

E-ISSN: 3108-4192

APSSHS

Academic Publications of Social Sciences and Humanities Studies

2024, Volume 4, Issue 2, Page No: 65-73

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

Journal of Applied Organizational Systems and Behavior

How Algorithmic Certainty Weakens Dissent, Experimentation, and Organizational Learning

Elena Petrova^{1*}, Ivan Georgiev¹, Nikolay Stoyanov², Petar Kolev³

1. Department of Algorithmic Certainty and Dissent, Faculty of Business, Sofia University, Sofia, Bulgaria.
2. Department of Organizational Learning and Experimentation, Faculty of Management, University of Plovdiv, Plovdiv, Bulgaria.
3. Department of Innovation and Risk, Faculty of Psychology, University of Ruse, Ruse, Bulgaria.

Abstract

Organizations increasingly rely on algorithmic systems because prediction appears to offer consistency, speed, scalability, and protection from human error. Yet predictive performance can acquire a broader organizational meaning: uncertain outputs may be enacted as settled answers, managerial authority may become attached to model recommendations, and disagreement may be reframed as resistance to objective evidence. This critical review examines how such algorithmic certainty may narrow inquiry, weaken the consequential influence of employee dissent, displace experimentation, and obstruct learning from error. The review integrates organizational, information-systems, work-design, voice, experimentation, and human–AI collaboration research. It evaluates convergent mechanisms, conflicting findings, alternative explanations, and level-of-analysis limitations rather than treating algorithms as uniformly controlling or employees as passive. The synthesis suggests that prediction becomes hazardous to learning when confidence cues, opaque evaluation, concentrated decision rights, and metric dependence reduce the legitimacy of questions, challenges, and locally generated alternatives. However, professional expertise, psychological safety, job autonomy, transparent uncertainty, and contestable decision processes may preserve inquiry and corrective action. The article's original contribution is a Certainty–Learning Framework that connects predictive outputs to organizational enactment, dissent, experimentation, error interpretation, and updating. The framework is a critical integrative synthesis, not a validated causal model. Longitudinal, comparative, multilevel, and intervention research is required to distinguish algorithm effects from pre-existing hierarchy, culture, task structure, and implementation choices. The implications remain context dependent.

Keywords: Algorithmic certainty, Employee dissent, Organizational learning, Experimentation, Algorithmic management, Critical review

How to cite this article: Petrova E, Georgiev I, Stoyanov N, Kolev P. How Algorithmic Certainty Weakens Dissent, Experimentation, and Organizational Learning. *J Appl Organ Syst Behav*. 2024;4(2):65-73. <https://doi.org/10.51847/ijMh6xn7N>

Received: 02 November 2023; **Revised:** 08 February 2024; **Accepted:** 13 February 2024

Corresponding author: Elena Petrova

E-mail ✉ elena.petrova@feb.uni-sofia.bg

Introduction

Algorithmic systems appeal to organizations because they promise to convert complexity into rankings, classifications, forecasts, and recommended actions. Their organizational significance, however, extends beyond computational accuracy. Algorithms can redistribute direction, evaluation, discipline, and expertise, producing new sites of control and contestation rather than merely automating isolated tasks [1–3]. The central problem is therefore not prediction itself, but the organizational process through which a provisional output becomes treated as an authoritative account of what is happening, what matters, and what should occur next. In this article, algorithmic certainty denotes this enacted condition: confidence in a prediction



© 2024 The Author(s).

Copyright CC BY-NC-SA 4.0

becomes embedded in routines, decision rights, evaluation systems, and expectations of compliance, even though the underlying model remains probabilistic and context dependent.

The prevailing framing is insufficient in two ways. First, debate often opposes automation to human judgment as though organizations must select one source of intelligence. The automation–augmentation paradox instead suggests that greater machine capability can increase the need for human interpretation, coordination, and exception handling. Second, rationality claims can conceal contestable assumptions about objectives, categories, acceptable errors, and whose knowledge is legitimate. When algorithmic recommendations are presented as technically settled, disagreement may be treated as noise, self-interest, or incompetence rather than information about model limits or local conditions [4]. Certainty thus becomes a property of organizational enactment, not an inherent property of the technology.

Trust further complicates the problem. Empirical research indicates that reliance on artificial intelligence depends on performance cues, process characteristics, purpose, embodiment, and user attributes; accuracy alone does not determine warranted trust [5]. High trust may support efficient coordination, but poorly calibrated trust can also discourage independent checking. Conversely, skepticism can preserve scrutiny but prevent useful augmentation. The critical question is not whether employees trust algorithms, but whether organizational arrangements allow them to question outputs, explain divergence, test alternatives, and influence consequential decisions.

This critical review examines how algorithmic certainty may weaken dissent, experimentation, and organizational learning. It makes three contributions. It distinguishes prediction from enacted certainty; it connects literatures that are usually treated separately, including regimes of knowing, employee voice, work design, experimentation, error learning, and human–AI delegation; and it develops an original Certainty–Learning Framework. The review does not claim that algorithms universally suppress voice or cause learning failure. It asks when prediction becomes organizational closure, which mechanisms are supported, where evidence conflicts, and what boundary conditions must constrain theory and governance.

Critical review method

The article used a transparent critical-review design intended to evaluate dominant assumptions, mechanisms, contradictions, and omissions rather than estimate a pooled effect. Critical reviews are suited to fields where constructs and levels of analysis are fragmented and where the contribution depends on problematization and integration rather than serial summary [6]. Review questions addressed how predictive outputs acquire authority, how inquiry and dissent are conditioned, how experimentation is displaced or preserved, and how these processes affect error interpretation and organizational updating.

Information sources comprised publisher and DOI-indexed journal searches, multidisciplinary scholarly-index and keyword searches, and backward and forward citation searching. Search concepts combined algorithmic control, algorithmic certainty, prediction, opacity, expert judgment, employee voice, dissent, experimentation, error learning, work design, delegation, and governance. Eligible sources were peer-reviewed Q1 journal articles with verifiable DOIs and a direct contribution to a planned construct, relationship, boundary, table cell, or framework component. Screening removed duplicates and excluded adjacent or redundant evidence lacking a distinct organizational-learning contribution.

The search log contained 52 records: 31 from publisher and DOI-indexed searches, 14 from multidisciplinary scholarly-index searches, and seven from citation-based searching. Four duplicates were removed, leaving 48 records screened. Seven were excluded at title and abstract. Forty-one reports were retrieved and assessed in full text; eight were excluded for substantive overlap, adjacent individual-judgment evidence without a direct organizational-learning claim, review redundancy, or diffuse unique contribution. The final synthesis included 33 sources. PRISMA-compatible reporting was used to make these transitions inspectable [7], while search-source and deduplication documentation followed search-reporting principles [8].

Figure 1 reports the verified flow of records through identification, duplicate removal, screening, full-text assessment, exclusion, and final inclusion for this critical review article.

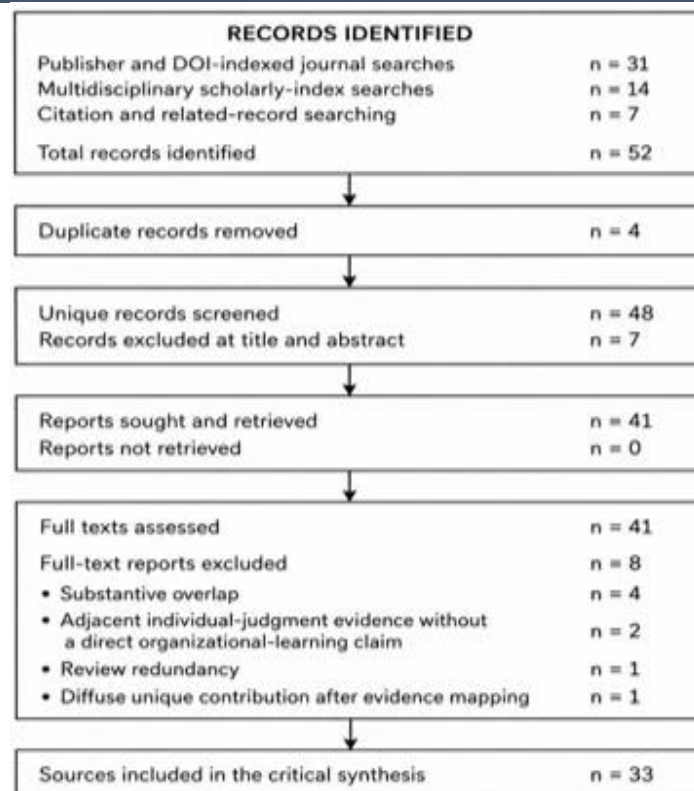


Figure 1. PRISMA Flow Diagram for Evidence Identification and Selection

Extraction recorded context, unit, design, focal construct, mechanism, moderator, outcome, limitation, and permitted claim. Appraisal emphasized direct claim fit, design–inference alignment, level-of-analysis discipline, and transferability. The synthesis compared convergence, conflict, and missing links, then developed an integrative argument while keeping evidence description, critical interpretation, and original propositions distinct. Contribution-led synthesis was used instead of study-by-study narration [9].

Prediction and the narrowing of inquiry

Prediction narrows inquiry when an algorithm’s categories and target variables become the boundaries of legitimate organizational attention. This narrowing is not equivalent to opacity. A transparent model can still privilege a limited objective, whereas an opaque model may be interrogated through strong professional routines. The relevant mechanism is interpretive closure: questions that fall outside the model’s representation become harder to formulate, justify, or escalate. Longitudinal qualitative evidence shows that introducing algorithms can change a regime of knowing by altering which forms of expertise and evidence are recognized in practice [10]. The finding supports a relational interpretation: algorithms do not simply replace knowledge but participate in negotiations over legitimacy. Black-boxing research similarly demonstrates that organizational practices can separate users from assumptions embedded in analytical technologies, changing who performs interpretive work and who can challenge an output [11]. Yet opacity does not produce one uniform response. Professionals may selectively engage, compare outputs with experiential cues, or decline engagement when model reasoning cannot be reconciled with critical judgment [12]. Expertise and discretion therefore moderate narrowing rather than merely delaying acceptance.

Experimental evidence adds a different form of support. Decision-makers can be influenced by erroneous artificial-intelligence recommendations, indicating that an output can anchor judgment even where independent assessment remains formally available [13]. Such susceptibility is consistent with inquiry narrowing but does not itself establish organizational closure. Alternative explanations include time pressure, interface framing, perceived institutional endorsement, and task-specific uncertainty. Furthermore, clinical and professional settings may contain stronger checking routines than platform work, while highly standardized jobs may offer fewer legitimate grounds for deviation.

The critical synthesis therefore treats narrowing as conditional and multilevel. At the individual level, it may appear as anchoring or reduced search. At the occupational level, it may alter the standing of experiential knowledge. At the organizational level, it may become embedded in metrics, escalation rules, and resource allocation. Evidence is strongest for changes in knowledge practices and judgment under particular conditions; it is weaker for long-term organizational learning outcomes.

Table 1 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for review design, evidence domains, and extraction framework for how algorithmic certainty weakens dissent, experimentation, and organizational learning.

Table 1. Review design, evidence domains, and extraction framework for How Algorithmic Certainty Weakens Dissent, Experimentation, and Organizational Learning.

Evidence domain or review construct	Review question	Eligible evidence type	Extraction fields	Appraisal or relevance criterion	Main limitation	Interpretation boundary
Enacted algorithmic certainty	How does prediction acquire authority?	Organizational theory, qualitative studies, empirical reviews	Framing, confidence cues, decision rights, trust	Distinguishes technical output from organizational enactment	Mostly indirect measurement	Certainty is proposed as enacted, not measured as a stable trait
Narrowing of inquiry	What knowledge becomes legitimate or excluded?	Longitudinal and comparative qualitative research; experiments	Knowledge regime, opacity, engagement, susceptibility	Identifies observable changes in questioning or judgment	Context and task heterogeneity	Influence on judgment does not prove organizational closure
Employee dissent	Does voice retain consequential influence?	Voice research; algorithmic-control field studies	Voice form, risk, power, response to challenge	Separates speaking opportunity from decision influence	Direct algorithm-specific voice measures remain scarce	Silence, compliance, and constrained agency are not interchangeable
Experimentation	Is variation authorized, displaced, or hidden?	Ethnography, experiments, innovation research	Trials, prototypes, workarounds, idea generation and selection	Distinguishes inquiry from unauthorized deviation	Evidence spans unlike settings	Experimentation is not inherently beneficial or safe
Error learning	Can anomalies challenge assumptions and routines?	Learning reviews, work-design research, delegation experiments	Feedback, autonomy, attribution, updating	Requires usable feedback and capacity for revision	Organizational updating is seldom observed directly	Error reduction is not equivalent to learning
Governance	Which arrangements preserve contestability?	Decision-structure, accountability, and human-centered design research	Appeal, override, audit, responsibility, escalation	Identifies mechanisms rather than abstract principles	Effectiveness evidence remains incomplete	Governance principles are not validated outcomes

Note. Evidence-derived entries are linked to approved references. “Interpretive closure,” “enacted algorithmic certainty,” and the cross-domain synthesis roles are original proposed constructs developed for this review.

Suppression of employee dissent

Employee dissent refers here to the expression of information, concern, or disagreement that challenges an endorsed course of action. Suppression is not limited to formal censorship. It also occurs when employees can speak but their input cannot alter classifications, priorities, or decisions. This distinction is essential because algorithmic systems may expand channels for reporting while simultaneously concentrating authority in metrics, rankings, or model-mediated escalation.

Field research in online work settings shows that opaque algorithmic control can create substantial power asymmetries while leaving workers with constrained forms of agency [14]. Workers interpret signals, adapt behavior, and sometimes resist, but they often lack access to evaluation criteria or meaningful negotiation over outcomes. Longitudinal evidence on opaque evaluations further indicates that workers may redirect effort toward inferred metrics, producing self-discipline around what they believe the system rewards [15]. Such reactivity can reduce overt challenge, yet it can also represent strategic adaptation rather than genuine agreement. The evidence therefore supports constrained influence more clearly than complete silence.

Voice research provides construct precision. Promotive voice proposes improvements, whereas prohibitive voice identifies harmful practices or risks; the two forms have distinguishable antecedents and consequences [16]. Algorithmic certainty may affect them differently. Promotive suggestions can be absorbed as optimization inputs, while prohibitive dissent questions whether the model, objective, or decision rule should govern at all. The latter poses a stronger challenge to institutionalized certainty and may therefore carry greater perceived risk.

Psychological safety is associated with speaking up and learning-related behavior, but much of this evidence is correlational and not algorithm specific [17]. Safety may encourage expression without ensuring that decision-makers treat dissent as consequential evidence. Conversely, low voice may reflect hierarchy, employment insecurity, occupational status, time pressure, or prior punitive climates rather than the algorithm itself. A defensible interpretation is therefore conditional:

algorithmic certainty may suppress dissent when opaque evaluation, concentrated decision rights, dependency, and weak appeal mechanisms combine. It may be less constraining where employees possess recognized expertise, protected escalation routes, collective support, and authority to pause or test a recommendation.

The organizational-learning consequence depends on what happens after voice. A concern that is recorded but cannot trigger review, experimentation, or revision is visible without being influential. Accordingly, this review treats consequential influence—not message frequency—as the critical bridge between dissent and learning.

Loss of experimentation

Experimentation is the deliberate production of informative variation through trials, prototypes, comparisons, or reversible deviations. Its loss should not be equated with the disappearance of all variation. Under tightly specified algorithmic routines, experimentation may move outside authorized workflows, become oriented toward improving a prescribed metric, or survive only as informal workaround activity. Ethnographic evidence from robotic surgery shows that when approved training arrangements impede access to practice, skill development can migrate into “shadow learning” that is organizationally necessary yet weakly governed [18]. This finding supports displacement, not a general claim that formal control eliminates learning.

Authority also matters differently across stages of innovation. Experimental research indicates that hierarchy can have distinct effects on idea generation and idea selection [19]. Applied to algorithmic certainty, employees may still generate alternatives but anticipate that only model-compatible proposals will survive selection. Conversely, hierarchy may accelerate selection without necessarily reducing initial ideation. The evidence therefore cautions against treating experimentation as a single outcome. Organizations can preserve the appearance of ideation while narrowing which hypotheses receive resources, authorization, or evaluation.

A randomized trial of entrepreneurial decision making found that explicit hypotheses and structured testing can improve decision discipline under uncertainty [20]. This evidence establishes the value of organized experimentation in its setting, not that algorithms generally prevent it. Research on design thinking likewise associates iteration with cultural supports for empathy, prototyping, and tolerance of incompleteness [21]. Pre-existing culture, regulation, task stakes, resource scarcity, and safety requirements may explain experimentation patterns independently of algorithmic systems. The critical mechanism is thus whether prediction closes questions before alternatives can be tested.

Evidence also conflicts over whether constraint always damages innovation. Standardized evaluation can focus scarce resources, prevent unsafe variation, and make comparisons more interpretable. Conversely, visible experimentation programs may become ceremonial when hypotheses, acceptable outcomes, and termination criteria are fixed in advance. An adequate assessment must therefore examine the decision rights surrounding trials: who can propose them, which alternatives are admissible, what counts as disconfirming evidence, and whether unsuccessful results can revise the governing model.

Table 2 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for cross-study findings, mechanisms, convergence, and conflict in how algorithmic certainty weakens dissent, experimentation, and organizational learning.

Table 2. Cross-study findings, mechanisms, convergence, and conflict in How Algorithmic Certainty Weakens Dissent, Experimentation, and Organizational Learning.

Construct, mechanism, or relationship	Representative context	Reported finding	Evidence design	Convergent evidence	Conflicting evidence	Methodological concern	Review interpretation
Formal constraint and shadow experimentation	Robotic-surgery learning	Restricted approved practice displaced learning into informal arrangements [18]	Comparative ethnography	Experimentation may persist outside formal authorization	Informal learning can preserve capability rather than weaken it	Specialized, high-skill setting	Certainty may relocate experimentation instead of eliminating it
Hierarchy across innovation stages	Group innovation tasks	Authority affected idea generation and selection differently [19]	Controlled experiments	Supports stage-specific constraint	Hierarchy may aid selection or coordination	Nonalgorithmic task context	Examine generation, authorization, selection, and retention separately
Hypothesis-driven testing	Entrepreneurial decisions	Structured experiments improved decision discipline under uncertainty [20]	Randomized trial	Supports testing as learning infrastructure	Effects depend on intervention and context	Transferability to established organizations	Algorithmic forecasts should remain testable hypotheses

Iterative culture	Design-thinking settings	Prototyping depends on	Integrative review	Culture conditions experimentation	Design practices are not uniformly effective	Heterogeneous studies and limited causal tests	Technology effects are conditional on enacted culture
		cultural support for incompleteness and iteration [21]					

Note. The cross-domain categories and review interpretations are original synthesis. The table does not rank evidence or imply that experimentation is universally beneficial.

Error avoidance and organizational learning failure

Error avoidance seeks reliability by preventing deviations, whereas learning from error requires anomalies to be noticed, discussed, interpreted, and used to revise assumptions or routines. These aims can complement each other, but certainty can make them conflict. A synthesis of failure research argues that learning requires opportunity, motivation, and ability; error data alone are insufficient. Digital work-design research further shows that autonomy, feedback, skill use, and task structure condition whether employees can interpret and act on information, while algorithmic management may redesign these characteristics in enabling or constraining ways [22–24]. The grouped evidence supports organizational conditions for learning, not a causal sequence applicable to every technology.

Algorithmic certainty can weaken those conditions when deviations are classified primarily as compliance problems. Employees may correct outputs without examining why the model, workflow, or environment produced the discrepancy. Learning then becomes first-order adjustment within accepted parameters rather than reconsideration of objectives or categories. Yet strict error controls may be legitimate in high-hazard work, and standardization can improve coordination. The relevant boundary is whether reliability arrangements preserve protected routes for anomaly escalation, causal inquiry, and controlled testing.

Human–AI delegation introduces another risk. Experiments on productive delegation show that joint performance depends on allocating tasks according to relative capabilities and cognitive limitations [25]. Poor calibration can generate automation bias, excessive human intervention, or fragmented responsibility. However, task performance in an experiment is not equivalent to organizational learning. Longitudinal research must determine whether organizations retain information about disagreements, overrides, near misses, and unsuccessful predictions and whether that information changes future models, routines, or decision rights.

Organizational learning failure should likewise be separated from an unfavorable outcome. A system may produce errors while supporting rapid detection and revision, or achieve stable short-term performance while suppressing information about changing conditions. The appropriate outcome is therefore not error frequency alone but the organization’s capacity to recognize surprises, preserve discrepant knowledge, investigate causes, and alter future action. This distinction prevents reliability improvements from being misrepresented as proof of adaptive learning.

Proposed certainty–learning framework

The proposed framework treats algorithmic certainty as a multilevel organizational process rather than a technical attribute. Decision-structure research shows that human and artificial-intelligence decision rights can be configured in different ways, with distinct implications for authority and control [26]. The framework therefore begins with predictive outputs and certainty cues but locates their consequences in managerial framing, institutional endorsement, metric dependence, and the allocation of rights to decide, challenge, pause, or revise.

Delegation is also reciprocal. Users delegate analysis or action to agentic systems, while systems can direct users through recommendations, prompts, rankings, and required responses [27]. This reciprocity distinguishes autonomy from responsibility transfer: retaining a nominal override does not establish meaningful control when employees lack information, time, protection, or authority to use it. Hybrid-intelligence theory offers a contrasting design logic in which human and machine capabilities are organized as complementary and mutually improving rather than substitutive [28]. Complementarity, however, cannot be presumed.

Meta-analytic evidence indicates that human–AI combinations are not uniformly superior and that outcomes depend on task characteristics, baseline capabilities, and combination design [29]. The framework accordingly proposes three connected but testable mechanisms: enacted certainty narrows inquiry, reduces the consequential influence of dissent, and constrains or displaces experimentation. These mechanisms may weaken anomaly interpretation and organizational updating. Expertise, psychological safety, work design, task stakes, transparency, contestability, slack, and accountability are proposed moderators. Feedback may also be recursive: learning can recalibrate certainty, while apparent success can reinforce it.

Figure 2 presents the original conceptual synthesis for certainty–learning framework, showing how the article’s principal constructs, mechanisms, uncertainties, and decision boundaries are connected.

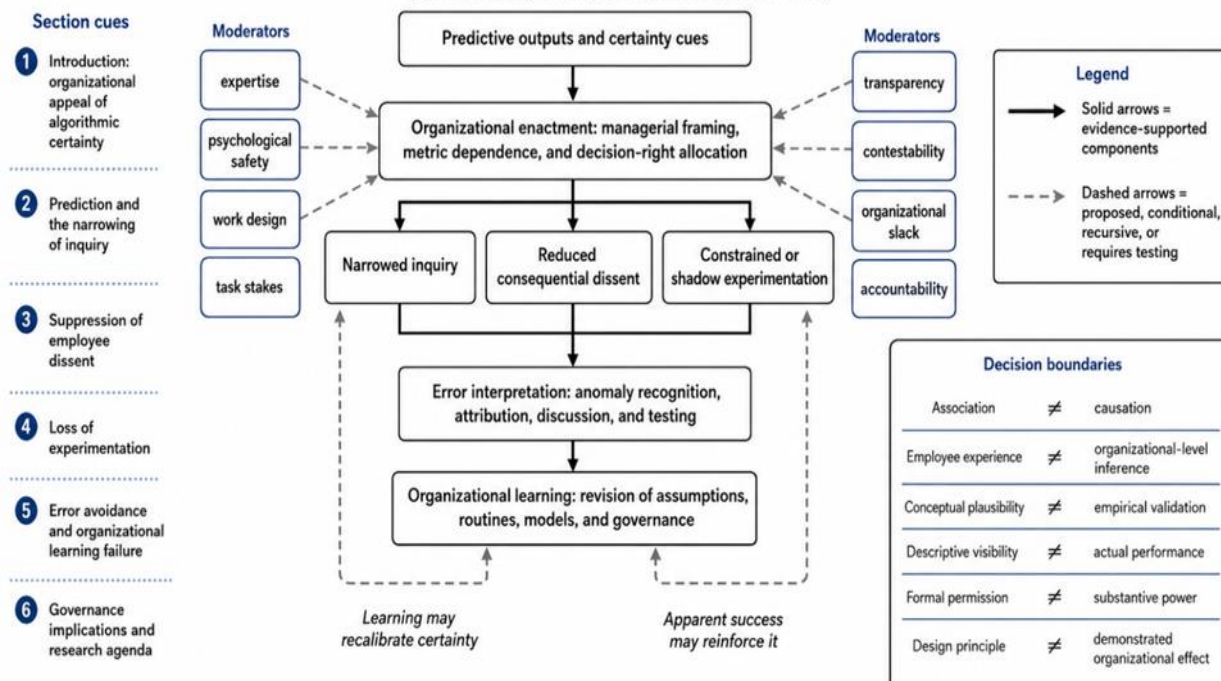


Figure 2. Certainty–Learning Framework

Table 3 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for evidence limitations, unresolved questions, and future research priorities.

Table 3. Evidence limitations, unresolved questions, and future research priorities.

Context or level	Moderator or boundary condition	Organizational implication	Alternative explanation	Transferability limit	Uncertainty	Research need
Individual judgment	Expertise, time pressure, confidence display	Preserve independent assessment and visible uncertainty	Anchoring, interface framing, institutional endorsement	Task experiments	Persistence over time is unknown	Repeated-measures field experiments
Employee and occupation	Psychological safety, status, appeal protection	Evaluate whether dissent changes decisions	Prior hierarchy or employment insecurity	Platform and nonalgorithmic voice evidence	Consequential influence is rarely measured	Multilevel studies linking challenge to revision
Work system	Autonomy, feedback, task stakes, slack	Record overrides, anomalies, and tests	Regulation, workload, or standardization	Conceptual and task-level evidence	Organizational updating remains underobserved	Longitudinal comparative case designs
Governance	Decision rights, transparency, accountability	Design contestable and reversible decision processes	Capability differences or implementation maturity	Cross-sector transfer is uncertain	No validated governance bundle	Intervention studies testing mechanisms and unintended effects

Note. Organizational implications and research priorities are original proposed synthesis. They are hypotheses for evaluation, not demonstrated prescriptions.

Governance implications and research agenda

Governance should address how certainty is produced and challenged, not only whether a model is accurate. AI management research identifies continuing problems of learning, autonomy, and opacity; accountability analysis locates responsibility in organizational actors and design choices; and ethics scholarship cautions that principles require operational institutions and procedures [30–32]. Accordingly, governance should specify who may question an output, what evidence triggers review, how disagreement is documented, and whether employees can pause consequential action without disproportionate penalty. Human-centered artificial intelligence emphasizes reliable operation alongside meaningful human control [33]. For this review, meaningful control requires more than an override button. It includes access to relevant reasoning and uncertainty, protected escalation, time and expertise for independent assessment, reversible trials where appropriate, and responsibility for responding to dissent. These are proposed design conditions rather than validated safeguards.

Research should measure the full process from prediction to updating. Longitudinal studies could compare organizations with different decision-right and appeal arrangements. Field experiments could vary uncertainty displays, dissent protections, or review triggers while monitoring challenge quality and unintended gaming. Multilevel designs should distinguish individual reliance, team voice, occupational authority, and organizational revision. Studies must also test alternative explanations,

including pre-existing hierarchy, selection into algorithmic work, regulation, workload, and task interdependence. The framework should be refined or rejected where these explanations better account for observed learning patterns. Measurement development is equally important. Studies should distinguish perceived certainty from displayed confidence, managerial endorsement, mandatory compliance, and actual predictive validity. Dissent measures should capture whether concerns are heard, investigated, and consequential, rather than counting submissions alone. Experimentation measures should separate authorized tests, informal workarounds, and routine variance. Learning should be evidenced through revision of assumptions, models, routines, or decision rights across time. These distinctions would allow competing mechanisms to be tested without inferring organization-level learning from individual attitudes or isolated task accuracy.

Conclusion

Algorithmic prediction does not weaken organizational learning simply by being computational, accurate, or opaque. The critical risk arises when organizations enact probabilistic outputs as settled authority and reorganize inquiry, dissent, and experimentation around that settlement. Across heterogeneous evidence, the most defensible conclusion is conditional: certainty may narrow legitimate questions, reduce the consequential influence of employee challenge, displace testing, and limit the interpretation of anomalies, but these relationships depend on expertise, hierarchy, psychological safety, work design, task stakes, and governance.

The Certainty–Learning Framework integrates these mechanisms while preserving distinctions between technical capability and legitimate authority, formal voice and substantive influence, deviation and experimentation, error prevention and error learning, and task performance and organizational updating. It is an original synthesis rather than a validated causal model. Its value lies in making assumptions and testable relationships visible. Evidence remains fragmented across occupations, technologies, methods, and levels of analysis. Comparative longitudinal and intervention research is therefore necessary before the framework can support general prescriptions. Organizations should treat certainty itself as an object of inquiry: not as a defect to eliminate, but as a condition whose sources, consequences, and limits require continuing examination.

Acknowledgments: None

Conflict of interest: None

Financial support: None

Ethics statement: None

References

1. 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
2. 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
3. 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
4. Lindebaum D, Vesa M, den Hond F. Insights from The Machine Stops to better understand rational assumptions in algorithmic decision making and its implications for organizations. *Acad Manag Rev.* 2020;45(1):247-63. doi:10.5465/amr.2018.0181
5. 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
6. Snyder H. Literature review as a research methodology: An overview and guidelines. *J Bus Res.* 2019;104:333-9. doi:10.1016/j.jbusres.2019.07.039
7. Page MJ, McKenzie JE, Bossuyt PM, Boutron I, Hoffmann TC, Mulrow CD, et al. The PRISMA 2020 statement: An updated guideline for reporting systematic reviews. *BMJ.* 2021;372. doi:10.1136/bmj.n71
8. Rethlefsen ML, Kirtley S, Waffenschmidt S, Ayala AP, Moher D, Page MJ, et al. PRISMA-S: An extension to the PRISMA statement for reporting literature searches in systematic reviews. *Syst Rev.* 2021;10(1):39. doi:10.1186/s13643-020-01542-z
9. Paul J, Criado AR. The art of writing literature review: What do we know and what do we need to know? *Int Bus Rev.* 2020;29(4):101717. doi:10.1016/j.ibusrev.2020.101717

10. 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
11. Anthony C. When knowledge work and analytical technologies collide: The practices and consequences of black boxing algorithmic technologies. *Adm Sci Q.* 2021;66(4):1173-212. doi:10.1177/00018392211016755
12. 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
13. Gaube S, Suresh H, Raue M, Merritt A, Berkowitz SJ, Lerner E, et al. Do as AI say: Susceptibility in deployment of clinical decision-aids. *NPJ Digit Med.* 2021;4(1):31. doi:10.1038/s41746-021-00385-9
14. Curchod C, Patriotta G, Cohen L, Neysen N. Working for an algorithm: Power asymmetries and agency in online work settings. *Adm Sci Q.* 2020;65(3):644-76. doi:10.1177/0001839219867024
15. Rahman HA. The invisible cage: Workers' reactivity to opaque algorithmic evaluations. *Adm Sci Q.* 2021;66(4):945-88. doi:10.1177/00018392211010118
16. Chamberlin M, Newton DW, LePine JA. A meta-analysis of voice and its promotive and prohibitive forms: Identification of key associations, distinctions, and future research directions. *Pers Psychol.* 2017;70(1):11-71. doi:10.1111/peps.12185
17. Frazier ML, Fainshmidt S, Klinger RL, Pezeshkan A, Vacheva V. Psychological safety: A meta-analytic review and extension. *Pers Psychol.* 2017;70(1):113-65. doi:10.1111/peps.12183
18. Beane M. Shadow learning: Building robotic surgical skill when approved means fail. *Adm Sci Q.* 2019;64(1):87-123. doi:10.1177/0001839217751692
19. Keum DD, See KE. The influence of hierarchy on idea generation and selection in the innovation process. *Organ Sci.* 2017;28(4):653-69. doi:10.1287/orsc.2017.1142
20. Camuffo A, Cordova A, Gambardella A, Spina C. A scientific approach to entrepreneurial decision making: Evidence from a randomized control trial. *Manage Sci.* 2020;66(2):564-86. doi:10.1287/mnsc.2018.3249
21. Elsbach KD, Stigliani I. Design thinking and organizational culture: A review and framework for future research. *J Manage.* 2018;44(6):2274-306. doi:10.1177/0149206317744252
22. Dahlin KB, Chuang YT, Roulet TJ. Opportunity, motivation, and ability to learn from failures and errors: Review, synthesis, and ways to move forward. *Acad Manag Ann.* 2018;12(1):252-77. doi:10.5465/annals.2016.0049
23. Parker SK, Grote G. Automation, algorithms, and beyond: Why work design matters more than ever in a digital world. *Appl Psychol.* 2022;71(4):1171-204. doi:10.1111/apps.12241
24. Parent-Rocheleau X, Parker SK. Algorithms as work designers: How algorithmic management influences the design of jobs. *Hum Resour Manag Rev.* 2022;32(3):100838. doi:10.1016/j.hrmr.2021.100838
25. Fügener A, Grahl J, Gupta A, Ketter W. Cognitive challenges in human-artificial intelligence collaboration: Investigating the path toward productive delegation. *Inf Syst Res.* 2022;33(2):678-96. doi:10.1287/isre.2021.1079
26. Shrestha YR, Ben-Menahem SM, von Krogh G. Organizational decision-making structures in the age of artificial intelligence. *Calif Manage Rev.* 2019;61(4):66-83. doi:10.1177/0008125619862257
27. 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
28. Dellermann D, Ebel P, Söllner M, Leimeister JM. Hybrid intelligence. *Bus Inf Syst Eng.* 2019;61(5):637-43. doi:10.1007/s12599-019-00595-2
29. Vaccaro M, Almaatouq A, Malone TW. When combinations of humans and AI are useful: A systematic review and meta-analysis. *Nat Hum Behav.* 2024;8(12):2293-303. doi:10.1038/s41562-024-02024-1
30. Berente N, Gu B, Recker J, Santhanam R. Managing artificial intelligence. *MIS Q.* 2021;45(3):1433-50. doi:10.25300/MISQ/2021/16274
31. Martin K. Ethical implications and accountability of algorithms. *J Bus Ethics.* 2019;160(4):835-50. doi:10.1007/s10551-018-3921-3
32. Mittelstadt B. Principles alone cannot guarantee ethical AI. *Nat Mach Intell.* 2019;1(11):501-7. doi:10.1038/s42256-019-0114-4
33. Shneiderman B. Human-centered artificial intelligence: Reliable, safe & trustworthy. *Int J Hum Comput Interact.* 2020;36(6):495-504. doi:10.1080/10447318.2020.1741118