

Human-AI Collaboration and Trust in Human Resource Management: Opportunities, Challenges and Framework for Fair Implementation

Nidhi Goel

Assistant Professor, Vivekananda School of Law and Legal Studies, Vivekananda Institute Of Professional Studies

Abstract

AI is being increasingly integrated into HRM, modifying recruitment, performance analysis, learning and development, and managerial decisions. This paper uses eighteen recent organizational psychology studies and focuses on information systems and human-computer interaction. It provides a synthesis on what triggers trust between employees and AI, the human-algorithm workload balance, and fairness concerns regarding AI in people management and assessment. The review concludes that trust in AI is separate from trust in human colleagues, that the process shapes trust and the delegation of trust can improve employees' performance and overall satisfaction, even when the trust is not complete. This review also notes that trust, explainability, accountability, and the varying perceptions of fairness limit the deployment of AI. Task Division, trust construction, varying perceptions of fairness, and the organizational structure are the main components of the Human-AI and HRM collaboration system, which this review seeks to examine. It also identifies research gaps.

Keywords: Human-AI Collaboration, Trust, Human Resource Management, Fairness, AI Adoption

Introduction

Managing human resources is about making decisions about people using people skills. In the last ten or so years, the need for people skills to make HR decisions has decreased because of developments in technology, particularly artificial intelligence (AI). AI can evaluate resumes, analyse the likelihood of various outcomes of employee-related decisions, and perform tasks that used to require a human resource manager or consultant (Madanchian et al., 2023) [6]. As reported in the studies reviewed here, most business leaders believe AI will change the way all businesses operate within the next several years, and businesses that employ AI in their HR departments will experience increased revenues and reduced costs (Xu et al., 2024) [12]. The studies also claim that the main factors in the success of AI in HRM are the systems' technical abilities and employee acceptance of the systems, appropriate task allocation to the most competent parties, and process outcomes that are perceived to be fair to all affected (Tinguely et al., 2023; Park et al., 2022) [8, 9].

Using eighteen peer-reviewed and practitioner-oriented studies from the years 2022 to 2025, this paper aims to consolidate existing research in HRM on Human-AI collaboration and trust. Instead of providing

a synthesis or summary of each of the studies, this paper is structured to provide responses to three questions common in the studies: (1) How is trust in AI formed and how is it different from trusting human coworkers? (2) To what extent, and how, should tasks and decisions be allocated to AI and humans? (3) In what ways do the design and fairness dilemmas be resolved in order for the organizations to be able to employ AI ethically? Then the paper proposes an integrative framework for fair implementation and points out the gaps left open by the reviewed literature.

Review of Literature

The Degree of AI Adoption in HRM

Before diving into the deep end of trust, task allocation, and fairness, it helps to look at where AI actually stands in HR today. According to a 2023 review by Madanchian and colleagues [6], the way companies use AI breaks down into three main areas. First, there are the reasons organizations adopt it in the first place—which usually comes down to technological readiness and perceived usefulness, though trust is ultimately the dealbreaker that decides whether a tool actually gets used. The second cluster concerns the applied use cases, e.g., algorithmic shortlisting of candidates in recruitment, ethical decision-making frameworks to guide algorithmic HR decisions, and AI systems in performance evaluation. The third cluster is about organizational implications. It contends that for AI-augmented HRM to succeed, there needs to be a capability framework that combines technical resources with non-technical resources such as human skills, leadership and organizational culture. This statement is reiterated in various forms throughout the rest of the literature reviewed here.

Nguyen and Elbanna (2025) [14] frame this fragmentation against the backdrop of the disorders of human-AI collaboration. They observe that the scattered nature of such research in management, information systems, and organizational studies makes it difficult to develop an integrated body of knowledge that is useful to both scholars and practitioners. They also make a distinction between human-AI augmentation and human-AI automation. The former refers to AI systems that offer support to human decision-making and improve productivity, whereas the latter refers to AI systems that are substitutes to human tasks. This distinction and the absence of a unified definition of “collaboration” in HRM literature are explained by the studies that were reviewed for this research. For example, Hemmer et al. (2023) [4] supported the idea of augmentation, whereas Madanchian et al. (2023) [6] described the automation of judgment via algorithmic shortlisting.

Trust as the Mediating Variable

Most sources suggest that adoption and effective use are driven by the extent to which AI-HRM tools are trusted, rather than by technical accuracy alone. Using a survey of 426 Chinese employees, Xu et al. applied the social exchange theory and leader-member exchange theory to find that initial trust in AI is built on cognitive and affective pathways and that a leader with high trustworthiness enhances employees’ trust in AI and their intention to adopt it (Xu et al., 2024) [12]. Interestingly, their study did not find a direct association between explaining how an AI system works and trust but that explanation moderates the association, reducing the dependence of employees on leader trust after they understand the technology on their own. This suggests that leadership and technical transparency are not complements, but rather partial substitutes in the formation of early stage trust.

Building on this work, Georganta and Ulfert (2024) [3] explored this question in a team context, where a new AI teammate was trusted relative to a new human teammate in a controlled experiment with 127

participants. They found that AI was perceived as less trustworthy and generated less affective interpersonal trust than human teammates, but cognitive interpersonal trust and trust-related behaviors were similar. This distinction is relevant for HRM practice: Employees seem to be able to rationally trust an AI colleague once competence and reliability have been demonstrated, but the emotional bond of a human working relationship is more difficult to reproduce. Ulfert et al. (2024) [10] operationalized this insight into a multidisciplinary theoretical framework of team trust in human-AI teams, arguing that trust in such teams cannot be reduced to a single construct derived from human-human trust research, but must take into account the AI's perceived competence, benevolence, and predictability separately from affective bonds, and must be studied at the team level rather than only at the individual level.

Wen et al. (2025) [11] investigated trust from the decision-weighting perspective and tested the Socio-Cognitive Model of Trust in two studies involving 111 and 210 managers, respectively. They found that trust in AI increases the importance managers place on AI in decision making about personnel selection and evaluation of performance and that willingness to collaborate with AI mediates this relationship. Importantly, in the second study, the more free will managers perceived AI agents to possess, the weaker was the relationship between trust and their willingness to collaborate with AI agents. This suggests that anthropomorphizing AI agency can, paradoxically, lead to the erosion rather than enhancement of collaborative trust.

Zárate-Torres et al. (2025) [15] present a leadership-focused description of maintaining trust and oversight with AI in organizational decision-making. They conducted a thematic study of leadership and AI and argue that human leaders act as mediators of AI and ethics/strategy. They place automated decisions in a "hybrid space of cooperation" where human judgment offers the real-world context, ethical governance stipulations constrain the AI-enhanced system, and cognitive flexibility adjusts the equilibrium of algorithms and human oversight. This gives leadership a distinctive role: where Xu et al. (2024) [12] placed leaders as the initial trust and adoption drivers, Zárate--Torres et al. place leaders as the ongoing governance function that ensures the human element is retained in the system long after the adoption process.

Human-AI Division of Labour

A second theme is the division of labor and decision-making between human and artificial agents. Tinguely et al. (2023) [9] developed a typology of HR-AI collaboration systems, distinguishing between routine and non-routine tasks, and tasks of low and high cognitive complexity. Automatable tasks are characterized by being low in complexity and routine. For high-stakes decisions for employees' careers or where AI explainability is limited, human judgment still applies to non-routine tasks with cognitive complexity. They say the rollout of AI must be guided by social acceptability, not just efficiency, as how employees perceive an organization's fairness depends on which jobs are offshored to algorithms.

Hemmer et al. (2023) [4] studied a related but distinct mechanism, AI delegation, in which the AI model itself decides whether to make a prediction or to pass a task instance to a human, depending on estimated confidence of both parties. They found, in a study with 196 participants, that task performance and satisfaction increased when AI delegated instances to humans, even when participants were not aware that such delegation was occurring; and that this improvement was driven by increased self-efficacy in humans to whom tasks were delegated. This finding challenges the assumption that transparency is always necessary for positive Human-AI collaboration results. The delegation mechanism in this study improved

outcomes even without disclosure, but the ethical implications of undisclosed delegation are an open question that the study does not address.

Kolbjørnsrud (2024) [5] identifies six principles to accommodate this division of labor regarding human-AI interaction within organizations. These are: addition (AI is used to supplement the work of teams), relevance (the most appropriate actors for the task are arranged), substitution (tasks are done by AI where it can accomplish them equally well or better than humans), diversity (collective problem-solving incorporates as many different viewpoints as possible), collaboration (interaction is designed so that humans and AI are in a complementary and collaborative relationship), and explanation (AI choices are rendered to a level that humans can act upon). While the efficiency of employees individually is often the focus of most discussions, Kolbjørnsrud believes that the end goal should be organizational intelligence. Leaders should purposefully design the blend of humans and digital resources. It should not be assumed that a natural division of labor will occur with the application of AI within the organization.

Lou, Lu, Raghu, and Zhang (2025) [18] expand this division-of-labour question by focusing on the development of learning and adaptive AI agents as intelligent, active participants in humans-AI collaborations. Using Team Situation Awareness theory, they identify two major problems within the scope of AI research: first, the ‘frustrating challenge’ of AI agents’ behavioral alignment, and second, the failure to recognize AI as a teammate rather than a tool. Their research agenda on human-AI collaboration progresses through the phases of formulation, coordination, maintenance, and training. Their agenda illustrates that shared mental models and trust are as critical for the sustainability of collaboration as the division of labor is. This is in line with Kolbjørnsrud’s (2024) [5] principles by arguing that the division of labor is not a design choice; rather, a persistent effort is required to maintain it, as AI agents increase their autonomy. (This study is a pre-print and has not completed peer review.)

Raftopoulos and Hamari (2023) [16] have reached a similar conclusion, but from the perspective of organizational value. Their review of human-AI collaboration literature outlines five positions that organizations must align to achieve sustainable business value from collaboration: strategic positioning, human engagement, organizational evolution, technology development, and intelligence building. Their review suggests that the internal microfoundations (structure, culture, skills) are more critical to sustainable value generation, and the ability to manage augmented collaboration, than the sophistication of the AI.

Fairness, Explainability and Stakeholder Tensions

Fairness is the third and perhaps the most contentious of the three themes. Park et al. (2022) [8] identify five tensions from their participatory design sessions. As part of their sessions, they included employees, human resource managers, and AI experts, among others. The five tensions are: competing fairness concerns, AI accuracy, algorithmic transparency, and explainability and perceptions of the tensions between productivity and the ‘inhumanity’ of reducing employee performance evaluations to an impersonal numerical score. Their design sessions showed that when stakeholders back the use of AI systems, there is often disagreement about what fairness actually is. For example, some employees wanted historical performance evaluations to be insulated from the prevailing market conditions; and some AI/business experts wanted to include some demographic information to enhance the precision of the prediction, something that was overwhelmingly rejected among the employee participants. These workshops gave design directions that showed that exercising human discretion in addition to AI

recommendations through the process of informed consent, is essential, and that the AI recommendations must not be taken and accepted as the ultimate and the unquestioned decision.

Abedin et al. (2022) [1] synthesized a special issue regarding the design and management of Human-AI interactions. They identified four major open challenges in the field: AI User Interface design; human-AI dialogue and collaboration; the challenge of explainability, responsibility, fairness, and bias; and human interaction with increasingly autonomous or “agentic” AI. They state that research in explainability is often developed by and for computer scientists, and there is an enduring gap between technically explainable models and the explanations that are accessible to the managers and staff that need to address and work with AI outputs. This gap is a driving force to the issues of fairness that Park et al. (2022) [8] highlighted because stakeholders can’t assess the fairness of an AI-mediated decision if they can’t understand the process of how that decision was made.

Jiang et al. (2023) [13] focus on AI-based context services and how to explain AI decisions. They believe that the “black box” of AI makes human and AI interaction a necessity. They believe this collaboration is a way to address and improve the uncertainty of AI algorithms and their outputs. Their framework segments AI-context systems into phases of context acquisition, interpretation, and application. Their framework shows that trust-related breakdowns could happen at each of the phases. This shows that the AI answer does not need to be the final output for trust breakdowns to occur. This extends the concepts mentioned by Park et al. (2022) [8] and Abedin et al. (2022) [1] through decades of design and trust within an AI system, and not through a singular design element.

Fenwick et al. (2024) [2] and Madanchian and Taherdoost (2025) [7] incorporate a broader perspective. With regard to AI-integrated HRM, Fenwick et al. identify three stages: phase one is technocratic, and phases two and three are integrated and fully embedded. They contend that ethical and human dimensions must be included starting from the technocratic phase, rather than retrofitted once AI is deeply embedded into HR systems. In their review and critique of the organizational and technological dimensions of AI adoption in HRM, Madanchian and Taherdoost cite the reluctance to change, data security, and integration costs as the major obstacles. On the other hand, strong digital leadership, a supportive organizational culture, and advances in natural language processing are factors in favor of AI adoption in HRM. Both articles importantly disregard that the use of AI in HRM is solely a technology-diffusion phenomenon. Instead, they stress that the employee experience of AI implementation in HRM as either empowering or disruptive is shaped by the power relationships and organizational culture.

Discussion and Theoretical Framework

The reviewed literature synthesizes these three themes and converges on an integrative framework for fair Human-AI collaboration in HRM built on four interdependent pillars.

Distribution of the task. As pointed out by Tinguely et al. (2023) [9] and Kolbjørnsrud (2024) [5], organizations should assign tasks based on the routineness, cognitive complexity and social acceptability of tasks, not only on efficiency. The Kolbjørnsrud substitution principle applies to low-complexity, routine HR tasks, such as scheduling or basic screening, while the Kolbjørnsrud diversity and collaboration principles should apply to higher-stakes, non-routine decisions, such as promotion or termination, in which the retention of multiple human perspectives guards against both algorithmic and individual human bias.

Trust. According to Xu et al. (2024) [12], Georganta and Ulfert (2024) [3], and Ulfert et al. (2024) [10], creating trust in AI-HRM systems requires both the cognitive (proven trust and consistency) and the affective (support from leadership and company-wide trust) elements. Furthermore, both affective and

cognitive elements are necessary, and neither elements will substitute the other. In the early phases of AI implementation, leaders' trust in AI is critical in generating company-wide trust in AI. It is important to note that even a well-designed AI system will not automatically generate trust within the organization.

Equitable outcomes for all stakeholders. Park et al. (2022) [8] show that fairness is not a homogenous construct. Rather, participatory design and implementation frameworks must accommodate multiple, and often conflicting, stakeholder perspectives, and even design fairness frameworks for AI systems. This pillar also includes the explainability requirement identified by Abedin et al. (2022) [1], where provided fairness explanations must be actionable for employees and managers alike.

Organizational design and human discretion. Wen et al. (2025) [11] and Hemmer et al. (2023) [4] recommended limiting the decisional power of highly autonomous AI agents and retaining mechanisms for delegating AI tasks under human supervision in the AI-augmented decision-making process. Fenwick et al. (2024) [2] cite this pillar and state that when embedding AI-Aided Human Resources Management (AI-HRM) systems, it is much easier to design the system with human-centered design in mind than to do so after the system has been fully embedded in organizational processes.

There are two literature contradictions to address. Hemmer et al. (2023) [4] argued that even with the lack of human notification, satisfaction and performance improved with the AI delegation. In turn, Park et al. (2022) [8] and Abedin et al. (2022) [1] considered transparency and explainability to be prerequisites of fairness. This would not be a contradiction, since Hemmer et al. (2023) studied the delegation of relatively low-stake tasks, while the examples given by the other authors were Human Resource Management (HRM) tasks with significant implications. However, it suggests that the importance of transparency may be context-dependent and valuable in low-stake tasks, while high-stake tasks may require a higher level of transparency. Finally, Xu et al. (2024) [12] argued that disclosing AI operations did not create trust, while Abedin et al. (2022) [1] positioned explainability at the core of accountability and fairness. This contradiction shows that while trust may not be created by transparency, the explanation may be necessary to create accountability and fairness.

Challenges and Fairness Issues

There are some common challenges across the reviewed studies that limit the practical implementation of the above framework. The first involves the accountability gap created by increasingly autonomous, or agentic, AI systems. Abedin et al. (2022) [1] point out that as AI systems become more autonomous, it is harder to determine who is responsible for a particular HR decision – the system designer, the deploying organization or the HR professional who acted on the AI's advice – especially when the AI's internal reasoning is not fully transparent.

Another challenge is the lack of consensus among different stakeholders on what constitutes a fair evaluation. An interesting and useful disagreement is reported by Park et al. (2022) [8], where employees in structurally disadvantaged roles believed that fairness meant evaluating their performance outside of business or market conditions that were beyond their control, while some AI and business experts argued that the addition of more context and even sensitive demographic information would make the evaluation more accurate. This is not a debate over implementation detail, it is a debate over two different underlying definitions of fairness: one procedural (i.e., was the process applied consistently) and one substantive (i.e., did the outcome take into account circumstances beyond the employee's control). Organizations adopting AI evaluation tools must decide explicitly which definition they are optimizing for.

A third problem is that the barriers to adoption are as much organizational and cultural as they are technical. Madanchian and Taherdoost (2025) [7] found that change aversion and organizational culture were at least as important as the accuracy of the underlying algorithms in determining whether AI-HRM tools were accepted. This was consistent with Xu et al.'s (2024) [12] finding that leadership trust, not technical capability, was the driver of initial adoption. This implies that organizations that focus primarily on improving the accuracy of AI, but do not consider change management and leadership engagement, are likely to realize a limited return on adoption.

A fourth challenge is associated with the emotional dimension of trust described by Georganta and Ulfert (2024) [3]. Even if AI teammates are perceived as competent and reliable, it seems hard to replicate the affective trust and interpersonal warmth of human colleagues. This might constrain the roles AI can credibly play in HRM functions that depend on employees feeling supported such as coaching, grievance handling or career counseling, even if AI is technically able to perform the underlying analytical task.

In addition to the aforementioned literature on trust and fairness, we can further analyze the psychological costs from collaborative Human-AI systems. When integration of technostress is considered, Xia (2023) [17] suggests that AI-enabled collaboration exert challenge and hindrance stressors at the same time and have positive and negative psychological effects simultaneously instead of having a strictly negative psychological effect. When employees have prior experience with AI, it is likely to increase their AI-driven satisfaction and decrease their technostress. Satisfied employees are likely to have a lower level of AI-induced technostress. Organizations likely have a ways to reduce the psychological AI adoption burden relying on the creation of positive, purposeful workplace opportunities. This result likely supports the barriers established by Madanchian and Taherdoost (2025) [7] and suggests a new well-being framework to the barriers, which seem largely cultural and structural.

Future Research Gaps

The reviewed literature leaves a number of gaps for future research. First, most studies in this group, such as Xu et al. (2024) [12] and Wen et al. (2025) [11], rely on samples from one country at a time, especially expatriate Chinese managers and employees. Thus, they exhibit low cross-cultural generality. There is a sharp absence of cross-cultural studies on trust formation in the context of AI-HRM. Second, the issue regarding the degree of disclosure as raised by Hemmer et al. (2023) [4] has not, to our knowledge, been examined in the field of HRM, and in particular, concerning high-risk scenarios (e.g., promotion or termination). Most studies on task delegation have centered on low-risk, passive, annotated activities, and the potential benefits of performance and satisfaction would be unknown to employees if delegated in a 'secret' manner, as they would likely perceive this as a fairness violation. Third, while Park et al. (2022) [8] and Abedin et al. (2022) [1] both call for stakeholder-centered and explainable design, none of the studies we reviewed test whether implementing their suggested design solutions actually changes downstream trust or fairness perceptions in a live organizational deployment; the field currently has more diagnostic research on tensions than evaluative research on solutions. Finally, longitudinal studies are scarce – most studies evaluate perceptions of trust or fairness at a single point in time (often immediately after the introduction of an AI), leaving open questions about how trust, task allocation, and perceptions of fairness develop as employees get prolonged experience working with AI systems in HRM.

Conclusion

The eighteen studies reviewed in this article collectively suggest that Human-AI collaboration in HRM is

not merely a question of deploying sufficiently accurate algorithms. Trust in AI is mediated by different cognitive and affective pathways and is strongly influenced by organizational leadership; the division of labor between humans and AI should be based on the routineness of the task, the cognitive complexity, and social acceptability, not just efficiency; and fairness is a contested multi-stakeholder construct, not a single design target. The integrative framework presented here based on task allocation, trust-building, stakeholder-centered fairness, and organizational design provides HR practitioners and researchers with a structure to think about these interdependent requirements simultaneously rather than separately. As AI systems in HRM become more autonomous and make more consequential decisions about people's careers, closing the gaps identified above, especially around cross-cultural generalizability, high-stakes disclosure, and longitudinal trust dynamics, will be crucial to ensure that Human-AI collaboration in HRM is not only efficient but also trusted and fair.

References

1. Abedin B.R, Meske C, Junglas I, Rabhi F, Motahari-Nezhad H.R, "Designing and Managing Human-AI Interactions", Information Systems Frontiers, 2022. <https://doi.org/10.1007/s10796-022-10313-1>
2. Fenwick A, Molnar G, Frangos P, "Revisiting the Role of HR in the Age of AI: Bringing Humans and Machines Closer Together in the Workplace", Frontiers in Artificial Intelligence, 2024, 6, 1272823. <https://doi.org/10.3389/frai.2023.1272823>
3. Georganta E, Ulfert A.S, "My Colleague Is an AI! Trust Differences between AI and Human Teammates", Team Performance Management, 2024, 30 (1/2), 23-37. <https://doi.org/10.1108/TPM-07-2023-0053>
4. Hemmer P, Westphal M, Schemmer M, Vetter S, Vössing M., Satzger G, "Human-AI Collaboration: The Effect of AI Delegation on Human Task Performance and Task Satisfaction", Proceedings of the 28th International Conference on Intelligent User Interfaces (IUI '23), ACM, 2023. <https://doi.org/10.1145/3581641.3584052>
5. Kolbjørnsrud V, "Designing the Intelligent Organization: Six Principles for Human-AI Collaboration", California Management Review, 2024, 66 (2), 44-64. <https://doi.org/10.1177/00081256231211020>
6. Madanchian M, Taherdoost H, Mohamed N, "AI-Based Human Resource Management Tools and Techniques: A Systematic Literature Review", Procedia Computer Science, 2023, 229, 367-377.
7. Madanchian M, Taherdoost H, "Barriers and Enablers of AI Adoption in Human Resource Management: A Critical Analysis of Organizational and Technological Factors", Information, 2025, 16 (1), 51-59. <https://doi.org/10.3390/info16010051>
8. Park H, Ahn D, Hosanagar K, Lee J, "Designing Fair AI in Human Resource Management: Understanding Tensions Surrounding Algorithmic Evaluation and Envisioning Stakeholder-Centered Solutions", Proceedings of the 2022 CHI Conference on Human Factors in Computing Systems (CHI '22), ACM, 2022. <https://doi.org/10.1145/3491102.3517672>
9. Tinguely P.N, Lee J, He V.F, "Designing Human Resource Management Systems in the Age of AI", Journal of Organization Design, 2023, 12, 263-269. <https://doi.org/10.1007/s41469-023-00153-x>
10. Ulfert A.S, Georganta E, Centeio Jorge C, Mehrotra S, Tielman M, "Shaping a Multidisciplinary Understanding of Team Trust in Human-AI Teams: A Theoretical Framework", European Journal of Work and Organizational Psychology, 2024, 33 (2), 158-171. <https://doi.org/10.1080/1359432X.2023.2200172>

11. Wen Y, Wang J, Chen X, “Trust and AI Weight: Human-AI Collaboration in Organizational Management Decision-Making”, *Frontiers in Organizational Psychology*, 2025, 3, 1419403. <https://doi.org/10.3389/forgp.2025.1419403>
12. Xu Y, Huang Y, Wang J, Zhou D, “How Do Employees Form Initial Trust in Artificial Intelligence: Hard to Explain but Leaders Help”, *Asia Pacific Journal of Human Resources*, 2024, 62, e12402. <https://doi.org/10.1111/1744-7941.12402>
13. Jiang N, Liu X, Liu H, Lim E.T.K, Tan C.W, Gu J, “Beyond AI-Powered Context-Aware Services: The Role of Human–AI Collaboration”, *Industrial Management & Data Systems*, 2023, 123(11), 2771-2802. <https://doi.org/10.1108/IMDS-03-2022-0152>
14. Nguyen T, Elbanna A, “Understanding Human-AI Augmentation in the Workplace: A Review and a Future Research Agenda”, *Information Systems Frontiers*, 2025. <https://doi.org/10.1007/s10796-025-10591-5>
15. Zárate-Torres R, Rey-Sarmiento C.F, Acosta-Prado J.C, Gómez-Cruz N.A, Rodríguez Castro D.Y, Camargo J, “Influence of Leadership on Human–Artificial Intelligence Collaboration”, *Behavioral Sciences*, 2025, 15(7), 873. <https://doi.org/10.3390/bs15070873>
16. Raftopoulos M, Hamari J, “Human-AI Collaboration in Organisations: A Literature Review on Enabling Value Creation”, *Proceedings of the 31st European Conference on Information Systems (ECIS 2023)*, Kristiansand, 2023.
17. Xia M, “Co-Working with AI Is a Double-Sword in Technostress? An Integrative Review of Human-AI Collaboration from a Holistic Process of Technostress”, *SHS Web of Conferences*, 2023, 155, 03022. <https://doi.org/10.1051/shsconf/202315503022>
18. Lou B, Lu T, Raghu T.S, Zhang Y, “Unraveling Human-AI Teaming: A Review and Outlook”, *arXiv:2504.05755*, 2025 <https://doi.org/10.48550/arXiv.2504.05755>