



E-ISSN: 3108-4192

APSSHS

Academic Publications of Social Sciences and Humanities Studies

2026, Volume 6, Issue 1, Page No: 124-132

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

## Journal of Applied Organizational Systems and Behavior

# Safety and Dissent Rights in Algorithmic Work Systems

Ivan Petrov<sup>1</sup>, Olga Ivanova<sup>1\*</sup>, Dmitry Smirnov<sup>2</sup>

1. Department of Safety and Dissent Rights, Faculty of Business, Lomonosov Moscow State University, Moscow, Russia.
2. Department of Algorithmic Work Systems, Faculty of Management, Saint Petersburg State University, Saint Petersburg, Russia.

### Abstract

Algorithmic work systems are commonly evaluated through technical properties such as accuracy, reliability, consistency, or predictive performance. These properties are important but incomplete when algorithms also allocate work, shape discretion, structure evaluation, influence consequential judgments, and redistribute authority between employees, managers, and computational systems. This Original Conceptual Framework Article develops an organizational-safety perspective centered on employees' capacity to identify, communicate, escalate, contest, and interrupt algorithm-mediated decisions without disproportionate penalty. The article integrates scholarship on algorithmic control, work design, employee voice and silence, safety communication, human-AI judgment, and organizational governance. Its principal proposed contribution is an Organizational Safety and Dissent Rights Framework that distinguishes protected voice from mere channel availability, escalation from initial reporting, human override authority from nominal human presence, and non-retaliation from formal permission to speak. The framework proposes that organizational safety depends on connected capacities for hazard recognition, protected reporting, responsive listening and escalation, usable human override, reasoned review and recourse, and learning. These relationships are conceptual propositions rather than validated causal effects. Their operation is expected to depend on hierarchy, employment precarity, task stakes, time pressure, expertise, algorithmic opacity, local voice climate, and institutional context. The framework therefore offers a bounded basis for theorizing and testing organizational safety beyond technical reliability without claiming universal applicability, implementation readiness, or demonstrated organizational effectiveness.

**Keywords:** Algorithmic work, Organizational safety, Employee voice, Incident escalation, Human override, Dissent rights

**How to cite this article:** Petrov I, Ivanova O, Smirnov D. Safety and Dissent Rights in Algorithmic Work Systems. J Appl Organ Syst Behav. 2026;6(1):124-132. <https://doi.org/10.51847/GQpYgY34wD>

**Received:** 17 July 2026; **Revised:** 04 October 2026; **Accepted:** 08 October 2026

**Corresponding author:** Ivan Petrov

**E-mail** ✉ [olga.ivanova@mgu.ru](mailto:olga.ivanova@mgu.ru)

### Introduction

Algorithmic work systems do more than generate predictions or recommendations. They can allocate tasks, evaluate workers, rank options, structure performance visibility, and shape the terms through which managers and employees contest decisions. Research on algorithmic control accordingly portrays algorithms as part of an organizational terrain in which control and worker responses are negotiated rather than as neutral technical instruments [1]. This matters for safety because a technically dependable system can still participate in work arrangements that suppress challenge, obscure responsibility, or make correction difficult. Organizational safety therefore cannot be inferred from system reliability alone.

The coexistence of discretion and control further complicates simple classifications. Evidence from platform work shows that workers may experience flexibility while remaining subject to intensive algorithmic direction and evaluation [2]. Visible choice is therefore an unreliable proxy for substantive control. Related work conceptualizes algorithms as work designers because their use can alter autonomy, demands, feedback, and other job characteristics [3]. For safety analysis, the relevant



© 2026 The Author(s).

Copyright CC BY-NC-SA 4.0

unit is consequently the combined work system: technology, surrounding authority structures, and the conditions under which employees can question system outputs.

The same reasoning applies to human–AI role allocation. Automation and augmentation are theoretically interdependent rather than mutually exclusive alternatives [4]. A system can automate portions of judgment while increasing the importance of human interpretation, exception handling, or accountability. Retaining a nominal human role thus does not establish that employees possess the information, competence, time, or authority needed to intervene.

Work-design scholarship similarly indicates that the consequences of digital technologies depend substantially on choices concerning jobs, discretion, coordination, and demands [5]. Building on this literature, the present article proposes that organizational safety in algorithmic work systems concerns the capacity to recognize possible harm, raise concerns, obtain consequential attention, escalate unresolved issues, exercise legitimate intervention authority, and learn without suppressing future dissent. The resulting framework is an original non-empirical synthesis rather than a validated causal model.

### *Conceptual and governance foundations*

Algorithmic management should not be treated as uniformly autonomy-reducing. Conceptual research emphasizes its duality: algorithms may constrain discretion while also enabling coordination, information access, or value creation, depending on organizational arrangements and worker interactions [6]. The safety problem is therefore not algorithmic control by definition, but whether credible mechanisms of correction remain available when algorithm-mediated authority becomes consequential. Delegation provides a second foundation. Agentic information systems can receive tasks or decision roles from humans and participate in arrangements in which action is reassigned back to human actors [7]. Delegation, however, is not identical to legitimate authority. The proposed framework consequently distinguishes technological capability from organizational authorization and automation of action from transfer of responsibility. A system may be capable of ranking or triggering action without resolving who is entitled to accept, interrupt, reverse, or answer for that action.

Artificial intelligence also introduces organizational challenges related to autonomy, learning, opacity, and governance that cannot always be treated as conventional information-technology problems [8]. Greater technical sophistication nevertheless does not guarantee procedural legitimacy. Research on algorithmic reductionism demonstrates that attempts to reduce human bias can still generate procedural-justice concerns where relevant information is excluded from decision processes [9]. Technical improvement and organizational legitimacy are therefore analytically distinct.

Formal governance must also be separated from enacted governance. Written policies may allocate review or override authority while everyday work makes those rights difficult to exercise. The framework treats formal rules as evidence of organizational design rather than proof of functioning safety capacity. Its governance claims remain bounded: existing evidence does not identify one universally optimal allocation of human and algorithmic authority.

### *Employee voice protection*

Employee voice protection begins with conceptual precision. Contemporary scholarship distinguishes voice from silence and shows that both depend partly on expectations concerning interpersonal and organizational consequences [10]. In algorithmic work, the presence of a reporting channel therefore cannot demonstrate protection. Employees may formally possess an avenue for concern while expecting that speaking will have little effect or attract undesirable consequences.

Voice is also heterogeneous. Meta-analytic evidence distinguishes promotive voice from prohibitive voice directed toward problematic practices or risks [11]. Safety-relevant dissent is closer to, but not identical with, prohibitive voice. A challenge to an algorithmic recommendation may concern error, procedural unfairness, inappropriate authority allocation, or uncertainty. The proposed framework therefore does not classify every disagreement as a safety report or presume that every challenge is substantively correct.

Voice and silence are not simple opposites, and perceived impact and psychological safety relate differently to them [12]. Protection consequently requires more than permission to speak. Even expression does not guarantee influence. Longitudinal evidence on voice cultivation shows that upward input may require continued organizational work before reaching implementation [13]. The framework therefore distinguishes voice opportunity from consequential voice: a concern may be raised and acknowledged while remaining practically inert.

**Table 1** organizes the conceptual distinctions required to separate technical reliability, voice, escalation, override authority, and dissent protection.

**Table 1.** Definitions, constructs, and conceptual boundaries for Safety and Dissent Rights in Algorithmic Work Systems

| Construct or component       | Working definition                       | Organizational problem addressed    | Conceptual boundary | Proposed function         | Key uncertainty                                       | Interpretation boundary                  |
|------------------------------|------------------------------------------|-------------------------------------|---------------------|---------------------------|-------------------------------------------------------|------------------------------------------|
| <b>Technical reliability</b> | Dependable technical operation or output | Malfunction or unstable performance | Excludes authority, | Baseline safety condition | Reliable systems may still enable unsafe arrangements | Reliability is not organizational safety |

|                                                           |                                                                                                                                    |                                                 |                                              |                                               |                                           |                                                                    |
|-----------------------------------------------------------|------------------------------------------------------------------------------------------------------------------------------------|-------------------------------------------------|----------------------------------------------|-----------------------------------------------|-------------------------------------------|--------------------------------------------------------------------|
|                                                           |                                                                                                                                    |                                                 | voice, and response                          |                                               |                                           |                                                                    |
| <b>Organizational safety beyond technical reliability</b> | Proposed capacity to surface, evaluate, interrupt, and learn from algorithm-mediated risk                                          | Failures around technically functioning systems | Broader than accuracy, uptime, or compliance | <b>Proposed umbrella construct</b>            | Empirical distinctiveness is untested     | Not a validated safety measure                                     |
| <b>Algorithmic control</b>                                | Algorithm-mediated allocation, monitoring, evaluation, or direction of work                                                        | Opaque or concentrated control                  | Distinct from automation itself              | Defines the governance context                | Effects may be enabling or constraining   | Control does not imply harm                                        |
| <b>Protected employee voice</b>                           | Capacity to raise good-faith concerns without disproportionate penalty                                                             | Hidden concerns and unchallenged decisions      | More than permission to speak                | Entry route for challenge                     | Formal and enacted protection may diverge | Channel availability does not prove protection                     |
| <b>Consequential voice</b>                                | Proposed condition in which a concern receives substantive consideration                                                           | Symbolic participation                          | Beyond initial expression                    | <b>Proposed bridge from voice to response</b> | Requires operationalization               | Consideration does not require agreement                           |
| <b>Responsive escalation</b>                              | Proposed routing of unresolved concerns to competent authority                                                                     | Locally contained concerns                      | Broader than reporting                       | <b>Proposed coordination mechanism</b>        | Appropriate depth depends on stakes       | Not every concern warrants escalation                              |
| <b>Human override authority</b>                           | Legitimate, usable authority to pause, modify, defer, reject, or reroute algorithm-mediated action                                 | Nominal human oversight                         | More than a human “in the loop”              | <b>Proposed intervention capability</b>       | Scope depends on competence and risk      | Human intervention is not inherently superior                      |
| <b>Dissent rights and non-retaliation</b>                 | Proposed standing to challenge in good faith without disproportionate social, status, workload, opportunity, or procedural penalty | Chilling effects and suppressed contestation    | Normative governance principle               | <b>Proposed sustaining protection</b>         | Legal and institutional forms vary        | Dissent may be wrong; not every adverse consequence is retaliation |

Note. “Proposed” elements are original conceptual contributions rather than empirically validated constructs or universally applicable governance requirements.

### Error reporting and incident escalation

Error reporting becomes organizationally meaningful only when a concern can travel beyond the initial act of speaking. In healthcare settings, safety-management approaches and speak-up-related climates are associated with willingness to raise concerns and with reported speaking or withholding behavior [14, 15]. These findings support attention to local reporting conditions, but their cross-sectional designs do not establish that climate causes safer outcomes. Willingness, actual voice, and completed organizational response must therefore remain separate stages.

Evidence from aviation accidents sharpens this distinction. Analysis of cockpit voice recordings found safety voice in nearly all examined accidents, demonstrating that speaking can occur without effective resolution [16]. The relevant organizational mechanism is therefore not voice alone but what follows it: whether the concern is heard, interpreted, accepted as actionable, routed appropriately, and connected to an actor able to intervene. This article labels that sequence responsive escalation. The label is an original synthesis; the underlying distinction between safety voice and safety listening is evidence-supported.

Safety-related speaking also has antecedent conditions. A proactive goal-regulation perspective connects safety voice with safety envisioning and work contexts supportive of proactive action [17]. In algorithmic work, recognition may be complicated by opacity, uncertainty, fragmented responsibility, or habitual reliance on system outputs.

The proposed escalation sequence is recognition → reporting → listening → routing → response → closure. Only parts of that chain have direct empirical support, and evidence from healthcare and aviation should not be generalized mechanically to every algorithmic workplace. Its conceptual value lies in locating failure points that may remain invisible when reporting-channel availability is used as the principal indicator of safety.

### Human override authority

Human override should be distinguished from simply keeping a person “in the loop.” Experimental research shows that people may reject imperfect algorithms when they cannot influence them, while even limited modification rights can increase willingness to rely on algorithmic forecasts [18]. Yet algorithm appreciation also occurs, with people sometimes preferring algorithmic over human judgment [19]. Reliance is therefore an unsuitable safety criterion by itself. The organizational question is whether an actor can appropriately challenge or interrupt a system when circumstances warrant.

Human involvement also creates possible failure modes. Research on human–AI collaboration indicates that combining human and artificial intelligence is not inherently superior; benefits and drawbacks depend on task allocation and interaction [20]. Professional engagement with opaque diagnostic AI likewise varies with expertise and capacity to interpret outputs

contextually [21]. The article therefore proposes human override authority as legitimate and practically usable capacity to pause, modify, defer, reject, or reroute algorithm-mediated action.

Override has ex ante, operational, and ex post dimensions. Governance must determine who possesses intervention authority, what information they can access, whether intervention can actually be executed, and how decisions are subsequently reviewed. Delegation theory reinforces the need to separate allocation of system agency from continuing human responsibility [7]. Authority without competence or epistemic access may be hollow, while human involvement without clear decision rights may reproduce ambiguity rather than oversight.

**Table 2** organizes the proposed mechanisms linking algorithmic control, voice, escalation, override, dissent, and organizational learning while retaining rival explanations and validation requirements.

**Table 2.** Proposed mechanisms, relationships, and organizational consequences in Safety and Dissent Rights in Algorithmic Work Systems

| Mechanism or relationship                          | Starting condition                                                       | Proposed process                                                  | Boundary condition                     | Possible consequence                         | Competing explanation                             | Empirical test required                            |
|----------------------------------------------------|--------------------------------------------------------------------------|-------------------------------------------------------------------|----------------------------------------|----------------------------------------------|---------------------------------------------------|----------------------------------------------------|
| <b>Algorithmic control → constrained challenge</b> | Algorithm structures work or evaluation                                  | Fewer credible opportunities to question outputs                  | Hierarchy, precarity, opacity          | Reduced challenge                            | General managerial control                        | Compare algorithmic and non-algorithmic control    |
| <b>Voice → consequential attention</b>             | Concern is expressed                                                     | Listening, acknowledgment, and follow-through                     | Receptiveness, workload, issue quality | Review or inaction                           | Better concerns are easier to act on              | Trace concerns from expression to response         |
| <b>Reporting → responsive escalation</b>           | Initial recipient cannot resolve concern                                 | Routing to competent authority                                    | Time pressure, stakes, reversibility   | Intervention or unresolved hazard            | Severe issues naturally escalate                  | Analyze routing, authority, and closure            |
| <b>Override authority → intervention</b>           | Output is judged questionable                                            | Authorized pause, modification, deferral, rejection, or rerouting | Expertise, opacity, workload           | Correction or unnecessary disruption         | Expertise rather than authority explains outcomes | Compare nominal and usable override                |
| <b>Dissent → social or resource consequences</b>   | Employee challenges an accepted decision                                 | Status, threat, and workload dynamics shape reception             | Power, relationship quality            | Support, ostracism, burden, or later silence | Communication style drives reaction               | Separate concern quality from subsequent treatment |
| <b>Non-retaliation → sustained dissent</b>         | Prior dissent occurred                                                   | Response signals future safety and expected impact                | Precarity, career dependence           | Continued voice or strategic silence         | General trust explains future voice               | Longitudinal post-dissent analysis                 |
| <b>Integrated configuration → learning</b>         | Voice, listening, escalation, override, recourse, and protection coexist | Concerns become traceable inputs to review and redesign           | Cross-functional coordination          | Greater learning visibility                  | Generic quality capability                        | Test configurations against simpler models         |

Note. Possible consequences are theoretically bounded outcomes, not demonstrated causal effects. No row specifies a validated universal sequence.

### *Dissent rights and non-retaliation*

Speaking up can create costs even when organizations formally encourage voice. Evidence indicates that social reception partly depends on perceived voice quality and that voicers can experience ostracism [22]. This does not make every adverse reaction retaliation, but it demonstrates why formal permission is incomplete protection. Employees may anticipate relational penalties without any explicit punitive policy.

Leader response is likewise conditioned by hierarchy. Research shows that power threat and leader–member relationship quality can shape support for challenging voice [23]. Such dynamics may become particularly consequential when dissent questions decisions managers have delegated to, defended through, or invested in algorithmic systems. The article therefore treats dissent rights as proposed organizational standing: a good-faith challenge should be eligible for consideration without requiring the employee first to prove that the algorithm is wrong.

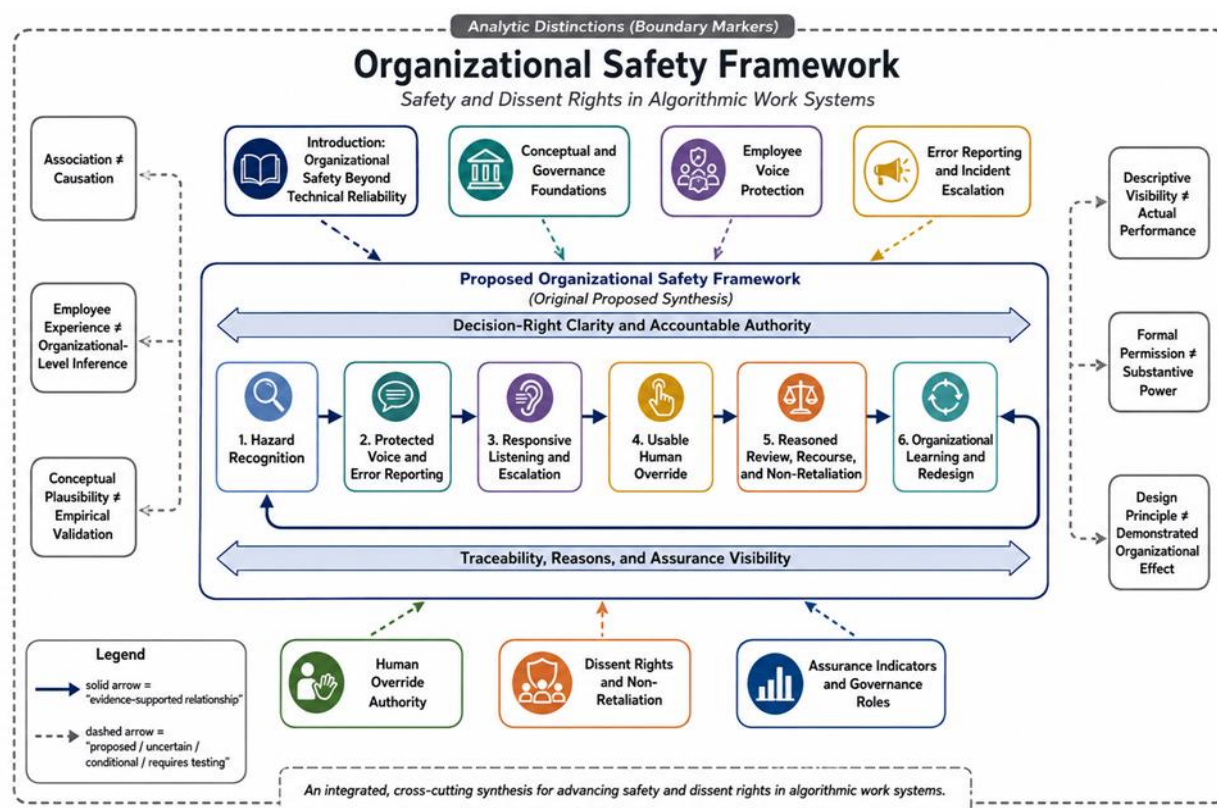
Costs may also be indirect. Speaking up can produce additional delegated work and resource depletion, generating regret even where the response appears supportive [24]. Non-retaliation should consequently extend analytically beyond formal punishment to disproportionate social, status, opportunity, workload, or procedural disadvantages that may chill future challenge. Legitimate performance evaluation nevertheless remains possible; not every unfavorable consequence is retaliatory.

The proposed dissent right is narrower than an unrestricted entitlement to oppose managerial decisions. It concerns good-faith challenges to algorithm-mediated action where safety, fairness, responsibility, or consequential work outcomes are plausibly

implicated. A defensible process should separate evaluation of the concern from treatment of the person raising it. An unsupported concern can therefore be rejected while maintaining a legitimate reporting process and providing reasons. The proposed right and non-retaliation obligation remain normative governance constructs rather than empirically established universal rights. Their form and effectiveness require testing across occupations, employment arrangements, institutional environments, and levels of algorithmic authority.

### Proposed organizational safety framework

Algorithmic coordination can become intertwined with organizational control, as research on platform matching demonstrates [25]. **Figure 1** presents the original conceptual synthesis for organizational safety framework, showing how the article's principal constructs, mechanisms, uncertainties, and decision boundaries are connected.



**Figure 1.** Organizational Safety Framework

The proposed framework connects six organizational capacities: hazard recognition; protected voice and error reporting; responsive listening and escalation; usable human override; reasoned review, recourse, and non-retaliation; and organizational learning and redesign. Two proposed infrastructures cross these capacities: decision-right clarity with accountable authority, and traceability with reason-giving and assurance visibility. The complete configuration and its cross-domain relationships are original propositions.

Research on gig work demonstrates that control and resistance may be co-constituted rather than independent [26]. Related evidence indicates that repeated worker choices can remain structurally confined by algorithmic arrangements [27]. Formal options therefore should not automatically be interpreted as substantive decision power. Surveillance research further shows that avoidance practices and managerial surveillance justifications can become mutually reinforcing [28]. The proposed framework incorporates this as a potential feedback process rather than assuming that stronger control necessarily produces safer compliance.

The framework is consequently cyclical rather than linear. Reasoned organizational response may influence later expectations concerning whether voice has impact; ignored or punitive responses may contribute to future withholding. Similarly, worker adaptation can provoke changes in control, producing further adaptation. Although evidence supports the plausibility of recursive control–resistance dynamics [26], it does not establish the direction or magnitude of the proposed safety-framework feedback loops.

Hierarchy, expertise, opacity, employment precarity, task stakes, reversibility, and time pressure are treated as boundary conditions rather than universal moderators. Their inclusion identifies conditions requiring empirical examination and prevents reporting-channel visibility from being interpreted as evidence that the complete safety architecture is functioning. **Table 3** organizes the proposed framework into falsifiable propositions, evaluation criteria, and practical governance responsibilities.

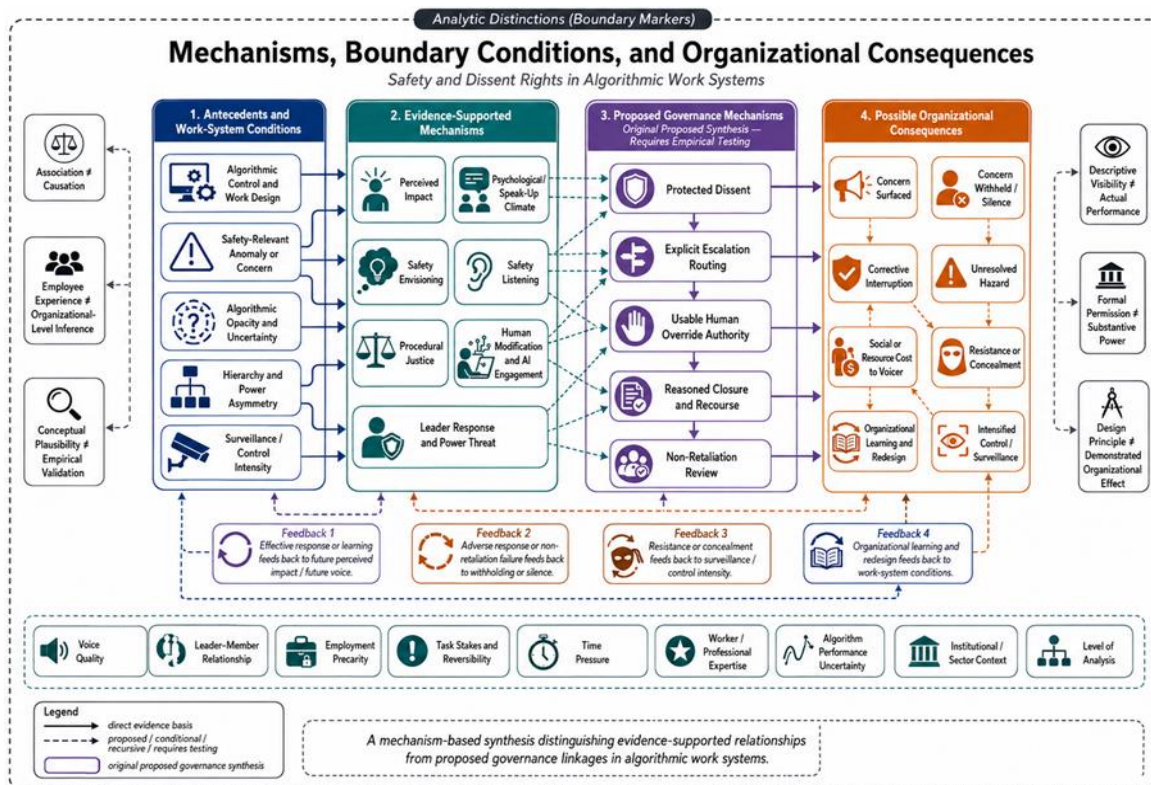
**Table 3.** Testable propositions, evaluation criteria, and practical responsibilities

| Proposition or evaluation criterion                                                                           | Primary unit            | Observable evidence                                                 | Evidence needed                       | Falsification condition                                    | Primary responsibility         | Misuse risk                        | Readiness boundary                                           |
|---------------------------------------------------------------------------------------------------------------|-------------------------|---------------------------------------------------------------------|---------------------------------------|------------------------------------------------------------|--------------------------------|------------------------------------|--------------------------------------------------------------|
| <b>P1. Protected voice without responsive listening is insufficient.</b>                                      | Employee–team–incident  | Concern is acknowledged, interpreted, and followed through          | Longitudinal incident records         | Reliable correction occurs without substantive listening   | Line management and safety     | Symbolic listening                 | Causal requirement unvalidated                               |
| <b>P2. Explicit escalation matters when initial recipients lack authority or expertise.</b>                   | Incident/process        | Traceable transfer to a named decision owner                        | Event logs and process tracing        | Ambiguous and explicit routing perform equivalently        | Operations, risk, system owner | Procedural overload                | Escalation depth is context-dependent                        |
| <b>P3. Override matters only when legitimate, executable, and competence-matched.</b>                         | Decision episode        | Authorized pause, modification, deferral, rejection, or rerouting   | Comparative field evidence            | Nominal presence performs equivalently to usable authority | System owner and operations    | Responsibility dumping             | No universal override threshold                              |
| <b>P4. Non-retaliation supports sustained dissent when protection is enacted.</b>                             | Employee/team over time | No disproportionate post-dissent penalty                            | Multi-source longitudinal evidence    | Formal policy alone predicts sustained voice               | HR, ethics, leadership         | Shielding poor performance         | Retaliation must be distinguished from legitimate management |
| <b>P5. Reasoned closure and traceability improve assurance visibility.</b>                                    | Incident/organization   | Rationale, owner, unresolved status, and recourse are documented    | Audit evidence linked to incidents    | Records add no useful governance information               | Assurance, risk, system owner  | Audit theatre                      | Visibility is not safety performance                         |
| <b>P6. Integrated capacities should explain organizational safety better than channel availability alone.</b> | Multi-level system      | Co-occurrence of voice, listening, override, recourse, and learning | Longitudinal/configurational research | Channel presence explains outcomes equally well            | Cross-functional governance    | Premature scoring or certification | Framework remains conceptual                                 |

Note. All propositions are original, testable conceptual claims. Responsibilities indicate plausible role ownership rather than universal prescriptions.

### *Assurance indicators and governance roles*

Digital high reliability depends on more than technically dependable components. Research on digital organizing emphasizes the organizational arrangements through which actors recognize limitations, frame anomalies, and compensate for what digital operations do not capture [29, 30]. **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

Assurance should focus on observable processes rather than an invented composite score. Candidate indicators include whether concerns receive acknowledgement, whether unresolved issues reach actors with relevant authority, whether override routes can actually be executed, whether decisions receive reasoned closure, whether recourse exists, and whether adverse consequences for good-faith dissent are examined. Safety voice climate is measurable as a distinct facet of perceived encouragement for safety-related speaking [31], but it should not be treated as a validated proxy for the complete organizational-safety framework.

Leading process evidence must also be distinguished from organizational outcomes. A recorded escalation or override can establish that a process occurred without proving that the underlying decision was correct or harm was prevented. Conversely, few recorded incidents can indicate safe operations, low exposure, weak detection, or suppressed reporting. Assurance therefore requires contextual interpretation rather than simple counts.

Formal availability and actual use may also diverge. Rare use of an override path can indicate safe operation, limited awareness, high procedural friction, or fear of consequences. Similarly, high reporting volume can reflect deteriorating conditions or a stronger reporting culture. Indicator interpretation therefore requires contextual comparison, qualitative follow-up, and attention to exposure. The framework consequently rejects fixed maturity rankings or universal thresholds at the conceptual stage.

Governance responsibilities should remain distributed rather than transferred to the employee who identifies a concern. System owners can maintain decision-right maps and technical traceability; line managers can receive and route operational concerns; safety, risk, or quality functions can examine escalation patterns; HR or employee-relations functions can examine protection and retaliation risks; technical owners can investigate system behavior; and appropriately independent assurance actors can examine whether formal controls operate in practice. These are proposed role allocations, not universally prescribed structures.

Governance independence is likewise a boundary rather than an assumed solution. Technical owners may possess expertise while having incentives to defend system performance; operational managers may understand context while facing throughput pressures; employee-relations functions may protect process while lacking technical authority. The framework therefore proposes complementary rather than singular ownership, calibrated to decision stakes and institutional context. Whether these arrangements improve detection, correction, trust, or learning remains an empirical question.

## Conclusion

Safety in algorithmic work systems cannot be inferred from technical reliability alone. When algorithms allocate tasks, shape evaluation, influence judgment, or mediate authority, organizational safety also depends on whether concerns can be

recognized, voiced, heard, escalated, acted upon, reviewed, and learned from without suppressing future challenge. This article proposes the Organizational Safety and Dissent Rights Framework to connect those capacities while preserving distinctions between voice and influence, human presence and usable override, formal permission and substantive dissent protection, and visibility and actual performance.

The framework is intentionally non-empirical. Its relationships, governance mechanisms, role allocations, and propositions require direct testing across occupations, employment arrangements, technological architectures, and institutional settings. Existing evidence supports constituent mechanisms and important boundaries but does not validate the integrated framework, prove causal effects, or establish a universally appropriate governance design. Its contribution is conceptual: to make organizational safety around algorithmic work more contestable, analytically differentiated, and empirically testable. Dissent is not presumed correct and override is not presumed beneficial; both become governed organizational processes whose safety value depends on context, authority, response, and evidence.

**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. Wood AJ, Graham M, Lehdonvirta V, Hjorth I. Good gig, bad gig: Autonomy and algorithmic control in the global gig economy. *Work Employ Soc.* 2019;33(1):56-75. doi:10.1177/0950017018785616
3. 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
4. Raisch S, Krakowski S. Artificial intelligence and management: The automation-augmentation paradox. *Acad Manage Rev.* 2021;46(1):192-210. doi:10.5465/amr.2018.0072
5. 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
6. Meijerink J, Bondarouk T. The duality of algorithmic management: Toward a research agenda on HRM algorithms, autonomy and value creation. *Hum Resour Manag Rev.* 2023;33(1):100876. doi:10.1016/j.hrmr.2021.100876
7. 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
8. Berente N, Gu B, Recker J, Santhanam R. Managing artificial intelligence. *MIS Q.* 2021;45(3):1433-50. doi:10.25300/MISQ/2021/16274
9. Newman DT, Fast NJ, Harmon DJ. When eliminating bias isn't fair: Algorithmic reductionism and procedural justice in human resource decisions. *Organ Behav Hum Decis Process.* 2020;160:149-67. doi:10.1016/j.obhdp.2020.03.008
10. 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
11. 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
12. Sherf EN, Parke MR, Isaakyan S. Distinguishing voice and silence at work: Unique relationships with perceived impact, psychological safety, and burnout. *Acad Manage J.* 2021;64(1):114-48. doi:10.5465/amj.2018.1428
13. Satterstrom P, Kerrissey M, DiBenigno J. The voice cultivation process: How team members can help upward voice live on to implementation. *Adm Sci Q.* 2021;66(2):380-425. doi:10.1177/0001839220962795
14. Alingh CW, van Wijngaarden JDH, van de Voorde K, Paauwe J, Huijsman R. Speaking up about patient safety concerns: The influence of safety management approaches and climate on nurses' willingness to speak up. *BMJ Qual Saf.* 2019;28(1):39-48. doi:10.1136/bmjqs-2017-007163
15. Schwappach DLB, Richard A. Speak up-related climate and its association with healthcare workers' speaking up and withholding voice behaviours: A cross-sectional survey in Switzerland. *BMJ Qual Saf.* 2018;27(10):827-35. doi:10.1136/bmjqs-2017-007388
16. Noort MC, Reader TW, Gillespie A. Safety voice and safety listening during aviation accidents: Cockpit voice recordings reveal that speaking-up to power is not enough. *Saf Sci.* 2021;139:105260. doi:10.1016/j.ssci.2021.105260

17. Curcuruto M, Strauss K, Axtell C, Griffin MA. Voicing for safety in the workplace: A proactive goal-regulation perspective. *Saf Sci.* 2020;131:104902. doi:10.1016/j.ssci.2020.104902
18. 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
19. 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
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. 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
22. Ng TWH, Wang M, Hsu DY, Su C. Voice quality and ostracism. *J Manage.* 2022;48(2):281-318. doi:10.1177/0149206320921230
23. Urbach T, Fay D. Leader member exchange in leaders' support for voice: Good relationships matter in situations of power threat. *Appl Psychol.* 2021;70(2):674-708. doi:10.1111/apps.12245
24. Newton DW, Sessions H, Lam CF, Welsh DT, Wu W. Loaded down from speaking up: A resource-based examination of voicer regret following supervisor delegation. *J Manage.* 2024;50(5):1911-38. doi:10.1177/01492063231163583
25. 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
26. Cameron LD, Rahman H. Expanding the locus of resistance: Understanding the co-constitution of control and resistance in the gig economy. *Organ Sci.* 2022;33(1):38-58. doi:10.1287/orsc.2021.1557
27. Cameron LD. The making of the "good bad" job: How algorithmic management manufactures consent through constant and confined choices. *Adm Sci Q.* 2024;69(2):458-514. doi:10.1177/00018392241236163
28. Anteby M, Chan CK. A self-fulfilling cycle of coercive surveillance: Workers' invisibility practices and managerial justification. *Organ Sci.* 2018;29(2):247-63. doi:10.1287/orsc.2017.1175
29. Salovaara A, Lyytinen K, Penttinen E. High reliability in digital organizing: Mindlessness, the frame problem, and digital operations. *MIS Q.* 2019;43(2):555-78. doi:10.25300/MISQ/2019/14577
30. Poláková-Kersten M, Khanagha S, van den Hooff B, Khapova SN. Digital transformation in high-reliability organizations: A longitudinal study of the micro-foundations of failure. *J Strateg Inf Syst.* 2023;32(1):101756. doi:10.1016/j.jsis.2023.101756
31. Mathisen GE, Tjora T. Safety voice climate: A psychometric evaluation and validation. *J Safety Res.* 2023;86:174-84. doi:10.1016/j.jsr.2023.05.008