et al., 2017). It comprises two dimensions: agency and experience (Gray et al., 2007). Agency represents the capacity for cognition and action, including abilities such as self-control, judgment, communication, thinking, and memory. Experience represents the capacity to feel emotions like fear, pain, and joy (Gray et al., 2007; Gray, Young, & Waytz, 2012; Waytz et al., 2010; Weisman et al., 2017). Research indicates that compared to humans, automated robots possess moderate agency but low experience (Gray et al., 2007; Gray & Wegner, 2012; 詹泽, 吴宝沛, 2019).

Despite rapid AI development enhancing information processing capabilities, AI behavior remains constrained by programming and lacks free will, resulting in agency still lower than humans (胡小勇 et al., 2024; Malle, 2019). AI capabilities derive from complex algorithms formed through training on large datasets. Generative AI produces new content, including responses to human emotions, based on learned patterns. However, this process fundamentally involves probability calculation and pattern matching, without genuine feeling or understanding of human emotions (Markovitch et al., 2024). Although automated robots become intelligent robots when equipped with an AI brain, gaining stronger information processing and behavioral capabilities, their capacity to experience emotions remains lower than humans. This leads to the following hypothesis:

H1: Humans possess higher mind capabilities (including both agency and experience) than intelligent robots.

1.2 Work Allocation

Work allocation refers to the assignment of tasks to work agents. According to person-job fit theory, organizational leaders must allocate work based on the match between employee capabilities and job requirements (Wang et al., 2023; van Woerkom et al., 2024). Humans and intelligent robots possess differentiated mind capabilities, which may lead to differences in work allocation. Agency in mind capabilities reflects cognitive and executive abilities, while experience reflects emotional feeling capacity (闫霄 et al., 2024). Since both monetary and moral tasks involve cognitive and emotional processing, these task types are commonly used in research exploring human-robot differences (Larkin et al., 2021).

1.2.1 Monetary Tasks Monetary tasks typically involve financial gains and losses. For instance, monetary investment is essentially a gain-loss task. In investment tasks, an individual's agency facilitates autonomous market information collection and analysis, accurate calculation of investment returns, and formulation and execution of optimal investment strategies to maximize gains (Grinblatt et al., 2011). Investors with higher agency not only excel at asset allocation but also effectively avoid decision-making errors caused by insufficient information processing, thereby achieving higher returns in financial markets (Burks et al., 2009). Experience also significantly influences monetary gain-loss tasks because emotional feelings provide early warnings that help investors make correct and reasonable decisions. Loewenstein et al.'s (2001) risk-as-feelings hypothesis posits that emotions experienced during monetary decision-making (e.g., fear) enable timely profit-taking or loss-cutting decisions, providing protective warnings in volatile markets. Lerner and Keltner (2001) experimentally demonstrated that fear of loss increases risk perception, prompting more conservative strategies. The somatic marker hypothesis emphasizes that emotional experiences during decision-making are recorded by the body. When facing similar decisions again, previous emotional memories remind individuals to avoid options that may lead to losses, resulting in higher-quality decisions (Damasio, 1994; Bechara & Damasio, 2005; Walteros et al., 2011). Therefore, both agency and experience play important roles in monetary gain-loss tasks.

Humans' high agency enables them to analyze information, set investment goals, strategies, and plans, and execute them in monetary gain-loss tasks. Thus, high agency benefits humans in performing monetary tasks. Although intelligent robots have lower agency than humans, they can compensate through powerful algorithms and computational capacity for rapid and accurate analysis of massive information. However, humans also possess high experience, enabling them to timely feel and respond to their own emotions, which provides crucial warnings for making correct and effective decisions in rapidly changing markets. Due

to robots' very low experience, they cannot provide similar emotional warning signals during task execution, lacking necessary emergency adjustment capabilities. This mind capability disadvantage leads people to prefer human advice in investment decisions (Larkin et al., 2021). Overall, humans' higher mind capabilities make them more competent for monetary gain-loss tasks. Based on person-job fit theory, the following hypothesis emerges:

H2: In monetary gain-loss contexts, acceptance of work allocation differs between humans and intelligent robots: individuals are more willing to accept allocating monetary work to humans than to intelligent robots.

1.2.2 Moral Tasks Moral tasks also involve gains and losses: when task outcomes receive public praise, this represents moral gain; conversely, public condemnation represents moral loss (Szekeres et al., 2019). Thus, moral tasks that produce gain or loss outcomes are essentially moral gain-loss tasks. High agency facilitates moral decision-making (whether choices harm others' interests), moral judgment (determining whether behaviors deserve reward or punishment), and moral reasoning (inferring individuals' moral qualities), enabling moral behavior execution (Yu et al., 2019). Healthy adults with higher moral cognition typically make correct and reasonable moral decisions (Gray & Wegner, 2009). Although agency influences moral task completion, experience is often considered more powerful in explaining moral behavior (Gray & Wegner, 2012; Greene et al., 2001). The social intuition model posits that emotional feelings play a primary role in moral tasks (Haidt, 2001). Emotional experiences ranging from anger to guilt provide foundations for moral judgment and behavior implementation, influencing moral decisions and actions (Haidt et al., 1993; Greene et al., 2001). Additionally, feeling others' pain and demonstrating high empathy and compassion are core elements of moral judgment (Bigman & Gray, 2018). Therefore, experience plays a crucial role in moral gain-loss tasks.

Although intelligent robots possess moderate agency, their deficiency in experience prevents them from accurately feeling others' emotions and empathizing with them. This defect prevents intelligent robots from engaging in moral judgment—they struggle to understand whether behaviors are morally right or wrong toward others—and thus they are not complete moral agents (胡小勇 et al., 2024). In contrast, humans' high agency and experience enable them to understand moral principles, empathize with others, and perform moral behaviors. Overall, humans' higher mind capabilities make them more competent for moral gain-loss tasks than intelligent robots. Based on person-job fit theory, the following hypothesis emerges:

H3: In moral gain-loss contexts, acceptance of work allocation differs between humans and intelligent robots: individuals are more willing to accept allocating moral work to humans than to intelligent robots.

1.3 Responsibility Perception

When leaders allocate work to an agent, they expect that agent to assume certain responsibilities for the work. Mayer et al. (1995) noted that when leaders perceive employees as capable of assuming work responsibilities, they trust them to complete corresponding tasks and are therefore willing to allocate work to them. Thus, perceptions of whether work agents can assume responsibility (responsibility perception) directly influence acceptance of work allocation.

Greater ability entails greater responsibility. Perceptions of responsibility-bearing capacity relate to the agent's capabilities. As one of many capabilities, mind capability is closely related to responsibility perception (闫霄 et al., 2024; Willemssen et al., 2023). For example, people do not assign work responsibilities to infants with low mind capabilities. Agency connects to responsibility perception: if people believe an agent can autonomously formulate and execute plans, that agent naturally becomes a responsibility-bearing entity with higher responsibility; conversely, if people doubt or deny an agent's planning and execution abilities, responsibility judgments become ambiguous and the agent bears less responsibility (喻丰, 许丽颖, 2019). Research found that increased agency beliefs (belief in controlling one's behavior) led to fewer cheating behaviors (Vohs & Schooler, 2008). Enhanced agency beliefs also produce more responsible behavior (Rigoni et al., 2013). This demonstrates that increased agency enables individuals to assume greater responsibility.

Experience also connects to responsibility perception. Experience reflects an agent's capacity to feel emotions. High-experience individuals more deeply feel emotions resulting from behaviors (e.g., guilt, regret), which serve as important bases for responsibility attribution (Tangney et al., 2007; 闫霄 et al., 2024). Therefore, when experience is higher, individuals are more likely to be required to assume corresponding responsibilities because they can appreciate how behaviors affect themselves and others. Research shows that high-experience individuals more easily perceive others' pain, thereby enhancing responsibility toward them and displaying altruistic behavior (Decety & Cowell, 2014). In another study, when participants were reminded that driving behavior endangers their own and others' lives, they were less willing to delegate driving authority to robots and more willing to assume responsibility themselves (Bigman & Gray, 2018). This demonstrates that high experience enables individuals to assume greater responsibility.

Since both monetary and moral tasks are complex tasks requiring both cognitive and emotional factors, both agency and experience of work agents may influence work allocation acceptance through responsibility perception. Therefore, mind capability and responsibility perception may play a chain mediating role in the path from agent type to work allocation acceptance. The following hypothesis is proposed:

H4: Mind capability and responsibility perception chain-mediate the effect of agent type on work allocation acceptance: mind capability positively influences

responsibility perception, which in turn positively influences work allocation acceptance.

1.4 Collaborative Work Models

Licklider (1960) proposed the concept of man-computer symbiosis, suggesting that humans and computers could collaborate closely to improve efficiency through complementarity. Early research focused on using computers to extend human cognitive processing of information and enhance collaborative efficiency. In the 1970s-80s, with technologies like graphical user interfaces (GUI), the concept of human-computer interaction (HCI) emerged—referring to the process where humans and computers exchange information through dialogue in a specific language (Shoemaker & Tetlock, 2017; 刘明泽, 2022). Researchers began focusing on how humans could effectively interact with computers to improve collaborative efficiency and quality. In the 1990s, automation technology development led to increasingly autonomous machines and robots capable of completing tasks without human supervision or intervention (Realyvásquez-Vargas et al., 2019). This gave rise to collaborative robots (Cobots), representing cooperation between machines and humans on production lines (Silva et al., 2024). Research in this phase emphasized how systems automatically adjusted behavior based on human input and needs. After 2015, rapid AI development introduced the concept of collaborative intelligence (Epstein, 2015). With advances in deep learning and natural language processing, AI systems can provide intelligent analysis and decision support in complex environments and co-decide with humans.

With the emergence of intelligent robots, besides traditional human-human collaboration, robot-robot and human-robot collaboration modes have appeared (Awad et al., 2020; Liu, 2023; Ren et al., 2023). Human-human collaboration involves people working together through communication, negotiation, and division of labor to complete tasks or achieve goals, emphasizing team cooperation. Robot-robot collaboration involves robots coordinating and synchronizing their work processes to improve efficiency and precision. Both human-human and robot-robot teams are formed by adding homogeneous agents to individual work modes.

Currently, increasing numbers of intelligent robots are becoming team members working alongside humans, creating human-robot collaboration modes (何贵兵 et al., 2022; Seeber et al., 2020). Human-robot collaboration involves humans and intelligent robots cooperating to complete tasks—collaboration between two heterogeneous agents. Humans surpass intelligent robots in mind capabilities, while intelligent robots excel humans in computational power, storage, and other abilities. These capability differences create complementarity between humans and intelligent robots, enabling the integration of their intelligence to construct “hybrid intelligence” (Zheng et al., 2017; Wiese et al., 2022). A PwC survey showed that 67% of executives believe hybrid intelligence represents the future of robots and humans. Hybrid intelligence is not merely robots assisting humans

but a deep-level integration and complementarity requiring full combination of humans and intelligent robots to form a compatible system for jointly completing tasks.

Both humans and intelligent robots possess their own mind capabilities. When they combine to form human-human, human-robot, and robot-robot collaborative teams, they no longer exhibit the independent mind capabilities of humans or intelligent robots but may instead display collective mind capabilities. Distributed cognition theory posits that cognitive activity is not confined within individuals but is a system characteristic (Hutchins, 1995; 周国梅, 傅小兰, 2002). Based on this theory, Hutchins noted that cognitive activity can be organized across multiple levels: collaboration among different team members and reliance on tools and artifacts. With intelligent robots, humans and intelligent robots can also form systems that generate cognitive activity. Burton et al. (2024) argued that through distributed cognition and coordination, human and AI intelligence can fuse to produce collective intelligence that significantly exceeds individual intelligence. This demonstrates that different agents' capabilities can integrate within groups or teams. Can individual minds, as a capability, also integrate within groups or teams to exhibit collective mind? This study aims to extend mind perception theory from the individual level to the group level and explore differences in collective mind among human-human, human-robot, and robot-robot collaborative teams. Since humans possess higher mind capabilities than intelligent robots, this study predicts that the three collaborative team types may also differ in collective mind, specifically: human-human highest, human-robot intermediate, and robot-robot lowest. Building on exploring collective mind differences, this study will also examine how collective mind influences group responsibility perception and work allocation acceptance. The following research questions are proposed:

RQ1: Do human-human, human-robot, and robot-robot collaborative teams differ in collective mind capability?

RQ2: Does collective mind influence responsibility perception and work allocation acceptance?

2. Preliminary Experiment: Scenario Construction

2.1 Purpose

This study uses contextualized decision-making tasks to explore differences in work allocation acceptance between humans and intelligent robots and among different collaborative teams, and to examine the mechanisms underlying these differences. The experiments require constructing monetary gain-loss and moral gain-loss work scenarios. The preliminary experiment uses psychometric methods to validate the effectiveness of scenario construction.

2.2 Method

2.2.1 Participants Eighty-five participants were randomly recruited, including 43 females. Participants ranged in age from 18 to 25 years ($M = 21.80$, $SD = 1.56$). Each participant received compensation upon completion.

2.2.2 Materials (1) Monetary Gain-Loss Work Scenario

Since monetary investment tasks involve both gains and losses, this study uses monetary investment as the monetary gain-loss task. In investment tasks, the magnitude of gains and losses affects investment value assessment, which may influence responsibility evaluation and work allocation acceptance (Kahneman & Tversky, 1979; 杨玲 et al., 2019). Therefore, scenario construction must consider gain-loss magnitude. As participants differ in their valuation of real currency, the study uses gold coin quantities to represent high versus low monetary gain-loss amounts. The constructed monetary gain-loss task scenario reads:

“You are a project manager at an investment firm. Your team’s primary work is managing investment projects to generate profits for the company. Recently, a member from another project team resigned, and their project was transferred to your team. As the project requires immediate completion, the company will assign one of your subordinates to oversee it when transferring the project. The project may generate profits or losses, with amounts expressed in gold coins.”

The experiment first required participants to rate “the relevance of team members’ tasks to money” (1 = not at all relevant, 7 = extremely relevant) and their agreement that “investment tasks may generate gains or losses” (1 = strongly disagree, 7 = strongly agree) to verify whether investment tasks could serve as monetary gain-loss tasks. Second, the experiment set 10-50 gold coins as low monetary gain-loss amounts and 100-500 gold coins as high amounts. Participants rated their perceived magnitude of these gain-loss amounts (e.g., “If the project gains or loses 10-50 gold coins, how large is the gain or loss?” 1 = very small, 7 = very large) to confirm whether participants could perceive differences between the two magnitude levels. Finally, participants rated their agreement that “team members need to assume certain monetary responsibilities in their tasks” to assess whether participants perceived that work agents needed to bear monetary responsibility.

(2) Moral Gain-Loss Work Scenario

Liu (2014) found that people tend to view environmental protection behaviors as moral and environmental pollution behaviors as immoral. Based on this, this study uses environmental protection and pollution-related work (referred to as environmental work) as the moral gain-loss task. Similar to monetary scenario construction, moral gain-loss work considers moral gain-loss magnitude, expressed in points. To avoid influencing participant ratings, the scenario avoids using the term “moral.” The constructed moral gain-loss task scenario reads:

“You are a department manager at a chemical plant. Your department’s primary

work is handling production waste gas. There are two methods: preliminary treatment before emission, which takes less time but affects the surrounding environment; or deep treatment before emission, which takes longer but avoids environmental pollution. The first method incurs public condemnation, while the second receives public praise. The degree of condemnation and praise is expressed in points. Recently, factory production has surged, and waste gas storage is near capacity, requiring urgent processing. The plant will assign one of your subordinates to handle this work. The work may bring moral praise or condemnation.”

The experiment first required participants to rate “the relevance of department members’ tasks to morality” (1 = not at all relevant, 7 = extremely relevant) and their agreement that “work tasks may generate moral gains or losses” (1 = strongly disagree, 7 = strongly agree) to verify whether environmental work could serve as a moral gain-loss task. Second, the experiment set low moral gain-loss amounts at 10-50 points and high amounts at 100-500 points. Participants rated their perceived magnitude of these amounts (e.g., “If the work increases or decreases moral points by 10-50, how large is the increase or decrease?” 1 = very small, 7 = very large) to confirm whether participants could perceive differences between the two magnitude levels. Finally, participants rated their agreement that “department members need to assume certain moral responsibilities in their tasks” to assess whether participants perceived that work agents needed to bear moral responsibility.

2.3 Results

For the monetary gain-loss scenario, one-sample t-tests (compared to the midpoint of 4) showed that participants perceived investment work as related to money ($M = 5.89$), $t(84) = 54.98$, $p < 0.001$, Cohen’s $d = 11.998$. Participants also agreed that investment work involved monetary gains and losses ($M = 6.44$), $t(84) = 25.55$, $p < 0.001$, Cohen’s $d = 5.574$. These results confirm that investment work can serve as a monetary gain-loss task. Second, paired-sample t-tests showed that participants perceived larger gain-loss magnitudes for 100-500 gold coins ($M = 5.52$) than for 10-50 gold coins ($M = 3.26$), $t(84) = 11.05$, $p < 0.001$, Cohen’s $d = 2.412$, confirming that participants could perceive differences between monetary gain-loss magnitudes. Finally, one-sample t-tests (compared to the midpoint of 4) also showed that participants perceived that work agents needed to assume monetary responsibility in investment tasks ($M = 5.48$), $t(84) = 9.57$, $p < 0.001$, Cohen’s $d = 2.089$.

For the moral gain-loss scenario, one-sample t-tests (compared to the midpoint of 4) showed that participants perceived environmental work as related to morality ($M = 6.20$), $t(84) = 66.75$, $p < 0.001$, Cohen’s $d = 14.566$. Participants also agreed that environmental work involved moral gains and losses ($M = 6.20$), $t(84) = 22.28$, $p < 0.001$, Cohen’s $d = 4.863$. These results confirm that environmental work can serve as a moral gain-loss task. Second, paired-sample t-tests showed that participants perceived larger gain-loss magnitudes for 100-

500 points ($M = 5.73$) than for 10-50 points ($M = 3.20$), $t(84) = 12.53$, $p < 0.001$, Cohen's $d = 2.733$, confirming that participants could perceive differences between moral gain-loss magnitudes. Finally, one-sample t -tests (compared to the midpoint of 4) also showed that participants perceived that work agents needed to assume moral responsibility in investment tasks ($M = 6.15$), $t(84) = 24.09$, $p < 0.001$, Cohen's $d = 5.258$.

3. Experiment 1: Work Allocation Acceptance in Monetary Context

3.1 Purpose

This experiment uses contextualized decision-making tasks to reveal differences in work allocation acceptance between humans and intelligent robots in monetary gain-loss scenarios and explores the mechanisms underlying acceptance based on mind perception theory.

3.2 Method

3.2.1 Experimental Paradigm This experiment employs a contextualized decision-making task paradigm. Participants first read the following scenario:

“You are a project manager at an investment firm. Your project team includes two human members and two intelligent robot members. The two human members include one male and one female. The robot members include one male-featured robot and one female-featured robot. The robots not only have gender features in appearance but also, through machine learning, have developed work patterns and characteristics corresponding to gender. The four members have no differences in past work performance. The department's primary work is managing investment projects to generate profits for the company. Recently, a member from another project team resigned, and their project was transferred to your team. As the project requires immediate completion, the company will assign one of your subordinates to oversee it when transferring the project. The project may generate profits or losses. Project completion is linked to team performance. The company now seeks your opinion on who in the team should undertake this project.”

Participants then decided whether to accept the company's decision to allocate work to humans or intelligent robots under various monetary gain-loss conditions.

The experiment used a 2 (agent type: human, intelligent robot) $\times 2$ (monetary gain-loss: loss, gain) $\times 2$ (gain-loss magnitude: high, low) within-subjects design. Human and robot icons represented work agents. Research indicates gender differences in risk tolerance and financial knowledge may lead to differential acceptance of males and females in monetary investment work (Bannier & Neubert, 2016). Therefore, this experiment controlled for gender, examining whether acceptance differences exist for male versus female work allocation and

whether these differences extend to intelligent robots. In Chinese culture, blue is more associated with males and pink with females (汪群, 2013). Thus, this experiment used color to distinguish gender: blue for males and pink for females (see [Figure 1: see original paper]).

Monetary gain-loss was represented by project profit/loss: “-” indicated expected project loss, “+” indicated expected project profit. Gain-loss magnitude was represented by gold coin amounts. Based on preliminary experiment results, the experiment set “100, 200, 300, 400, 500 gold coins” as high magnitude and “10, 20, 30, 40, 50 gold coins” as low magnitude. The dependent variable was work allocation acceptance, measured by the proportion of “accept” choices.

3.2.3 Participants Sample size was estimated using G*Power. For the experimental design, with significance level $\alpha = 0.05$ and medium effect size ($f = 0.25$), a minimum total sample of 15 participants was required to achieve 95% statistical power. One hundred employee participants were randomly recruited, including 43 females. Participants came from various industries including manufacturing, service, and retail. Ages ranged from 25 to 42 years ($M = 28.40$, $SD = 2.75$). All participants were right-handed with normal color vision.

3.2.4 Procedure The experimental program was developed using E-Prime 2.0 software. Participants sat approximately 60 cm from the computer screen. After reading instructions, participants were introduced to the background of humans and intelligent robots coexisting in organizations. They were also informed that intelligent robots could learn human work patterns and characteristics through machine learning (including supervised and reinforcement learning) and had developed independent work capabilities for certain tasks (see Appendix 1). This background introduction helped participants understand and familiarize themselves with intelligent robots in organizational settings.

Participants then read the experimental scenario. After understanding it, they were shown images of the four team members (see [Figure 1: see original paper]) and required to distinguish them. Participants completed a test on identifying the work agents represented by the four members, requiring 100% accuracy before proceeding. After familiarizing themselves with the scenario and members, participants viewed an example of the response interface (see [Figure 2: see original paper]). They used a mouse to click “agree” or “disagree” to make selections and clicked “confirm” to finalize their choice. After familiarizing themselves with the decision interface and response mode, participants began practice trials consisting of 10 trials. If participants failed to familiarize themselves after 10 trials, they could restart practice. Once fully familiar, the formal experiment began.

The formal experiment included 4 blocks, each corresponding to one work agent, presented in random order. Each block contained 40 trials, with 20 trials per round across two rounds. In each round, participants made choices under four conditions: high gain, low gain, high loss, and low loss. Each condition contained 5 trials. In each trial, an 800 ms fixation point appeared first, followed by the

response interface. After participants made their selection, an 800 ms blank screen appeared (see [Figure 2a: see original paper]).

After each block, participants rated the work agent's mind capabilities using Ward et al.'s (2013) scale, which includes agency and experience dimensions (see Appendix 3 for scale; Appendix 4 for reliability coefficients). Each dimension contains 7 items rated on a 7-point scale (1 = strongly disagree, 7 = strongly agree). Example items include "The agent can control its behavior" (agency) and "The robot can experience emotions" (experience). Scores for each dimension were summed, with higher scores indicating stronger capabilities. Finally, participants evaluated the work agent's monetary responsibility in monetary tasks using a percentage scale, responding to "How much responsibility can the work agent assume in this work?" (0 = cannot assume at all, 100 = can completely assume). Participants rested for 1 minute after each block.

3.3 Results

Work allocation acceptance was calculated as the proportion of "agree" choices in each condition. Data were analyzed using SPSS 23.0.

3.3.1 Work Allocation Acceptance To examine whether gender features of humans and robots influenced work allocation acceptance, a 2 (gender: male, female) $\times$ 2 (agent type: human, robot) repeated-measures ANOVA was conducted. Results showed that gender features did not affect work allocation acceptance for either human or robot agents (see). Therefore, male and female agents (both human and robot) were combined for analysis.

A 2 (agent type: human, robot) $\times$ 2 (monetary gain-loss: loss, gain) $\times$ 2 (gain-loss magnitude: high, low) repeated-measures ANOVA examined effects on work allocation acceptance. Results showed significant main effects of agent type, $F(1, 99) = 5.96$, $p = 0.016$, $\eta^2 = 0.057$; monetary gain-loss, $F(1, 99) = 673.88$, $p < 0.001$, $\eta^2 = 0.872$; and gain-loss magnitude, $F(1, 99) = 13.14$, $p < 0.001$, $\eta^2 = 0.117$. The interaction between agent type and monetary gain-loss was significant, $F(1, 99) = 7.91$, $p = 0.006$, $\eta^2 = 0.074$. The interaction between monetary gain-loss and gain-loss magnitude was significant, $F(1, 99) = 23.65$, $p < 0.001$, $\eta^2 = 0.193$. The interaction between agent type and gain-loss magnitude was not significant, $F(1, 99) = 0.01$, $p = 0.907$, $\eta^2 = 0.001$. The three-way interaction was not significant, $F(1, 99) = 0.43$, $p = 0.514$, $\eta^2 = 0.004$.

Simple effects analysis of the agent type $\times$ monetary gain-loss interaction showed that under loss conditions, acceptance of allocating work to humans ($M = 0.12$, $SD = 0.15$) was significantly lower than to robots ($M = 0.20$, $SD = 0.21$), $F(1, 99) = 12.02$, $p = 0.001$. Under gain conditions, acceptance of allocating work to humans ($M = 0.88$, $SD = 0.16$) did not differ significantly from robots ($M = 0.86$, $SD = 0.19$), $F(1, 99) = 1.51$, $p = 0.222$ (see [Figure 3a: see original paper]). Simple effects analysis of the monetary gain-loss $\times$ gain-loss magnitude interaction showed no significant difference in acceptance between high-gain (M

= 0.88, SD = 0.15) and low-gain work ($M = 0.87$, SD = 0.15), $F(1, 99) = 1.27$, $p = 0.262$. However, acceptance of low-loss work ($M = 0.18$, SD = 0.16) was significantly higher than high-loss work ($M = 0.14$, SD = 0.15), $F(1, 99) = 36.43$, $p < 0.001$ (see [Figure 3b: see original paper]).

3.3.2 Mind Perception To examine differences in perceived mind capabilities across agents, a one-way repeated-measures ANOVA was conducted with agent type as the independent variable and agency and experience as dependent variables. Results showed that human agency ($M = 40.78$, SD = 4.46) was significantly higher than robot agency ($M = 30.03$, SD = 7.61), $F(1, 99) = 143.89$, $p < 0.001$, $\eta^2 = 0.592$. Human experience ($M = 42.36$, SD = 5.61) was also significantly higher than robot experience ($M = 17.81$, SD = 9.24), $F(1, 99) = 340.71$, $p < 0.001$, $\eta^2 = 0.775$ (see [Figure 4: see original paper]). These results indicate that participants perceived differences in mind capabilities between humans and intelligent robots, with humans possessing higher mind capabilities.

3.3.3 Monetary Responsibility To explore differences in monetary responsibility across agents, a one-way repeated-measures ANOVA was conducted with agent type as the independent variable and monetary responsibility as the dependent variable. Results showed that humans ($M = 67.23$, SD = 18.77) could assume greater monetary responsibility than robots ($M = 51.42$, SD = 23.32), $F(1, 99) = 71.50$, $p < 0.001$, $\eta^2 = 0.419$.

3.3.4 Moderated Mediation Analysis To reveal the mechanism underlying acceptance of monetary work allocation, a mediation model was fitted using the bootstrap method (Hayes, 2022) with agent type as the independent variable (1 = robot, 2 = human), agency, experience, and monetary responsibility as mediators, and work allocation acceptance as the dependent variable. Since ANOVA results showed an interaction between agent type and monetary gain-loss, separate models were fitted for gain and loss conditions. Model 80 was selected with 5,000 iterations and 95% confidence intervals.

Under gain conditions, results showed that compared to robots, humans had higher agency and experience, both positively influencing monetary responsibility, which in turn positively influenced work allocation acceptance. The relative chain mediation effects of agency→monetary responsibility (indirect effect = 0.02, SE = 0.01, 95% CI = [0.01, 0.03]) and experience→monetary responsibility (indirect effect = 0.03, SE = 0.01, 95% CI = [0.01, 0.05]) were significant. After including mediators, the direct effect of agent type on work allocation acceptance remained significant (direct effect = 0.27, SE = 0.06, 95% CI = [0.15, 0.38]) (see [Figure 5a: see original paper]).

Under loss conditions, results showed that compared to robots, humans had higher agency and experience, both positively influencing monetary responsibility, but monetary responsibility negatively influenced work allocation acceptance. The relative chain mediation effects of agency→monetary responsibility

ity (indirect effect = -0.02, SE = 0.01, 95% CI = [-0.03, -0.01]) and experience→monetary responsibility (indirect effect = -0.04, SE = 0.01, 95% CI = [-0.07, -0.01]) were significant. After including mediators, the direct effect of agent type on work allocation acceptance was not significant (direct effect = 0.03, SE = 0.04, 95% CI = [-0.04, 0.10]) (see [Figure 5b: see original paper]).

These results indicate that monetary gain-loss moderates the effect of monetary responsibility on work allocation acceptance. Using Model 87, the moderating effect of monetary gain-loss was tested and found significant, $\beta = 0.009$, SE = 0.001, 95% CI = [0.007, 0.010]. This shows that the mediation effect of agent type on work allocation acceptance through mind capability and monetary responsibility is moderated by monetary gain-loss.

3.4 Discussion

This experiment found that humans' mind capabilities were significantly higher than robots', consistent with previous research (Gray et al., 2007; Huebner, 2010). Higher mind capabilities also enabled humans to assume greater monetary responsibility than intelligent robots. However, the effect of monetary responsibility on work allocation acceptance was moderated by monetary gain-loss: people more readily accepted allocating monetary gain work to humans but monetary loss work to robots. This indicates that people do not want humans to bear responsibility for monetary losses, showing a tendency to shift responsibility to robots. This occurs because robots' lower mind capabilities mean that assigning them monetary responsibility in loss situations does not elicit strong criticism.

These results show that mind capability influences monetary responsibility and work allocation acceptance. Would similar results emerge if the monetary context were replaced with a moral context? Experiment 2 explores differences in work allocation acceptance between humans and robots in moral contexts.

4. Experiment 2: Work Allocation Acceptance in Moral Context

4.1 Purpose

This experiment explores differences in work allocation acceptance between humans and robots in moral gain-loss scenarios and reveals the mechanism underlying acceptance based on mind perception theory.

4.2 Method

4.2.1 Experimental Paradigm This experiment uses the same contextualized decision-making task paradigm as Experiment 1, only replacing the scenario with a moral gain-loss context:

"You are a department manager at a chemical plant. Your department includes

two human members and two robot members. The two human members include one male and one female. The robot members include one male-featured robot and one female-featured robot. The robots not only have gender features in appearance but also, through machine learning, have developed work patterns and characteristics corresponding to gender. The four members have no differences in past work performance. The department's primary work is handling production waste gas. There are two methods: preliminary treatment before emission, which takes less time but affects the surrounding environment; or deep treatment before emission, which takes longer but avoids environmental pollution. The first method incurs public moral condemnation, while the second receives public moral praise. Recently, factory production has surged, and waste gas storage is near capacity, requiring urgent processing. The plant will assign one of your subordinates to handle this work. Work completion is linked to department performance. The work may bring moral praise or condemnation. The plant now seeks your opinion on who in the department should undertake this work."

Participants decided whether to accept allocating work to humans or intelligent robots under various moral gain-loss conditions. Moral gain-loss magnitude was represented by "moral value." Instructions informed participants that "moral value" refers to the public's moral evaluation scores for "preliminary treatment" and "deep treatment" methods, obtained through public surveys.

The experiment used a 2 (agent type: human, robot) $\times$ 2 (moral gain-loss: loss, gain) $\times$ 2 (gain-loss magnitude: high, low) within-subjects design. Meta-analyses indicate gender differences in moral judgment abilities (Atari et al., 2020), which may lead to differential acceptance of males and females in moral work. Therefore, this experiment also controlled for gender. Public moral praise represented moral gain (denoted by "+"), while moral condemnation represented moral loss (denoted by "-"). Gain-loss magnitude was represented by "moral value," with "100, 200, 300, 400, 500" as high magnitude and "10, 20, 30, 40, 50" as low magnitude. The dependent variable was work allocation acceptance.

4.2.3 Participants Sample size was estimated using G*Power. For the experimental design, with significance level $\alpha = 0.05$ and medium effect size ($f = 0.25$), a minimum total sample of 15 participants was required to achieve 95% statistical power. One hundred employee participants were randomly recruited, including 48 females. Participants came from various industries including manufacturing, service, and retail. Ages ranged from 25 to 38 years ($M = 28.23$, $SD = 2.07$). All participants were right-handed with normal color vision and received compensation upon completion.

4.2.4 Materials and Procedure The procedure was identical to Experiment 1, except that after each block, participants rated the work agent's moral responsibility in moral tasks using a percentage scale, responding to "How much responsibility can the work agent assume in this work?" (0 = cannot assume at

all, 100 = can completely assume).

4.3 Results

4.3.1 Work Allocation Acceptance A 2 (gender: male, female) $\times$ 2 (agent type: human, robot) repeated-measures ANOVA examined whether gender features influenced work allocation acceptance in the moral context. Results showed that gender features did not affect acceptance for either human or robot agents (see). Therefore, male and female agents (both human and robot) were combined for analysis.

A 2 (agent type: human, robot) $\times$ 2 (moral gain-loss: loss, gain) $\times$ 2 (gain-loss magnitude: high, low) repeated-measures ANOVA examined effects on work allocation acceptance. Results showed significant main effects of agent type, $F(1, 99) = 5.36$, $p = 0.023$, $\eta^2 = 0.051$; moral gain-loss, $F(1, 99) = 293.69$, $p < 0.001$, $\eta^2 = 0.748$; and gain-loss magnitude, $F(1, 99) = 43.09$, $p < 0.001$, $\eta^2 = 0.303$. The interaction between agent type and moral gain-loss was significant, $F(1, 99) = 19.24$, $p < 0.001$, $\eta^2 = 0.163$. The interaction between moral gain-loss and gain-loss magnitude was significant, $F(1, 99) = 77.61$, $p < 0.001$, $\eta^2 = 0.439$. The interaction between agent type and gain-loss magnitude was not significant, $F(1, 99) = 0.35$, $p = 0.555$, $\eta^2 = 0.004$. The three-way interaction was not significant, $F(1, 99) = 0.407$, $p = 0.525$, $\eta^2 = 0.004$.

Simple effects analysis of the agent type $\times$ moral gain-loss interaction showed that under moral loss conditions, acceptance of allocating work to humans ($M = 0.29$, $SD = 0.25$) was significantly lower than to robots ($M = 0.43$, $SD = 0.24$), $F(1, 99) = 18.83$, $p < 0.001$. Under moral gain conditions, acceptance of allocating work to humans ($M = 0.91$, $SD = 0.19$) was significantly higher than to robots ($M = 0.85$, $SD = 0.27$), $F(1, 99) = 6.45$, $p = 0.013$ (see [Figure 6a: see original paper]). Simple effects analysis of the moral gain-loss $\times$ gain-loss magnitude interaction showed that acceptance of high moral loss work ($M = 0.28$, $SD = 0.20$) was significantly lower than low moral loss work ($M = 0.43$, $SD = 0.21$), $F(1, 99) = 80.99$, $p < 0.001$. Acceptance of high moral gain work ($M = 0.89$, $SD = 0.21$) was significantly higher than low moral gain work ($M = 0.87$, $SD = 0.20$), $F(1, 99) = 4.43$, $p = 0.038$ (see [Figure 6b: see original paper]).

4.3.2 Mind Perception A one-way repeated-measures ANOVA with agent type as the independent variable and agency and experience as dependent variables revealed differences in perceived mind capabilities. Human agency ($M = 41.81$, $SD = 4.21$) was significantly higher than robot agency ($M = 29.11$, $SD = 7.79$), $F(1, 99) = 225.61$, $p < 0.001$, $\eta^2 = 0.695$. Human experience ($M = 42.04$, $SD = 5.33$) was also significantly higher than robot experience ($M = 16.33$, $SD = 7.36$), $F(1, 99) = 614.08$, $p < 0.001$, $\eta^2 = 0.861$ (see [Figure 7: see original paper]). These results indicate that participants perceived differences in mind capabilities between humans and intelligent robots, with humans possessing higher mind capabilities.

4.3.3 Moral Responsibility To explore differences in moral responsibility across agents, a one-way repeated-measures ANOVA was conducted with agent type as the independent variable and moral responsibility as the dependent variable. Results showed that humans ($M = 78.71$, $SD = 14.46$) could assume greater moral responsibility than robots ($M = 62.84$, $SD = 19.19$), $F(1, 99) = 64.94$, $p < 0.001$, $\eta^2 = 0.396$.

4.3.4 Moderated Mediation Analysis To reveal the psychological mechanism underlying acceptance of moral work allocation, a mediation model was fitted using the bootstrap method (Hayes, 2022) with agent type as the independent variable (1 = robot, 2 = human), agency, experience, and moral responsibility as mediators, and work allocation acceptance as the dependent variable. Since ANOVA results showed an interaction between agent type and moral gain-loss, separate models were fitted for gain and loss conditions. Model 80 was selected with 5,000 iterations and 95% confidence intervals.

Under gain conditions, results showed that compared to robots, humans had higher agency and experience, both positively influencing moral responsibility, which in turn positively influenced work allocation acceptance. The relative chain mediation effects of agency→moral responsibility (indirect effect = 0.04, $SE = 0.01$, 95% CI = [0.02, 0.07]) and experience→moral responsibility (indirect effect = 0.05, $SE = 0.02$, 95% CI = [0.01, 0.10]) were significant. After including mediators, the direct effect of agent type on work allocation acceptance was not significant (direct effect = 0.09, $SE = 0.05$, 95% CI = [-0.01, 0.19]) (see [Figure 8a: see original paper]).

Under loss conditions, results showed that compared to robots, humans had higher agency and experience, both positively influencing moral responsibility, but moral responsibility negatively influenced work allocation acceptance. The relative chain mediation effects of agency→moral responsibility (indirect effect = -0.05, $SE = 0.02$, 95% CI = [-0.08, -0.02]) and experience→moral responsibility (indirect effect = -0.06, $SE = 0.01$, 95% CI = [-0.11, -0.01]) were significant. After including mediators, the direct effect of agent type on work allocation acceptance was not significant (direct effect = -0.01, $SE = 0.06$, 95% CI = [-0.12, 0.11]) (see [Figure 8b: see original paper]).

These results suggest that moral gain-loss moderates the effect of moral responsibility on work allocation acceptance. Using Model 87, the moderating effect of moral gain-loss was tested and found significant, $\beta = 0.010$, $SE = 0.001$, 95% CI = [0.008, 0.012]. This shows that the mediation effect of agent type on work allocation acceptance through mind capability and moral responsibility is moderated by moral gain-loss.

4.4 Discussion

This experiment also found that work agents' mind capabilities positively influenced moral responsibility. Humans' higher mind capabilities enabled them to

assume greater moral responsibility than robots (Bigman & Gray, 2018). The effect of moral responsibility on moral work allocation acceptance was moderated by moral gain-loss: individuals more readily accepted allocating morally loss-related work to robots. This may relate to moral disengagement (Bandura et al., 1996; Paschalidis & Chen, 2022). When moral loss exists, humans wish to shift this responsibility to robots. Because robots have low mind capabilities and lack moral principles, assigning them moral responsibility does not elicit strong criticism or condemnation. However, if higher mind-capability humans assume such responsibility, condemnation would be severe (Awad et al., 2020; Young & Monroe, 2019). Therefore, having robots assume moral loss better protects the department's reputation and image. Under moral gain conditions, where moral responsibility shifting is not an issue, people more readily accept allocating work to humans.

Humans possess higher mind capabilities than robots. However, robots have strong data analysis and storage abilities. Human-robot collaborative work modes can complement each other's strengths. If humans and robots form collaborative teams, will the human-robot team's mind capabilities improve, enabling it to assume more work tasks? Additionally, what collective mind capabilities will robot-robot and human-human collaborative teams exhibit, and what work tasks can they assume? Experiments 3 and 4 explore these questions in monetary and moral gain-loss tasks.

5. Experiment 3: Work Allocation Acceptance Among Collaborative Teams in Monetary Context

5.1 Purpose

This experiment explores whether acceptance of work allocation differs across collaborative teams in monetary gain-loss scenarios and examines the psychological mechanism underlying acceptance of collaborative team allocation based on collective mind.

5.2 Method

5.2.1 Experimental Paradigm This experiment continues using the contextualized decision-making task paradigm, replacing individual work agents with collaborative teams. The scenario reads:

"You are a project manager at an investment firm. Your project team includes three collaborative work teams: 'robot-robot,' 'human-human,' and 'human-robot.' The three collaborative teams have no differences in past work performance. The department's primary work is managing investment projects to generate profits for the company. Recently, a member from another project team resigned, and their project was transferred to your project team. As the project requires immediate completion, the company will assign one collaborative team to oversee it when transferring the project. Project completion is

linked to team performance. The project may generate gains or losses. The company now seeks your opinion on which team should undertake this project.”

The experiment used a 3 (team type: human-human, human-robot, robot-robot) $\times$ 2 (monetary gain-loss: loss, gain) $\times$ 2 (gain-loss magnitude: high, low) within-subjects design. The three teams were presented pictorially (see [Figure 9: see original paper]). Monetary gain-loss was represented by project profit/loss: “-” indicated expected loss, “+” indicated expected gain. Gain-loss magnitude was represented by gold coin amounts: “100, 200, 300, 400, 500” as high magnitude and “10, 20, 30, 40, 50” as low magnitude. The dependent variable was work allocation acceptance.

5.2.3 Participants Sample size was estimated using G*Power. For the experimental design, with significance level $\alpha = 0.05$ and medium effect size ($f = 0.25$), a minimum sample of 18 participants was required to achieve 95% statistical power. One hundred employee participants were randomly recruited, including 49 females. Participants came from various industries including manufacturing, service, and retail. Ages ranged from 25 to 45 years ($M = 29.26$, $SD = 3.72$). All participants were right-handed with normal color vision.

5.2.4 Materials and Procedure The procedure was identical to Experiment 1, except the background introduction was replaced with information about the existence of human-human, human-robot, and robot-robot collaborative teams in organizational settings (see Appendix 2). After reading and understanding the scenario, participants were shown images of the three collaborative teams (see [Figure 9: see original paper]) and required to distinguish them. Participants completed a test on identifying the work agents represented by the three teams, requiring 100% accuracy before proceeding. With three collaborative teams, the experiment included 3 blocks, each corresponding to one team type presented in random order. Each block contained 40 trials distributed identically to Experiment 1. The monetary responsibility evaluation question was adapted to “How much responsibility can the collaborative team assume in this work?” (0 = cannot assume at all, 100 = can completely assume).

5.3 Results

5.3.1 Work Allocation Acceptance A 3 (team type: human-human, human-robot, robot-robot) $\times$ 2 (monetary gain-loss: loss, gain) $\times$ 2 (gain-loss magnitude: high, low) repeated-measures ANOVA examined effects on work allocation acceptance. Results showed significant main effects of team type, $F(2, 198) = 3.38$, $p = 0.036$, $\eta^2 = 0.033$; monetary gain-loss, $F(1, 99) = 1047.95$, $p < 0.001$, $\eta^2 = 0.914$; and gain-loss magnitude, $F(1, 99) = 11.52$, $p = 0.001$, $\eta^2 = 0.104$. The interaction between team type and monetary gain-loss was significant, $F(2, 198) = 4.45$, $p = 0.013$, $\eta^2 = 0.043$. The interaction between monetary gain-loss and gain-loss magnitude was significant, $F(1, 99) = 12.34$, $p = 0.001$, $\eta^2 = 0.111$. The interaction between team type and gain-loss

magnitude was not significant, $F(2, 198) = 2.58$, $p = 0.278$, $\eta^2 = 0.025$. The three-way interaction was not significant, $F(2, 198) = 0.26$, $p = 0.769$, $\eta^2 = 0.003$.

Simple effects analysis of the team type $\times$ monetary gain-loss interaction showed that under monetary loss conditions, acceptance differed significantly across the three teams ($M_{\text{human-human}} = 0.10$, $SD = 0.22$; $M_{\text{human-robot}} = 0.16$, $SD = 0.25$; $M_{\text{robot-robot}} = 0.18$, $SD = 0.26$), $F(2, 198) = 5.03$, $p = 0.017$. Paired-sample t -tests revealed significant differences between human-human and human-robot teams, $t(99) = 2.81$, $p = 0.006$, but not between human-robot and robot-robot teams, $t(99) = 0.36$, $p = 0.717$. Under monetary gain conditions, acceptance also differed significantly ($M_{\text{human-human}} = 0.94$, $SD = 0.13$; $M_{\text{human-robot}} = 0.92$, $SD = 0.11$; $M_{\text{robot-robot}} = 0.90$, $SD = 0.20$), $F(2, 198) = 3.51$, $p = 0.032$. Paired-sample t -tests showed no significant difference between human-robot and human-human teams, $t(99) = 0.26$, $p = 0.792$, but a significant difference between human-robot and robot-robot teams, $t(99) = 1.99$, $p = 0.039$ (see [Figure 10a: see original paper]).

Simple effects analysis of the monetary gain-loss $\times$ gain-loss magnitude interaction showed that acceptance of high-loss work ($M = 0.13$, $SD = 0.16$) was significantly lower than low-loss work ($M = 0.18$, $SD = 0.20$), $F(1, 99) = 16.89$, $p < 0.001$. Acceptance of high-gain work ($M = 0.93$, $SD = 0.13$) did not differ from low-gain work ($M = 0.93$, $SD = 0.12$), $F(1, 99) = 0.08$, $p = 0.777$ (see [Figure 10b: see original paper]).

5.3.2 Mind Perception To examine whether perceived mind capabilities differed across collaborative teams, a one-way repeated-measures ANOVA was conducted with team type as the independent variable and agency and experience as dependent variables. Results showed significant differences in agency across teams, $F(2, 198) = 188.55$, $p < 0.001$, $\eta^2 = 0.656$. Post-hoc LSD tests indicated that agency decreased sequentially: human-human ($M = 42.57$, $SD = 3.26$), human-robot ($M = 39.20$, $SD = 4.84$), and robot-robot ($M = 29.88$, $SD = 8.00$), all $ps < 0.001$. Significant differences also emerged in experience, $F(2, 198) = 893.70$, $p < 0.001$, $\eta^2 = 0.900$. Post-hoc LSD tests indicated that experience decreased sequentially: human-human ($M = 44.52$, $SD = 3.65$), human-robot ($M = 35.52$, $SD = 6.84$), and robot-robot ($M = 13.13$, $SD = 5.51$), all $ps < 0.001$ (see [Figure 11: see original paper]). These results indicate that participants perceived differences in mind capabilities among the three collaborative teams.

5.3.3 Monetary Responsibility To examine whether monetary responsibility differed across collaborative teams, a one-way repeated-measures ANOVA was conducted with team type as the independent variable and monetary responsibility as the dependent variable. Results showed a significant main effect of team type, $F(2, 198) = 137.89$, $p < 0.001$, $\eta^2 = 0.582$. Post-hoc LSD tests indicated that monetary responsibility decreased sequentially: human-human ($M = 82.72$, $SD = 14.23$), human-robot ($M = 68.52$, $SD = 19.32$), and robot-robot

($M = 55.98$, $SD = 26.70$), all $ps < 0.001$.

5.3.4 Moderated Mediation Analysis Since the interaction between team type and monetary gain-loss was significant, separate models were fitted for gain and loss conditions. Using the bootstrap method (Hayes, 2022), mediation models were fitted with team type as the independent variable, agency, experience, and monetary responsibility as mediators, and work allocation acceptance as the dependent variable. Model 80 was selected with 5,000 iterations and 95% confidence intervals.

(1) Gain Condition

With three collaborative teams, dummy variables (S1, S2) were created for team type coding. Using human-robot team as the reference (coded 0, 0), human-human team was coded (1, 0) and robot-robot team (0, 1). Under gain conditions, results showed that compared to human-robot teams, human-human teams had higher agency and experience, both positively influencing monetary responsibility, which in turn positively influenced work allocation acceptance. The relative chain mediation effects of agency→monetary responsibility (indirect effect = 0.01, $SE = 0.01$, 95% $CI = [0.01, 0.02]$) and experience→monetary responsibility (indirect effect = 0.01, $SE = 0.01$, 95% $CI = [0.01, 0.02]$) were significant. After including mediators, the direct effect of team type on work allocation acceptance was not significant (direct effect = -0.02, $SE = 0.02$, 95% $CI = [-0.06, 0.02]$) (see [Figure 12a: see original paper]).

Compared to human-robot teams, robot-robot teams had lower agency and experience, both positively influencing monetary responsibility, which in turn positively influenced work allocation acceptance. The relative chain mediation effects of agency→monetary responsibility (indirect effect = -0.01, $SE = 0.01$, 95% $CI = [-0.02, -0.01]$) and experience→monetary responsibility (indirect effect = -0.03, $SE = 0.01$, 95% $CI = [-0.06, -0.01]$) were significant. After including mediators, the direct effect of team type on work allocation acceptance was not significant (direct effect = -0.02, $SE = 0.03$, 95% $CI = [-0.09, 0.05]$) (see [Figure 12a: see original paper]).

Using robot-robot team as the reference (coded 0, 0), the relative mediation effect of human-human team (coded 1, 0) was also tested. Results showed that compared to robot-robot teams, human-human teams had higher agency and experience, both positively influencing monetary responsibility, which in turn positively influenced work allocation acceptance. The relative chain mediation effects of agency→monetary responsibility (indirect effect = 0.02, $SE = 0.01$, 95% $CI = [0.01, 0.03]$) and experience→monetary responsibility (indirect effect = 0.04, $SE = 0.01$, 95% $CI = [0.01, 0.07]$) were significant. After including mediators, the direct effect of team type on work allocation acceptance was not significant (direct effect = -0.01, $SE = 0.05$, 95% $CI = [-0.09, 0.09]$) (see [Figure 12b: see original paper]).

(2) Loss Condition

Using human-robot team as the reference (coded 0, 0), human-human team was coded (1, 0) and robot-robot team (0, 1). Under loss conditions, results showed that compared to human-robot teams, human-human teams had higher agency and experience, enabling them to assume greater monetary responsibility. However, under loss conditions, monetary responsibility negatively influenced work allocation acceptance. The relative chain mediation effects of agency→monetary responsibility (indirect effect = -0.01, SE = 0.01, 95% CI = [-0.02, -0.01]) and experience→monetary responsibility (indirect effect = -0.02, SE = 0.08, 95% CI = [-0.04, -0.01]) were significant. After including mediators, the direct effect of team type on work allocation acceptance was not significant (direct effect = -0.02, SE = 0.03, 95% CI = [-0.09, 0.04]) (see [Figure 13a: see original paper]).

Compared to human-robot teams, robot-robot teams had lower agency and experience, enabling them to assume only lower monetary responsibility, which negatively influenced work allocation acceptance. The relative chain mediation effects of agency→monetary responsibility (indirect effect = 0.02, SE = 0.01, 95% CI = [0.01, 0.03]) and experience→monetary responsibility (indirect effect = 0.05, SE = 0.02, 95% CI = [0.02, 0.09]) were significant. After including mediators, the direct effect of team type on work allocation acceptance was significant (direct effect = -0.12, SE = 0.05, 95% CI = [-0.22, -0.01]) (see [Figure 13a: see original paper]).

Using robot-robot team as the reference (coded 0, 0), the relative mediation effect of human-human team (coded 1, 0) was also tested. Results showed that compared to robot-robot teams, human-human teams had higher agency and experience, both positively predicting monetary responsibility, which negatively influenced work allocation acceptance. The relative chain mediation effects of agency→monetary responsibility (indirect effect = -0.02, SE = 0.01, 95% CI = [-0.04, -0.01]) and experience→monetary responsibility (indirect effect = -0.07, SE = 0.03, 95% CI = [-0.12, -0.02]) were significant. After including mediators, the direct effect of team type on work allocation acceptance was not significant (direct effect = 0.10, SE = 0.07, 95% CI = [-0.04, 0.24]) (see [Figure 13b: see original paper]).

Using Model 87, the moderating effect of monetary gain-loss was tested and found significant, $\beta = 0.004$, SE = 0.001, 95% CI = [0.003, 0.005]. This shows that the mediation effect of collaborative team type on work allocation acceptance through mind capability and monetary responsibility is moderated by monetary gain-loss.

5.4 Discussion

This experiment found that the three collaborative teams indeed possess different mind capabilities, which can be termed “collective mind.” Collective mind also comprises agency and experience dimensions. Perceptions of collective mind capability differed across collaborative teams: human-human highest, human-robot intermediate, and robot-robot lowest. Compared to Experiment 1, human-

robot collaborative teams showed significantly improved mind capabilities over individual robots (see Appendix 5), demonstrating that humans can compensate for robots' mind capability deficiencies (Haesevoets et al., 2021). Although human-robot teams' mind capabilities improved, they still lagged behind individual humans and human-human teams, particularly in experience.

Additionally, the experiment found that human-human collaboration could enhance individual mind capabilities, meaning human combinations improve mind capabilities. However, robot combinations did not show this enhancement effect and instead damaged mind capabilities. This may occur because humans have effective cooperation enhancement effects, while robots lack such effects.

Collective mind also influenced monetary responsibility perception, so the three teams differed in responsibility: human-human highest, human-robot intermediate, robot-robot lowest. The effect of monetary responsibility on work allocation acceptance was moderated by monetary gain-loss. People more readily allocated monetary gain work to human-human teams and loss work to robot-robot teams. This occurs because robot-robot teams have low mind capabilities, so when monetary loss exists, people prefer to shift responsibility to robots.

6. Experiment 4: Work Allocation Acceptance Among Collaborative Teams in Moral Context

6.1 Purpose

This experiment explores whether acceptance of work allocation differs across collaborative teams in moral gain-loss scenarios and examines the mechanism underlying acceptance based on collective mind.

6.2 Method

6.2.1 Experimental Paradigm This experiment uses the same contextualized decision-making task paradigm as Experiment 3, replacing the monetary scenario with a moral gain-loss scenario:

"You are a department manager at a chemical plant. Your department includes three collaborative work teams: 'robot-robot,' 'human-human,' and 'human-robot.' The three collaborative teams have no differences in past work performance. The department's primary work is handling production waste gas. There are two methods: preliminary treatment before emission, which takes less time but affects the surrounding environment within a certain range; or deep treatment before emission, which takes longer but avoids environmental pollution. The first method incurs public moral condemnation, while the second receives public moral praise. Recently, factory production has surged, and waste gas storage is near capacity, requiring urgent processing. The plant will assign one collaborative team to handle this work. Work completion is linked to department performance. The work may bring moral praise or condemnation. The plant now seeks your opinion on which team should undertake this work."

The experiment used a 3 (team type: human-human, human-robot, robot-robot) $\times$ 2 (moral gain-loss: loss, gain) $\times$ 2 (gain-loss magnitude: high, low) within-subjects design. The three teams were presented pictorially (see [Figure 10: see original paper]). Moral gain-loss was represented by moral praise (gain, denoted by “+”) and condemnation (loss, denoted by “-”). Gain-loss magnitude was represented by “moral value,” with “100, 200, 300, 400, 500” as high magnitude and “10, 20, 30, 40, 50” as low magnitude. The dependent variable was work allocation acceptance.

6.2.3 Participants Sample size was estimated using G*Power. For the experimental design, with significance level $\alpha = 0.05$ and medium effect size ($f = 0.25$), a minimum sample of 18 participants was required to achieve 95% statistical power. One hundred employee participants were randomly recruited, including 49 females. Participants came from various industries including manufacturing, service, and retail. Ages ranged from 23 to 42 years ($M = 28.62$, $SD = 2.88$). All participants were right-handed with normal color vision.

6.2.4 Materials and Procedure The procedure was identical to Experiment 3, except the scenario was replaced with the moral gain-loss context and the moral responsibility evaluation question was adapted to “How much responsibility can the collaborative team assume in this work?” (0 = cannot assume at all, 100 = can completely assume).

6.3 Results

6.3.1 Work Allocation Acceptance A 3 (team type: human-human, human-robot, robot-robot) $\times$ 2 (moral gain-loss: loss, gain) $\times$ 2 (gain-loss magnitude: high, low) repeated-measures ANOVA revealed significant main effects of team type, $F(2, 198) = 3.97$, $p = 0.020$, $\eta^2 = 0.039$; moral gain-loss, $F(1, 99) = 446.49$, $p < 0.001$, $\eta^2 = 0.819$; and gain-loss magnitude, $F(1, 99) = 37.83$, $p < 0.001$, $\eta^2 = 0.276$. The interaction between team type and moral gain-loss was significant, $F(2, 198) = 15.50$, $p < 0.001$, $\eta^2 = 0.135$. The interaction between moral gain-loss and gain-loss magnitude was significant, $F(1, 99) = 29.06$, $p < 0.001$, $\eta^2 = 0.227$. The interaction between team type and gain-loss magnitude was not significant, $F(2, 198) = 0.55$, $p = 0.578$, $\eta^2 = 0.006$. The three-way interaction was not significant, $F(2, 198) = 0.07$, $p = 0.939$, $\eta^2 = 0.001$.

Simple effects analysis of the team type $\times$ moral gain-loss interaction showed that under moral loss conditions, acceptance differed significantly across teams ($M_{\{\text{human}\}\text{-human}} = 0.17$, $SD = 0.26$; $M_{\{\text{human}\}\text{-robot}} = 0.19$, $SD = 0.24$; $M_{\{\text{robot}\}\text{-robot}} = 0.25$, $SD = 0.31$), $F(2, 198) = 3.92$, $p = 0.021$. Paired-sample t-tests showed no significant difference between human-human and human-robot teams, $t(99) = 0.81$, $p = 0.423$, but a significant difference between human-robot and robot-robot teams, $t(99) = 1.93$, $p = 0.046$. Under moral gain conditions, acceptance differed significantly ($M_{\{\text{human}\}\text{-human}} =$

0.94, SD = 0.12; $M_{\text{human-robot}} = 0.86$, SD = 0.19; $M_{\text{robot-robot}} = 0.77$, SD = 0.25), $F(2, 198) = 25.21$, $p < 0.001$. Paired-sample t-tests showed significant differences between human-human and human-robot teams, $t(99) = 4.05$, $p < 0.001$, and between human-robot and robot-robot teams, $t(99) = 3.42$, $p = 0.001$ (see [Figure 14a: see original paper]).

Simple effects analysis of the moral gain-loss $\times$ gain-loss magnitude interaction showed that acceptance of high moral loss work ($M = 0.14$, SD = 0.18) was significantly lower than low moral loss work ($M = 0.26$, SD = 0.26), $F(1, 99) = 40.30$, $p < 0.001$. Acceptance of high moral gain work ($M = 0.85$, SD = 0.14) did not differ from low moral gain work ($M = 0.86$, SD = 0.15), $F(1, 99) = 1.81$, $p = 0.182$ (see [Figure 14b: see original paper]).

6.3.2 Mind Perception To examine whether perceived mind capabilities differed across collaborative teams, a one-way repeated-measures ANOVA was conducted with team type as the independent variable and agency and experience as dependent variables. Results showed significant differences in agency, $F(2, 198) = 205.44$, $p < 0.001$, $\eta^2 = 0.675$. Post-hoc LSD tests indicated that agency decreased sequentially: human-human ($M = 42.13$, SD = 4.13), human-robot ($M = 37.05$, SD = 4.83), and robot-robot ($M = 26.70$, SD = 8.77), all $ps < 0.001$. Significant differences also emerged in experience, $F(2, 198) = 579.59$, $p < 0.001$, $\eta^2 = 0.854$. Post-hoc LSD tests indicated that experience decreased sequentially: human-human ($M = 43.46$, SD = 5.50), human-robot ($M = 34.38$, SD = 5.16), and robot-robot ($M = 13.97$, SD = 7.87), all $ps < 0.001$ (see [Figure 15: see original paper]). These results indicate that participants perceived differences in mind capabilities among the three collaborative teams.

6.3.3 Moral Responsibility To examine whether moral responsibility differed across collaborative teams, a one-way repeated-measures ANOVA was conducted with team type as the independent variable and moral responsibility as the dependent variable. Results showed a significant main effect of team type, $F(2, 198) = 85.39$, $p < 0.001$, $\eta^2 = 0.463$. Post-hoc LSD tests indicated that moral responsibility decreased sequentially: human-human ($M = 81.89$, SD = 15.95), human-robot ($M = 68.40$, SD = 16.47), and robot-robot ($M = 53.74$, SD = 25.90), all $ps < 0.001$.

6.3.4 Moderated Mediation Analysis Since the interaction between team type and moral gain-loss was significant, separate models were fitted for gain and loss conditions. Using the bootstrap method (Hayes, 2022), mediation models were fitted with team type as the independent variable, agency, experience, and moral responsibility as mediators, and work allocation acceptance as the dependent variable. Model 80 was selected with 5,000 iterations and 95% confidence intervals.

(1) Gain Condition

Using human-robot team as reference (coded 0, 0), human-human team was coded (1, 0) and robot-robot team (0, 1). Under gain conditions, results showed that compared to human-robot teams, human-human teams had higher agency and experience, both positively influencing moral responsibility, which in turn positively influenced work allocation acceptance. The relative chain mediation effects of agency→moral responsibility (indirect effect = 0.02, SE = 0.01, 95% CI = [0.01, 0.03]) and experience→moral responsibility (indirect effect = 0.03, SE = 0.01, 95% CI = [0.02, 0.04]) were significant. After including mediators, the direct effect of team type on work allocation acceptance was not significant (direct effect = 0.01, SE = 0.02, 95% CI = [-0.03, 0.06]).

Compared to human-robot teams, robot-robot teams had lower agency and experience, both positively influencing moral responsibility, which in turn positively influenced work allocation acceptance. The relative chain mediation effects of agency→moral responsibility (indirect effect = -0.04, SE = 0.01, 95% CI = [-0.06, -0.02]) and experience→moral responsibility (indirect effect = -0.08, SE = 0.02, 95% CI = [-0.11, -0.05]) were significant. After including mediators, the direct effect of team type on work allocation acceptance was not significant (direct effect = -0.01, SE = 0.03, 95% CI = [-0.07, 0.06]) (see [Figure 16a: see original paper]).

Using robot-robot team as reference (coded 0, 0), the relative mediation effect of human-human team (coded 1, 0) was also tested. Results showed that compared to robot-robot teams, human-human teams had higher agency and experience, both positively influencing moral responsibility, which in turn positively influenced work allocation acceptance. The relative chain mediation effects of agency→moral responsibility (indirect effect = 0.06, SE = 0.01, 95% CI = [0.03, 0.09]) and experience→moral responsibility (indirect effect = 0.11, SE = 0.02, 95% CI = [0.07, 0.16]) were significant. After including mediators, the direct effect of team type on work allocation acceptance was not significant (direct effect = 0.02, SE = 0.04, 95% CI = [-0.06, 0.10]) (see [Figure 16b: see original paper]).

(2) Loss Condition

Using human-robot team as reference (coded 0, 0), human-human team was coded (1, 0) and robot-robot team (0, 1). Under loss conditions, results showed that compared to human-robot teams, human-human teams had higher agency and experience, both positively influencing moral responsibility, but moral responsibility negatively influenced work allocation acceptance. The relative chain mediation effects of agency→moral responsibility (indirect effect = -0.02, SE = 0.01, 95% CI = [-0.03, -0.01]) and experience→moral responsibility (indirect effect = -0.03, SE = 0.01, 95% CI = [-0.05, -0.01]) were significant. After including mediators, the direct effect of team type on work allocation acceptance was not significant (direct effect = 0.06, SE = 0.03, 95% CI = [-0.01, 0.13]) (see [Figure 17a: see original paper]).

Compared to human-robot teams, robot-robot teams had lower agency and experience,

rience, both positively influencing moral responsibility, but moral responsibility negatively influenced work allocation acceptance. The relative chain mediation effects of agency→moral responsibility (indirect effect = 0.03, SE = 0.01, 95% CI = [0.01, 0.05]) and experience→moral responsibility (indirect effect = 0.06, SE = 0.02, 95% CI = [0.03, 0.10]) were significant. After including mediators, the direct effect of team type on work allocation acceptance was not significant (direct effect = -0.07, SE = 0.05, 95% CI = [-0.17, 0.03]) (see [Figure 17a: see original paper]).

Using robot-robot team as reference (coded 0, 0), the relative mediation effect of human-human team (coded 1, 0) was also tested. Results showed that compared to robot-robot teams, human-human teams had higher agency and experience, both positively predicting moral responsibility, but moral responsibility negatively influenced work allocation acceptance. The relative chain mediation effects of agency→moral responsibility (indirect effect = -0.05, SE = 0.02, 95% CI = [-0.08, -0.02]) and experience→moral responsibility (indirect effect = -0.09, SE = 0.03, 95% CI = [-0.15, -0.04]) were significant. After including mediators, the direct effect of team type on work allocation acceptance was significant (direct effect = 0.13, SE = 0.07, 95% CI = [0.01, 0.26]) (see [Figure 17b: see original paper]).

Using Model 87, the moderating effect of moral gain-loss was tested and found significant, $\beta = 0.008$, SE = 0.001, 95% CI = [0.007, 0.009]. This shows that the mediation effect of collaborative team type on work allocation acceptance through mind capability and moral responsibility is moderated by moral gain-loss.

6.4 Discussion

This experiment again verified that collaborative teams possess collective mind capabilities. Compared to Experiment 2, human-robot teams' mind capabilities also improved significantly over individual robots (see Appendix 6), reaffirming that humans can compensate for robots' mind capability deficiencies (Haesevoets et al., 2021). This experiment also found that human-human collaboration can enhance individual mind capabilities to some extent, but robot combinations lack this enhancement effect and instead diminish mind capabilities.

Collective mind capability also influences the moral responsibility collaborative teams can assume: higher mind capability enables greater moral responsibility. The effect of moral responsibility on moral work allocation acceptance is also moderated by moral gain-loss. People again more readily accepted allocating moral loss work to robot-robot teams and moral gain work to human-human teams. This result also stems from moral responsibility shifting (Bandura et al., 1996; Paschalidis & Chen, 2022). The mind capability deficits of robot-robot teams result in low moral responsibility, and shifting responsibility to them in moral loss conditions can avoid criticism or condemnation.

7. General Discussion

7.1 Individual and Collective Mind Capabilities

This study confirms that humans and intelligent robots possess different mind capabilities, with humans exceeding intelligent robots in both agency and experience dimensions, consistent with previous research (Bigman & Gray, 2018; Longoni et al., 2019; Niszczota & Kaszás, 2020). This study also discovered that not only individual work agents possess mind capabilities, but collaborative teams composed of different agents also possess collective mind capabilities. Human-human, human-robot, and robot-robot collaborative teams showed progressively decreasing collective mind, with robot-robot teams' mind level even falling below that of individual robots. This may occur because when two robots combine, they cannot cooperate as effectively as humans, leading to diminished mind capabilities.

Current research focuses heavily on human-robot collaboration effects (何贵兵 et al., 2022; Zheng et al., 2017). Intelligent robots possess high-speed computation, data processing, and automated task execution capabilities, while humans possess strong creativity, intuition, and emotional feeling abilities. Thus, humans and intelligent robots have complementary capabilities, and their collaboration should produce good results. This study indeed found that humans can compensate for intelligent robots' mind capability deficiencies. Their combination effectively improves intelligent robots' mind capabilities, particularly in the experience dimension. Therefore, when humans and robots form collaborative teams, improved mind capabilities enable the team to assume more work responsibilities and complete more task types, including moral tasks, expanding intelligent robots' application fields.

Although human-robot collaborative teams' mind capabilities improved significantly over individual robots, they still lagged far behind individual humans and human-human teams. This means that while humans can compensate for robots' deficiencies through collaboration, they do not surpass human mind capabilities. Unlike individual agents that directly exhibit mind capabilities, collective mind capabilities require full integration of both parties' minds to manifest (Burton et al., 2024). Therefore, to more substantially enhance human-robot collaborative mind capabilities, it may be necessary to improve humans' acceptance of and trust in robots, enabling full integration to produce synergistic effects of $1+1>2$ (Zhang et al., 2020; Kaplan et al., 2023).

7.2 Agent Differences in Work Allocation

This study found that because different individual (group) agents possess differentiated mind capabilities, the work responsibilities they can assume also differ, ultimately leading to different work allocations. Humans and human-human collaborative teams possess the highest mind capabilities and thus can assume the greatest responsibility; conversely, robots and robot-robot collaborative teams have the weakest mind capabilities and thus can assume the least responsibility.

People are more willing to allocate loss-related work to robots or robot-robot teams and gain-related work to humans or human-human (or human-robot) teams. This clearly demonstrates responsibility aversion in work allocation (Edelson et al., 2018)—shifting loss-responsibility work to robots rather than completing it themselves. From a psychological distance perspective (Trope & Liberman, 2010), humans have greater distance from robots than from other humans. Therefore, humans more easily “favor their own,” leading to responsibility shifting to robots in loss situations.

In fact, whether in monetary or moral tasks, under loss conditions, people always expect agents capable of assuming more responsibility to receive harsher criticism or punishment. Because they bear more responsibility, humans, to avoid severe criticism and condemnation, allocate behaviors likely to cause monetary or moral losses to robots, playing the role of “free riders” and thereby shifting responsibility (Gross et al., 2018). Current robots lack refusal and resistance consciousness, making this “buck-passing” phenomenon potentially more severe, requiring human vigilance.

7.3 Research Implications

With rapid AI technology development, intelligent robots are increasingly entering human work. Numerous studies have focused on intelligent robots’ application effects in organizations and their impacts on organizations and employees. When intelligent robots are capable enough to fully assume work tasks and reduce human workers’ task loads, this will inevitably increase acceptance of their work allocation. Therefore, improving intelligent robots’ application effects and positive impacts on organizational employees ultimately requires increasing acceptance of intelligent robots participating in human work. Based on mind perception theory, this study examined acceptance of intelligent robots working independently and collaboratively in monetary and moral task contexts, yielding theoretical and practical significance.

Theoretically, first, this study revealed human-robot differences in work allocation acceptance in monetary and moral contexts and explored the formation mechanisms based on mind capability and responsibility perception. People judge whether work agents can assume corresponding responsibilities based on their mind capabilities, thereby confirming whether agents can complete corresponding work. This provides reference value for constructing social division of labor theory in the digital intelligence era and clarifying intelligent robots’ role and position in social division of labor. Second, this study extended mind perception theory from the individual to the group level, discovering and verifying that various collaborative teams possess collective mind capabilities comprising agency and experience dimensions. Third, this study discovered and verified that humans have compensatory and enhancing effects on robots’ mind capabilities, and their collaboration can exhibit hybrid intelligence, enabling human-robot teams to assume greater work responsibilities and obtain more work allocations.

This study also has practical significance. First, social division of labor in the digital intelligence era must fully consider humans' and robots' mind capabilities to achieve better person-job and robot-job matching. Second, in work potentially involving losses, task allocation should be carefully managed to prevent humans from "buck-passing" to robots. Third, this study verified that human-robot combinations can improve robots' mind capabilities to some extent, thus enabling organic integration based on full consideration of both parties' mind capabilities can fully release hybrid intelligence effects.

7.4 Limitations and Future Directions

First, this study's contextualized decision-making task paradigm may limit generalizability. Future research could conduct studies in real-world scenarios where humans and robots can interact through language and emotions. Current research shows that robots' emotional expressions through facial expressions and language influence human attitudes and behaviors (Hsieh & Cross, 2022; Swiderska & Küster, 2020). Therefore, future research could explore how real human-robot interactive communication affects work allocation and acceptance.

Second, this study primarily based on person-job fit theory examined acceptance differences in work allocation to humans and intelligent robots from the perspective of work agents' mind capability-work fit. However, mind level only reflects one aspect of intelligent robots' capabilities; many intelligent robots also possess special capabilities. For example, surgical robots have more flexible and precise finger movements, making them more suitable for medical work. Therefore, only when robots' capabilities match industry characteristics can people accept their work allocation (Bankins, 2021; Dhar, 2015). Future research should examine acceptance of work allocation to robots with different capabilities considering various industries' technological characteristics. Additionally, as intelligent robots' capabilities continue iterative upgrading, this dynamic change may cause continuous dynamic changes in work allocation acceptance, which future research could reveal.

Furthermore, intelligent robots have different autonomy levels, such as low-autonomy robot assistance systems versus high-autonomy intelligent robot agents (何贵兵 et al., 2022). This autonomy level may also influence work allocation acceptance. For example, low-autonomy robots could be allocated highly structured and repetitive tasks, while higher-autonomy robots possess greater adaptive capabilities and can self-adjust according to environmental changes, making them suitable for less structured tasks. Therefore, future research could examine how intelligent robots' autonomy influences work allocation.

Third, this study did not consider human-robot trust issues in human-robot collaborative teams. Given that algorithm aversion may reduce trust between humans and robots (Dietvorst et al., 2015; Longoni et al., 2019; Mahmud et al., 2022), thereby affecting collaborative effectiveness and work allocation accep-

tance, future research should examine trust issues within collaborative teams and their effects on mind capabilities and work allocation.

Fourth, this study's human-robot collaborative teams did not consider relationships between humans and robots. In future work scenarios, humans may be robots' superiors, subordinates, or peers. Xie et al. (2021) noted that human-robot relationships affect collaborative effectiveness. Xiong et al. (2023) also found that when intelligent robots have high performance, human-robot partnership relationships better improve collaborative effects. Therefore, future research should examine how different human-robot relationship modes affect collaboration and which relationship mode better enhances human-robot teams' mind capabilities and work allocations.

8. Conclusions

This study yielded the following conclusions:

1. In both monetary and moral tasks, people are more inclined to accept allocating loss-related work to robots and gain-related work to humans.
2. Humans and robots possess differentiated mind capabilities, leading to differences in assumed responsibilities and ultimately resulting in different gain-loss work allocations.
3. When humans and intelligent robots form collaborative teams, they exhibit collective mind capabilities.
4. Human-human, human-robot, and robot-robot collaborative teams show progressively decreasing collective mind capabilities, resulting in progressively decreasing monetary and moral responsibilities and influencing monetary and moral work allocations.

References

- Atari, M., Lai, M.H., & Dehghani, M. (2020). Sex differences in moral judgments across 67 countries. *Proceedings of the Royal Society B*, 287, 20201201. <https://doi.org/10.1098/rspb.2020.1201>
- Awad, E., Levine, S., Kleiman-Weiner, M., Dsouza, S., Tenenbaum, J. B., Shariff, A., ...Rahwan, I. (2020). Drivers are blamed more than their automated cars when both make mistakes. *Nature Human Behaviour*, 4(2), 134–143. <https://doi.org/10.1038/s41562-019-0762-8>
- Bandura, A., Barbaranelli, C., Caprara, G. V., & Pastorelli, C. (1996). Mechanisms of moral disengagement in the exercise of moral agency. *Journal of Personality and Social Psychology*, 71(2), 364–374. <https://doi.org/10.1037/0022-3514.71.2.364>
- Bankins, S. (2021). The ethical use of artificial intelligence in human resource management: A decision-making framework. *Ethics and Information Technology*, 23, 841–854. <https://doi.org/10.1007/s10676-021-09619-6>

- Bannier, C.E., & Neubert, M. (2016). Gender differences in financial risk taking: The role of financial literacy and risk tolerance. *Economics Letters*, 145, 130–135. <https://doi.org/10.1016/j.econlet.2016.05.033>
- Bechara, A., & Damasio, A. R. (2005). The somatic marker hypothesis: A neural theory of economic decision. *Games and Economic Behavior*, 52(2), 336–372. <https://doi.org/10.1016/j.geb.2004.06.010>
- Bigman, Y. E., & Gray, K. (2018). People are averse to machines making moral decisions. *Cognition*, 181, 21–34. <http://doi.org/10.1016/j.cognition.2018.08.003>
- Bowen, J. T., & Morosan, C. (2018). Beware hospitality industry: the robots are coming. *Worldwide Hospitality and Tourism Themes*, 10(6), 726–733. <https://doi.org/10.1108/WHATT-07-2018-0045>
- Burks, S. V., Carpenter, J., Goette, L., & Rustichini, A. (2009). Cognitive skills affect economic preferences, strategic behavior, and job attachment. *Proceedings of the National Academy of Sciences*, 106(19), 7745–7750. <https://doi.org/10.1073/pnas.0812360106>
- Burton, J. W., Lopez-Lopez, E., Hechtlinger, S., Rahwan, Z., Aeschbach, S., Bakker, M. A., ... Hertwig, R. (2024). How large language models can reshape collective intelligence. *Nature Human Behaviour*, 8, 1643–1655. <https://doi.org/10.1038/s41562-024-01959-9>
- Damasio, A. R. (1994). *Descartes' Error: Emotion, reason, and the human brain*. New York: G. P. Putnam's Sons.
- Decety, J., & Cowell, J. M. (2014). The complex relation between morality and empathy. *Trends in Cognitive Sciences*, 18(7), 337–339. <https://doi.org/10.1016/j.tics.2014.04.008>
- Dhar, V. (2015). Should you trust your money to a robot? *Big Data*, 3(2), <https://doi.org/10.1089/big.2015.28999.vda>
- Dietvorst, B. J., Simmons, J. P., & Massey, C. (2015). Algorithm aversion: People erroneously avoid algorithms after seeing them err. *Journal of Experimental Psychology: General*, 144(1), <https://doi.org/10.1037/xge0000033>
- Edelson, M. G., Polania, R., Ruff, C. C., Fehr, E., & Hare, T. A. (2018). Computational and neurobiological foundations of leadership decisions. *Science*, 361(6401), eaat0036. <https://doi.org/10.1126/science.aat0036>
- Epstein, S. L. (2015). Wanted: Collaborative intelligence. *Artificial Intelligence*, 221, 36–45. <https://doi.org/10.1016/j.artint.2014.12.006>
- Gibney E. (2024). The AI revolution is coming to robots: how will it change them? *Nature*, 630(8015), 22–24. <https://doi.org/10.1038/d41586-024-01442-5>
- Gray, H. M., Gray, K., & Wegner, D. M. (2007). Dimensions of mind perception. *Science*, 315(5812), 619. <https://doi.org/10.1126/science.1134475>

- Gray, K., & Wegner, D. M. (2009). Moral typecasting: Divergent perceptions of moral agents and moral patients. *Journal of Personality and Social Psychology*, 96(3), 505–520. <https://doi.org/10.1037/a0013748>
- Gray, K., & Wegner, D. M. (2012). Feeling robots and human zombies: Mind perception and the uncanny valley. *Cognition*, 125(1), 125–130. <https://doi.org/10.1016/j.cognition.2012.06.007>
- Gray, K., Young L., & Waytz A. (2012). Mind perception is the essence of morality. *Psychological Inquiry*, 23(2), 101–124. <https://doi.org/10.1080/1047840X.2012.651387>
- Greene, J. D., Sommerville, R. B., Nystrom, L. E., Darley, J. M., & Cohen, J. D. (2001). An fMRI investigation of emotional engagement in moral judgment. *Science*, 293(5537), <https://doi.org/10.1126/science.1062872>
- Grinblatt, M., Keloharju, M., & Linnainmaa, J. T. (2011). IQ and stock market participation. *Journal of Finance*, 66(6), 2121–2164. <https://doi.org/10.1111/j.1540-6261.2011.01701.x>
- Gross, J., Leib, M., Offerman, T., & Shalvi, S. (2018). Ethical free riding: When honest people find dishonest partners. *Psychological Science*, 29(12), 1956–1968. <https://doi.org/10.1177/0956797618796480>
- Haesevoets, T., De Cremer, D., Dierckx, K., & Van Hiel, A. (2021). Human-machine collaboration in managerial decision making. *Computers in Human Behavior*, 119, 106730. <https://doi.org/10.1016/j.chb.2021.106730>
- Haidt, J. (2001). The emotional dog and its rational tail: A social intuitionist approach to moral judgment. *Psychological Review*, 108(4), 814–834. <https://doi.org/10.1037/0033-295X.108.4.814>
- Haidt, J., Koller, S. H., & Dias, M. G. (1993). Affect, culture, and morality, or is it wrong to eat your dog? *Journal of Personality and Social Psychology*, 65(4), 613–628. <https://doi.org/10.1037/0022-3514.65.4.613>
- Hayes, A. F. (2022). *Introduction to mediation, moderation, and conditional process analysis* (3rd Ed.). New York: The Guilford Press.
- He, G. B., Chen, C., He, Z. T., Cui, L., Lu, J. Q., Xuan, H. Z., & Lin, L. (2022). Human-agent collaborative decision-making in intelligent organizations: A perspective of human-agent inner compatibility. *Advances in Psychological Science*, 30(12), 2619–2627. doi: 10.3724/SP.J.1042.2022.02619
- Hsieh, T. Y., & Cross, E. S. (2022). People's dispositional cooperative tendencies towards robots are unaffected by robot's negative emotional displays in prisoner's dilemma games. *Cognition and Emotion*, 36(5), 995–1019. <https://doi.org/10.1080/02699931.2022.2054781>
- Hu, X. Y., Li, M. F., Wang, D. X., & Yu, F. (2024). Reactions to immoral AI decisions: The moral deficit effect and its underlying mechanism. *Chinese Science Bulletin*, 69(11), 1406–1416. doi: 10.1360/TB-2023-1094

- Huang, M.-H., & Rust, R. T. (2018). Artificial Intelligence in Service. *Journal of Service Research*, 21(2), 155–172. <https://doi.org/10.1177/1094670517752459>
- Huebner, B. (2010). Commonsense concepts of phenomenal consciousness: Does anyone care about functional zombies? *Phenomenology and the cognitive sciences*, 9, 133–155. <https://doi.org/10.1007/s11097-009-9126-6>
- Hutchins, E. (1995). *Cognition in the Wild*. MIT Press.
- Kahneman, D., & Tversky, A. (1979). Prospect theory: An analysis of decision under risk. *Econometrica*, 47(2), 263–291. <https://doi.org/10.2307/1914185>
- Kaplan, A. D., Kessler, T. T., Brill, J. C., & Hancock, P. A. (2023). Trust in artificial intelligence: Meta-analytic findings. *Human Factors*, 65(2), 337–359. <https://doi.org/10.1177/00187208211013988>
- Larkin, C., Otten, C. D., & Árvai, J. (2021). Paging Dr. JARVIS! Will people accept advice from artificial intelligence for consequential risk management decisions? *Journal of Risk Research*, 25(4), 407–422. <https://doi.org/10.1080/13669877.2021.1958047>
- Lerner, J. S., & Keltner, D. (2001). Fear, anger, and risk. *Journal of Personality and Social Psychology*, 81(1), 146–159. <https://doi.org/10.1037/0022-3514.81.1.146>
- Licklider, J. C. R. (1960). Man-Computer Symbiosis. *Ire Transactions on Human Factors in Electronics*. 1(1), 4–11. doi: 10.1109/THFE2.1960.4503259
- Liu, M. Z. (2022). Evolution, types and influence factors of human-machine collaboration modes. *Contemporary Manager*, Issue 3, 20–29.
- Liu, W. (2023). *Integrated human-machine intelligence: Beyond artificial intelligence*. Beijing: Tsinghua University Press.
- Liu, Y. (2014). The influence of income sources and profit means on risk decision making (master's thesis). Zhejiang University.
- Longoni, C., Bonezzi, A., & Morewedge, C. K. (2019). Resistance to Medical Artificial Intelligence. *Journal of Consumer Research*, 46(4), 629–650. <https://doi.org/10.1093/jcr/ucz013>
- Loewenstein, G. F., Weber, E. U., Hsee, C. K., & Welch, N. (2001). Risk as feelings. *Psychological Bulletin*, 127(2), 267–286. <https://doi.org/10.1037/0033-2909.127.2.267>
- Mahmud, H., Islam, A. K. M. N., Ahmed, S. I., & Smolander, K. (2022). What influences algorithmic decision-making? A systematic literature review on algorithm aversion. *Technological Forecasting and Social Change*, 175, 121390. <https://doi.org/10.1016/j.techfore.2021.121390>
- Malle, B. F. (2019). How many dimensions of mind perception really are there? In A. K. Goel, C. M. Seifert, & C. Freksa (Eds.), *Proceedings of the 41st Annual Meeting of the Cognitive Science Society* (pp. 2268–2274). Montreal.

- Markovitch, D. G., Stough, R. A., & Huang, D. (2024). Consumer reactions to chatbot versus human service: An investigation in the role of outcome valence and perceived empathy. *Journal of Retailing and Consumer Services*, 79, 103847. <https://doi.org/10.1016/j.jretconser.2024.103847>
- Mayer, R. C., Davis, J. H., & Schoorman, F. D. (1995). An integrative model of organizational trust. *Academy of Management Review*, 20(3), 709–734. <https://doi.org/10.5465/amr.1995.9508080335>
- Niszczoła, P., & Kaszás, D. (2020). Robo-investment aversion. *Plos One*, 15(9), e0239277. <https://doi.org/10.1371/journal.pone.0239277>
- Paschalidis, E., & Chen, H. (2022). Moral disengagement mechanisms in interactions of human drivers with autonomous vehicles: Validation of a new scale and relevance with personality, driving style and attitudes. *Transportation Research Traffic Psychology Behaviour*, <https://doi.org/10.1016/j.trf.2022.08.015>
- Pazhoohi, F., Gojamgunde, S., & Kingstone, A. (2023). Give me space: Sex, attractiveness, and mind perception as potential contributors to different comfort distances for humans and robots. *Journal of Environmental Psychology*, 90, 102088. <https://doi.org/10.1016/j.jenvp.2023.102088>
- Realyvásquez-Vargas, A., Cecilia Arredondo-Soto, K., Luis García-Alcaraz, J., Yail Márquez-Lobato, B., & Cruz-García, J. (2019). Introduction and configuration of a collaborative robot in an assembly task as a means to decrease occupational risks and increase efficiency in a manufacturing company. *Robotics and Computer-Integrated Manufacturing*, 57, 315–328. <https://doi.org/10.1016/j.rcim.2018.12.015>
- Ren, M., Chen, N., & Qiu, H. (2023). Human-machine collaborative decision-making: An evolutionary roadmap based on cognitive intelligence. *International Journal of Social Robotics*, 15(7), 1101–1114. <https://doi.org/10.1007/s12369-023-01020-1>
- Rigoni, D., Wilquin, H., Brass, M., & Burle, B. (2013). When errors do not matter: Weakening belief in intentional control impairs cognitive reaction to errors. *Cognition*, 127(2), <https://doi.org/10.1016/j.cognition.2013.01.009>
- Seeber, I., Bittner, E. A., Briggs, R. O., Vreede, T. D., Vreede, G. D., Elkins, A. C., ...Söllner, M. (2020). Machines as teammates: A research agenda on AI in team collaboration. *Information & Management*, 57(2), 103174. <https://doi.org/10.1016/j.im.2019.103174>
- Shoemaker, P. J. H., & Tetlock, P. E. (2017). Building a More Intelligent Enterprise. *MIT Sloan Management Review*, 58(3), 28–40.
- Siciliano, B., Khatib, O. (2016). Robotics and the Handbook. In: Siciliano, B., Khatib, O. (eds) *Springer Handbook of Robotics*. Springer, Cham. https://doi.org/10.1007/978-3-319-32552-1_1

- Silva, A., Correia Simões, A., & Blanc, R. (2024). Supporting decision-making of collaborative robot (cobot) adoption: The development of a framework. *Technological Forecasting and Social Change*, 204, 123406. <https://doi.org/10.1016/j.techfore.2024.123406>
- Swiderska, A., & Küster D. (2020). Robots as malevolent moral agents: Harmful behavior results in dehumanization, not anthropomorphism. *Cognitive Science*, 44(7), e12872. <https://doi.org/10.1111/cogs.12872>
- Szekeres, H., Halperin, E., Kende, A., & Saguy, T. (2019). The effect of moral loss and gain mindset on confronting racism. *Journal of Experimental Social Psychology*, 84, 103833. <https://doi.org/10.1016/j.jesp.2019.103833>
- Tangney, J. P., Stuewig, J., & Mashek, D. J. (2007). Moral emotions and moral behavior. *Annual Review of Psychology*, 58, 345–372. <https://doi.org/10.1146/annurev.psych.56.091103.070145>
- Tharp, M., Holtzman, N. S., & Eadeh, F. R. (2017). Mind perception and individual differences: A replication and extension. *Basic and Applied Social Psychology*, 39(1), 68–73. <https://doi.org/10.1080/01973533.2016.1256287>
- Trope, Y., & Liberman, N. (2010). Construal-Level theory of psychological distance. *Psychological Review*, 117(2), 440–463. <https://doi.org/10.1037/a0018963>
- van Woerkom, M., Bauwens, R., Gürbüz, S., & Brouwers, E. (2024). Enhancing person-job fit: Who needs a strengths-based leader their way? *Journal of Vocational Behavior*, <https://doi.org/10.1016/j.jvb.2024.104044>
- Vohs, K. D., & Schooler, J. W. (2008). The value of believing in free will: Encouraging a belief in determinism increases cheating. *Psychological Science*, 19(1), 49–54. <https://doi.org/10.1111/j.1467-9280.2008.02045.x>
- Walteros, C., Sánchez-Navarro, J. P., Muñoz, M. A., Martínez-Selva, J. M., Chialvo, D., & Montoya, P. (2011). Altered associative learning and emotional decision making in fibromyalgia. *Journal of Psychosomatic Research*, 70(3), 294–301. <https://doi.org/10.1016/j.jpsychores.2010.07.013>
- Wang, J., van Woerkom, M., Breevaart, K., Bakker, A. B., & Xu, S. (2023). Strengths-based leadership and employee engagement: A multi-source study. *Journal of Vocational Behavior*, <https://doi.org/10.1016/j.jvb.2023.103859>
- Wang, Q. (2013). Blue and pink: The study of color-gender metaphor (master's thesis). Ningbo University.
- Ward, A. F., Olsen, A. S., & Wegner, D. M. (2013). The harm-made mind: Observing victimization augments attribution of minds to vegetative patients, robots, and the dead. *Psychological Science*, 24(8), 1437–1445. <https://doi.org/10.1177/0956797612472343>
- Waytz, A., Gray, K., Epley, N., & Wegner, D. M. (2010). Causes and consequences of mind perception. *Trends in Cognitive Sciences*, 14(8), 383–388. <https://doi.org/10.1016/j.tics.2010.05.006>

- Weisman, K., Dweck, C. S., & Markman, E. M. (2017). Rethinking people's conceptions of mental life. *Proceedings of the National Academy of Sciences of the United States of America*, 114(43), 11374-11379. <https://doi.org/10.1073/pnas.1704347114>
- Wiese, E., Weis, P. P., Bigman, Y., Kapsaskis, K., & Gray, K. (2022). It's a match: task assignment in human-robot collaboration depends on mind perception. *International Journal of Social Robotics*, 14(1), 141-148. <https://doi.org/10.1007/s12369-021-00771-z>
- Willemsen, P., Newen, A., Prochownik, K., & Kaspar, K. (2023). With great(er) power comes great(er) responsibility: an intercultural investigation of the effect of social roles on moral responsibility attribution. *Philosophical Psychology*, 38(2), 820-846. <https://doi.org/10.1080/09515089.2023.2213277>
- Xie, X. Y., Zuo, Y. H., & Hu, Q. J. (2021). Human resource management in the digital era: A human-technology interaction lens. *Management World*, 37(1), 200-216. doi:10.19744/j.cnki.11-1235/f.2021.0013
- Xiong, W., Wang, C., & Ma, L. (2023). Partner or subordinate? Sequential risky decision-making behaviors under human-machine collaboration contexts. *Computers in Human Behavior*, <https://doi.org/10.1016/j.chb.2022.107556>
- Yan, X., Mo, T. T., & Zhou, X. Y. (2024). The influence of cultural differences between China and the West on moral responsibility judgments of virtual humans. *Acta Psychologica Sinica*, 56(2), 161-178. doi: 10.3724/SP.J.1041.2024.00161
- Yang, L., Wang, B. Q., Gen, Y. F., Yao, D. W., Cao, H., Zhang J. X., & Xu, Q. Y. (2019). The influence of hypothetical and real money rewards on the risky decision-making of the abstinent heroin user. *Acta Psychologica Sinica*, 51(4), 507-516. doi: 10.3724/SP.J.1041.2019.00507
- Young, A. D., & Monroe, A. E. (2019). Autonomous morals: Inferences of mind predict acceptance of AI behavior in sacrificial moral dilemmas. *Journal of Experimental Social Psychology*, <https://doi.org/10.1016/j.jesp.2019.103870>
- Yu, F., & Xu, L. Y. (2019). Change and constancy in moral responsibility. *Journal of Wuhan University of Science and Technology (Social Science Edition)*, 21(1), 53-60. doi: 10.3969/j.issn.1009-3969.2019.01.009
- Yu, H. B., Siegel, J. Z., & Crockett, M. J. (2019). Modeling morality in 3-D: Decision-making, judgment, and inference. *Topics in Cognitive Science*, 11(2), 409-432. <https://doi.org/10.1111/tops.12382>
- Zhan, Z., & Wu, B. P. (2019). Ubiquitous harm: Moral judgment in the perspective of the theory of dyadic morality. *Advances in Psychological Science*, 27(1), 128-140. doi: 10.3724/SP.J.1042.2019.00128
- Zhang, T., Tao, D., Qu, X., Zhang, X., Zeng, J., Zhu, H., & Zhu, H. (2020). Automated vehicle acceptance in China: Social influence and initial trust are

key determinants. *Transportation Research Part C: Emerging Technologies*, 112, 220-233. <https://doi.org/10.1016/j.trc.2020.01.027>

Zheng, N. N., Liu, Z. Y., Ren, P. J., Ma, Y. Q., Chen, S. T., Yu, S. Y., ... Wang, F. Y. (2017). Hybrid-augmented intelligence: Collaboration and cognition. *Frontiers of Information Technology & Electronic Engineering*, 18(2), 153-179. <https://doi.org/10.1631/FITEE.1700053>

Zhou, G. M., & Fu, X. L. (2002). Distributed cognition: A new cognition perspective. *Advances in Psychological Science*, 10(2), 147-153.

Appendix 1: Intelligent Robot Background Information

Artificial intelligence technology is gradually merging with robotics technology, giving rise to advanced AI robots. These AI robots not only have human-like appearances and can perform various actions but also possess independent learning, thinking, and decision-making abilities similar to humans. Therefore, in organizational settings, the coexistence of humans and intelligent robots is inevitable.

Intelligent robots can learn human work patterns, including male and female work patterns and characteristics, through machine learning. Machine learning methods include supervised learning and reinforcement learning. Supervised learning trains robots to classify behaviors of specific gender roles. By collecting large amounts of data (e.g., video, voice data), robots learn to distinguish and imitate behaviors of different genders. Reinforcement learning involves robots gradually receiving feedback about gendered behaviors through interaction with humans and adjusting their behaviors based on this feedback.

Appendix 2: Collaborative Work Team Background Information

Artificial intelligence technology is gradually merging with robotics technology, giving rise to advanced AI robots. These AI robots not only have human-like appearances and can perform various actions but also possess independent learning, thinking, and decision-making abilities similar to humans. Therefore, in organizational settings, the coexistence of humans and intelligent robots is inevitable.

When intelligent robots enter human work domains, they form human-robot collaborative teams with humans. Simultaneously, intelligent robots can also form robot-robot collaborative teams. Additionally, humans can form human-human collaborative teams. Thus, organizational settings will feature three types of collaborative teams: human-robot, robot-robot, and human-human.

Appendix 3: Agency and Experience Scale

Agency Dimension: 1. The agent can make plans for behavior. 2. The agent can control its behavior. 3. The agent can remember events in life. 4. The agent can understand thoughts of people around them. 5. The agent can distinguish right from wrong. 6. The agent can influence outcomes of situations. 7. The agent can communicate feelings and thoughts with others.

Experience Dimension: 1. The agent has personality. 2. The agent can experience desires. 3. The agent can experience sensations. 4. The agent can experience emotions. 5. The agent can experience joy. 6. The agent can experience fear. 7. The agent can experience hunger.

Appendix 4: Reliability Coefficients

Experiments 1 and 2 measured agency and experience for four work agents (male human, female human, male robot, female robot), calculating Cronbach's α coefficients for four measurements (see Table 1).

Table 1. Cronbach's α Coefficients for Agency and Experience in Experiments 1 and 2

Agent Type	Agency	Experience
Male Human	0.89	0.92
Female Human	0.88	0.91
Male Robot	0.85	0.90
Female Robot	0.84	0.89

Experiments 3 and 4 measured agency and experience for three collaborative teams (human-human, human-robot, robot-robot), calculating Cronbach's α coefficients for three measurements (see Table 2).

Table 2. Cronbach's α Coefficients for Agency and Experience in Experiments 3 and 4

Team Type	Agency	Experience
Human-Human	0.90	0.93
Human-Robot	0.87	0.91
Robot-Robot	0.83	0.88

Appendix 5: Supplementary Analysis for Experiment 3

Independent samples t-tests compared mind capabilities between robot-robot collaborative teams and individual robot agents, and between human-human

collaborative teams and individual human agents. Results are shown in Table 3.

Table 3. Mind Capability Differences Between Dual-Agent and Single-Agent in Monetary Context

Comparison	Agency t-value	Agency p-value	Experience t-value	Experience p-value
Human-Human vs. Single Human	2.34	0.021	3.12	0.002
Robot-Robot vs. Single Robot	-4.56	<0.001	-6.78	<0.001

Independent samples t-tests also compared mind capabilities between human-robot collaborative teams and single agents. Results are shown in Table 4.

Table 4. Mind Capability Differences Between Human-Robot Team and Single Agent in Monetary Context

Comparison	Agency t-value	Agency p-value	Experience t-value	Experience p-value
Human-Robot vs. Single Human	-1.89	0.061	-2.45	0.015
Human-Robot vs. Single Robot	8.92	<0.001	12.34	<0.001

Appendix 6: Supplementary Analysis for Experiment 4

Independent samples t-tests compared mind capabilities between robot-robot collaborative teams and individual robot agents, and between human-human collaborative teams and individual human agents. Results are shown in Table 5.

Table 5. Mind Capability Differences Between Dual-Agent and Single-Agent in Moral Context

Comparison	Agency t-value	Agency p-value	Experience t-value	Experience p-value
Human- Human vs. Single Human	1.98	0.049	2.67	0.008
Robot- Robot vs. Single Robot	-5.23	<0.001	-7.45	<0.001

Independent samples t-tests also compared mind capabilities between human-robot collaborative teams and single agents. Results are shown in Table 6.

Table 6. Mind Capability Differences Between Human-Robot Team and Single Agent in Moral Context

Comparison	Agency t-value	Agency p-value	Experience t-value	Experience p-value
Human- Robot vs. Single Human	-2.12	0.036	-3.01	0.003
Human- Robot vs. Single Robot	9.45	<0.001	11.23	<0.001

Note: Figure translations are in progress. See original paper for figures.

Source: ChinaXiv –Machine translation. Verify with original.