

# Psychological Safety at Scale: Measuring and Building Trust in Modern Organizations

## Abstract and Executive Summary

### Abstract

Psychological safety—commonly defined as a shared belief that a team is safe for interpersonal risk taking—has become one of the most empirically grounded explanations for why some teams learn faster, surface problems earlier, and innovate more reliably than others. [1] Yet most organizations still struggle to “scale” psychological safety beyond pockets of high-performing teams because voice and silence are shaped by power dynamics, fear of reputation damage, and leaders’ moment-to-moment responses to risk taking. [2]

This whitepaper synthesizes foundational scholarship in organizational behavior and team learning, plus evidence from large organizations and applied research, to propose a privacy-first approach to measuring and improving psychological safety at scale. [3] It introduces a productization model for **RUDY AI** that uses (1) opt-in safety signals, (2) manager coaching recommendations grounded in research, and (3) team-level aggregation to create actionable insights without workplace surveillance. [4]

### Executive Summary

Psychological safety is best understood as an **enabler of learning behavior**—help seeking, asking questions, discussing errors, and experimentation—which then drives team performance outcomes. [5] Meta-analytic evidence spanning over 100 samples and tens of thousands of individuals supports psychological safety’s broad association with performance-relevant outcomes and clarifies that it operates at both individual and team levels. [6]

Employees often withhold ideas or concerns not because they “lack engagement,” but because speaking up can trigger predictable social and career risks (e.g., humiliation, being labeled difficult, or retaliation), especially under status and authority gradients. [7] These dynamics are amplified when managers experience upward voice as implicit criticism and react defensively. [8]

At scale, measurement must prioritize **trustworthiness**: collecting only what is needed, providing clear purpose limitation, and avoiding monitoring that can cause self-protection and reduce experimentation. [9] The recommended approach is a hybrid system: - team-level psychological safety measurement using validated items and short, recurring pulses [10]

- manager coaching interventions that reliably increase safety signals (e.g., inclusive invitation, consultative input-seeking, learning framing) [11]

- causal evaluation using cluster randomized A/B tests or stepped-wedge rollouts, tied to measurable KPIs (innovation throughput, retention risk, conflict quality). [12]

## Literature Review

### Psychological safety as a team-level climate for learning

Across the modern psychological safety literature, the most durable formulation is the team-level one: psychological safety is a **shared belief** about interpersonal consequences—whether speaking up will lead to embarrassment, rejection, or punishment—rather than a purely individual personality trait. [13] In field research on 51 work teams, psychological safety was associated with learning behavior, and learning behavior mediated the relationship between psychological safety and performance—supporting the argument that safety is a mechanism that enables the behaviors through which teams adapt. [14]

Critically, psychological safety is not synonymous with cohesion or “being nice.” Research distinguishes psychologically safe climates from cohesiveness because cohesion can suppress dissent (e.g., groupthink), whereas psychological safety supports productive disagreement and candor in service of shared goals. [15]

Psychological safety is also related to, but distinct from, interpersonal trust. Trust emphasizes expectations about others’ future actions and willingness to be vulnerable across time; psychological safety emphasizes the **immediate interpersonal risk calculus** of whether it is safe to ask a question, admit an error, or challenge a decision in the moment. [16]

### Team learning and innovation as outcomes of speaking up

Learning in teams requires interpersonal risk: surfacing errors, questioning assumptions, and experimenting with new practices. [5] In organizational contexts where performance depends on coordination, psychological safety helps shift attention from self-protection to problem solving—supporting early detection and correction of issues and enabling the iterative learning cycles that innovation demands. [16]

Applied evidence from Google[17]’s Google re:Work[18] synthesis of Project Aristotle identifies psychological safety as the most important of five team dynamics linked to team effectiveness, emphasizing that members must feel safe taking interpersonal risks such as asking “basic” questions, raising concerns, or offering new ideas. [19]

At the firm level, research in the Journal of Organizational Behavior argues that innovation efforts produce stronger performance returns when accompanied by climates including psychological safety; “innovation is not enough” without the climate conditions that allow people to raise problems and adapt practices. [20]

## Why employees don't speak up: fear, futility, and power

The employee silence literature frames silence as a systematic organizational problem, not an individual deficiency. Morrison and Milliken define “organizational silence” as widespread withholding of information about problems and issues, driven by contextual forces and shared sensemaking that speaking up is unwise; they argue this silence impedes organizational change and development. [21] A later integrative review explains that upward voice is the voluntary communication of suggestions or concerns to someone higher in the organization, and silence withholds potentially useful information—often with measurable costs to decision quality and organizational learning. [22]

A key driver is **fear of negative consequences**. Kish-Gephart and colleagues describe fear-based silence as a pervasive response to perceived threat cues—social and professional—leading employees to withhold ideas or concerns even when speaking up could help the organization. [23] In qualitative research on employee silence, Milliken, Morrison, and Hewlin highlight how employees develop “implicit theories” about the risks of speaking up, including image damage; they also connect how shared beliefs about danger and futility can spread through observation and social contagion. [24]

Power dynamics and managerial threat perceptions matter. Work on managers' perspectives suggests that improvement-oriented voice can feel personally threatening to managers because it implicitly critiques their competence or decisions; managers may therefore welcome some forms of input while resisting voice that challenges the status quo. [25] Consistent with this, Detert and Edmondson's “implicit voice theories” research emphasizes that employees hold taken-for-granted beliefs about when and why speaking up is risky or inappropriate—beliefs that can produce self-censorship even for pro-organizational suggestions. [26]

## Psychological safety and conflict quality

Innovation requires disagreement, but not all conflict is productive. A meta-analysis distinguishes task conflict (substantive disagreements) from relationship conflict (interpersonal friction) and finds relationship conflict is strongly negatively associated with team performance and satisfaction. [27] Psychological safety changes what conflict “does” in a team: evidence from 117 project teams indicates that psychological safety climate moderates the relationship between task conflict and performance, such that task conflict is positively associated with performance under conditions of high psychological safety. [28]

This connects directly to modern misconceptions: psychological safety does not avoid “hard truths.” Instead, it enables candid, constructive debate in service of improvement—and it is built interaction by interaction rather than mandated by policy. [29]

## Theoretical Framework

### A multilevel causal chain from safety to innovation

The most evidence-aligned model for business measurement treats psychological safety as a multilevel social climate that enables learning behavior, which then drives distal outcomes such as performance and innovation. [30] At a practical level, this implies psychological safety is not the “end metric”; it is a **leading indicator** for the behavioral mechanisms that produce innovation output: idea sharing, experimentation, coordination, and error detection. [31]

This view is reinforced by meta-analytic and systematic review work that maps psychological safety’s nomological network: antecedents include leadership and team climate variables, and outcomes include performance-relevant behaviors. [6] For scaling purposes, the key implication is that interventions must target both (a) interpersonal climate and (b) the structural and leadership conditions that shape it. [32]

### Power dynamics and risk as “friction” in the learning system

From a systems perspective, modern organizations often ask for “innovation behavior” while simultaneously creating conditions that punish the interpersonal risks required for learning—especially in hierarchical contexts where voice can be interpreted as criticism. [33] Fear-based silence and implicit “rules” of self-censorship can be viewed as endogenous frictions that reduce information flow in precisely the areas where leaders need bad news early and creative alternatives surfaced quickly. [34]

Psychological safety functions as a mechanism that reduces this friction by shifting the perceived payoff of voice: it lowers the expected interpersonal cost of speaking up. [5] The result is not “comfort,” but increased participation in hard conversations—help seeking, error admission, and critical questioning—behaviors empirically linked to learning and performance. [35]

### Scaling without surveillance

Scaling psychological safety introduces a paradox: measurement systems can unintentionally undermine safety if they increase workers’ sense of being watched—triggering self-protection rather than experimentation. [36] Bernstein’s “transparency paradox” argues that high observability can counterintuitively reduce performance by inducing self-protection and reducing local experimentation—precisely the behaviors psychological safety is meant to unlock. [37]

Therefore, psychological safety at scale must be operationalized using **privacy-preserving measurement and action loops**—collecting only what teams voluntarily contribute, tightly limiting interpretation to team-level patterns, and using the data primarily to support learning and leadership development rather than evaluation. [9]

# Productization: How RUDY AI Implements Psychological Safety at Scale

## Design principles for a privacy-first, non-surveillance system

RUDY AI's product thesis is that psychological safety improves when teams experience (1) low interpersonal fear, (2) visible leader inclusiveness, and (3) reliable action-following measurement—without exposing individuals to monitoring risk. [38] The system design is guided by four principles:

First, **autonomy and consent**: making participation opt-in supports autonomy, a core contributor to motivation and well-being in self-determination theory, and is also foundational to trust in measurement systems. [39] Second, **data minimization**: collect the smallest set of signals necessary to estimate team-level climate and to choose coaching interventions. [40] Third, **team-level focus**: psychological safety is a team climate construct, and teams exhibit shared perceptions shaped by common interactions and the same leader context. [5] Fourth, **purpose limitation**: measurement should support learning and improvement; otherwise, measurement itself can amplify self-protection. [41]

## Opt-in safety signals

RUDY AI operationalizes psychological safety through voluntary, low-burden “safety pulses” anchored in validated constructs. One model is a short recurring pulse derived from Edmondson's team psychological safety items (e.g., “It is safe to take a risk on this team,” “Members of this team are able to bring up problems and tough issues”), presented as an opt-in check-in after key moments such as team meetings, retrospectives, launches, or incident reviews. [10]

Signals are designed to reduce fear-based nonresponse: - **Private response** (employee completes on their own device). - **Time-bounded prompts** (e.g., 30 seconds, not a “survey project”). - **Non-identifying aggregation** (no individual dashboard, no ranking of employees). - **Explicit anti-surveillance UX** (clear descriptions of what is and is not collected). [42]

In practice, the system focuses on three opt-in signal categories aligned with the literature:

- **Interpersonal risk perception** (classic psychological safety items). [10]
- **Voice friction** (self-reported “I held back a concern because...” categories aligned to fear-based silence mechanisms). [43]
- **Learning behavior opportunity** (“We discussed an error openly,” “We tested a new approach,” “We asked for help”). [5]

## Manager coaching recommendations grounded in evidence

RUDY AI translates team-level signals into manager coaching behaviors that are repeatedly supported as antecedents of psychological safety.

Leader behaviors matter because team members attend especially to leaders' responses when interpreting interpersonal risk. In the foundational team study, leader coaching and context support are proposed antecedents of psychological safety, with leaders' supportive and non-defensive reactions shaping whether teams engage in risky learning behaviors. [5]

Evidence-based coaching recommendations include:

**Inclusive invitation and appreciation of contributions.** Nembhard and Edmondson define leader inclusiveness as words and deeds that invite and appreciate others' contributions and propose that it helps teams overcome inhibiting status differences—directly addressing power dynamics that suppress speaking up. [44]

**Consultative and supportive leadership routines.** A McKinsey survey and modeling approach reports that authoritative leadership behaviors are detrimental to psychological safety, while consultative and supportive behaviors promote psychological safety, partly through creating a positive team climate; consultative leadership includes soliciting input and considering the team's views, while supportive leadership involves showing concern for team members as individuals. [45]

**Learning-oriented framing and leader fallibility.** Google re:Work's guidance—grounded in research and internal findings—recommends that leaders frame work as a learning problem, acknowledge fallibility, and model curiosity by asking questions. [19]

**Conflict hygiene: enabling productive task conflict without relational harm.** Because psychological safety supports the benefits of task conflict, manager coaching can include structured debate practices (e.g., “surface two alternatives,” “ask for a disconfirming view”) designed to keep disagreements substantive rather than personal. [46]

RUDY AI surfaces these as specific behavioral prompts and micro-habits (e.g., “ask the quietest voice first,” “thank and restate dissent,” “name the learning goal”) rather than generic “be supportive” guidance, because psychological safety is built interaction by interaction. [29]

## Team-level aggregation and organizational insights

To scale responsibly, RUDY AI aggregates signals at the team level and only reports results when privacy thresholds are met. This aligns with the construct's team-level nature—teams develop shared perceptions because members are subject to similar interactions, manager behaviors, and resource access. [5]

The aggregation model is designed to support action without surveillance: - **Minimum n thresholds:** no reporting for small teams or low-response windows. - **Time windowing:** trends displayed over weeks/months to reduce any perception of “tracking an individual moment.” - **No individual inference:** no attempt to identify who is “safe” or “unsafe,” and no “voice scoring” of employees. [42]

At the organizational level, RUDY AI reports only de-identified distributions across teams (e.g., “percent of teams improving,” “teams at risk”) to support leadership development investment and resourcing—consistent with evidence that leadership development at scale influences psychological safety through leader behaviors and team climate. [47]

## Hypotheses, Measurement Models, KPI Mapping, and Experiments

### Hypotheses and measurable constructs

The hypotheses below translate psychological safety theory into testable statements for RUDY AI deployments. They are framed at the team level unless otherwise stated.

**H1: RUDY AI increases team psychological safety.** Teams with RUDY AI (opt-in pulses + manager coaching) will show higher psychological safety scores over time relative to control teams. [48]

**H2: The effect is mediated by manager behaviors and team climate.** Improvements in psychological safety will be mediated by increases in inclusive/consultative/supportive leader behaviors and positive team climate indicators. [49]

**H3: Psychological safety increases learning behavior and voice.** Improvements in psychological safety will predict increased learning behavior (help seeking, discussing errors, experimentation) and increased upward voice, consistent with foundational and review evidence. [50]

**H4: Innovation output improves through learning behavior.** Innovation throughput (e.g., experiments completed, ideas implemented) will improve, with learning behavior serving as a mediator between psychological safety and innovation outcomes. [51]

**H5: Retention improves via psychological safety.** Teams that increase psychological safety will show lower turnover intention and/or lower voluntary turnover risk, consistent with evidence that work group psychological safety associates with intent to remain and can be a strong predictor of turnover intention. [52]

**H6: Conflict becomes more productive.** Teams with increased psychological safety will show better conflict quality—task conflict more likely to correlate with performance, and reduced costly relationship conflict dynamics—consistent with psychological safety’s moderating role and conflict meta-analytic findings. [46]

### Measurement models for psychological safety at scale

A privacy-first measurement approach should treat psychological safety as a latent climate variable estimated via a combination of validated survey indicators and lightweight contextual signals, with clear separation between **measurement** and **interpretation**.

**Survey-based climate estimation.** Edmondson’s team psychological safety scale is widely used and can be implemented as a periodic pulse. Google re:Work reproduces the

seven items and positions them as a practical team measure (e.g., items about mistakes being held against someone, ability to raise problems, and safety in taking risks). [10]

**Team aggregation logic.** Because psychological safety is theorized and empirically treated as team-level in much of the foundational work, RUDY AI aggregates individual responses to a team score and focuses interpretation on team patterns, not individuals. [30]

**Measurement validity and limitations.** Systematic review work emphasizes that psychological safety research spans levels of analysis and that measurement challenges remain; therefore, organizations should explicitly document how they define psychological safety, how often they measure it, and how they separate psychological safety from related constructs. [53]

**Non-surveillance proxies.** RUDY AI should avoid automated inference from message content in collaboration tools as a default design choice, because increased observability can produce self-protective behavior and reduce experimentation. [41] Instead, “proxy” measures should be opt-in and employee-controlled (e.g., meeting reflections, voluntary tagging of moments where the employee withheld input, or structured check-ins after incidents). [54]

## Linking psychological safety to innovation output

The innovation linkage must be defined in measurable terms. Edmondson’s team learning model provides a direct template: psychological safety enables learning behavior, and learning behavior mediates performance outcomes. [5] Translating “performance” to innovation in modern organizations implies defining innovation outcomes as observable outputs of learning cycles, not as abstract creativity sentiment.

A practical linkage model uses three levels of innovation KPIs:

**Input KPIs (idea and concern flow).** Examples include number of improvement suggestions submitted, number of risks raised, and rate of cross-boundary coordination requests. This aligns with the view that voice and information sharing are central to organizational learning and process improvement. [55]

**Throughput KPIs (experimentation).** Examples include number of experiments launched and completed, cycle time from idea to test, and rate of documented retrospectives/postmortems. These map to learning behaviors such as experimentation and discussion of error. [5]

**Output KPIs (implemented innovation).** Examples include number of improvements shipped, process innovations adopted, or quality improvements embedded into practice. Firm-level evidence suggests that innovations produce better outcomes when supported by climates including psychological safety, underscoring that innovation capability depends on the interpersonal climate around risk taking. [56]

## KPI mapping for innovation, retention, and conflict reduction

To connect psychological safety to business outcomes without overclaiming, KPI mapping should distinguish between (a) leading indicators (climate + behavior) and (b) lagging indicators (organizational outcomes), and explicitly test mediation.

**Innovation KPIs.** Psychological safety is expected to increase innovation-relevant learning behaviors and contribute to team effectiveness in applied organizational settings. [57]

Candidate KPI relationships include: - Increased idea sharing and experimentation (leading) driven by reduced fear and increased learning behaviors. [58]

- Improved innovation implementation outcomes when psychological safety climate supports the interpersonal risk involved in challenging existing routines. [59]

**Retention KPIs.** Evidence from a large statewide study in public child welfare indicates that work group psychological safety had the strongest effect on turnover intention and was positively associated with intent to remain employed. [52] Organizationally, this can be operationalized as: - Voluntary attrition rate (lagging) - Turnover intention from periodic surveys (leading) - Regretted loss rate (lagging), evaluated cautiously due to definitional variance across firms. [52]

**Conflict reduction and conflict quality KPIs.** Relationship conflict is reliably costly for teams, negatively correlating with performance and satisfaction. [27] Psychological safety, however, can enable the benefits of task conflict—making disagreement productive rather than personal. [60] Candidate KPIs include: - Reduction in relationship conflict (leading), measured via validated conflict scales where appropriate and feasible. [27]

- Increase in “productive task conflict” behaviors (leading), such as voiced alternatives during decision-making. [28]

- Downstream reductions in costly escalations (lagging), such as formal grievances or repeated rework due to unraised risks, treated as context-specific and tested empirically rather than assumed. [61]

## Experimental design: A/B tests and stepped-wedge rollouts

To meet a publication-quality standard, the system should be evaluated with designs capable of supporting causal claims, not just correlations.

**Cluster randomized A/B tests (team-level randomization).** Because psychological safety is a team climate and interventions are delivered to managers and team routines, randomizing at the team level reduces contamination. Cluster randomized trial methods provide guidance on design and analysis when treatment is assigned to groups. [62]

A recommended A/B design for RUDY AI: - **Units:** teams (or stable workgroups) - **Arms:** (A) RUDY AI + manager coaching prompts + opt-in pulses; (B) measurement-only or business-as-usual - **Pre-period:** baseline survey and KPI capture (4–8 weeks) - **Intervention:** 8–16 weeks (or longer for lagging outcomes like attrition) - **Primary endpoints:** change in team psychological safety (survey), change in learning behavior proxy metrics - **Secondary**

**endpoints:** innovation throughput metrics; near-term conflict quality measures; retention intention where available [63]

Field experiment best practices emphasize clear outcome definitions, random assignment integrity, and interpretation aligned to design. [64]

**Stepped-wedge rollout (when randomization constraints exist).** If an organization requires that all teams eventually receive the intervention, stepped-wedge designs allow phased rollouts while preserving causal leverage, with established methods for analysis. [65]

**Mediation testing.** Because the core theory predicts effects via learning behavior, experiments should pre-specify mediation pathways (e.g., psychological safety → learning behavior → innovation throughput) rather than treating psychological safety as a standalone “culture score.” [5]

## Ethical Considerations

### Privacy-first data practices

A psychological safety system must earn trust by being structurally incapable of becoming a surveillance tool. Bernstein’s evidence on the transparency paradox underscores that increased observability can induce self-protection and reduce localized experimentation—outcomes directly opposed to the goals of psychological safety and innovation. [37] Therefore, ethical implementation requires: - tight limits on data collection and retention - aggregation by default - separation from performance evaluation systems. [42]

### Non-surveillance guardrails and anti-retaliation alignment

Because fear of negative consequences is a primary driver of silence, any system that could be perceived as identifying individuals increases the risk of backlash and reduced speaking up. [66] RUDY AI should prohibit individual-level reporting and should provide clear internal governance statements aligned to “speak-up” ethics and psychologically safe learning norms. [67]

### Fairness, bias, and the limits of measurement

NIST’s AI Risk Management Framework emphasizes that AI is socio-technical, with risks emerging from human and societal dynamics; it also describes trustworthy AI as privacy-enhanced and fair with harmful biases managed. [68] For psychological safety measurement, this means: - treating metrics as decision aids, not ground truth - monitoring for differential impacts (e.g., whether some demographic groups report systematically lower safety) - ensuring interventions do not “average away” marginalized experiences through aggregation without follow-up qualitative inquiry. [69]

## Change management and motivational integrity

Opt-in measurement supports autonomy and reduces coercion risk. Self-determination theory research emphasizes that employees' motivation and well-being depend strongly on whether motivation is autonomous versus controlled. [70] Organizations should therefore implement RUDY AI as a learning-support tool with clear value to employees (improving meeting norms, reducing blame responses, creating safer error discussions), not as a compliance mechanism. [29]

## References

- Baer, M., & Frese, M. (2003). Innovation is not enough: Climates for initiative and psychological safety, process innovations, and firm performance. *Journal of Organizational Behavior*, 24(1), 45–68. doi:10.1002/job.179 [20]
- Bernstein, E. S. (2012). The transparency paradox: A role for privacy in organizational learning and operational control. *Administrative Science Quarterly*, 57(2), 181–216. [37]
- Bradley, B. H., Postlethwaite, B. E., Klotz, A. C., Hamdani, M. R., & Brown, K. G. (2012). Reaping the benefits of task conflict in teams: The critical role of team psychological safety climate. *Journal of Applied Psychology*, 97(1), 151–158. doi:10.1037/a0024200 [28]
- Deci, E. L., Olafsen, A. H., & Ryan, R. M. (2017). Self-determination theory in work organizations: The state of a science. *Annual Review of Organizational Psychology and Organizational Behavior*, 4, 19–43. doi:10.1146/annurev-orgpsych-032516-113108 [71]
- De Dreu, C. K. W., & Weingart, L. R. (2003). Task versus relationship conflict, team performance, and team member satisfaction: A meta-analysis. *Journal of Applied Psychology*, 88(4), 741–749. [27]
- Detert, J. R., & Edmondson, A. C. (2011). Implicit voice theories: Taken-for-granted rules of self-censorship at work. *Academy of Management Journal*, 54(3), 461–488. [26]
- Edmondson, A. (1999). Psychological safety and learning behavior in work teams. *Administrative Science Quarterly*, 44(2), 350–383. [14]
- Edmondson, A. C. (2002). *Managing the risk of learning: Psychological safety in work teams* (Working Paper No. 02-062). Harvard Business School. [72]
- Edmondson, A. C., & Lei, Z. (2014). Psychological safety: The history, renaissance, and future of an interpersonal construct. *Annual Review of Organizational Psychology and Organizational Behavior*, 1, 23–43. doi:10.1146/annurev-orgpsych-031413-091305 [73]
- Frazier, M. L., Fainshmidt, S., Klinger, R. L., Pezeshkan, A., & Vracheva, V. (2017). Psychological safety: A meta-analytic review and extension. *Personnel Psychology*, 70(1), 113–165. doi:10.1111/peps.12183 [74]

Gerber, A. S., & Green, D. P. (2012). *Field experiments: Design, analysis, and interpretation*. W. W. Norton. [75]

Google re:Work. Guides: Understand team effectiveness (Project Aristotle). [19]

Hayes, R. J., & Moulton, L. H. (2017). *Cluster randomised trials* (2nd ed.). Taylor & Francis / CRC Press. [62]

Hussey, M. A., & Hughes, J. P. (2007). Design and analysis of stepped wedge cluster randomized trials. *Contemporary Clinical Trials*, 28(2), 182–191. doi:10.1016/j.cct.2006.05.007 [76]

Kish-Gephart, J. J., Detert, J. R., Treviño, L. K., & Edmondson, A. C. (2009). Silenced by fear: The nature, sources, and consequences of fear at work. *Research in Organizational Behavior*, 29, 163–193. doi:10.1016/j.riob.2009.07.002 [23]

Kruzich, J. M., Mienko, J. A., & Courtney, M. E. (2014). Individual and work group influences on turnover intention among public child welfare workers: The effects of work group psychological safety. *Children and Youth Services Review*, 42, 20–27. doi:10.1016/j.chilyouth.2014.03.005 [52]

McKinsey & Company. (2021). Psychological safety and the critical role of leadership development. [77]

MIT Sloan Management Review. (2023). Proven tactics for improving teams' psychological safety. [78]

Milliken, F. J., Morrison, E. W., & Hewlin, P. F. (2003). An exploratory study of employee silence: Issues that employees don't communicate upward and why. *Journal of Management Studies*, 40(6), 1453–1476. doi:10.1111/1467-6486.00387 [79]

Morrison, E. W. (2014). Employee voice and silence. *Annual Review of Organizational Psychology and Organizational Behavior*, 1, 173–197. doi:10.1146/annurev-orgpsych-031413-091328 [22]

Morrison, E. W., & Milliken, F. J. (2000). Organizational silence: A barrier to change and development in a pluralistic world. *Academy of Management Review*, 25(4), 706–725. [21]

National Institute of Standards and Technology. (2023). *Artificial Intelligence Risk Management Framework (AI RMF 1.0)* (NIST AI 100-1). [68]

Newman, A., Donohue, R., & Eva, N. (2017). Psychological safety: A systematic review of the literature. *Human Resource Management Review*, 27(3), 521–535. doi:10.1016/j.hrmr.2017.01.001 [80]

---

[1] [3] [5] [13] [14] [15] [30] [31] [32] [35] [38] [48] [50] [51] [57] [58] [63]

[https://web.mit.edu/curhan/www/docs/Articles/15341\\_Readings/Group\\_Performance/Edmondson%20Psychological%20safety.pdf](https://web.mit.edu/curhan/www/docs/Articles/15341_Readings/Group_Performance/Edmondson%20Psychological%20safety.pdf)

[https://web.mit.edu/curhan/www/docs/Articles/15341\\_Readings/Group\\_Performance/Edmondson%20Psychological%20safety.pdf](https://web.mit.edu/curhan/www/docs/Articles/15341_Readings/Group_Performance/Edmondson%20Psychological%20safety.pdf)

[2] [21] [61]

<https://openurl.ebsco.com/contentitem/doi%3A10.5465%2FAMR.2000.3707697?id=ebsco%3Aedb%3A3707697&sid=ebsco%3Aplink%3Asitemap>

<https://openurl.ebsco.com/contentitem/doi%3A10.5465%2FAMR.2000.3707697?id=ebsco%3Aedb%3A3707697&sid=ebsco%3Aplink%3Asitemap>

[4] [10] [11] [19] <https://rework.withgoogle.com/intl/en/guides/understanding-team-effectiveness>

<https://rework.withgoogle.com/intl/en/guides/understanding-team-effectiveness>

[6] [74] [https://digitalcommons.odu.edu/management\\_fac\\_pubs/13/](https://digitalcommons.odu.edu/management_fac_pubs/13/)

[https://digitalcommons.odu.edu/management\\_fac\\_pubs/13/](https://digitalcommons.odu.edu/management_fac_pubs/13/)

[7] [23] [34] [43] [66]

<https://www.sciencedirect.com/science/article/abs/pii/S019130850900015X>

<https://www.sciencedirect.com/science/article/abs/pii/S019130850900015X>

[8] [25] [33]

[https://msbfile03.usc.edu/digitalmeasures/nathanaf/intellcont/managing\\_to\\_stay\\_in\\_the\\_dark-1.pdf](https://msbfile03.usc.edu/digitalmeasures/nathanaf/intellcont/managing_to_stay_in_the_dark-1.pdf)

[https://msbfile03.usc.edu/digitalmeasures/nathanaf/intellcont/managing\\_to\\_stay\\_in\\_the\\_dark-1.pdf](https://msbfile03.usc.edu/digitalmeasures/nathanaf/intellcont/managing_to_stay_in_the_dark-1.pdf)

[9] [36] [37] [41] [42]

[https://www.hbs.edu/ris/Publication%20Files/Bernstein\\_TransparencyParadox\\_ASQ\\_June2012\\_cdaaee20-3a45-4a07-8867-9761a5d4b5e8.pdf](https://www.hbs.edu/ris/Publication%20Files/Bernstein_TransparencyParadox_ASQ_June2012_cdaaee20-3a45-4a07-8867-9761a5d4b5e8.pdf)

[https://www.hbs.edu/ris/Publication%20Files/Bernstein\\_TransparencyParadox\\_ASQ\\_June2012\\_cdaaee20-3a45-4a07-8867-9761a5d4b5e8.pdf](https://www.hbs.edu/ris/Publication%20Files/Bernstein_TransparencyParadox_ASQ_June2012_cdaaee20-3a45-4a07-8867-9761a5d4b5e8.pdf)

[12] [64] [75] <https://wnorton.com/books/9780393979954>

<https://wnorton.com/books/9780393979954>

[16] [59] [72] [https://www.hbs.edu/ris/Publication%20Files/02-062\\_0b5726a8-443d-4629-9e75-736679b870fc.pdf](https://www.hbs.edu/ris/Publication%20Files/02-062_0b5726a8-443d-4629-9e75-736679b870fc.pdf)

[https://www.hbs.edu/ris/Publication%20Files/02-062\\_0b5726a8-443d-4629-9e75-736679b870fc.pdf](https://www.hbs.edu/ris/Publication%20Files/02-062_0b5726a8-443d-4629-9e75-736679b870fc.pdf)

[17] [24] [79] [https://homepages.se.edu/cvonbergen/files/2012/12/AN-EXPLORATORY-STUDY-OF-EMPLOYEE-SILENCE\\_ISSUES-THAT-EMPLOYEES-DONT-COMMUNICATE-UPWARD-AND-WHY.pdf](https://homepages.se.edu/cvonbergen/files/2012/12/AN-EXPLORATORY-STUDY-OF-EMPLOYEE-SILENCE_ISSUES-THAT-EMPLOYEES-DONT-COMMUNICATE-UPWARD-AND-WHY.pdf)

[https://homepages.se.edu/cvonbergen/files/2012/12/AN-EXPLORATORY-STUDY-OF-EMPLOYEE-SILENCE\\_ISSUES-THAT-EMPLOYEES-DONT-COMMUNICATE-UPWARD-AND-WHY.pdf](https://homepages.se.edu/cvonbergen/files/2012/12/AN-EXPLORATORY-STUDY-OF-EMPLOYEE-SILENCE_ISSUES-THAT-EMPLOYEES-DONT-COMMUNICATE-UPWARD-AND-WHY.pdf)

[18] [20] [56] <https://onlinelibrary.wiley.com/doi/10.1002/job.179>

<https://onlinelibrary.wiley.com/doi/10.1002/job.179>

[22] [55] <https://www.annualreviews.org/content/journals/10.1146/annurev-orgpsych-031413-091328>

<https://www.annualreviews.org/content/journals/10.1146/annurev-orgpsych-031413-091328>

[26] <https://journals.aom.org/doi/10.5465/amj.2011.61967925>

<https://journals.aom.org/doi/10.5465/amj.2011.61967925>

[27]  
[https://web.mit.edu/curhan/www/docs/Articles/15341\\_Readings/Negotiation\\_and\\_Conflict\\_Management/De\\_Dreu\\_Weingart\\_Task-conflict\\_Meta-analysis.pdf](https://web.mit.edu/curhan/www/docs/Articles/15341_Readings/Negotiation_and_Conflict_Management/De_Dreu_Weingart_Task-conflict_Meta-analysis.pdf)

[https://web.mit.edu/curhan/www/docs/Articles/15341\\_Readings/Negotiation\\_and\\_Conflict\\_Management/De\\_Dreu\\_Weingart\\_Task-conflict\\_Meta-analysis.pdf](https://web.mit.edu/curhan/www/docs/Articles/15341_Readings/Negotiation_and_Conflict_Management/De_Dreu_Weingart_Task-conflict_Meta-analysis.pdf)

[28] [46] [60] <https://pubmed.ncbi.nlm.nih.gov/21728397/>

<https://pubmed.ncbi.nlm.nih.gov/21728397/>

[29] [54] <https://hsph.harvard.edu/news/debunking-misconceptions-about-workplace-psychological-safety/>

<https://hsph.harvard.edu/news/debunking-misconceptions-about-workplace-psychological-safety/>

[39] [70] [71] [https://selfdeterminationtheory.org/wp-content/uploads/2017/03/2017\\_DeciOlafsenRyan\\_annurev-orgpsych.pdf](https://selfdeterminationtheory.org/wp-content/uploads/2017/03/2017_DeciOlafsenRyan_annurev-orgpsych.pdf)

[https://selfdeterminationtheory.org/wp-content/uploads/2017/03/2017\\_DeciOlafsenRyan\\_annurev-orgpsych.pdf](https://selfdeterminationtheory.org/wp-content/uploads/2017/03/2017_DeciOlafsenRyan_annurev-orgpsych.pdf)

[40] [68] <https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>

<https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>

[44] [49] <https://www.utsouthwestern.edu/employees/leadership-programs/leadership-foundations/making-it-safe.pdf>

<https://www.utsouthwestern.edu/employees/leadership-programs/leadership-foundations/making-it-safe.pdf>

[45] [47] [77]

<https://www.mckinsey.com/~/media/McKinsey/Business%20Functions/Organization/Our%20Insights/Psychological%20safety%20and%20the%20critical%20role%20of%20leadership%20development/Psychological-safety-and-the-critical-role-of-leadership-development-final.pdf>

<https://www.mckinsey.com/~/media/McKinsey/Business%20Functions/Organization/Our%20Insights/Psychological%20safety%20and%20the%20critical%20role%20of%20leadership%20development/Psychological-safety-and-the-critical-role-of-leadership-development-final.pdf>

[52] <https://www.sciencedirect.com/science/article/abs/pii/S019074091400098X>

<https://www.sciencedirect.com/science/article/abs/pii/S019074091400098X>

[53] [80] <https://www.sciencedirect.com/science/article/pii/S1053482217300013>

<https://www.sciencedirect.com/science/article/pii/S1053482217300013>

[62] <https://www.taylorfrancis.com/books/mono/10.4324/9781315370286/cluster-randomised-trials-richard-hayes-lawrence-moulton>

<https://www.taylorfrancis.com/books/mono/10.4324/9781315370286/cluster-randomised-trials-richard-hayes-lawrence-moulton>

[65] [76] <https://pubmed.ncbi.nlm.nih.gov/16829207/>

<https://pubmed.ncbi.nlm.nih.gov/16829207/>

[67] <https://shop.sloanreview.mit.edu/fostering-ethical-conduct-through-psychological-safety>

<https://shop.sloanreview.mit.edu/fostering-ethical-conduct-through-psychological-safety>

[69] [73] <https://www.annualreviews.org/content/journals/10.1146/annurev-orgpsych-031413-091305>

<https://www.annualreviews.org/content/journals/10.1146/annurev-orgpsych-031413-091305>

[78] <https://shop.sloanreview.mit.edu/proven-tactics-for-improving-teams-psychological-safety>

<https://shop.sloanreview.mit.edu/proven-tactics-for-improving-teams-psychological-safety>