



Cognitive Monoculture and the Individual-Team Reliance Paradox: A Formal Model of Shared AI Dependence in Organizational Decision Making

Kwan Hong TAN¹

¹ *Singapore University of Social Sciences, 463 Clementi Road, Singapore 599494.*

Article Info

Article History:

Published: 20 August 2026

Publication Issue:

*Volume 3, Issue 8
August-2026*

Page Number:

465-480

Corresponding Author:

Kwan Hong TAN

Abstract:

Generative artificial intelligence can improve the quality and speed of individual knowledge work while creating an overlooked organizational risk: when many employees rely on the same model, their judgments cease to be statistically independent. This paper develops a formal theory of cognitive monoculture as common-mode dependence in human-artificial intelligence decision systems. An analytical model represents each employee's judgment as a weighted combination of an idiosyncratic human estimate and advice from one of a set of artificial intelligence sources. The model derives the team-level error variance, an individual-optimal reliance level, a team-optimal reliance level, a common-model dominance threshold, and an effective independent judgment count. The central result is an individual-team reliance paradox: a reliance level that minimizes each person's expected error can increase the error of an aggregated team decision because common model error does not diversify away. Under an illustrative baseline, a twenty-person team using a single shared model more than doubles its team-level error when all members adopt the individually optimal reliance level, while its nominal twenty judgments collapse to approximately 1.31 effective independent judgments. Diversifying model sources mitigates the effect only when cross-model errors are imperfectly correlated. The framework therefore shifts artificial intelligence governance from counting human reviewers to measuring the independence architecture of their judgments.

Keywords: artificial intelligence governance, cognitive monoculture, collective intelligence, correlated error, decision diversity, human-AI collaboration

1. Introduction

Generative artificial intelligence has moved rapidly from an experimental interface to an ordinary layer in professional work. Controlled evidence shows that generative systems can raise productivity and improve output quality on many writing-intensive tasks [1]. In organizations, this value proposition is usually framed at the level of the individual worker: a manager, analyst, engineer, researcher, or service employee becomes faster or more capable when paired with artificial intelligence. Research on human-AI symbiosis and the automation-augmentation paradox similarly emphasizes how algorithmic capabilities can complement human judgment [2], [3]. Organizational design research has therefore concentrated on allocating decision rights between people and algorithms and on identifying conditions under which one or the other should lead [4].

The unit of analysis changes, however, when many people use the same artificial intelligence system. A team does not gain robustness merely because several human names appear in a review chain. If those reviewers independently consult the same model, accept similar suggestions, or begin from a common generated draft, part of their error becomes shared. The organization may preserve nominal human oversight while losing the statistical independence that makes multiple reviewers, committees, expert panels, and distributed teams valuable. This paper calls that condition cognitive monoculture: a reduction in the effective diversity of organizational judgment caused by common dependence on one or a small number of correlated artificial intelligence sources.

This problem is not reducible to conventional automation bias. Classic human-factors research showed that automation can produce new forms of vulnerability, misuse, overreliance, and monitoring failure [5]-[8]. Those mechanisms primarily concern the quality of an individual's relationship with automation. The present problem is relational and organizational. Even when every employee is rationally using a model that is more accurate than the employee acting alone, the shared model can make employee errors correlate. The resulting system may become less resilient to model-specific bias, hallucination, framing, omission, stale knowledge, or systematic misclassification precisely because many people are consulting the same epistemic source.

The risk is increasingly plausible because recent evidence suggests that generative artificial intelligence can improve individual performance while compressing collective variation. Theoretical work on diversity and collective problem solving has long shown why heterogeneous perspectives can improve group performance [9]. Experimental work demonstrates that social influence can undermine the wisdom of crowds by reducing the diversity and independence of estimates [10]. More recently, generative artificial intelligence has been shown to enhance individual creativity while reducing the collective diversity of outputs [11], and controlled studies of large-language-model-assisted ideation and open-ended survey responses report related homogenization effects [12], [13]. These findings motivate a broader question: what happens when homogenization affects not only style or ideas, but the error structure of consequential organizational decisions?

The paper answers this question with a parsimonious formal model. It makes four contributions. First, it defines cognitive monoculture in probabilistic rather than rhetorical terms, as common-mode dependence among otherwise distinct human judgments. Second, it proves an individual-team reliance paradox: the artificial intelligence weight that minimizes an individual's error is generally higher than the weight that minimizes the error of a team average when the same model is shared. Third, it derives a common-model dominance threshold and an effective independent judgment count that translate hidden dependence into auditable quantities. Fourth, it formalizes model-source diversification and shows that adding providers or model families is useful only to the extent that their errors are genuinely less correlated. These results connect human factors, collective intelligence, organizational design, statistical decision theory, and artificial intelligence governance.

The analysis is deliberately model-based rather than empirical. Numerical values are illustrative, not universal prescriptions. The purpose is to identify mechanisms and thresholds that can later be calibrated in specific settings such as investment committees, medical review teams, procurement boards, safety assurance, compliance, academic evaluation, and strategic planning. This distinction is important: the contribution is not a claim that a particular percentage of artificial intelligence reliance is always unsafe. It is a demonstration that organizational accuracy and independence can move in different directions, and that governance based only on individual productivity, reviewer headcount, or model accuracy can therefore be systematically misleading.

2. Literature Review and Theoretical Gap

2.1 Automation reliance and human oversight

Bainbridge's classic account of the ironies of automation established that automating a task can leave humans responsible for the hardest residual situations while simultaneously depriving them of the practice needed to handle those situations [5]. Parasuraman and Riley distinguished appropriate use from misuse, disuse, and abuse of automation [6], while later work on levels of automation emphasized that automation design redistributes information acquisition, analysis, decision selection, and action functions between humans and machines [7]. Lee and See subsequently framed trust calibration as a central condition for appropriate reliance [8]. Together, these traditions imply that human oversight is not meaningful merely because a human remains formally present.

Generative artificial intelligence intensifies this problem because the system is often not a bounded control device but a general-purpose cognitive intermediary. It can frame a problem, propose alternatives, draft an argument, summarize evidence, rank options, and generate the language through which a later reviewer evaluates the decision. Reliance can

therefore enter upstream, before a reviewer experiences the artifact as an object of independent judgment. Rinta-Kahila et al. show how cognitive automation may contribute to skill erosion through reinforcing loops of reliance, complacency, and reduced mindful conduct [14]. Such mechanisms matter for cognitive monoculture because weakly exercised human expertise makes independent challenge less likely even when an organization retains nominal review steps.

2.2 Collective intelligence requires useful independence

Collective intelligence depends not simply on adding people, but on combining information that is at least partly nonredundant. Hong and Page demonstrate, under formal conditions, that diverse problem solvers can outperform collections of individually high-ability problem solvers because diversity expands the search repertoire of the group [9]. Lorenz et al. provide an empirical complement: social influence can reduce diversity without necessarily improving accuracy, thereby undermining the wisdom-of-crowds effect [10]. The implication for artificial intelligence adoption is direct. A tool can increase the quality of every individual's starting point while simultaneously shrinking the independence of the errors that remain.

This distinction between competence and dependence is central. Suppose twenty analysts each make an estimate. If their errors are largely independent, averaging can cancel much of the idiosyncratic noise. If all twenty anchor on the same artificial intelligence recommendation, the number of visible reviewers remains twenty, but a nontrivial part of the final error is repeated twenty times. The organization has acquired a common-mode component analogous to correlated failure in engineering or correlated assets in portfolio risk. The key governance variable is therefore not only how accurate each decision maker is, but how much independent information survives after shared artificial intelligence use.

2.3 Generative artificial intelligence and homogenization

Emerging generative artificial intelligence research makes this concern empirically concrete. Doshi and Hauser found that access to generative artificial intelligence increased assessed creativity for individual short stories but made the set of stories more similar to one another [11]. Anderson, Shah, and Kreminski similarly reported that large-language-model assistance in creative ideation could increase output quantity and detail while decreasing semantic distinctiveness across users [12]. Zhang, Xu, and Alvero found artificial-intelligence-mediated homogenization in open-ended survey responses, creating a methodological threat when researchers assume that respondents independently generated their language [13]. These studies differ in task and outcome, but all expose the same structural tension: individual uplift can coexist with collective convergence.

Most homogenization studies focus on semantic or creative diversity rather than decision error. Yet the organizational consequence may be more serious when common suggestions affect classification, forecasting, risk assessment, legal reasoning, hiring, procurement, strategy, or safety review. In those contexts, similarity is not merely aesthetic. It changes the covariance structure of judgment. A shared blind spot can survive multiple review stages because each reviewer independently reproduces the same model-mediated premise.

2.4 Governance gap

Contemporary governance frameworks clearly recognize the need for lifecycle risk management, measurement, oversight, and accountability. The National Institute of Standards and Technology Artificial Intelligence Risk Management Framework organizes governance around the functions Govern, Map, Measure, and Manage [15], while its generative artificial intelligence profile adds risk-specific considerations for generative systems [16]. ISO/IEC 42001 embeds artificial intelligence in a management-system cycle [17]. The European Union Artificial Intelligence Act includes human oversight among the requirements applicable to high-risk systems [18]. These instruments are important, but they do not by themselves make a multi-person review independent. An organization can comply with a human-oversight requirement while every human in the chain consults the same model.

A related prior contribution by Tan formalizes how artificial intelligence controls can lose effectiveness over time and therefore require risk-sensitive revalidation [19]. The present paper addresses a different axis of assurance. Governance half-life concerns temporal durability: whether a control remains valid as systems and environments change. Cognitive monoculture concerns cross-sectional dependence: whether multiple contemporaneous judgments are genuinely independent enough to diversify error. Together, the two problems suggest that trustworthy artificial intelligence governance requires both evidence renewal over time and independence across reviewers or epistemic sources.

The unresolved research gap is therefore a quantitative one. Existing literatures provide strong reasons to care about human reliance, diversity, homogenization, and oversight, but there is no simple organizational model that simultaneously links team size, human error correlation, artificial intelligence reliance, model-source diversity, cross-model correlation, and the effective number of independent judgments. The next section develops such a model.

3. Method and Analytical Framework

3.1 Research design

The study uses analytical theory building supported by deterministic sensitivity analysis and Monte Carlo verification. No human participants, personal data, proprietary organizational data, or externally supplied datasets are used. The model is intentionally minimal so that the effect of shared artificial intelligence dependence can be isolated. The analysis begins with an unbiased stochastic case, derives closed-form results, then examines model-source diversification and systematic-bias extensions. Numerical scenarios are illustrative and are chosen to expose comparative statics rather than to estimate population parameters.

The baseline considers a team of human decision makers estimating an unknown scalar target. A scalar representation is sufficient because many multidimensional organizational choices can be locally interpreted as estimating a score, expected value, risk level, probability, ranking utility, or latent decision criterion. The same logic extends to vector outcomes by replacing variances with covariance matrices, but the scalar case yields transparent thresholds suitable for governance use.

3.2 Human and artificial intelligence judgments

Let the unaided judgment of an arbitrary human decision maker be the true target plus an error term:

$$h_i = \theta + \varepsilon_i \quad (1)$$

Human errors are assumed exchangeable, with a common individual error variance and a common pairwise error correlation. The variance of the average unaided human error is therefore:

$$A_N = \sigma_h^2 [\rho_h + (1 - \rho_h)/N] \quad (2)$$

Equation (2) makes the aggregation advantage explicit. If human errors are independent, the variance falls in inverse proportion to team size. If human errors are positively correlated, the variance approaches a nonzero floor as team size increases. Human teams are therefore never assumed to possess perfect independence; the model asks how shared artificial intelligence adds further dependence to whatever correlation already exists.

Each decision maker combines an unaided human judgment with advice from an assigned artificial intelligence source. Let the reliance weight range from zero to one:

$$y_i = (1 - r)h_i + ra_{g(i)} \quad (3)$$

The baseline uses a common organizational reliance level. In practice, reliance can represent explicit numerical weighting, de facto anchoring, the fraction of a decision pathway delegated to a model, or an empirically estimated influence coefficient. Artificial intelligence source errors are characterized by a single-source error variance and a

pairwise cross-source error correlation. When distinct sources are allocated evenly across the team, the variance of their average contribution is:

$$B_M = \sigma_A^2 [\rho_A + (1 - \rho_A)/M] \quad (4)$$

With one artificial intelligence source, Equation (4) reduces to the variance of a single source because every team member inherits the same model error. Adding distinct sources reduces the aggregated artificial intelligence variance only when cross-model errors are imperfectly correlated. If cross-source error correlation is perfect, adding models changes names but not the common-mode risk. This captures the practical distinction between vendor diversity and genuine error diversity.

3.3 Team error, optimal reliance, and the reliance paradox

Assume initially that human and artificial intelligence errors are mutually independent and unbiased. The mean squared error of the team average is then equal to its variance:

$$V(r) = (1 - r)^2 A_N + r^2 B_M \quad (5)$$

Differentiating Equation (5) yields the team-optimal reliance level:

$$r_{team}^* = A_N / (A_N + B_M) \quad (6)$$

By contrast, minimizing the mean squared error of a single assisted individual yields:

$$r_{ind}^* = \sigma_h^2 / (\sigma_h^2 + \sigma_A^2) \quad (7)$$

Equations (6) and (7) generate the individual-team reliance paradox. For any team with more than one member and less-than-perfect human error correlation, aggregation makes the human variance term smaller than the variance of one unaided human. With one shared artificial intelligence source, the artificial intelligence variance term does not receive an equivalent diversification benefit. Consequently, team-optimal reliance is strictly lower than individual-optimal reliance. The reliance level that is individually accuracy-maximizing is therefore not generally team-optimal. Each person has an incentive to exploit the model's individual advantage, but simultaneous use converts that advantage into common dependence. The individual decision maker does not internalize the correlation imposed on the team aggregate.

3.4 Common-model dominance and harmful-reliance thresholds

A useful diagnostic is the point at which the artificial intelligence component contributes more variance to the team average than the residual human component. Setting the two terms of Equation (5) equal gives the common-model dominance threshold:

$$r_{dom} = \sqrt{A_N} / (\sqrt{A_N} + \sqrt{B_M}) \quad (8)$$

Above the dominance threshold in Equation (8), more than one-half of the team's stochastic error variance originates in the shared artificial intelligence layer. A second threshold answers a stricter question: when does the assisted team become worse than the unaided team? Solving Equation (5) against the no-assistance baseline yields a positive harmful-reliance boundary equal to twice the aggregated human variance divided by the sum of the aggregated human and artificial intelligence variance terms, when that boundary lies within the unit interval. The harmful-reliance boundary is performance-based, while the dominance threshold is architecture-based. A team can become common-mode dominated before its measured error is worse than baseline.

3.5 Effective independent judgment count

Reviewer headcount can conceal dependence. To translate the covariance architecture into an intuitive governance quantity, define the effective independent judgment count as the number of equally accurate independent judgments

that would produce the same variance of the mean. Let the numerator in Equation (9) represent the variance of one assisted individual's judgment. Then:

$$N_{eff} = D(r) / V(r) \quad (9)$$

When all twenty reviewers rely completely on one shared model, the effective independent judgment count tends to one even though the nominal headcount remains twenty. With several distinct but correlated model sources, the limiting value is greater than one but can remain far below both the number of reviewers and the number of nominal models. The measure therefore separates procedural multiplicity from epistemic multiplicity.

3.6 Systematic bias extension

The unbiased core is conservative with respect to common-mode risk because systematic error does not diversify at all. If both the average human judgment and the artificial intelligence layer may be biased, the team mean squared error becomes:

$$MSE(r) = [(1 - r)b_h + rb_A]^2 + (1 - r)^2A_N + r^2B_M \quad (10)$$

Equation (10) shows why correlated variance is only part of the problem. A shared artificial intelligence bias, such as systematic omission of a regulatory exception or consistent underestimation of a rare hazard, survives every reviewer who inherits it. The closed-form results below use zero bias to isolate the reliance mechanism; nonzero common bias would generally strengthen the case for independent challenge.

Table 1. Model notation and governance interpretation

Symbol	Definition	Governance interpretation
N	Number of human decision makers	Nominal human review headcount
M	Number of substantively distinct AI sources	Distinct model families or independently governed model sources
r	AI reliance weight, 0 to 1	Influence of AI advice on each final judgment
σ_h^2	Variance of unaided human error	Individual human uncertainty
ρ_h	Pairwise correlation of human error	Pre-existing dependence among human judgments
σ_A^2	Variance of one AI source error	Uncertainty of the model advice
ρ_A	Pairwise correlation of AI-source error	Degree to which model sources fail alike
A_N	Variance of the mean unaided human error	Residual human error after aggregation

Symbol	Definition	Governance interpretation
B_M	Variance of the mean AI-source error	Residual common-model error after source diversification
N_{eff}	Effective independent judgment count	Independence-equivalent size of the review system

3.7 Baseline parameterization and verification

The illustrative baseline sets individual human error variance to 1.00, pairwise human error correlation to 0.05, and single-source artificial intelligence error variance to 0.36. Thus the artificial intelligence source is substantially more accurate than one unaided human judgment in variance terms, while human reviewers retain modest positive correlation. The primary scenario uses one artificial intelligence source. Diversification scenarios set pairwise cross-source error correlation to 0.25, meaning distinct artificial intelligence sources share some common error but are far from perfectly correlated. Team sizes of 5, 20, and 100 illustrate small, medium, and large review structures.

The closed-form calculations were independently checked with 100,000-trial Gaussian Monte Carlo simulations using a fixed random seed. Across eight validation cases spanning one and four artificial intelligence sources and reliance values from zero to the individual optimum, the maximum relative difference between simulated and analytical team mean squared error was 0.55%. The simulations therefore serve as a numerical verification of the algebra rather than as an empirical estimate of organizational behavior.

4. Results and Analysis

4.1 The individual-team reliance paradox

Under the baseline variances, the individually optimal artificial intelligence reliance is 0.735. An isolated decision maker can therefore place roughly three-quarters of the weight on the more accurate artificial intelligence signal and reduce personal expected error substantially. The team optimum is very different because the same model error is repeated across the team. For a twenty-person team using one shared source, the aggregated unaided human error variance is 0.0975 and team-optimal reliance falls to 0.213. Figure 1 displays this divergence. The two optimization problems are not equivalent: the individual minimizes personal error variance, whereas the team must also price the correlation externality generated by shared model use.

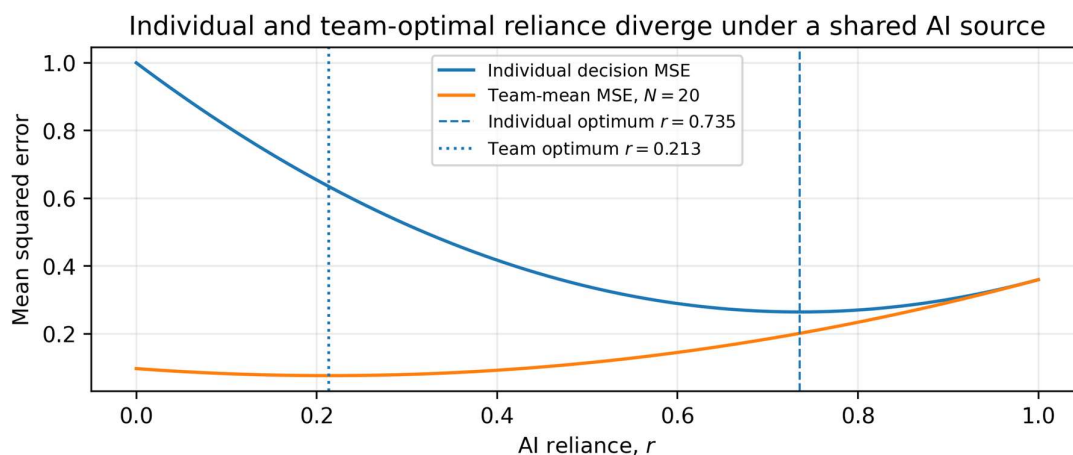


Figure 1. Individual and team-optimal artificial intelligence reliance under the baseline shared-model scenario. The individual optimum occurs at reliance 0.735, while the twenty-person team optimum occurs at reliance 0.213.

The paradox becomes visible when the organization allows every member to adopt the individual optimum. For a twenty-person team, the unaided team mean squared error is 0.0975. At the individually optimal reliance level of 0.735 with a shared model, the team mean squared error rises to approximately 0.2015, or 2.07 times the unaided value. No individual is behaving irrationally under the individual objective. The failure comes from composing individually rational decisions into a correlated system.

Table 2. Shared-model results as team size increases

N	A_N	Team-optimal r	Dominance r	MSE at individual optimum / no-AI MSE	N_{eff} : no AI to individual optimum
5	0.2400	0.400	0.449	0.88x	4.17 -> 1.25
20	0.0975	0.213	0.342	2.07x	10.26 -> 1.31
100	0.0595	0.142	0.289	3.34x	16.81 -> 1.33

Table 2 reveals a team-size inversion. A five-person team still benefits slightly when all members use the individual-optimal reliance level; its assisted team error is about 0.88 times the unaided level. At twenty people, the same behavior becomes harmful. At one hundred people, the team error is about 3.34 times the unaided level. The reason is not that artificial intelligence becomes less accurate in larger teams. It is that the diversified human component becomes more accurate as team size grows, while the single model's common error does not shrink. Consequently, a large team has more collective intelligence to lose when independent judgments are collapsed onto a shared source.

4.2 Common-model dominance occurs at lower reliance in larger teams

Figure 2 plots team mean squared error against reliance for teams of 5, 20, and 100 people. The curves converge as reliance approaches one because the human aggregation advantage disappears. The common-model dominance threshold falls from approximately 0.449 for a five-person team to 0.342 for a twenty-person team and 0.289 for a one-hundred-person team. Larger teams can therefore become common-model dominated at lower artificial intelligence reliance. This result is counterintuitive if headcount is treated as a proxy for robustness: increasing the number of reviewers can make the organization more sensitive to a common epistemic dependency because the counterfactual value of human diversification is larger.

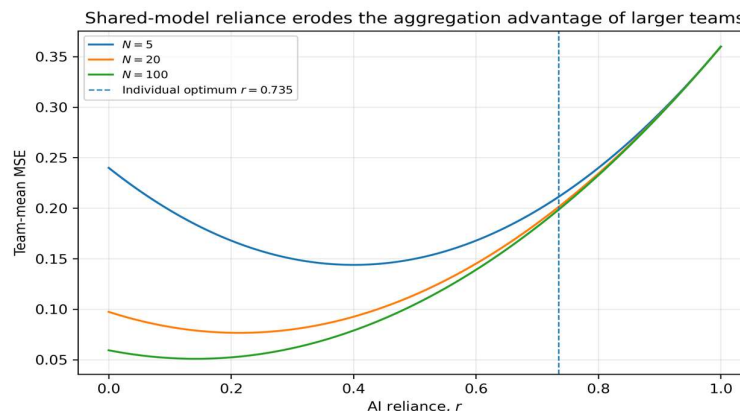


Figure 2. Team mean squared error across artificial intelligence reliance levels for different team sizes using one shared model. Larger teams have lower error at low reliance but lose that advantage as common-model dependence increases.

The distinction between the dominance threshold and the harmful-reliance threshold is operationally useful. At the dominance threshold, shared artificial intelligence already contributes more than half of the stochastic error variance, yet total error may still be below the no-artificial-intelligence baseline. An organization that monitors only average accuracy may therefore conclude that the system is healthy even as resilience deteriorates. A subsequent model update, domain shift, or systematic bias can then affect many reviewers simultaneously because the independence buffer has already been consumed.

4.3 Effective independence collapses faster than reviewer headcount suggests

The effective independent judgment count exposes this hidden loss. With pairwise human error correlation of 0.05, twenty unaided reviewers correspond to about 10.26 independent judgments rather than twenty because human errors are already mildly correlated. When all twenty use the single shared model at the individually optimal reliance level, the effective count falls to approximately 1.31. The organization still sees twenty reviewers, twenty approvals, and perhaps twenty separate signatures, yet the error variance of the system resembles that of barely more than one independent assisted judgment.

This observation sharpens the meaning of cognitive monoculture. The concept does not require identical wording, identical beliefs, or complete deference to artificial intelligence. It is enough that a common source explains a large share of residual error. A review process can look socially diverse and remain statistically concentrated.

4.4 Model-source diversification helps only when errors are genuinely different

A natural mitigation is to diversify the artificial intelligence layer. Table 3 evaluates a twenty-person team while distributing reviewers evenly across increasing numbers of sources whose pairwise error correlation is 0.25. Source diversification reduces the aggregated artificial intelligence variance, raises the team-optimal reliance level, raises the common-model dominance threshold, and increases the effective independent judgment count. Four sources are sufficient in the illustrative baseline to make the team slightly better than the unaided team even when everyone adopts the individual-optimal reliance level. Nevertheless, effective independence remains only 2.88 judgments, far below the unaided value of 10.26. Accuracy recovery and independence recovery are therefore distinct governance objectives.

Table 3. Sensitivity to artificial intelligence source diversification for a twenty-person team with cross-source error correlation of 0.25

M	B_M	Team-optimal r	Dominance r	MSE at individual optimum / no-AI MSE	N_{eff} at individual optimum
1	0.3600	0.213	0.342	2.07x	1.31
2	0.2250	0.302	0.397	1.32x	2.06
4	0.1575	0.382	0.440	0.94x	2.88
10	0.1170	0.455	0.477	0.72x	3.78

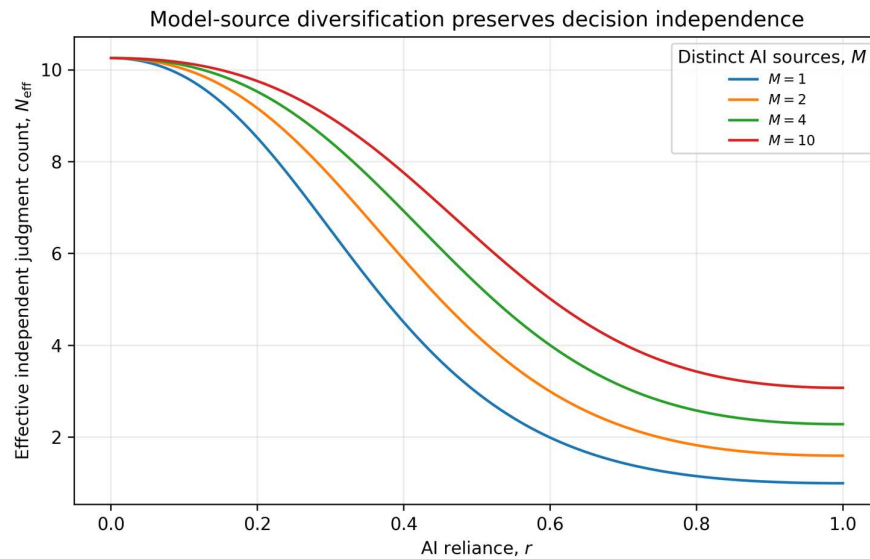


Figure 3. Effective independent judgment count for a twenty-person team under increasing artificial intelligence reliance and different numbers of distinct model sources. Cross-model error correlation is fixed at 0.25.

Figure 3 also shows diminishing returns. As the number of sources grows, the common artificial intelligence variance approaches a floor determined by the single-source variance multiplied by cross-source error correlation rather than approaching zero. In other words, diversifying vendors or model families cannot eliminate shared errors that arise from common training data, common benchmarks, similar retrieval sources, shared industry assumptions, or correlated design incentives. When cross-source error correlation is high, nominal multi-model architecture can still behave like a monoculture. Governance should therefore test cross-model disagreement empirically rather than infer independence from branding or procurement diversity.

4.5 Boundary conditions and extensions

The paradox weakens when unaided human errors are already highly correlated. As human error correlation approaches one, additional team members add little independent information, so aggregation provides little variance reduction and team-optimal reliance converges toward the individual optimum. In such settings, a good artificial intelligence system can legitimately substitute for redundant human consensus. The framework therefore does not romanticize human diversity; it values error diversity that is informative rather than merely demographic or procedural.

The paradox also weakens when the diversified artificial intelligence variance is smaller than the aggregated human variance. In that case, heavy model reliance may remain team-optimal. Conversely, systematic model bias, domain shift, or adversarial manipulation strengthens common-mode risk because those components do not average away. Equation (10) makes this explicit. An organization should therefore calibrate the model using empirical error decompositions rather than adopt the illustrative parameters directly.

Finally, reliance need not be homogeneous. If employees use different weights, prompts, retrieval systems, or escalation rules, the covariance structure becomes heterogeneous. The same principle remains: team robustness depends on the weighted covariance matrix of all human and artificial intelligence contributions. The scalar framework is thus a governance approximation for a more general portfolio-of-judgments problem.

5. Discussion

5.1 A new unit of analysis for human-AI governance

The central theoretical shift is from the human-in-the-loop to the independence-of-the-loop. Most oversight architectures ask whether a human reviews, approves, can intervene, or remains accountable. Those questions are necessary but incomplete. If the human's evidentiary basis, framing, counterarguments, and proposed conclusion are all supplied by the same model used by upstream colleagues, formal human control may coexist with epistemic dependence. The relevant unit of analysis is not the reviewer in isolation but the correlation network through which judgments are produced.

This view also clarifies why generative artificial intelligence differs from many earlier decision-support tools. A conventional calculator may standardize arithmetic while leaving problem framing and interpretation dispersed. A general-purpose language model can influence problem definition, evidence retrieval, option generation, causal explanation, writing, and critique within the same workflow. The same foundation model can therefore become an epistemic infrastructure spanning multiple organizational roles. Once that occurs, model-specific error is no longer local. It becomes a potential common-mode organizational exposure.

5.2 The reliance externality

The individual-team reliance paradox can be understood as a reliance externality. An employee chooses artificial intelligence assistance because it improves personal expected performance. The private optimization does not price the reduction in the marginal informational value of other reviewers using the same source. As adoption diffuses, each additional user's decision may become more accurate in isolation but less additive to collective review. This resembles diversification problems in finance and reliability engineering: the quality of each component is not sufficient to determine portfolio resilience when component errors are correlated.

The externality creates a managerial dilemma. Policies that encourage uniform use of the best-performing approved model can improve consistency, productivity, security, and procurement simplicity. The same standardization can reduce epistemic redundancy. Conversely, allowing every team to use different tools may increase diversity but create data leakage, audit, cost, and quality-control problems. The appropriate response is not uncontrolled tool proliferation. It is deliberate diversity at designated challenge points, with governance over which models, retrieval sources, prompts, and human pathways are permitted.

5.3 From reviewer count to independence architecture

The model suggests that reviewer count should be treated as a weak governance indicator unless paired with an independence measure. A four-eye principle, committee vote, second-line risk review, or senior sign-off adds little resilience if all participants inherit the same artificial intelligence premise. The effective independent judgment count offers a conceptual replacement: ask how many independent judgments the process is statistically equivalent to after shared tool use. Estimating this measure in practice requires repeated decisions with known outcomes or proxy disagreement data, but even qualitative assessment can reveal obvious concentration, such as a single model generating the first draft, the counterargument, and the final review checklist.

5.4 Governance design principles

Four design principles follow. First, preserve pre-model judgment for high-consequence decisions. At least one reviewer should record an initial assessment before seeing the shared model output when independent judgment has material value. Second, diversify artificial intelligence sources at challenge points rather than indiscriminately. A separate model family, independently curated retrieval corpus, or non-model analytical method can act as a red-team source. Third, separate generation from verification. A model used to propose an answer should not automatically be the only model used to critique that answer. Fourth, monitor disagreement and convergence. Sudden collapse in dispersion after a model rollout can be a risk signal even when average performance improves.

These principles complement rather than replace current governance frameworks. The National Institute of Standards and Technology emphasizes governance, measurement, and management [15], [16]; ISO/IEC 42001 requires a systematic management process [17]; and the European Union Artificial Intelligence Act requires human oversight for covered high-risk systems [18]. The present framework adds a measurable question that can sit inside those processes: how independent are the humans and artificial intelligence sources that constitute the oversight chain?

Table 4. Governance controls for cognitive monoculture risk

Control	Risk addressed	Measurable implementation
Pre-AI judgment capture	Anchoring and shared framing	Require at least one reviewer to record an initial view before model exposure; compare pre- and post-AI judgments.
Model-source challenge	Single-model common error	Route selected high-impact decisions to a substantively different model family or independently governed analytical source.
Generation-verification separation	Self-confirming AI review	Do not rely exclusively on the same model instance or model family to generate and verify a consequential output.
Disagreement telemetry	Invisible convergence	Track dispersion, reversal rates, cross-model disagreement, and changes in reviewer variance after AI deployment.
Independence stress test	False comfort from reviewer count	Estimate the effective independent judgment count or a proxy using historical errors, correlated overrides, or controlled parallel-review exercises.
Human capability maintenance	Reviewers unable to challenge	Use periodic no-AI exercises, calibration tasks, and domain-specific challenge drills for critical roles.
Escalation by common-mode exposure	High-impact shared blind spots	Increase independent challenge when reliance, model concentration, or cross-source correlation is high.

5.5 Implications for organizational design and policy

For managers, the results caution against equating standardization with robustness. Standardized artificial intelligence can be valuable for routine work where consistency is the objective. In decisions where error diversification is valuable, however, organizations should identify a minimum independence architecture in the same way that critical processes specify segregation of duties. One useful policy distinction is between production use and challenge use: production teams may share a preferred model for efficiency, while designated reviewers use independent evidence sources or alternative models before final approval.

For boards and risk committees, the framework suggests adding model concentration and epistemic concentration to artificial intelligence inventories. A conventional inventory records which business units use which systems. An

independence inventory asks which decisions are exposed to the same model family, retrieval corpus, prompt library, vendor, or synthetic training feedback. Concentration at any of these layers can create common-mode risk even when applications appear operationally separate.

For regulators and standard setters, the model offers a way to make human oversight more substantive without prescribing a universal number of reviewers. Requirements could ask organizations to demonstrate that oversight for high-consequence decisions is not fully dependent on the same artificial intelligence source being overseen. This could be evidenced through independent verification, alternative analytical methods, pre-model judgment capture, model diversity testing, or measured disagreement. The principle is proportionality: the greater the consequence and common dependence, the stronger the independence requirement.

5.6 Research propositions

Proposition 1: In multi-reviewer tasks, greater shared artificial intelligence reliance will increase pairwise correlation among final judgment errors, controlling for individual model-assisted accuracy.

Proposition 2: The reliance level that maximizes individual accuracy will exceed the reliance level that minimizes team-aggregate error when reviewers share an artificial intelligence source and unaided human errors are not perfectly correlated.

Proposition 3: The performance penalty from individually optimized shared-model reliance will increase with team size when additional humans contribute partially independent information.

Proposition 4: Model-source diversification will improve team robustness in proportion to the reduction in cross-model error correlation, not simply the number of models procured.

Proposition 5: Effective independent judgment count will predict resilience to model-specific shocks more strongly than nominal reviewer headcount.

Proposition 6: Workflows that separate generation and verification across independent human or artificial intelligence sources will exhibit lower common-mode failure rates than workflows that use the same model for both functions, holding average model accuracy constant.

6. Limitations and Future Research

The framework is intentionally stylized. First, it uses a scalar outcome and exchangeable error structures. Real decisions may be multidimensional, path-dependent, and strategically interactive. Future work should estimate full covariance matrices across roles, models, retrieval sources, and decision stages. Second, reliance is represented as a weight. In practice, human responses to artificial intelligence include selective adoption, anchoring, verification, counterfactual reasoning, and task delegation that may require nonlinear or state-dependent models.

Third, the illustrative parameters are not empirical estimates. Field studies should measure human error correlation before and after generative artificial intelligence deployment, then estimate whether observed reviewer headcount overstates effective independence. Suitable settings include forecasting tournaments, medical second opinions, credit committees, safety cases, procurement review, legal research, and academic peer evaluation. Controlled experiments could randomly assign reviewers to shared versus diversified models while holding model quality approximately constant.

Fourth, artificial intelligence sources are treated as having stable pairwise correlation. In practice, cross-model correlation may vary by domain and failure type. Models from different providers can still share public training corpora, reinforcement preferences, benchmark incentives, retrieval sources, or common factual misconceptions. Future research should therefore develop task-specific model correlation matrices and distinguish ordinary output disagreement from disagreement on rare, high-impact errors.

Fifth, cognitive diversity has normative as well as statistical dimensions. The effective independent judgment count measures error independence, not viewpoint fairness, cultural representation, legitimacy, or participation. A statistically diverse system can still be socially exclusionary, and a socially diverse committee can still share the same artificial intelligence dependency. These dimensions should be measured separately rather than collapsed into one index.

Finally, the framework does not imply that human disagreement is inherently good. Excessive dispersion can reflect noise, poor expertise, or inconsistent standards. The objective is calibrated independence: enough common infrastructure to capture the benefits of accurate artificial intelligence, but enough independent evidence and reasoning to prevent one model-specific error from propagating through the entire decision system.

7. Conclusion

Artificial intelligence can make individuals better while making groups more correlated. That possibility creates a governance problem that is easy to miss because conventional indicators move in the wrong direction: productivity rises, individual error falls, outputs become more consistent, and every formal reviewer remains in place. Yet the effective number of independent judgments can collapse. This paper formalizes that condition as cognitive monoculture and derives the individual-team reliance paradox that produces it.

The analytical result is simple but consequential. When many people rely on the same artificial intelligence source, common model error does not diversify away. The reliance level that is optimal for one person is therefore generally too high for an aggregating team. Larger teams can be especially exposed because their counterfactual unaided error is more strongly diversified. Model-source diversification can restore robustness, but only when sources fail differently; adding highly correlated models creates the appearance rather than the substance of diversity.

The practical implication is that human oversight should be evaluated as an independence architecture. Organizations should complement reviewer counts with measures of model concentration, disagreement, cross-source correlation, generation-verification separation, and effective independent judgment. The purpose is not to minimize artificial intelligence use. It is to preserve the error-correcting value of human and technological plurality while retaining the productivity gains that make artificial intelligence attractive in the first place. A trustworthy organization is not one in which many people approve the same machine-mediated answer. It is one in which consequential decisions remain contestable by genuinely independent sources of judgment.

Acknowledgments

The author acknowledges no external assistance that meets the journal criteria for acknowledgment.

Author Contributions

Kwan Hong TAN: Conceptualization; methodology; formal analysis; simulation design; visualization; writing: original draft; writing: review and editing.

Funding Details

This research received no external funding.

Disclosure Statement / Conflicts of Interest

The author declares no financial or non-financial conflict of interest relevant to this work.

Data and Reproducibility Statement

No empirical dataset was used. All equations, assumptions, parameter values, and numerical scenarios required to reproduce the analytical results are reported in the manuscript. The Monte Carlo verification used synthetic Gaussian draws only.

Ethics Statement

Ethics committee review was not applicable because the study did not involve human participants, animals, personal data, or confidential organizational data.

References

- [1] S. Noy and W. Zhang, "Experimental evidence on the productivity effects of generative artificial intelligence," *Science*, vol. 381, no. 6654, pp. 187-192, 2023, doi: 10.1126/science.adh2586.
- [2] M. H. Jarrahi, "Artificial intelligence and the future of work: Human-AI symbiosis in organizational decision making," *Business Horizons*, vol. 61, no. 4, pp. 577-586, 2018, doi: 10.1016/j.bushor.2018.03.007.
- [3] S. Raisch and S. Krakowski, "Artificial Intelligence and Management: The Automation-Augmentation Paradox," *Academy of Management Review*, vol. 46, no. 1, pp. 192-210, 2021, doi: 10.5465/amr.2018.0072.
- [4] Y. R. Shrestha, S. M. Ben-Menahem, and G. von Krogh, "Organizational Decision-Making Structures in the Age of Artificial Intelligence," *California Management Review*, vol. 61, no. 4, pp. 66-83, 2019, doi: 10.1177/0008125619862257.
- [5] L. Bainbridge, "Ironies of automation," *Automatica*, vol. 19, no. 6, pp. 775-779, 1983, doi: 10.1016/0005-1098(83)90046-8.
- [6] R. Parasuraman and V. Riley, "Humans and Automation: Use, Misuse, Disuse, Abuse," *Human Factors*, vol. 39, no. 2, pp. 230-253, 1997, doi: 10.1518/001872097778543886.
- [7] R. Parasuraman, T. B. Sheridan, and C. D. Wickens, "A model for types and levels of human interaction with automation," *IEEE Trans. Syst., Man, Cybern. A, Syst. Humans*, vol. 30, no. 3, pp. 286-297, 2000, doi: 10.1109/3468.844354.
- [8] J. D. Lee and K. A. See, "Trust in Automation: Designing for Appropriate Reliance," *Human Factors*, vol. 46, no. 1, pp. 50-80, 2004, doi: 10.1518/hfes.46.1.50_30392.
- [9] L. Hong and S. E. Page, "Groups of diverse problem solvers can outperform groups of high-ability problem solvers," *Proc. Natl. Acad. Sci. U.S.A.*, vol. 101, no. 46, pp. 16385-16389, 2004, doi: 10.1073/pnas.0403723101.
- [10] J. Lorenz, H. Rauhut, F. Schweitzer, and D. Helbing, "How social influence can undermine the wisdom of crowd effect," *Proc. Natl. Acad. Sci. U.S.A.*, vol. 108, no. 22, pp. 9020-9025, 2011, doi: 10.1073/pnas.1008636108.
- [11] A. R. Doshi and O. P. Hauser, "Generative AI enhances individual creativity but reduces the collective diversity of novel content," *Science Advances*, vol. 10, no. 28, art. eadn5290, 2024, doi: 10.1126/sciadv.adn5290.
- [12] B. R. Anderson, J. H. Shah, and M. Kreminski, "Homogenization Effects of Large Language Models on Human Creative Ideation," in *Proc. 16th ACM Conf. Creativity & Cognition*, pp. 413-425, 2024, doi: 10.1145/3635636.3656204.
- [13] S. Zhang, J. Xu, and A. J. Alvero, "Generative AI Meets Open-Ended Survey Responses: Research Participant Use of AI and Homogenization," *Sociological Methods & Research*, vol. 54, no. 3, pp. 1197-1242, 2025, doi: 10.1177/00491241251327130.
- [14] T. Rinta-Kahila, E. Penttinen, A. Salovaara, W. Soliman, and J. Ruissalo, "The Vicious Circles of Skill Erosion: A Case Study of Cognitive Automation," *J. Assoc. Inf. Syst.*, vol. 24, no. 5, pp. 1378-1412, 2023, doi: 10.17705/1jais.00829.
- [15] E. Tabassi, *Artificial Intelligence Risk Management Framework (AI RMF 1.0)*, NIST AI 100-1. Gaithersburg, MD, USA: National Institute of Standards and Technology, 2023, doi: 10.6028/NIST.AI.100-1.

- [16] C. Autio et al., *Artificial Intelligence Risk Management Framework: Generative Artificial Intelligence Profile*, NIST AI 600-1. Gaithersburg, MD, USA: National Institute of Standards and Technology, 2024, doi: 10.6028/NIST.AI.600-1.
- [17] ISO and IEC, *ISO/IEC 42001:2023 Information technology - Artificial intelligence - Management system*, Geneva, Switzerland: ISO, 2023.
- [18] European Parliament and Council of the European Union, “Regulation (EU) 2024/1689 (Artificial Intelligence Act),” *Official Journal of the European Union*, OJ L, 2024/1689, 12 July 2024.
- [19] K. H. TAN, “The Governance Half-Life of Agentic Artificial Intelligence Controls: A Dynamic Assurance Framework for Control Decay, Residual Risk, and Revalidation,” *Int. J. Web Multidiscip. Stud.*, vol. 3, no. 7, pp. 217-229, 2026, doi: 10.71366/ijwos03072654084.