

E-ISSN: 3108-4192

APSSHs

Academic Publications of Social Sciences and Humanities Studies

2025, Volume 5, Issue 1, Page No: 72-80

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

Journal of Applied Organizational Systems and Behavior

How Predictive Systems Suppress Dissent and Organizational Learning

Nguyen Thanh Huy^{1*}, Pham Quang Minh¹, Le Thi Bich², Tran Van Nam³

1. Department of Predictive Systems and Dissent, Faculty of Business, Vietnam National University of Agriculture, Hanoi, Vietnam.
2. Department of Organizational Learning Suppression, Faculty of Management, Can Tho University, Can Tho, Vietnam.
3. Department of Work Innovation, Faculty of Psychology, Hue University, Hue, Vietnam.

Abstract

Predictive systems increasingly participate in organizational judgments about performance, risk, allocation, diagnosis, and future action. Their organizational significance, however, extends beyond predictive capability. When predictive output is treated as unusually authoritative, employees may face social, procedural, or epistemic costs for challenging it, potentially weakening the flow of discrepant local knowledge that organizations require for correction and learning. This Original Propositional Model Article develops a bounded theoretical account of that possibility. It distinguishes technological capability from legitimate authority, voice from silence, psychological safety from consequential influence, and algorithmic confidence from the proposed social signal of algorithmic certainty. Building on these distinctions, the article introduces the proposed constructs of predictive authority, algorithmic certainty signaling, defensive conformity, and anomalous-knowledge suppression, and integrates them in a proposed Dissent-Suppression Model. The model theorizes that predictive authority may reduce consequential dissent when authoritative framing, interpersonal risk, responsibility diffusion, opacity, or dependence make contradiction costly, thereby limiting the entry of anomalous knowledge into organizational review and weakening specified learning processes. The article further develops falsifiable propositions and conditional countermeasures centered on contestability, independent judgment, protected anomaly escalation, and learning-preservation review. These mechanisms are not presented as empirically validated, universal, or deployment-ready. Their operation should vary with task characteristics, expertise, psychological safety, system quality, work design, hierarchy, and level of analysis. The contribution is a testable organizational theory of when predictive systems may become epistemic authority structures rather than neutral decision aids.

Keywords: Predictive authority and organizational dissent, Conceptual foundations, Error avoidance and defensive conformity, Suppression of anomalous knowledge, Propositions and organizational countermeasures, Predictive

How to cite this article: Huy NT, Minh PQ, Bich LT, Nam TV. How Predictive Systems Suppress Dissent and Organizational Learning. J Appl Organ Syst Behav. 2025;5(1):72-80. <https://doi.org/10.51847/Gqw2ou6CCh>

Received: 19 August 2025; **Revised:** 05 October 2025; **Accepted:** 07 October 2025

Corresponding author: Nguyen Thanh Huy

E-mail ✉ huy.nguyen@vnua.edu.vn

Introduction

Predictive systems are increasingly embedded in organizational processes that rank, recommend, forecast, classify, or flag possible future states. Their influence is therefore not exhausted by technical accuracy. Algorithmic systems can also reorganize managerial control by redistributing observation, evaluation, and direction across digital infrastructures [1]. Learning algorithms further alter organizing because their outputs may be black-boxed, continuously updated, anticipatory, and embedded in data-intensive regimes of knowing [2]. The organizational problem addressed here begins when such outputs become socially consequential not merely as information but as cues about what can legitimately be questioned. A technically assistive system can therefore acquire organizational authority without any formal declaration that its prediction is binding.



© 2025 The Author(s).

Copyright CC BY-NC-SA 4.0

Authority matters because algorithm-mediated settings can create asymmetries between those evaluated by systems and those able to design, interpret, or enforce them, even while workers retain forms of agency [3]. This suggests that predictive influence should be analyzed as a sociotechnical relationship rather than a property located inside an algorithm. The central issue is not whether employees always obey predictions, but whether organizational arrangements make contradiction more costly, less credible, or less consequential. Such conditions may be especially important where predictive outputs are integrated into performance evaluation, risk management, professional judgment, or managerial escalation and where employees must decide whether locally held knowledge is worth voicing against an apparently authoritative forecast.

That problem intersects with, but is not reducible to, the established literature on employee voice and silence. Voice concerns discretionary expression intended to improve organizational functioning, while silence has distinct antecedents and should not be inferred simply from the absence of speaking [4]. Organizational dissent therefore requires attention both to whether employees can express disagreement and to whether that disagreement can affect consequential judgment. A formal channel for objection may coexist with practical futility, interpersonal risk, or expectations that predictive output will prevail. Accordingly, this article separates voice opportunity from consequential influence and treats non-voice as theoretically ambiguous rather than as automatic evidence of coercion.

The knowledge consequences may extend beyond a single decision. Field research on algorithmic technology adoption shows that organizational actors can change their practices and symbolic responses as a new regime of knowing becomes institutionalized [5]. Yet existing research does not establish a general causal sequence from predictive systems to silence and then to learning loss. As an Original Propositional Model Article, this article therefore makes a narrower contribution: it develops a model in which predictive authority, certainty signaling, defensive conformity, anomalous-knowledge attenuation, and learning loss are proposed relationships that require direct testing. The model is designed to explain conditions under which prediction-mediated authority may suppress correction-relevant dissent, while preserving alternative explanations such as rational deference, algorithm aversion, low-quality human judgment, workload, and pre-existing hierarchy.

Conceptual foundations

The first conceptual requirement is to distinguish what predictive systems can do from what organizations authorize them to mean. AI may automate some activities while augmenting others, and these logics can coexist rather than forming a simple replacement continuum [6]. Likewise, digital technologies do not determine work outcomes independently of job and organizational design; the allocation of discretion, monitoring, task interdependence, and human influence shapes how technological systems enter work [7]. Predictive authority is therefore proposed here as an enacted organizational condition in which predictive output is treated as unusually legitimate, difficult to contest, or presumptively correct. It is not synonymous with model accuracy, automation, managerial authority, or employee trust.

A second requirement concerns interpersonal risk. Psychological safety refers to a context in which interpersonal risk taking is perceived as more acceptable, but it should not be treated as equivalent to speaking up itself [8]. In the present model, psychological safety is a boundary condition that may alter whether disagreement with a predictive output is expressed, withheld, or reformulated. It cannot establish whether voiced disagreement is consequential. The model therefore separates the opportunity to speak, the willingness to speak, and the organizational capacity of speech to affect a decision. These distinctions also guard against treating formal “human-in-the-loop” arrangements as evidence that meaningful human influence actually exists.

Third, voice and silence require separate treatment. Evidence indicates that voice and silence are distinguishable constructs, with perceived impact and psychological safety relating to them in different ways [9]. This matters because reduced dissent can arise through several pathways: employees may remain silent, speak ambiguously, defer to another actor, challenge privately rather than publicly, or voice a concern that is heard but ignored. The proposed model consequently defines consequential dissent as expressed disagreement or discrepant information that has a credible route into organizational review. Individual behavior is not automatically an organizational learning outcome; aggregation mechanisms are required before suppressed or discounted employee knowledge can be theorized as a loss at team or organizational level.

Algorithmic certainty as a social signal

The proposed construct of algorithmic certainty signaling addresses how predictive output is socially interpreted rather than how statistically certain it actually is. Empirical research on trust in AI shows that reliance varies with system characteristics, interaction design, reliability, transparency, and related cues rather than reflecting an invariant tendency to trust algorithms [10]. Experimental work also demonstrates algorithm appreciation under some conditions: people may place greater weight on algorithmic than human judgment, although expertise and self-comparison can alter that pattern [11]. These findings make it plausible that predictive output can acquire elevated epistemic weight, but they do not establish that such weight is universal or that it suppresses workplace dissent.

Task characteristics provide an important counterweight. Algorithm aversion is task dependent, with relative acceptance of algorithmic judgment varying according to perceived task objectivity [12]. The proposed model therefore treats certainty

signaling as conditional. A risk score presented as definitive, embedded in mandatory workflow, and backed by managerial enforcement may communicate greater epistemic settledness than the same score presented as provisional evidence open to challenge. Conversely, subjective tasks, recognized model limitations, professional expertise, or explicit uncertainty may reduce that signal. The construct concerns organizational presentation and enactment; it does not represent statistical confidence, calibration, predictive validity, or a validated psychological scale.

Responses to algorithmic advice are conditional on characteristics of the technology, decision maker, and task rather than reflecting uniformly high or low algorithmic authority [10–12]. The original theoretical step is to propose that these cues can become a social signal about the legitimacy of contradiction. Where an output is treated as already settled, dissent may appear to require stronger justification than agreement, shifting the burden of proof toward the employee. This asymmetry remains to be tested directly. It could also be offset by high expertise, visible model error, strong psychological safety, or organizational norms that reward anomaly detection rather than conformity.

Table 1 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for definitions, constructs, and conceptual boundaries for how predictive systems suppress dissent and organizational learning.

Table 1. Definitions, constructs, and conceptual boundaries for How Predictive Systems Suppress Dissent and Organizational Learning.

Construct or Component	Definition	Problem Addressed	Distinction From Adjacent Concepts	Proposed Role	Uncertainty	Boundary Statement
Predictive System	Technology producing forecasts, classifications, recommendations, or risk estimates used in work.	Separates prediction from authority.	Not identical to managerial control.	Antecedent infrastructure.	Systems vary in purpose and accuracy.	Use alone does not imply suppression.
Algorithmic Control	Digital structuring of monitoring, evaluation, direction, or matching.	Locates organizational power around technology.	Broader than prediction.	Context for enacted authority.	Control may coexist with worker agency.	Does not prove compliance.
Predictive Authority	Proposed condition in which predictive output is treated as unusually legitimate or difficult to contest.	Explains why advisory output may shape dissent.	Not model accuracy or managerial rank.	Proposed upstream construct.	Requires direct measurement and validation.	Must be separated from rational deference.
Voice And Silence	Distinct patterns of expression and withholding relevant to organizational improvement.	Prevents equating absence of speech with coercion.	Voice opportunity differs from consequential influence.	Established behavioral domain.	Motives for silence vary.	Non-voice alone cannot demonstrate suppression.
Psychological Safety	Perceived acceptability of interpersonal risk taking.	Identifies a context shaping willingness to challenge.	Not synonymous with voice.	Boundary condition.	May not ensure influence after speaking.	Enabling, not sufficient.
Algorithmic Certainty Signal	Proposed perception that predictive output is epistemically settled or unusually costly to challenge.	Connects prediction presentation to dissent incentives.	Not statistical confidence or calibration.	Proposed mediator.	Construct remains unvalidated.	Task and expertise conditions may reverse effects.

Error avoidance and defensive conformity

The next proposed mechanism concerns how employees manage the perceived costs of being wrong against an authoritative prediction. Psychological safety has a well-established multilevel nomological network relevant to interpersonal-risk behavior [13]. In prediction-mediated work, that evidence supports treating interpersonal climate as a boundary condition, not as proof that algorithms create fear. Defensive conformity is proposed here as public alignment with predictive output partly because contradiction appears to carry greater interpersonal, reputational, or responsibility risk than agreement. Such conformity may be strategic even when private judgment remains uncertain, and it should therefore be distinguished from genuine belief in the model.

Voice can also be inhibited when responsibility for speaking becomes diffuse. Research on the voice bystander effect shows that information redundancy can reduce felt individual responsibility to voice, especially when others appear able to access managers [14]. A predictive system may create an analogous but untested dynamic if employees infer that the organization, model developers, data specialists, or other reviewers have already considered the relevant evidence. The model does not

assume that this occurs automatically. Rather, it proposes a testable extension: authoritative predictive infrastructure may make contradictory knowledge feel redundant, already adjudicated, or somebody else's responsibility to escalate.

Speaking up can itself carry social consequences. Research on prohibitive voice shows that employees may experience anxiety through social uncertainty after raising concerns [15]. Applied cautiously, this evidence suggests that error avoidance may concern not only technical correctness but also the social consequences of contradicting a decision path backed by predictive authority. The proposed pathway is therefore conditional: certainty signaling may increase defensive conformity where psychological safety is low, perceived impact is weak, hierarchy is salient, or disagreement is associated with reputational exposure. No approved source directly demonstrates that predictive authority causes defensive conformity. That relationship remains a falsifiable proposition rather than an established organizational effect.

Suppression of anomalous knowledge

The proposed model becomes organizationally consequential only if reduced dissent changes what information reaches decision processes. Research on professional engagement with opaque AI shows that experts do not respond uniformly to algorithmic recommendations; instead, they develop different ways of engaging with or distancing themselves from opaque outputs when making critical judgments [16]. This supports treating opacity as a boundary condition rather than assuming that predictive systems simply override expertise. The relevant question is whether contradictory local knowledge can still enter consequential review when employees cannot readily interrogate how a prediction was produced.

Organizational practice can also create black boxes socially. Longitudinal research on analytical technologies shows that validating practices may gradually alter how experts relate to analytical outputs and can contribute to black-boxing even when the technology itself is not literally inaccessible [17]. The proposed construct of anomalous-knowledge suppression therefore refers to attenuation of discrepant local, experiential, contextual, or professional information before it can challenge or revise a prediction-mediated decision. Suppression here does not imply intentional censorship. It may involve withholding, discounting, routinized deference, reduced investigation, or organizational procedures that make contradictory evidence difficult to translate into action.

Dependence provides a further pathway. Ethnographic evidence shows that professionals may reconcile themselves to practices they regard as methodologically deficient when occupational dependence constrains their position, sometimes lowering the standards they are willing to articulate [18]. The Dissent-Suppression Model proposes that predictive authority may interact with similar dependence structures, making disagreement less consequential even where formal speaking opportunities remain. This relationship is proposed, not established. Expertise, model failure, professional norms, protected escalation rights, or managerial receptiveness may instead amplify anomalous knowledge and improve correction.

Table 2 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for proposed mechanisms, relationships, and organizational consequences in how predictive systems suppress dissent and organizational learning.

Table 2. Proposed mechanisms, relationships, and organizational consequences in How Predictive Systems Suppress Dissent and Organizational Learning.

Mechanism or Relationship	Antecedent	Mediator or Process	Moderator	Expected Consequence	Supporting Evidence	Rival Explanation	Validation Requirement
Safety-Conditioned Speaking	Interpersonal risk	Willingness to challenge	Psychological safety	More or less dissent	Established psychological-safety research	Personality or leadership style	Longitudinal or experimental separation of climate and voice
Responsibility Diffusion	Shared or apparently redundant knowledge	Reduced felt obligation to speak	Access to managers; role clarity	Lower voice probability	Voice-bystander evidence	Rational belief others will act	Direct manipulation in prediction-mediated work
Social-Cost Anticipation	Contradictory or prohibitive voice	Anxiety and social uncertainty	Hierarchy; peer norms	Defensive conformity	Voice-consequence evidence	General conflict avoidance	Measure predicted social cost before dissent
Opacity-Conditioned Engagement	Opaque predictive output	Reduced interrogability	Expertise; explainability	Less effective challenge	Professional AI field evidence	Appropriate trust in superior prediction	Compare equivalent tasks with varying opacity
Social Black-Boxing	Repeated validation routines	Increasing taken-for-grantedness	Professional autonomy	Reduced anomaly investigation	Longitudinal analytical-work evidence	Efficiency-driven routinization	Trace whether black-boxing precedes reduced challenge

Predictive Authority → Anomaly Attenuation	Proposed predictive authority	Proposed defensive conformity	Safety, dependence, task, system quality	Proposed reduction in consequential anomaly transmission	Cross-domain synthesis	Rational deference or poor human information	Direct test with accuracy held constant
---	--------------------------------------	--------------------------------------	--	---	------------------------	--	---

Organizational learning loss

Organizational learning should not be treated as a single undifferentiated outcome. Established research distinguishes processes such as search, knowledge creation, retention, retrieval, and transfer [19]. Accordingly, the proposed construct of organizational learning loss refers only to reduced opportunities for specified processes of correction, updating, retention, or experiential learning. It does not imply generalized organizational decline, poorer financial performance, or universal deskilling. The theoretical claim is narrower: if discrepant knowledge fails to reach consequential review, organizations may lose information that would otherwise challenge assumptions, revise routines, or expose prediction errors.

Technology can alter the opportunities through which learning occurs. Research on robotic surgery shows that technology-mediated restructuring of work can restrict conventional participation through which novices develop expertise, prompting alternative forms of learning [20]. Although this setting differs from predictive decision systems, it demonstrates why learning consequences must be traced through changed participation structures rather than assumed from technology adoption itself. Prediction-mediated work may similarly alter who encounters anomalies, who can interpret them, and whose judgment is considered legitimate.

More directly, research on clinical decision-support adoption has found that experiential-learning relationships can become flatter after algorithmic recommendation tools are introduced, with task difficulty and task variety conditioning those associations [21]. That evidence does not measure dissent and therefore cannot validate the proposed mediation mechanism. It does, however, establish that algorithmic recommendation systems can interact with experiential learning in organizational settings. The Dissent-Suppression Model consequently proposes that anomaly attenuation may weaken learning where contradictory experience would otherwise trigger investigation or updating. This prediction remains conditional on task structure, system quality, knowledge distribution, and whether anomalies are captured elsewhere in the organization.

Proposed dissent-suppression model

Algorithmic authority is embedded in work arrangements rather than confined to technical interfaces. Research on platform management shows that matching and control functions can jointly structure worker experience and organizational tensions [22]. Work-design research similarly indicates that algorithmic management reshapes demands, resources, discretion, and human influence rather than producing one invariant employment outcome [23]. These literatures support locating predictive authority inside a sociotechnical configuration of workflow, hierarchy, decision rights, and accountability.

The proposed Dissent-Suppression Model integrates these foundations into a bounded process model. It begins with organizational embedding of a predictive system and the proposed emergence of predictive authority. That authority may generate an algorithmic certainty signal, which may increase the relative burden of contradicting predictive output. Under conditions of low psychological safety, weak perceived impact, hierarchy, dependence, redundancy, or opacity, employees may respond through defensive conformity. Where such conformity prevents discrepant information from entering consequential review, anomalous-knowledge suppression may follow. Reduced anomaly exposure may then weaken specified processes of updating, correction, or experiential learning.

Cross-level inference is deliberately constrained. AI-related processes differ across individual, team, and organizational levels, and individual perceptions cannot be treated automatically as organizational outcomes [24]. Human-AI collaboration research also shows that productive delegation depends on task allocation and metaknowledge rather than a generic requirement to retain humans in the process [25]. The proposed model is therefore conditional on work-design architecture, level of analysis, and the allocation of judgment between humans and AI [23–25]. Organization-level learning loss should be inferred only where aggregation mechanisms prevent discrepant knowledge from changing collective decisions, routines, or stored knowledge.

Figure 1 presents the original conceptual synthesis for dissent-suppression model, showing how the article's principal constructs, mechanisms, uncertainties, and decision boundaries are connected.

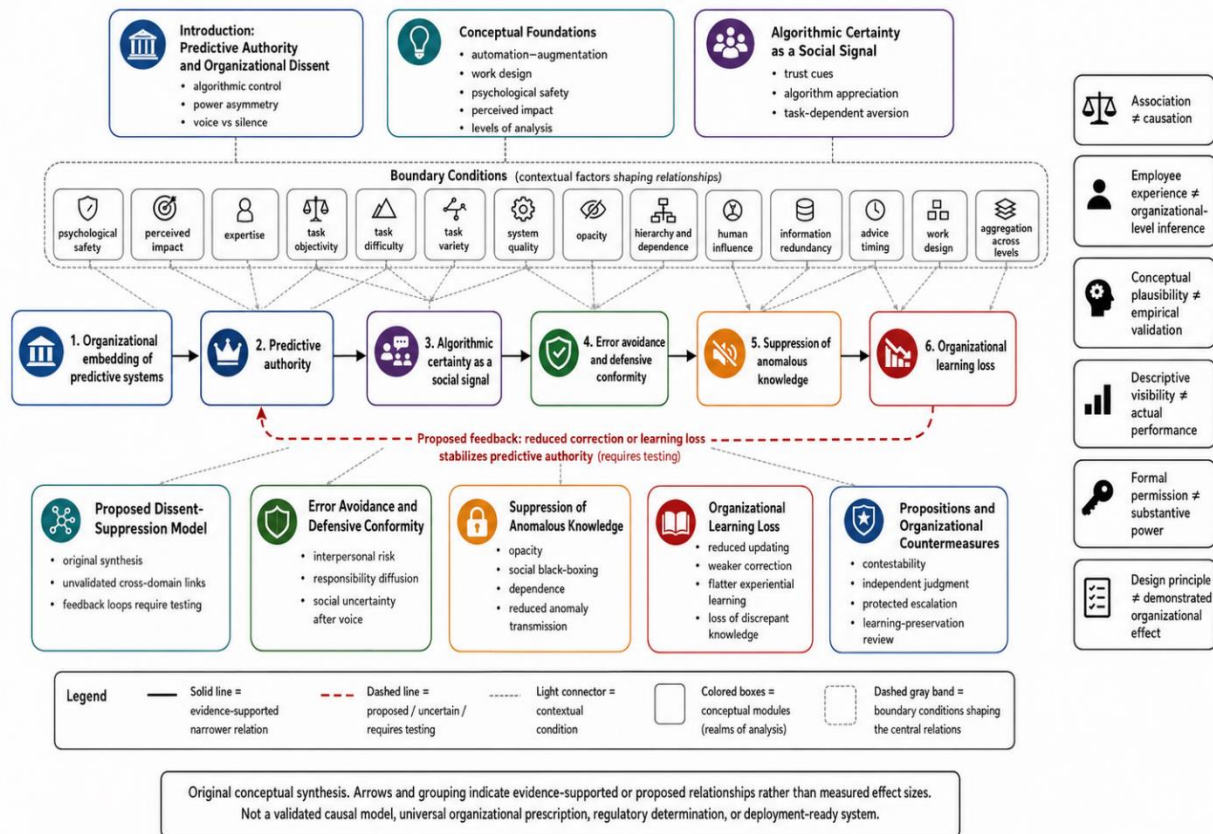


Figure 1. Dissent-Suppression Model

The figure is an original conceptual synthesis developed for “How Predictive Systems Suppress Dissent and Organizational Learning.” Arrows and grouping indicate evidence-supported or proposed relationships rather than measured effect sizes. The figure does not represent a validated causal model, universal organizational prescription, regulatory determination, or deployment-ready system.

Table 3 organizes the evidence, constructs, relationships, uncertainties, and boundary conditions required for testable propositions, evaluation criteria, and practical responsibilities.

Table 3. Testable propositions, evaluation criteria, and practical responsibilities.

Proposition or Criterion	Unit And Context	Observable Indicator	Evidence Required	Falsification Condition	Organizational Responsibility	Misuse Risk	Readiness Boundary
P1 Proposed Predictive-Authority Relationship	Employee decision context	Willingness to challenge equivalent outputs	Controlled comparison holding accuracy constant	No relationship or greater dissent under stronger authority cues	Preserve legitimate challenge routes	Treating correlation as coercion	Not validated
P2 Proposed Certainty-Signal Relationship	Human-prediction interaction	Perceived epistemic settledness	Manipulation of presentation and enactment	No change in dissent-related judgment	Communicate uncertainty appropriately	Confusing displayed confidence with model calibration	Construct requires validation
P3 Proposed Defensive-Conformity Mechanism	Employee/team	Public agreement despite private disagreement	Temporally ordered mediation evidence	Costs do not predict conformity	Protect good-faith contradiction	Assuming every agreement is defensive	Mechanism unproven
P4 Proposed Anomaly-Transmission Relationship	Team/decision process	Contradictory information reaching review	Traced disagreement-to-decision records	Dissent reduction does not reduce anomaly transmission	Maintain reviewable anomaly channels	Incentivizing low-quality objections	Test information quality
P5 Proposed Learning Relationship	Team/organization	Updating, correction, or revised routines	Longitudinal or experimental learning measures	Lower anomaly transmission does not affect learning	Preserve correction and feedback pathways	Equating learning with performance	Context-specific

P6 Proposed Contestability Criterion	Human-AI workflow	Effective ability to alter consequential decisions	Comparative intervention evidence	Contestability fails to improve calibrated challenge	Assign meaningful decision rights	Symbolic human-in-loop compliance	No universal safeguard
P7 Proposed Multilevel Criterion	Individual to organization	Aggregation of dissent into collective revision	Multilevel longitudinal evidence	Individual conformity does not affect collective learning	Specify escalation and aggregation mechanisms	Cross-level overclaiming	Organizational inference requires aggregation evidence

Propositions and organizational countermeasures

The proposed countermeasures follow from the model's uncertainty rather than from proof that any one safeguard works universally. Experimental evidence indicates that allowing people even limited modification of algorithmic forecasts can increase willingness to use algorithmic advice [26]. This finding suggests that perceived influence matters, but it does not establish that nominal adjustment rights preserve dissent. A meaningful contestability mechanism would require access to relevant evidence, protection from retaliation, and a credible route for altering a consequential decision rather than merely editing an inconsequential recommendation.

Human involvement is similarly conditional. Experimental research with public employees indicates that procedural responses to algorithmic managerial decisions depend partly on human involvement and task characteristics [27]. The implication is not that adding a manager automatically legitimizes predictive systems. Human review may be substantive, ceremonial, or even reinforce deference if reviewers lack expertise or authority. Proposed organizational responsibilities should therefore concern actual decision rights, escalation pathways, and accountability for considering contradictory information.

Decision sequencing provides another possible intervention point. Experimental work on AI-assisted diagnosis indicates that receiving AI advice after an independent initial judgment can alter information processing and engagement with contradictory recommendations relative to receiving advice earlier [28]. This supports testing independent-first judgment where task characteristics make it appropriate, but it does not establish a universally superior sequence. Evidence on user influence, human involvement, and advice sequencing indicates that organizational countermeasures should be evaluated as conditional design choices rather than universal safeguards [26–28].

Figure 2 presents the original conceptual synthesis for mechanisms, boundary conditions, and organizational consequences, showing how the article's principal constructs, mechanisms, uncertainties, and decision boundaries are connected.

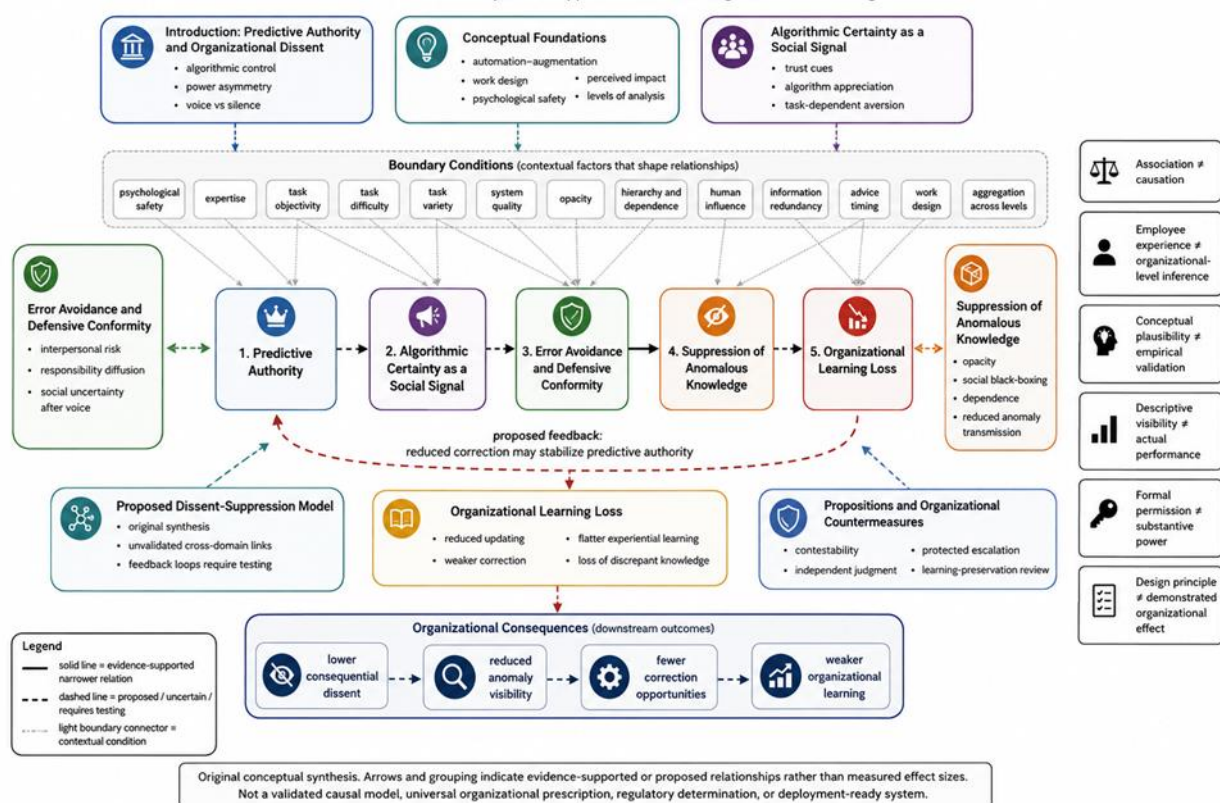


Figure 2. Mechanisms, Boundary Conditions, and Organizational Consequences

Conclusion

This article develops an original propositional explanation for how predictive systems may become organizational authority structures rather than remaining neutral decision aids. Its central contribution is the proposed Dissent-Suppression Model, which links predictive authority and algorithmic certainty signaling to defensive conformity, attenuation of anomalous knowledge, and reduced opportunities for organizational learning under specified conditions.

The model deliberately stops short of claiming that predictive systems generally suppress dissent. Predictive output may deserve deference, employees may hold inaccurate contradictory beliefs, and well-designed systems may improve rather than inhibit correction. Psychological safety, expertise, task characteristics, human influence, opacity, hierarchy, information redundancy, system quality, and decision sequencing may alter or reverse the proposed relationships. Individual silence also cannot establish organization-level learning loss without demonstrated aggregation mechanisms.

The framework is therefore best understood as a falsifiable research architecture. Its propositions identify what future longitudinal, experimental, qualitative, and multilevel studies would need to observe—and what findings could challenge the model. The practical implication is correspondingly bounded: organizations adopting predictive systems should examine not only prediction performance but also whether contradictory knowledge can remain visible, examinable, and consequential. Whether such arrangements improve learning remains an empirical question rather than an established organizational effect.

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. 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
3. 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
4. Morrison EW. Employee voice and silence: Taking stock a decade later. *Annu Rev Organ Psychol Organ Behav.* 2023;10:79-107. doi:10.1146/annurev-orgpsych-120920-054654
5. 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
6. 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
7. 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
8. Newman A, Donohue R, Eva N. Psychological safety: A systematic review of the literature. *Hum Resour Manag Rev.* 2017;27(3):521-35. doi:10.1016/j.hrmr.2017.01.001
9. Sherf EN, Parke MR, Isaakyan S. Distinguishing voice and silence at work: Unique relationships with perceived impact, psychological safety, and burnout. *Acad Manag J.* 2021;64(1):114-48. doi:10.5465/amj.2018.1428
10. 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
11. 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
12. Castelo N, Bos MW, Lehmann DR. Task-dependent algorithm aversion. *J Mark Res.* 2019;56(5):809-25. doi:10.1177/0022243719851788
13. Frazier ML, Fainshmidt S, Klinger RL, Pezeshkan A, Vracheva V. Psychological safety: A meta-analytic review and extension. *Pers Psychol.* 2017;70(1):113-65. doi:10.1111/peps.12183
14. Hussain I, Shu R, Tangirala S, Ekkirala S. The voice bystander effect: How information redundancy inhibits employee voice. *Acad Manag J.* 2019;62(3):828-49. doi:10.5465/amj.2017.0245
15. Welsh DT, Outlaw R, Newton DW, Baer MD. The social aftershocks of voice: An investigation of employees' affective and interpersonal reactions after speaking up. *Acad Manag J.* 2022;65(6):2034-57. doi:10.5465/amj.2019.1187

16. 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
17. 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
18. Stice-Lusvardi R, Hinds PJ, Valentine M. Legitimizing illegitimate practices: How data analysts compromised their standards to promote quantification. *Organ Sci.* 2024;35(2):432-52. doi:10.1287/orsc.2022.1655
19. Argote L, Lee S, Park J. Organizational learning processes and outcomes: Major findings and future research directions. *Manag Sci.* 2021;67(9):5399-429. doi:10.1287/mnsc.2020.3693
20. Beane M. Shadow learning: Building robotic surgical skill when approved means fail. *Adm Sci Q.* 2019;64(1):87-123. doi:10.1177/0001839217751692
21. Sundaresan S, Guler I. Algorithmic recommendation tools and experiential learning in clinical care. *Organ Sci.* 2025;36(5):1786-802. doi:10.1287/orsc.2022.16738
22. Möhlmann M, Zalmanson L, Henfridsson O, Gregory RW. Algorithmic management of work on online labor platforms: When matching meets control. *MIS Q.* 2021;45(4):1999-2022. doi:10.25300/MISQ/2021/15333
23. Parent-Rochelleau 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
24. Bankins S, Ocampo AC, 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
25. Fügner 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. Dietvorst BJ, Simmons JP, Massey C. Overcoming algorithm aversion: People will use imperfect algorithms if they can (even slightly) modify them. *Manag Sci.* 2018;64(3):1155-70. doi:10.1287/mnsc.2016.2643
27. Nagtegaal R. The impact of using algorithms for managerial decisions on public employees' procedural justice. *Gov Inf Q.* 2021;38(1):101536. doi:10.1016/j.giq.2020.101536
28. Yin J, Ngiam KY, Tan SSL, Teo HH. Designing AI-based work processes: How the timing of AI advice affects diagnostic decision making. *Manag Sci.* 2025;71(11):9361-83. doi:10.1287/mnsc.2022.01454