

E-ISSN: 3108-4192

APSSHs

Academic Publications of Social Sciences and Humanities Studies

2026, Volume 6, Issue 1, Page No: 67-75

Available online at: <https://apsshs.com/>

Journal of Applied Organizational Systems and Behavior

Artificial Coworkers, Emotional Attachment, and Identity Projection at Work

Emily Johnson^{1*}, Robert Smith¹, Laura Brown², Kevin Miller³

1. Department of Artificial Coworkers and Emotional Attachment, Faculty of Business, University of Toronto, Toronto, Canada.
2. Department of Identity Projection and Work, Faculty of Management, McGill University, Montreal, Canada.
3. Department of Human-AI Interaction, Faculty of Psychology, University of Guelph, Guelph, Canada.

Abstract

Advanced artificial intelligence is increasingly incorporated into interdependent work in ways that blur the familiar distinction between organizational tool and social collaborator. Existing scholarship explains human–AI teaming, algorithmic control, anthropomorphism, social presence, trust, and professional-role identity, yet these literatures remain fragmented when workers begin to interpret AI as a relationally consequential participant in work. This Original Conceptual Theory Article develops the proposed Theory of Artificial-Coworker Relations (TACR) to explain how AI may become construed along a tool–artificial-coworker boundary and how that construal may connect social-relational and identity processes. TACR proposes that AI affordances, work-role conditions, and human or team characteristics shape relational object construal; that social presence and repeated interaction may support workplace emotional attachment; and that work-role overlap, comparison, substitution, or complementarity may make AI interaction a site of proposed identity projection. Boundary management concerning role definition, decision rights, responsibility, reliance, disclosure, override, and escalation is theorized to condition these processes. The theory distinguishes anthropomorphism from social presence, trust from attachment, professional-role identity from identity projection, and technological agency from legitimate authority. Its propositions are conditional and falsifiable rather than claims of validated causality. The framework is bounded by uneven cross-context evidence, especially the transfer of attachment findings from social-chatbot settings to work. Validation requires longitudinal workplace research, discriminant construct development, experimental manipulation of relational cues and authority boundaries, and multilevel tests capable of challenging the theory.

Keywords: Artificial coworkers, Emotional attachment, Identity projection at work, Artificial coworkers as relational objects, Identity projection, Theory of artificial-coworker relations

How to cite this article: Johnson E, Smith R, Brown L, Miller K. Artificial Coworkers, Emotional Attachment, and Identity Projection at Work . J Appl Organ Syst Behav. 2026;6(1):67-75. <https://doi.org/10.51847/Vf2hulDu1X>

Received: 22 April 2026; **Revised:** 04 July 2026; **Accepted:** 07 July 2026

Corresponding author: Emily Johnson

E-mail ✉ emily.johnson@rotman.utoronto.ca

Introduction

Artificial intelligence is no longer confined to backstage computation. In many work systems, it now participates in recommendation, drafting, classification, coordination, monitoring, and decision support, altering how expertise and organizational action are distributed. Recent organization scholarship treats advanced AI not only as computational technology but also as a participant in changing organizational coordination and team collaboration [1, 2]. This development creates a conceptual problem: the categories of “tool” and “coworker” are increasingly insufficient when an AI system is recurrently involved in interdependent tasks yet remains outside ordinary assumptions about employment, personhood, responsibility, and authority.



© 2026 The Author(s).

Copyright CC BY-NC-SA 4.0

The relational problem is not reducible to whether an AI appears human. Workers may trust an AI for technical reasons without feeling socially connected to it, while socially expressive systems may evoke interpersonal responses without receiving legitimate decision authority. Research on human trust in AI shows that cognitive and emotional forms of trust are shaped by characteristics of the technology and its representation [3]. Trust therefore provides one bridge between technical use and relational interpretation, but it should not be treated as equivalent to attachment, identification, friendship, or organizational membership.

A second limitation of prevailing accounts is that they often position automation and augmentation as competing outcomes rather than interdependent organizational dynamics. Management theory instead shows that automation and augmentation can coexist and recursively shape one another [4]. The meaning of an AI system may therefore depend on whether it substitutes for valued tasks, complements judgment, evaluates performance, coordinates activity, or moves among these roles over time. These functional positions can also become socially meaningful when workers interpret what the system is doing in relation to their own role.

A third limitation concerns power. Algorithms can restructure surveillance, direction, evaluation, and control, generating new forms of contestation at work [5]. A system perceived as cooperative may simultaneously operate within a managerial architecture of monitoring or authority. Accordingly, this article develops an original, non-empirical theory of artificial-coworker relations rather than assuming that “coworker” is a factual status. It defines an artificial coworker as a proposed analytic category for an AI system recurrently involved in interdependent work and construed by workers as participating in coordination. The article integrates theoretical foundations, emotional attachment and social presence, identity projection, and boundary management. Its contribution is a proposed explanation of how relational object construal may emerge, how social-relational and identity pathways may interact, and where interpretation must stop short of personhood, causal proof, or organizational decision readiness.

Theoretical foundations

The first foundation concerns agentic artifacts. Information-systems theory distinguishes technologies that merely execute commands from agentic artifacts that can receive delegated tasks and participate in appraisal, distribution, and coordination [6]. This supports treating some AI systems as relationally consequential work participants, but functional agency does not establish consciousness, moral responsibility, employment status, or legitimate organizational authority. TACR therefore uses agency as a condition that may make relational interpretation possible, not as evidence that an AI literally becomes a colleague. The second foundation concerns anthropomorphism and communicative framing. Experimental research shows that anthropomorphic design cues and agency framing can influence users’ perceptions of conversational agents, including perceived social presence and evaluations of the agent [7]. Yet “humanlikeness” is not a unitary property. Visual cues, identity cues, and conversational interactivity can have distinguishable and interacting effects on perceived humanness [8]. A worker may consequently interpret an AI as socially responsive without judging it visually humanlike, or may perceive a polished persona without developing a durable relational orientation.

The third foundation is contextual contingency. Meta-analytic evidence indicates that anthropomorphism does not produce uniform responses across AI systems or service conditions; effects depend on characteristics of the technology and interaction context [9]. This matters theoretically because a coworker-like interpretation cannot be inferred from interface design alone. Task interdependence, reliability, role framing, familiarity, and organizational norms may provide alternative explanations for apparently relational behavior.

These foundations establish the minimum conditions for TACR while also limiting it. Agentic capability can support delegation, anthropomorphic and interactive cues can shape social interpretation, and contextual conditions can alter these effects. None of these findings demonstrates emotional attachment, identity projection, or coworker status by itself. The proposed theory therefore treats relational object construal as an emergent interpretation formed at the intersection of technological affordances, work roles, and user meaning-making rather than as a direct consequence of one design feature.

Emotional attachment and social presence

Social presence refers here to the experienced social salience or availability of an interaction partner, whereas emotional attachment denotes a stronger affective bond involving relational investment or motivation to maintain the relationship. The distinction is essential. Research on human–chatbot relationships shows that self-disclosure, trust, and repeated interaction can accompany the development of meaningful relational bonds with a conversational agent [10]. Such evidence demonstrates that artificial agents can become relational objects for users, but it does not establish that analogous bonds occur in employment settings or that social presence automatically becomes attachment.

Across social and companion-chatbot research, repeated conversation, self-disclosure, perceived relational qualities, and interaction intensity have been associated with relationship formation or attachment [10-12]. The temporal dimension is especially important for TACR because a worker’s interpretation of an AI may change through repeated exposure rather than

being fixed at first use. An initially instrumental system may become familiar, predictable, and socially salient; conversely, apparent warmth may weaken when the system performs poorly, becomes intrusive, or is reframed as an instrument of managerial control. These possibilities make attachment a process claim rather than a stable design outcome.

Workplace-oriented evidence strengthens the case for distinguishing social presence from anthropomorphism. Research on human–AI collaboration indicates that social presence can matter beyond anthropomorphic design and may shape collaborative motivation or willingness to depend on an AI partner [13]. This provides a direct organizational bridge to TACR: workers may experience an AI as socially consequential even without mistaking it for a human being. Nevertheless, social presence remains analytically different from trust, reliance, affection, or emotional dependence.

TACR therefore proposes, rather than assumes, a workplace social-relational pathway. Perceived social presence may make repeated interaction more relationally meaningful; repeated relational interaction may in turn create conditions under which emotional attachment or relational-maintenance motivation becomes possible. The pathway is conditional on familiarity, understandability, continuity, task context, individual social motivation, and organizational framing. Alternative explanations—including usability, novelty, performance, loneliness, or routine exposure—must be measured rather than absorbed into “attachment.” Evidence from companion-chatbot relationships establishes plausibility, whereas the workplace extension requires longitudinal testing that can separate temporal association from causal mechanism.

Identity projection

If attachment concerns the relationship to the AI, identity projection concerns the relationship between the AI and the worker’s interpretation of self. Professional-role identity can be affected when AI takes over decision activities previously central to an employee’s occupational role [14]. Such findings establish that AI can become identity-relevant because it alters the meaning, boundaries, or status of valued work. They do not, however, establish that workers project identity onto AI. TACR therefore introduces identity projection as an original proposed construct rather than a renamed empirical finding.

Organizational research also shows that algorithmic change can alter regimes of knowing and provoke symbolic responses as actors defend, reinterpret, or renegotiate established forms of expertise [15]. This indicates that reactions to AI may concern more than task efficiency. Workers can respond to what the technology implies about whose knowledge counts, what competence means, and which activities sustain professional legitimacy. These dynamics provide an organizational context in which AI may become a mirror, comparison target, threat, or complement to the work self.

Technology-driven work redesign can also reshape participation in learning. Research on robotic surgery shows how altered access to practice can push learners toward informal or “shadow” forms of skill development when approved participation becomes insufficient [16]. Although this evidence does not measure identity projection, it illustrates how technology can restructure opportunities through which competence and occupational self-understanding are maintained. Identity effects may therefore arise indirectly through changing participation, expertise, and role visibility rather than through social attachment to the technology.

A functional-identity perspective further suggests that identity consequences depend on what AI does in relation to workers—whether it substitutes for, complements, evaluates, or reorganizes valued functions [17]. TACR builds on this premise by proposing identity projection as the process through which workers use AI interaction, comparison, attribution, or role overlap as material for externalizing, defending, testing, or revising work-related self-understandings. Proposed identity projection is distinct from anthropomorphism because it concerns the worker’s self-meaning rather than the system’s humanlikeness; it is distinct from attachment because affective bonding is not required.

Table 1 clarifies how the central relational and identity constructs in the proposed theory differ conceptually, where they intersect, and where interpretive boundaries must be maintained.

Table 1. Conceptual architecture of artificial-coworker relations at work

Conceptual domain	Core meaning	Primary question	What it is not	Proposed function in TACR	Principal uncertainty	Interpretive boundary
Artificial coworker	AI interpreted as participating recurrently in interdependent work	When does AI become relationally consequential rather than merely instrumental?	Legal employee, person, or autonomous organizational actor	Focal relational category	Construal may vary across tasks and time	Perceived coworker status does not create formal membership
Relational object construal	Worker interpretation of AI along a tool–coworker continuum	What meaning does the worker assign to the AI’s role?	Objective property of technology	Proposed central organizing construct	Stability, reversibility, and variation remain unknown	Interpretation must be separated from capability
Social presence	Experienced social salience of AI during interaction	Does the AI feel socially available or responsive?	Attachment, trust, friendship, or authority	Social-relational pathway condition	Presence may reflect design,	Social presence does not establish

					familiarity, or task structure	emotional bonding
Emotional attachment	Affective investment in maintaining an AI relationship	Does continued interaction become emotionally meaningful?	Routine use, satisfaction, reliance, or preference	Proposed relational development process at work	Workplace formation and persistence remain unvalidated	Attachment cannot be inferred from use frequency
Professional-role identity	Work-related self-understanding tied to valued roles and competence	What does the worker believe their role means?	General self-esteem or technology attitude	Identity-pathway foundation	Responses may differ by occupation and role centrality	Identity change need not involve attachment
Identity projection	Use of AI comparison, overlap, or interaction to interpret the work self	What does AI imply about who the worker is becoming?	Anthropomorphism or attributing personality to AI	Original proposed TACR construct	Measurement and discriminant validity remain undeveloped	Conceptual plausibility is not empirical validation
Boundary calibration	Ongoing adjustment of role, reliance, authority, responsibility, override, and escalation	Where should the human–AI relationship stop?	Fixed technical permission setting	Proposed cross-cutting governance mechanism	Effective forms may be context dependent	Governance principles are not demonstrated interventions
Relational-authority asymmetry	Social influence exceeds formally legitimate AI authority	Can relational standing become stronger than formal decision rights?	Proof of manipulation or dependence	Original proposed risk condition	Prevalence and consequences are unknown	Social influence is not legitimate power

Table note. Bolded TACR elements identify original proposed constructs or mechanisms introduced by this article. The table is conceptual rather than classificatory: categories may overlap, shift over time, or be absent in particular human–AI work relationships.

Boundary management between tool and colleague

The distinction between tool and colleague is unlikely to remain stable when AI participates repeatedly in consequential work. Longitudinal evidence on human-in-the-loop configurations shows that organizational arrangements around algorithms can be repeatedly reconfigured as tasks, expertise, and system capabilities change [18]. TACR therefore proposes boundary management as a continuing process rather than a one-time classification. Workers and organizations may recalibrate what an AI is permitted to do, what remains a human decision, how much reliance is appropriate, and whether the system is described socially or instrumentally. Such calibration concerns enacted work relations; it does not convert an artificial system into an organizational member.

Boundary management becomes particularly important when AI recommendations are opaque or difficult to reconcile with professional knowledge. Research on consequential diagnostic work shows that professionals respond differently to opacity and remain actively involved in deciding when and how to engage with AI recommendations [19]. This supports a distinction between technological capability and legitimate decisional authority. An AI may produce a recommendation that influences work without possessing responsibility for accepting, contextualizing, contesting, or escalating that recommendation. Expertise, task stakes, institutional norms, and access to alternative evidence remain plausible explanations for variation in human engagement.

Human involvement also creates collective trade-offs. Evidence from human–AI decision settings indicates that AI advice can improve individual accuracy while reducing the distinctiveness of human knowledge under some configurations [20]. Relationally comfortable collaboration should therefore not be presumed beneficial merely because users experience less friction. A coworker-like construal might facilitate cooperation, but it could also increase convergence, inhibit independent checking, or make advice feel socially compelling. TACR treats these consequences as contingent possibilities rather than expected universal outcomes.

The proposed boundary-calibration mechanism consequently spans role definition, decision rights, reliance, responsibility, disclosure, human override, and escalation. Agentic-artifact theory supplies a basis for analyzing how functions can be delegated without requiring human equivalence [6]. TACR extends that reasoning by proposing that relational interpretation itself becomes an object of governance: organizations may need to distinguish how an AI is socially framed from what it is formally authorized to decide. Whether explicit calibration reduces inappropriate reliance or relational confusion remains unvalidated and should be tested rather than asserted.

Proposed theory of artificial-coworker relations

TACR integrates dyadic, team, and system perspectives. Experimental work shows that cooperation with an AI teammate depends on characteristics such as performance and user expertise; team research indicates that trust varies with human–AI team composition; and organization theory emphasizes that AI collaboration is embedded in wider sociotechnical arrangements [21–23]. Together, these streams imply that an artificial coworker cannot be understood solely from interface properties or nominal labels. Relational interpretation is situated within performance histories, team structures, work roles, and organizational systems.

Figure 1 presents the original conceptual synthesis for theory of artificial-coworker relations, showing how the article’s principal constructs, mechanisms, uncertainties, and decision boundaries are connected.

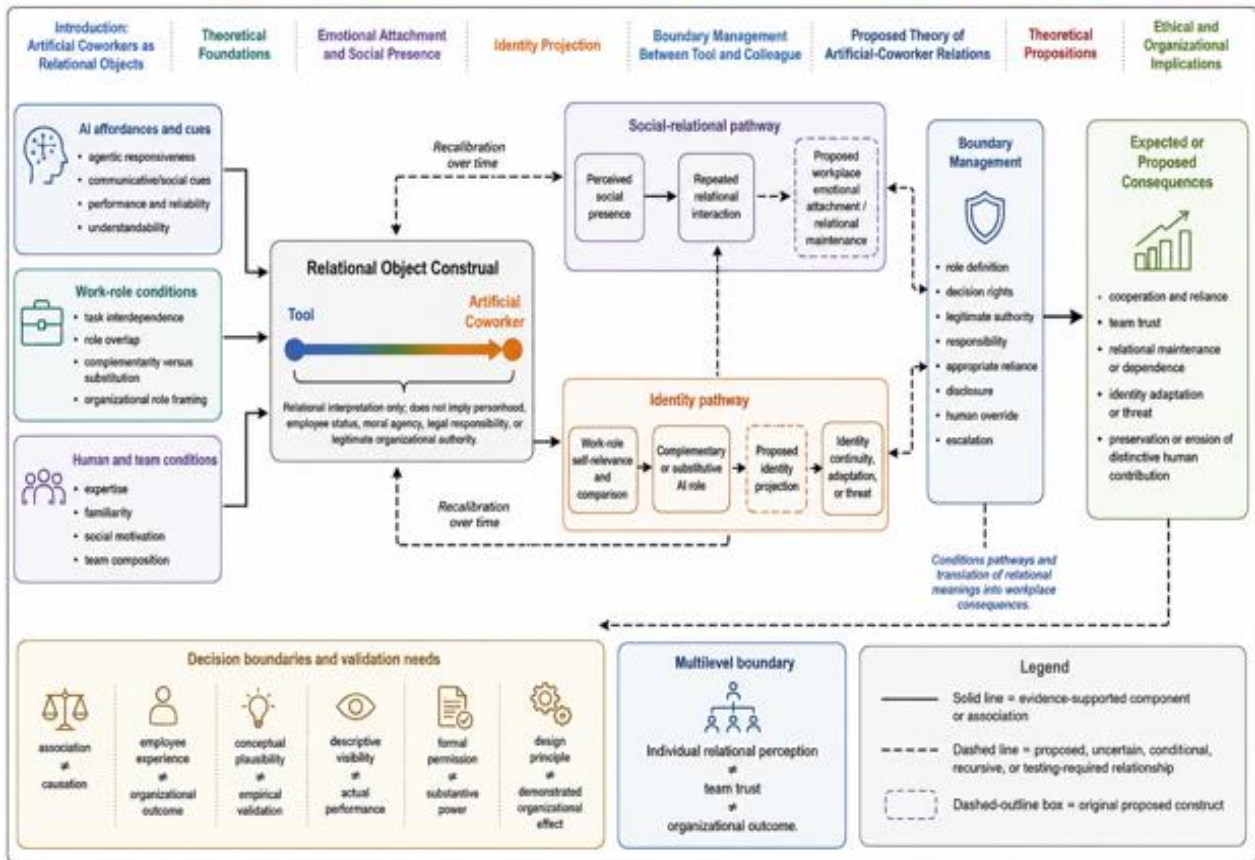


Figure 1. Theory of Artificial-Coworker Relations

A further boundary concerns variation among AI forms. Contemporary organization theory distinguishes predictive, generative, agentic, and embodied forms of AI because they redistribute work and agency differently [24]. TACR therefore predicts no generic “AI effect.” Systems that merely supply occasional predictions should be less likely to acquire coworker-like relational meaning than systems engaged in recurrent communication, task coordination, or adaptive action, other conditions equal. This comparison remains a theoretical expectation, not a demonstrated hierarchy.

At the center of TACR is the proposed construct of relational object construal: workers interpret an AI along a continuum from instrument to artificial coworker while retaining the possibility of mixed or unstable meanings. Delegation theory helps explain why an agentic system may become consequential to role allocation [6], but TACR adds the proposition that workers also interpret what that functional participation means socially. Construal is therefore neither an objective property of the system nor a claim about consciousness. It is a worker-level interpretation embedded in organizational arrangements.

From this center, TACR proposes two analytically distinct pathways. The social-relational pathway links perceived social presence and repeated interaction to the possibility of emotional attachment or relational-maintenance motivation. Workplace evidence establishes the relevance of social presence in human–AI collaboration [13], while the attachment step remains proposed. The identity pathway links task overlap, complementarity, substitution, and comparison with proposed identity projection. The pathways may intersect—for example, an identity-threatening system could still be socially engaging—but neither is necessary for the other.

TACR further proposes feedback. Performance histories, breakdowns, role changes, and organizational reframing may revise relational object construal over time. Cooperation with an AI teammate has been shown to vary with performance and expertise [21], supporting performance as one empirical anchor without establishing the proposed recursive process. Boundary

calibration is theorized to condition whether relational meaning becomes reliance, authority attribution, or dependence. Individual construal must remain distinct from team trust and organizational outcomes.

Table 2 translates the proposed Theory of Artificial-Coworker Relations into a process architecture that separates antecedent conditions, relational mechanisms, possible consequences, rival explanations, and requirements for empirical testing.

Table 2. Process architecture and testable boundaries of the proposed Theory of Artificial-Coworker Relations

Theory stage	Triggering condition	Proposed relational process	Possible workplace manifestation	Competing interpretation	Critical boundary	Empirical question
1. Relational activation	Recurrent interdependence, responsiveness, or coordination with AI	AI becomes salient as more than a passive instrument	Workers address, consult, or anticipate AI as a recurring work participant	High usage may reflect superior utility	Use intensity is not evidence of relational construal	Does construal explain behavior beyond usefulness and performance?
2. Social interpretation	Socially interpretable interaction and repeated communication	Perceived social presence develops	Interaction feels reciprocal, responsive, or socially consequential	Familiarity or fluent design may produce similar experience	Social presence is distinct from attachment	Does social presence persist after controlling novelty and interface quality?
3. Relational deepening	Continuity, predictability, self-disclosure, or personalization	Proposed emotional attachment or relational-maintenance motivation	Workers prefer continuity with the same AI or experience disruption when it changes	Habit, convenience, or functional dependence	Attachment must be measured directly	Do affective bonds predict behavior beyond habit and utility?
4. Identity engagement	AI overlaps with valued tasks, expertise, or professional roles	Proposed identity projection	Workers compare themselves with AI or reinterpret competence and role meaning	Job insecurity or ordinary technological change	Projection must be distinguished from threat alone	Does AI-related self-interpretation form a distinct construct?
5. Boundary negotiation	Ambiguous roles, opaque recommendations, or overlapping decision rights	Proposed boundary calibration	Workers redefine when to rely, override, disclose, contest, or escalate	Formal procedures may explain behavior	Relational meaning does not confer authority	Do explicit boundaries moderate reliance independently of trust?
6. Relational-authority divergence	Strong relational salience with unclear accountability	Proposed relational-authority asymmetry	Socially influential AI is treated as more authoritative than warranted	High perceived competence may explain deference	Influence and legitimate authority remain separate	When does social influence exceed assigned decision rights?
7. Consequence formation	Repeated relational and identity processes	Cooperation, dependence, identity adaptation, resistance, or knowledge convergence may emerge	Checking behavior, trust, contribution, and role perceptions change	Task difficulty, incentives, workload, or management pressure	Individual experience cannot establish organizational performance	Which consequences persist across occupations and technologies?
8. Recalibration over time	Performance changes, failures, redesign, role change, or reframing	Relational object construal is revised	AI moves toward the tool or coworker end of the continuum	Learning about system accuracy alone	Temporal change does not prove recursion	Do relational meanings predict later boundary changes and vice versa?

Table note. All directional processes identified as proposed remain theoretical relationships requiring empirical testing. Possible workplace manifestations are illustrative theoretical expectations rather than demonstrated outcomes, effect estimates, thresholds, or implementation claims.

Theoretical propositions

TACR's propositions convert the conceptual architecture into falsifiable claims while preserving contextual limits. First, task characteristics should constrain whether functional participation becomes relationally consequential. Experimental research demonstrates task-dependent algorithm aversion, with algorithm acceptance varying according to features such as perceived task subjectivity [25]. This does not establish an attachment mechanism, but it indicates that workers' responses to AI should

not be modeled independently of what the task represents. TACR therefore expects relational construal to vary across work that is objective, interpretive, identity-laden, or socially interdependent.

Second, retained human agency should affect how coworker-like interpretation translates into reliance. Experiments show that allowing people even limited ability to modify algorithmic forecasts can increase willingness to use an imperfect algorithm [26]. TACR extends this finding cautiously: override and modification rights may help separate relational comfort from unconditional deference because workers retain a recognized capacity to intervene. The proposition concerns boundary conditions on reliance, not a claim that user control always improves decisions or eliminates automation bias.

Third, human response to algorithms is not inherently aversive. Experimental evidence also documents algorithm appreciation, while showing that willingness to follow algorithmic advice varies with expertise and comparison conditions [27]. This competing pattern is important because artificial-coworker relations could emerge under positive, negative, or ambivalent evaluations. TACR does not require workers to like AI; it requires only that the system becomes sufficiently consequential to work and self-interpretation for a relational meaning to be constructed.

Proposition 1—Relational-Construal Proposition. Recurrent interdependence with AI possessing agentic and socially interpretable affordances will increase the likelihood of coworker-oriented relational object construal, conditional on role framing, perceived capability, and task context.

Proposition 2—Social-Presence/Attachment Proposition. Greater social presence across repeated work interaction will increase the likelihood of proposed emotional attachment or relational-maintenance motivation, conditional on familiarity, understandability, continuity, and social motivation. Social presence provides the empirical anchor [13]; the workplace attachment transition remains proposed.

Proposition 3—Identity-Projection Proposition. Greater overlap between AI activity and identity-relevant work roles will increase the likelihood that AI becomes a site of proposed identity projection, conditional on whether its role is complementary, substitutive, evaluative, or controlling. Professional-role identity evidence provides an adjacent foundation [14].

Proposition 4—Boundary-Calibration Proposition. Explicit differentiation of decision rights, responsibility, override, and escalation will condition how relational construal translates into reliance. Evidence that human–AI work configurations evolve provides an empirical anchor for the boundary premise [18], while the specific moderating relationship remains proposed.

Proposition 5—Relational-Authority Asymmetry Proposition. When strong relational interpretation coexists with weakly specified human accountability, social influence may exceed the AI's legitimate organizational authority. Each proposition requires empirical designs capable of rejection, temporal-order testing, construct discrimination, and control of alternative explanations rather than confirmatory illustration alone.

Ethical and organizational implications

Relational AI introduces ethical questions that are not captured by accuracy or productivity alone. Qualitative research on a social chatbot documents reported harms associated with emotional dependence and perceived reciprocal relationships [28]. These findings cannot establish workplace prevalence, causation, or equivalent employee responses, but they identify a credible boundary condition for organizational theory. If work systems are deliberately made socially responsive, organizations should at least treat emotional dependence, disclosure, and vulnerability as possible relational consequences rather than assuming that anthropomorphic engagement is harmless.

Accountability creates a separate constraint. Ethical analysis emphasizes that algorithms are value-laden artifacts embedded in human and organizational choices and that responsibility cannot simply disappear when decision functions are delegated to technology [29]. TACR therefore separates relational presence from legitimate authority. A worker may experience an AI as supportive, familiar, or coworker-like while the organization remains responsible for defining decision rights, review obligations, escalation routes, and the limits of machine recommendations.

The implications are also multilevel. Organizational AI research spans individual reactions, team processes, work design, control, and organization-level arrangements [30]. An employee's attachment to an AI does not demonstrate team trust; team trust does not establish improved performance; and widespread use does not demonstrate legitimate authority. Research and governance should preserve these levels rather than treating visible adoption or favorable attitudes as evidence of organizational benefit.

Accordingly, TACR identifies role disclosure, responsibility allocation, relational-design scrutiny, appropriate override, and escalation as proposed governance principles rather than validated interventions. Social presence may be useful for collaboration [13], yet socially persuasive design may also complicate authority boundaries. Organizations should therefore distinguish the question “Does the system feel like a colleague?” from “What may it legitimately decide?” and “Who remains answerable?” The theory does not prescribe one universal boundary; it specifies questions whose answers should depend on task stakes, workforce context, technology type, and empirically demonstrated consequences.

Conclusion

Artificial intelligence can participate in work in ways that are simultaneously functional, social, symbolic, and identity-relevant. The proposed Theory of Artificial-Coworker Relations explains this possibility without asserting that AI systems are literally coworkers. Its central claim is that workers may construe an AI along a tool–artificial-coworker boundary and that this relational object construal may organize two distinct pathways: a social-relational pathway involving social presence and possible emotional attachment, and an identity pathway involving role comparison and proposed identity projection. Boundary calibration is proposed to connect these interpretations to reliance, authority, responsibility, override, and escalation.

The theory's contribution is integrative and conditional. It separates anthropomorphism from social presence, social presence from attachment, trust from relational dependence, professional-role identity from proposed identity projection, and functional technological agency from legitimate organizational authority.

TACR remains an original conceptual theory requiring empirical challenge. Its distinctive constructs need discriminant validation, its proposed pathways require longitudinal and experimental tests, and its multilevel implications require evidence that does not infer organizational consequences from individual perception. Artificial-coworker relations should therefore be understood as a researchable organizational possibility—not a validated causal model, universal description of human–AI work, or ready-made managerial prescription.

Acknowledgments: None

Conflict of interest: None

Financial support: None

Ethics statement: None

References

1. Faraj S, Pachidi S, Sayegh K. Working and organizing in the age of the learning algorithm. *Inf Organ.* 2018;28(1):62-70. doi:10.1016/j.infoandorg.2018.02.005
2. Seeber I, Bittner EAC, Briggs RO, de Vreede T, de Vreede GJ, Elkins A, et al. Machines as teammates: A research agenda on AI in team collaboration. *Inf Manag.* 2020;57(2):103174. doi:10.1016/j.im.2019.103174
3. Glikson E, Woolley AW. Human trust in artificial intelligence: Review of empirical research. *Acad Manag Ann.* 2020;14(2):627-60. doi:10.5465/annals.2018.0057
4. Raisch S, Krakowski S. Artificial intelligence and management: The automation-augmentation paradox. *Acad Manag Rev.* 2021;46(1):192-210. doi:10.5465/amr.2018.0072
5. Kellogg KC, Valentine MA, Christin A. Algorithms at work: The new contested terrain of control. *Acad Manag Ann.* 2020;14(1):366-410. doi:10.5465/annals.2018.0174
6. Baird A, Maruping LM. The next generation of research on IS use: A theoretical framework of delegation to and from agentic IS artifacts. *MIS Q.* 2021;45(1):315-41. doi:10.25300/MISQ/2021/15882
7. Araujo TB. Living up to the chatbot hype: The influence of anthropomorphic design cues and communicative agency framing on conversational agent and company perceptions. *Comput Human Behav.* 2018;85:183-9. doi:10.1016/j.chb.2018.03.051
8. Go E, Sundar SS. Humanizing chatbots: The effects of visual, identity and conversational cues on humanness perceptions. *Comput Human Behav.* 2019;97:304-16. doi:10.1016/j.chb.2019.01.020
9. Blut M, Wang C, Wunderlich NV, Brock C. Understanding anthropomorphism in service provision: A meta-analysis of physical robots, chatbots, and other AI. *J Acad Mark Sci.* 2021;49(4):632-58. doi:10.1007/s11747-020-00762-y
10. Skjuve M, Følstad A, Fostervold KI, Brandtzaeg PB. My chatbot companion - a study of human-chatbot relationships. *Int J Hum Comput Stud.* 2021;149:102601. doi:10.1016/j.ijhcs.2021.102601
11. Croes EAJ, Antheunis ML. Can we be friends with Mitsuku? A longitudinal study on the process of relationship formation between humans and a social chatbot. *J Soc Pers Relat.* 2021;38(1):279-300. doi:10.1177/0265407520959463
12. Pentina I, Hancock T, Xie T. Exploring relationship development with social chatbots: A mixed-method study of Replika. *Comput Human Behav.* 2023;140:107600. doi:10.1016/j.chb.2022.107600
13. Siemon D, Elshan E, de Vreede T, Ebel P, de Vreede GJ. Beyond anthropomorphism: Social presence in human-AI collaboration processes. *J Manag Stud.* 2026;63(2):515-60. doi:10.1111/joms.70000

14. Strich F, Mayer AS, Fiedler M. What do I do in a world of artificial intelligence? Investigating the impact of substitutive decision-making AI systems on employees' professional role identity. *J Assoc Inf Syst.* 2021;22(2):304-24. doi:10.17705/1jais.00663
15. Pachidi S, Berends H, Faraj S, Huysman M. Make way for the algorithms: Symbolic actions and change in a regime of knowing. *Organ Sci.* 2021;32(1):18-41. doi:10.1287/orsc.2020.1377
16. Beane M. Shadow learning: Building robotic surgical skill when approved means fail. *Adm Sci Q.* 2019;64(1):87-123. doi:10.1177/0001839217751692
17. Selenko E, Bankins S, Shoss M, Warburton J, Restubog SLD. Artificial intelligence and the future of work: A functional-identity perspective. *Curr Dir Psychol Sci.* 2022;31(3):272-9. doi:10.1177/09637214221091823
18. Grønsund TR, Aanestad M. Augmenting the algorithm: Emerging human-in-the-loop work configurations. *J Strateg Inf Syst.* 2020;29(2):101614. doi:10.1016/j.jsis.2020.101614
19. Lebovitz S, Lifshitz-Assaf H, Levina N. To engage or not to engage with AI for critical judgments: How professionals deal with opacity when using AI for medical diagnosis. *Organ Sci.* 2022;33(1):126-48. doi:10.1287/orsc.2021.1549
20. Fügener A, Grahl J, Gupta A, Ketter W. Will humans-in-the-loop become Borgs? Merits and pitfalls of working with AI. *MIS Q.* 2021;45(3):1527-56. doi:10.25300/MISQ/2021/16553
21. Zhang G, Chong L, Kotovsky K, Cagan J. Trust in an AI versus a human teammate: The effects of teammate identity and performance on human-AI cooperation. *Comput Human Behav.* 2023;139:107536. doi:10.1016/j.chb.2022.107536
22. Georganta E, Ulfert AS. Would you trust an AI team member? Team trust in human-AI teams. *J Occup Organ Psychol.* 2024;97(3):1212-41. doi:10.1111/joop.12504
23. Anthony C, Bechky BA, Fayard AL. "Collaborating" with AI: Taking a system view to explore the future of work. *Organ Sci.* 2023;34(5):1672-94. doi:10.1287/orsc.2022.1651
24. Chalmers D, Hunt R, Pachidi S, Potočník K, Townsend DM. The acceleration of artificial intelligence: Rethinking organization and work in an era of rapid technological change. *J Manag Stud.* 2026;63(2):285-314. doi:10.1111/joms.70063
25. Castelo N, Bos MW, Lehmann DR. Task-dependent algorithm aversion. *J Mark Res.* 2019;56(5):809-25. doi:10.1177/0022243719851788
26. Dietvorst BJ, Simmons JP, Massey C. Overcoming algorithm aversion: People will use imperfect algorithms if they can (even slightly) modify them. *Manage Sci.* 2018;64(3):1155-70. doi:10.1287/mnsc.2016.2643
27. Logg JM, Minson JA, Moore DA. Algorithm appreciation: People prefer algorithmic to human judgment. *Organ Behav Hum Decis Process.* 2019;151:90-103. doi:10.1016/j.obhdp.2018.12.005
28. Laestadius L, Bishop A, Gonzalez M, Illenčík D, Campos-Castillo C. Too human and not human enough: A grounded theory analysis of mental health harms from emotional dependence on the social chatbot Replika. *New Media Soc.* 2024;26(10):5923-41. doi:10.1177/14614448221142007
29. Martin K. Ethical implications and accountability of algorithms. *J Bus Ethics.* 2019;160(4):835-50. doi:10.1007/s10551-018-3921-3
30. Bankins S, Ocampo ACG, Marrone M, Restubog SLD, Woo SE. A multilevel review of artificial intelligence in organizations: Implications for organizational behavior research and practice. *J Organ Behav.* 2024;45(2):159-82. doi:10.1002/job.2735