

From Micro Signals to Macro Intelligence: Building Adaptive Organizations with AI

Executive summary

This whitepaper treats adaptation as a learning problem, not a tooling problem. The core insight from organizational learning research is that firms do not improve simply by collecting more data; they improve when they reliably convert experience into changed routines, changed assumptions, and changed policies. Classic work by Barbara Levitt[1] and James G. March[2] frames organizational learning as routine-based, history-dependent, and target-oriented, while George P. Huber[3] decomposes learning into knowledge acquisition, distribution, interpretation, and memory. David A. Garvin[4] adds the managerial requirement that learning be measurable, experimental, and shared across the firm. Amy Edmondson[5] shows that speaking up and reporting errors depend on psychological safety, and March's exploration/exploitation work explains why organizations struggle to balance efficiency with innovation. [6]

Organizations therefore fail to learn for recurring, structurally predictable reasons: leaders protect prior beliefs rather than test them; people underweight others' information; hierarchies suppress dissent and prematurely stop search; incentives reward local outputs instead of system outcomes; narrow metrics invite gaming; and governance often lacks explicit risk tolerance, end-user feedback channels, and disciplined review cadences. Chris Argyris[7] called this defensive reasoning; John D. Sterman[8] showed that managers routinely misperceive delayed and nonlinear feedback; HBS and MIT research shows that hierarchy, weak error analysis, and poorly aligned incentives can turn learning systems into performance traps rather than capability builders. [9]

AI matters because it can shorten the sensing-to-decision loop. In an industry-unspecified enterprise, aggregated micro signals from behavioral telemetry, operational metrics, and qualitative inputs can be turned into policy proposals for staffing, training, workflow thresholds, escalation rules, nudges, and standard operating procedures. The strongest evidence is not that "more monitoring" is inherently valuable, but that targeted monitoring becomes valuable when managers use it to identify bottlenecks and respond with specific interventions such as refresher training, staffing changes, or policy redesign. MIT research similarly argues that AI value emerges at the workflow level, that worker voice shapes adoption, and that advanced AI maturity requires accessible data, transparent outcomes, acceptable-use policies, and explicit human oversight. [10]

The practical implication is a product architecture: capture signals, normalize them into an organizational memory, infer likely causes and interventions, route proposals through human and governance review, deploy changes in controlled experiments, and feed measured outcomes back into the policy registry. This is not only a data platform. It is a closed-loop management system, supported by experiment infrastructure, KPIs, review

forums, and change management. Research on experimentation shows that controlled experiments create a virtuous cycle of hypothesis generation and improvement; institutional memory across experiments compounds learning; and sequential and bandit methods extend the toolset when decisions must be made earlier or traffic must be allocated adaptively. [11]

Organizational learning and feedback loops

Organizational learning theory converges on a useful engineering specification for adaptive enterprises. Learning begins with **acquisition** of signals, but it only becomes organizational when those signals are **distributed, interpreted**, and stored as **memory** that shapes future action. In practical terms, this means an adaptive organization must instrument experience, circulate evidence across functions, convert evidence into interpretable hypotheses, and preserve decisions and outcomes in reusable artifacts such as playbooks, experiment reports, and policy registries. HBS research emphasizes that learning organizations pair systematic problem solving with experimentation and knowledge sharing, rather than treating knowledge as detached from action. [12]

Feedback loops determine whether the organization compounds or destroys learning. March’s exploration/exploitation framework shows the tension between trying new options and optimizing known ones. Argyris’s double-loop logic shows that real learning requires not only correcting errors inside current rules, but occasionally questioning the rules themselves. Edmondson’s work shows that people report weak signals, errors, and dissenting interpretations only when interpersonal risk is low. System dynamics research then explains why many firms still miss obvious lessons: delays, feedback complexity, and short-term performance pressure make “working harder” look better than investing in capability, creating a capability trap in which improvement work is perpetually deferred. [13]

The table below translates the theory into design requirements for an AI-enabled adaptive organization. It is a synthesis grounded in HBS, MIT, and seminal organizational learning research. [14]

Theoretical lens	Core proposition	Product design implication
Organizational learning as routines and targets	Organizations encode experience into routines; goals shape what gets noticed	Maintain versioned policy registries, target hierarchies, and explicit change logs
Acquisition, distribution, interpretation, memory	Learning fails if signals are sensed but not shared or remembered	Build event pipelines, shared dashboards, review forums, and durable experiment memory
Exploration vs. exploitation	Firms over-optimize current routines and underinvest in	Run a portfolio of fixed-horizon experiments, bandits, and

Theoretical lens	Core proposition	Product design implication
Double-loop learning	novel search Firms must question assumptions, not only fix deviations	periodic exploration budgets Require retrospectives to document “what assumption changed,” not just “what metric moved”
Psychological safety	Weak signals emerge only when people can speak without punishment	Add low-friction worker voice, appeals, and exception reporting into the measurement system
Capability trap	Short-term pressure crowds out improvement work	Reserve time, staffing, and funding specifically for learning and experimentation
Interactive management systems	Metrics create value when they trigger discussion and adaptation, not passive reporting	Use operating reviews that challenge assumptions and revise action plans on a fixed cadence

Why organizations fail to learn

The most important failure modes are not mysterious. They are cognitive, structural, cultural, incentive-driven, measurement-related, and governance-related, and they reinforce one another. HBS and MIT work shows that people often defend prior beliefs, discount outside evidence, and rationalize failures instead of analyzing them; that hierarchy can suppress upward information flow and prematurely stop search; that workers respond to the incentives embedded in measurement systems; and that AI systems require explicit context-setting, risk tolerances, user feedback mechanisms, and ongoing monitoring if they are to remain trustworthy after deployment. [15]

The synthesis below shows the common causes and the corresponding design responses.

Cause	What goes wrong	Evidence-based interpretation	Design response
Cognitive	Leaders defend prior views, overtrust their own information, misread delayed feedback, and escalate failing paths	Defensive reasoning, underweighting of others’ information, and misperceptions of feedback distort learning	Require hypothesis logs, pre-mortems, after-action reviews, and evidence-linked decision memos
Structural	Silos and hierarchy slow or suppress signal flow and stop search too early	Upward information gets filtered; hierarchy can trap firms on local peaks	Create cross-functional review cells, delegated authority, and direct signal channels to

Cause	What goes wrong	Evidence-based interpretation	Design response
Cultural	Errors are hidden, dissent is muted, and learning becomes reputationally risky	Low psychological safety reduces error reporting and experimentation	decision forums Protect reporting channels, distinguish blameworthy from intelligent failure, and normalize retrospective discussion
Incentive	Employees optimize the measured local score, not the system outcome	High-powered explicit incentives can produce gaming, deception, or narrow effort allocation	Use balanced KPI sets, guardrails, and rewards for learning quality, not just output volume
Data and measurement	The firm tracks what is easy, not what matters; metrics are misinterpreted	Incomplete objective measures, poor OEC design, and metric pitfalls yield wrong conclusions	Use an overall evaluation criterion plus diagnostic metrics, refresh measures periodically, and validate causal assumptions
Governance	AI use lacks risk tolerance, context definitions, appeals, and post-deployment review	Systems drift beyond their intended setting; feedback from users is absent or ignored	Define use context, risk appetite, approval rights, monitoring rules, audit trails, and escalation paths

AI can either deepen or reduce these failures. It deepens them when it scales poor incentives, centralizes opaque decisions, or turns surveillance into resistance. It reduces them when it broadens sensing, captures worker voice, speeds falsification of bad assumptions, and formalizes safe experimentation inside a governance framework. MIT research on AI work redesign is explicit that frontline workers often have little say in design decisions and that this fuels resistance, avoidance, and failed pilots; at the same time, MIT research on human-AI collaboration shows that value comes from redesigning the workflow and learning from randomized comparisons of human-only, AI-only, and human-AI modes. [16]

From micro signals to macro intelligence

Micro signals are the smallest usable observations of organizational behavior and performance. In an industry-unspecified enterprise, they usually appear in three families.

Behavioral telemetry includes workflow clicks, approval paths, queue traversals, overrides, handoffs, dwell time, search behavior, and model-accept/reject events. **Operational metrics** include cycle time, throughput, rework, defects, waste, service levels, staffing coverage, cost, and incident rates. **Qualitative inputs** include survey text, worker comments, case notes, escalation narratives, postmortems, and manager observations. HBS studies of performance monitoring show that fine-grained operational telemetry becomes valuable when it reveals bottlenecks and triggers targeted managerial action; HBS work on values-based assessment and subjective evaluation shows why qualitative inputs are necessary complements to objective measures; MIT research adds that frontline voice is essential in AI design and deployment. [17]

Those signals become macro intelligence only after four transformations. First, they are **contextualized**: events are joined to teams, workflows, customer segments, policy versions, and risk classes. Second, they are **aggregated** into patterns: recurring bottlenecks, exception clusters, quality regressions, drift segments, or voice themes. Third, they are **attributed**: the system estimates whether the pattern reflects noise, workload mix, policy failure, training gaps, or model degradation. Fourth, they are **actionized**: the platform proposes a manageable change such as a new threshold, a revised escalation path, refresher training, a UI nudge, a staffing rule, or an experiment. The HBS monitoring evidence is especially relevant here: better monitoring created value when managers used it to identify bottlenecks and respond with staffing adjustments or retraining, not when monitoring existed in isolation. [18]

The table below is a cross-industry product view of how signal classes map to policy moves.

Signal class	Example micro signal	Typical inference	Example policy change
Behavioral telemetry	Repeated overrides of AI recommendations in one workflow step	Recommendation logic mismatched to context; poor trust calibration	Raise confidence threshold, add explanation, or reroute to human review
Operational metrics	One queue shows rising cycle time and rework while volume mix is stable	Local bottleneck or capability decay	Add refresher training, staffing change, or revised handoff rule
Qualitative inputs	Recurrent comments that a rule forces unnecessary escalation	Formal policy conflicts with frontline reality	Rewrite SOP, create exception pathway, or change escalation criteria
Combined signals	Higher adoption but worse downstream quality	Local gain is harming the system objective	Reinstate guardrail, redesign OEC, or reverse rollout

Signal class	Example micro signal	Typical inference	Example policy change
Combined signals over time	Same intervention works in one segment and fails in another	Heterogeneous treatment effect	Segment-specific policy or model routing

The entity model below shows how micro evidence should be stored so that learning compounds over time. This structure operationalizes acquisition, interpretation, and memory, while ensuring that every proposed policy change is linked to signals, experiments, and outcomes. [19]

erDiagram

```

SIGNAL }o--|| ENTITY : "about"
QUAL_NOTE }o--|| ENTITY : "describes"
SIGNAL ||--o{ FEATURE : "transformed_into"
QUAL_NOTE ||--o{ FEATURE : "summarized_into"
FEATURE ||--o{ MODEL_SCORE : "feeds"
MODEL_SCORE ||--o{ RECOMMENDATION : "supports"
RECOMMENDATION }o--|| POLICY_VERSION : "proposes_change_to"
POLICY_VERSION ||--o{ EXPERIMENT : "validated_by"
EXPERIMENT ||--o{ OUTCOME : "produces"
OUTCOME }o--|| KPI : "updates"
KPI ||--o{ REVIEW : "triggers"
REVIEW ||--o{ POLICY_VERSION : "approves_or_revises"

```

In this design, **macro intelligence** is not a dashboard. It is a risk-adjusted, evidence-linked recommendation that answers five questions at once: *what is changing, where, for whom, why now, and what policy move is likely to improve the system without violating guardrails*. That last clause is essential because MIT research shows that human-AI combinations create the most value when the workflow itself is redesigned and monitored continuously, not when AI is simply inserted into a single task. [20]

Reference architecture

The architecture below operationalizes Huber’s learning cycle, Kaplan’s closed-loop management logic, MIT CISR’s AI maturity guidance, and the NIST risk-management lifecycle. Its purpose is to transform distributed observations into governed policy change. [21]

flowchart LR

```

subgraph Sources
    A1[Behavioral telemetry]
    A2[Operational systems]
    A3[Qualitative inputs]
    A4[External context and benchmarks]
end

```

```

subgraph Ingestion

```

- B1[Event bus and CDC]
- B2[Batch ETL and API pulls]
- B3[Schema registry and data contracts]
- B4[PII minimization and access controls]

end

subgraph Storage

- C1[Raw object store]
- C2[Lakehouse or warehouse]
- C3[Feature store]
- C4[Document and vector store]
- C5[Experiment registry]
- C6[Policy registry and audit log]

end

subgraph Analytics

- D1[Data quality and observability]
- D2[Anomaly and bottleneck detection]
- D3[Causal inference and uplift models]
- D4[Forecasting and simulation]
- D5[LLM evidence synthesis]
- D6[Policy scoring and prioritization]

end

subgraph HumanLoop

- E1[Ops analyst]
- E2[Domain owner]
- E3[Risk and legal review]
- E4[Frontline review cell]
- E5[Executive operating review]

end

subgraph Action

- F1[New SOP or playbook]
- F2[Threshold or routing rule]
- F3[Training and coaching]
- F4[Staffing or resource allocation]
- F5[UI nudge or product change]
- F6[Experiment or staged rollout]

end

subgraph Governance

- G1[Risk taxonomy]
- G2[Use-context documentation]
- G3[Fairness, privacy, security checks]
- G4[Approval workflow]
- G5[Rollback and incident response]

end

```
subgraph Monitoring
  H1[KPI dashboard and OEC]
  H2[Drift and degradation monitor]
  H3[Appeals and user feedback]
  H4[Postmortem and retrospective]
  H5[Model and policy refresh]
end
```

```
A1 --> B1
A2 --> B1
A3 --> B2
A4 --> B2
B1 --> B3 --> B4 --> C1
B2 --> B3
C1 --> C2 --> C3
C2 --> C4
C2 --> C5
C2 --> C6
C3 --> D1 --> D2
C3 --> D3
C2 --> D4
C4 --> D5
D2 --> D6
D3 --> D6
D4 --> D6
D5 --> D6
D6 --> E1 --> E2 --> E3 --> E4 --> E5
E5 --> F6
F6 --> F1
F6 --> F2
F6 --> F3
F6 --> F4
F6 --> F5
G1 --> E3
G2 --> E3
G3 --> E3
G4 --> E5
G5 --> F6
F1 --> H1
F2 --> H1
F3 --> H1
F4 --> H1
F5 --> H1
F6 --> H1
H1 --> H2 --> H5 --> C6
H3 --> H4 --> C6
H1 --> C5
H4 --> D5
```

The guiding principle is simple: every recommendation must be reproducible back to evidence, every deployment must be measurable, and every measured outcome must update organizational memory. That principle matters because MIT CISR's maturity model ties superior AI performance to cumulative capability-building, transparent dashboards, and enterprise-wide ways of working rather than isolated pilots. NIST makes the same point in governance language: organizations should govern, map, measure, and manage risk across the AI lifecycle, document context and risk tolerance, and integrate end-user feedback into evaluation. [22]

The annotated diagram below shows the control logic of the learning loop.

flowchart TD

```
A[Acquire<br/>Capture event, metric, and narrative data]
B[Interpret<br/>Normalize, join, segment, estimate causes]
C[Decide<br/>Rank policy options by value, risk, and confidence]
D[Act<br/>Experiment, canary, rollout, train, or revise workflow]
E[Learn<br/>Measure outcomes, document lessons, refresh models]
```

```
A --> B --> C --> D --> E --> A
```

```
N1[[Data contracts, lineage, access control]]
N2[[Human review for high-risk changes]]
N3[[Guardrails for quality, fairness, privacy, and safety]]
N4[[Institutional memory: experiment and policy registry]]
```

```
N1 -.-> A
N2 -.-> C
N3 -.-> D
N4 -.-> E
```

Data pipelines and storage

Pipelines should separate **raw capture**, **analytical storage**, and **decision-ready features**. The raw layer preserves original events and text. The analytical layer standardizes entities, joins operational context, and computes lagged and segment-level metrics. The decision layer stores reusable features, embeddings for qualitative retrieval, experiment metadata, and versioned policy objects. This separation is important because organizational memory must preserve both the evidence and the interpretation history, while experiment registries allow patterns to generalize across tests rather than disappearing into local notebooks and presentations. [23]

Models and decisioning

A practical stack usually needs four model classes. The first is **observational**: anomaly detection, drift monitors, and bottleneck detectors surface unusual patterns. The second is **causal**: uplift models, quasi-experimental designs, and A/B analyses estimate whether a proposed change should work and for whom. The third is **predictive**: demand, workload,

and failure forecasting support resource planning. The fourth is **interpretive**: LLMs summarize qualitative evidence, cluster themes, and draft decision memos, but they should not be the sole arbiters of policy. MIT research on human-AI collaboration suggests the highest returns come when AI handles repetitive, data-rich subtasks and humans retain contextual and judgment-heavy work; NIST adds that model outputs should be explained, validated, documented, and interpreted within context. [24]

Human-in-the-loop and governance

Human review is not a fallback. It is part of the system design. MIT's AI work redesign research argues that frontline workers need voice in design and deployment, while NIST specifies that end-user feedback and appeal mechanisms should be integrated into evaluation metrics and that governance spans management, boards, domain experts, auditors, and other AI actors. In practice, low-risk interventions can move quickly through delegated review, while high-risk interventions should require documented use context, defined risk tolerance, fairness/privacy/security checks, and explicit signoff. [25]

Deployment and learning cadence

Deployment should follow the reversibility of the policy. For reversible changes such as prompts, thresholds, and nudges, use **shadow mode**, **canary release**, and **ringed rollout** with automatic rollback. For materially consequential changes, keep the AI in recommend-only mode until confirmatory evidence and governance review are complete. The operating cadence should mirror Kaplan's interactive management logic: frequent reviews, explicit discussion of assumptions, and alignment of reporting, target setting, initiative management, and strategy updates. MIT CISR's maturity work similarly emphasizes acceptable-use policies, metrics, and transparent dashboards before industrialization. [26]

Measurement and experimentation

A learning system fails if it does not know what "better" means. The most robust design pattern is to define an **overall evaluation criterion** that represents the system objective, then pair it with diagnostic and guardrail metrics. Experimentation research argues that teams should align early on what they are optimizing; Microsoft's experimentation literature warns that poor metric interpretation can lead to costly mistakes; and Kaplan's Balanced Scorecard work shows that measurement becomes powerful when it is used interactively to question assumptions and redirect action. Institutional memory across experiments then allows the organization to detect repeatable patterns rather than relearn the same lesson. [27]

The KPI mapping below is an illustrative, cross-industry starting point because the industry is unspecified. It should be customized to regulation, workflow criticality, labor model, and customer economics. The thresholds are intentionally practical rather than universal. [28]

KPI	Data sources	Measurement frequency	Owner	Example threshold
Learning cycle time	Event logs, issue tracker, policy registry	Weekly	COO or Product Ops	Median time from signal to policy decision > 30 days triggers escalation
Experiment velocity	Experiment registry, decision memos	Monthly	Product or Ops excellence lead	Fewer than 4 completed experiments per quarter in a priority domain
Policy adoption rate	Workflow logs, config audits	Weekly	Domain process owner	Adoption < 80% four weeks after rollout
Exception or escalation rate	Case system, service desk, workflow engine	Daily/weekly	Operations leader	> 10% above baseline for two consecutive periods
Cycle time	ERP, CRM, ticketing, workflow telemetry	Daily/weekly	Process owner	p95 cycle time worsens by > 5%
Quality or rework rate	QA audits, returns, defect logs	Daily/weekly	Quality leader	> 2 standard deviations above baseline
Customer outcome metric	CSAT, retention, service completion, claims	Weekly/monthly	CX or business GM	Any statistically credible decline on primary customer guardrail
Employee voice coverage	Pulse surveys, feedback forms, meeting participation	Monthly/quarterly	HR or People Analytics	Response rate < 60% or speak-up index declines materially
Model and data drift	Feature store, model monitor, shadow evaluations	Daily	Data science or MLOps	PSI > 0.2 or validation metric drops > 5%
Risk incident rate	Appeals, audit findings, privacy/security	Daily/monthly	Risk and Compliance	Any severe incident or sustained

KPI	Data sources	Measurement frequency	Owner	Example threshold
	logs			increase in moderate incidents
Realized economic value	Finance, productivity model, revenue/cost data	Monthly/quarterly	CFO plus business sponsor	Benefits below plan for 2 review cycles

Three experiment families belong in the operating model. **Fixed-horizon A/B tests** are best for interpretable, confirmatory policy changes. **Multi-armed bandits** are best for low-risk repeated choices where traffic can be allocated adaptively. **Sequential tests** are best when the organization must monitor continuously and still preserve inferential validity. HBS research ties A/B adoption to learning and improved startup performance; experimentation research emphasizes statistical power, sample size, and clear evaluation criteria; Thompson sampling formalizes adaptive arm allocation; and always-valid inference addresses the continuous-monitoring problem. Ron Kohavi[29], William R. Thompson[30], and later sequential work by Abraham Wald[31] and Johari and coauthors are the foundational references here. [32]

For planning, the standard approximations are still useful:

$$n_{\text{per arm, binary}} \approx \frac{[z_{1-\alpha/2}\sqrt{2\bar{p}(1-\bar{p})} + z_{1-\beta}\sqrt{p_1(1-p_1) + p_2(1-p_2)}]^2}{(p_2 - p_1)^2}$$

$$n_{\text{per arm, continuous}} \approx 2 \left(\frac{(z_{1-\alpha/2} + z_{1-\beta})\sigma}{\delta} \right)^2$$

The templates below use 80% power and 5% two-sided significance for fixed-horizon planning. The numerical sample sizes are illustrative calculations for this report, not universal prescriptions. Their role is to show the order of magnitude and the analysis logic. [33]

Design	Use when	Sample hypothesis	Unit and randomization	Illustrative sample size	Primary metric	Guardrails	Analysis plan
Fixed-horizon A/B	Clear policy change; interpretability matters	“A revised exception-handling SOP	Case or workflow instance; equal randomization	1,570 per arm assuming SD = 5 hours and	Mean cycle time	Quality, rework, customer score	Difference in means or ANCOVA/CUPED with pre-period baseline;

Design	Use when	Sample hypothesis	Unit and randomization	Illustrative sample size	Primary metric	Guardrails	Analysis plan
		reduces mean cycle time from 5.0 hours to 4.5 hours”		MDE = 0.5 hours			segment effects only if pre-specified
Multi-armed bandit	Repeated, low-risk choices among several nudges or prompts	“One of three coaching messages raises policy adoption from 20% to at least 23%”	User or team-week; burn-in then adaptive allocation	Burn-in 500 per arm, then adaptive allocation up to 9,000 total exposures	Adoption rate	Error rate, opt-out rate, fairness by segment	Thompson sampling with exploration floor; monitor posterior regret; confirm the winner with a short fixed-horizon A/B before permanent rollout
Sequential test	Need early stopping with continuous monitoring	“AI triage raises correct routing from 55% to 58%”	Case or ticket; equal allocation with interim looks	Maximum 4,286 per arm; stop earlier for successes or futility	Correct-routing rate	Escalations, downstream rework, severe incidents	Always-valid p-values or confidence sequences; interim looks every 1,000 cases; decision at success, harm, futility, or cap

Two operating rules matter more than the statistical motif. First, every experiment should have a **predeclared OEC**, explicit guardrails, a defined randomization unit, a stopping rule, and a named decision owner. Second, all experiment metadata and narrative interpretation should be written to a registry so that the enterprise can reuse what it learned. Without that memory, the organization accumulates data but not capability. [34]

Implementation and change management

Implementation should proceed as staged organizational redesign, not as a one-time AI deployment. MIT CISR’s maturity model suggests a progression from experimentation and preparation, to pilots and capability building, to enterprise ways of working. HBS change-management research adds that success depends not only on culture and motivation, but also on review cadence, team quality, commitment, and the extra work imposed on staff. Kaplan’s work further argues that alignment and performance review should be built into an ongoing governance cycle. [35]

```

timeline
  title Adaptive organization rollout
  Foundation : Define business outcomes and risk appetite
              : Select one priority workflow
              : Create data contracts and policy taxonomy
  Instrumentation : Capture telemetry, operational metrics, and worker
voice
                  : Establish KPI baselines and dashboards
                  : Create experiment and policy registries
  Pilot : Run recommend-only models
        : Launch first A/B tests
        : Hold weekly operating reviews and frontline retrospectives
  Scale : Add policy scoring and automated rollout paths
        : Expand to training, staffing, and routing decisions
        : Implement drift, fairness, and incident monitoring
  Institutionalize : Rebalance incentives toward system outcomes and
learning quality
                  : Quarterly strategy and threshold refresh
                  : Reuse experiment memory across domains
```

A practical implementation sequence for an industry-unspecified enterprise is shown below.

Phase	Objective	Key outputs	Exit criteria
Foundation	Define the learning problem	North-star outcome, use contexts, risk appetite, decision rights, target workflow	Clear sponsor, governance owner, and policy scope
Instrumentation	Make the workflow observable	Event schema, telemetry coverage, baseline metrics, feedback channels, data	Data completeness and trusted baseline dashboard

Phase	Objective	Key outputs	Exit criteria
Pilot	Prove one closed loop	quality checks Cause hypotheses, review forum, experiment templates, recommend-only models, first tests	One policy change validated and documented
Scale	Generalize beyond one workflow	Shared feature store, policy registry, experiment memory, rollout controls, MLOps	Reuse across at least three policy classes
Institutionalize	Change the operating model	Incentive updates, quarterly threshold review, learning SLAs, portfolio governance	Learning cadence embedded in normal management reviews

Change management should explicitly handle both the hard and soft sides. On the hard side, the DICE logic is highly relevant: keep milestone reviews frequent, staff high-integrity project teams, visibly signal senior commitment, and cap the additional work imposed on employees. On the soft side, MIT's work redesign research implies that worker resistance is not noise to be managed away; it is often valid signal about workflow design, job quality, and local knowledge. That means frontline co-design, clear communication about why the change exists, and visible response to feedback are not optional. [\[36\]](#)

Failure modes and mitigations

The table below summarizes the most likely implementation failures and the most direct mitigations.

Failure mode	Why it happens	Mitigation
Telemetry without action	Dashboards are built, but no operating forum converts signals into policy	Create a weekly review cell with named owners, SLAs, and decision templates
Local optimization	Teams optimize adoption or speed while harming quality or customer outcomes	Use an OEC plus guardrails and refresh KPI sets periodically
Capability trap	Performance pressure crowds out experiments and improvement work	Reserve budget, staff time, and calendar capacity for learning work
Frontline resistance	AI is experienced as surveillance or imposed redesign	Co-design with workers, show how feedback changed policy, keep appeals channels open
Hierarchical veto	Senior judgment overrides evidence or suppresses upward signal flow	Require evidence-linked memos, delegated authority, and experiment default rules

Failure mode	Why it happens	Mitigation
Qualitative summarization error	LLM outputs compress nuance or hallucinate causes	Use retrieval-backed summaries, source linking, and human signoff for consequential decisions
Governance bottleneck	Every change needs heavyweight review, so nothing ships	Define risk tiers and preapproved control paths for low-risk interventions
Model drift and context shift	A model or policy works in one setting and later degrades	Monitor drift, segment outcomes, run shadow evaluations, and enforce rollback rules
Incentive gaming	Employees respond to narrow targets rather than system goals	Blend objective and qualitative measures; reward learning quality and downstream outcomes
Experiment amnesia	The firm reruns old tests because memory is fragmented	Maintain a searchable experiment registry and policy history

Open questions and limitations

Because the industry is unspecified, the KPI thresholds, sample-size examples, and policy classes in this paper are intentionally generic. Highly regulated domains such as healthcare, finance, critical infrastructure, or labor-sensitive environments will require sector-specific controls, documentation, and legal review beyond the cross-industry baseline provided here. The human-AI boundary should also be tightened as the stakes of error rise; in such settings, recommend-only and confirmatory experimentation should dominate full automation until domain evidence is stronger. [37]

References

- Levitt, Barbara, and James G. March. “Organizational Learning.” *Annual Review of Sociology* 14 (1988). [38]
- Huber, George P. “Organizational Learning: The Contributing Processes and the Literatures.” *Organization Science* 2, no. 1 (1991). [39]
- Garvin, David A. “Building a Learning Organization.” *Harvard Business Review* (1993). [40]
- Garvin, David A. “Learning in Action.” *Working Knowledge*, Harvard Business School[41] (2000). [42]
- Garvin, David A. “Managing to Learn: How Companies Can Turn Knowledge into Action.” *Working Knowledge*, HBS (2000). [43]
- Argyris, Chris. “Teaching Smart People How to Learn.” *Harvard Business Review* (1991). [44]
- Edmondson, Amy C. “Psychological Safety and Learning Behavior in Work Teams.” *Administrative Science Quarterly* 44, no. 2 (1999). [45]

- Edmondson, Amy C., and Mark D. Cannon. “The Hard Work of Failure Analysis.” *Working Knowledge*, HBS (2005). [46]
- Adhvaryu, Achyuta, Anant Nyshadham, and Jorge Tamayo. “An Anatomy of Performance Monitoring.” HBS Working Paper 22-066 (2024). [47]
- Ghosh, Sayan, Rembrand Koning, and coauthors. “The Effects of Hierarchy on Learning and Performance in Experimentation.” HBS working paper (2020). [48]
- Koning, Rembrand, and coauthors. “Digital Experimentation and Startup Performance” and related HBS research on A/B testing and organizational learning (2020–2021). [49]
- Kaplan, Robert S. “Conceptual Foundations of the Balanced Scorecard.” HBS working paper (2010). [50]
- Kaplan, Robert S. “Managing Alignment as a Process.” *Working Knowledge*, HBS (2006). [51]
- Deller, Carolyn, and Mária Sandino. “In Search of Organizational Alignment Using a 360-Degree Assessment System.” HBS conference paper (2019). [52]
- Weill, Peter, Stephanie L. Woerner, and Ina M. Sebastian. “Building Enterprise AI Maturity.” *MIT CISR* (2024). [53]
- Wixom, Barbara H., Cynthia M. Beath, and coauthors. “AI Is Everybody’s Business.” *MIT CISR* (2024). [54]
- “What’s Your Company’s AI Maturity Level?” and “Leading the AI-Driven Organization.” *MIT Sloan School of Management* (2024–2025). MIT Sloan School of Management [55] [56]
- “The AI Work Redesign Playbook.” MIT Sloan Institute for Work and Employment Research (2025). [57]
- Malone, Thomas W., Abdullah Almaatouq, Michelle Vaccaro, and coauthors. “When Humans and AI Work Best Together — and When Each Is Better Alone.” MIT Sloan (2024). [58]
- Burnham, Kristin. “How AI Is Reshaping Workflows and Redefining Jobs.” MIT Sloan (2026). [59]
- Sterman, John D. *Business Dynamics* and related publications on system dynamics and feedback complexity. MIT faculty materials. [60]
- Repenning, Nelson P., and John D. Sterman. “Nobody Ever Gets Credit for Fixing Problems That Never Happened.” *California Management Review* 43, no. 4 (2001). [61]
- Conlon, John J., and coauthors. “Failing to Learn from Others.” MIT Economics working paper (2022). [62]
- “Artificial Intelligence Risk Management Framework.” National Institute of Standards and Technology [63] (NIST AI RMF 1.0, 2023). [64]
- Kohavi, Ron, Roger Longbotham, Dan Sommerfield, and Randal M. Henne. “Controlled Experiments on the Web: Survey and Practical Guide.” *Data Mining and Knowledge Discovery* 18, no. 1 (2009). [65]

- Kohavi, Ron, Randal M. Henne, and Dan Sommerfield. “Practical Guide to Controlled Experiments on the Web: Listen to Your Customers, Not to the HiPPO.” KDD (2007). [66]
- Dmitriev, Pavel, Somit Gupta, Dong Woo Kim, and Garnet Vaz. “A Dirty Dozen: Twelve Common Metric Interpretation Pitfalls in Online Controlled Experiments.” KDD (2017). [67]
- Johari, Ramesh, Leo Pekelis, and David J. Walsh. “Always Valid Inference: Continuous Monitoring of A/B Tests.” *Operations Research* (2021/2022). [68]
- Thompson, William R. “On the Likelihood That One Unknown Probability Exceeds Another in View of the Evidence of Two Samples.” *Biometrika* 25, no. 3–4 (1933). [69]
- Wald, Abraham. “Sequential Tests of Statistical Hypotheses.” *Annals of Mathematical Statistics* 16, no. 2 (1945). [70]
- Manso, Gustavo. “Motivating Innovation.” MIT working paper / *Journal of Finance* lineage (2011). [71]
- Sirkin, Harold L., Perry Keenan, and Alan Jackson. “The Hard Side of Change Management.” *Harvard Business Review* (2005). [72]

[1] [3] [10] [17] [18] [47] [55] https://www.hbs.edu/ris/Publication%20Files/22-066_3c5abc63-bac0-4c34-b53b-3bda4d14e75e.pdf

https://www.hbs.edu/ris/Publication%20Files/22-066_3c5abc63-bac0-4c34-b53b-3bda4d14e75e.pdf

[2] [45] <https://journals.sagepub.com/doi/10.2307/2666999>

<https://journals.sagepub.com/doi/10.2307/2666999>

[4] [70] <https://projecteuclid.org/journals/annals-of-mathematical-statistics/volume-16/issue-2/Sequential-Tests-of-Statistical-Hypotheses/10.1214/aoms/1177731118.full>

<https://projecteuclid.org/journals/annals-of-mathematical-statistics/volume-16/issue-2/Sequential-Tests-of-Statistical-Hypotheses/10.1214/aoms/1177731118.full>

[5] [32] https://www.hbs.edu/ris/Publication%20Files/20-018_0e319556-b28a-4528-bfa1-51db6d5831b8.pdf

https://www.hbs.edu/ris/Publication%20Files/20-018_0e319556-b28a-4528-bfa1-51db6d5831b8.pdf

[6] [14] [38]

<https://www.annualreviews.org/content/journals/10.1146/annurev.so.14.080188.001535>

<https://www.annualreviews.org/content/journals/10.1146/annurev.so.14.080188.001535>

[7] [62]

<https://economics.mit.edu/sites/default/files/publications/Conlon%20et%20al.%20--%20Failing%20to%20Learn%20from%20Others%20%28wit.pdf>

<https://economics.mit.edu/sites/default/files/publications/Conlon%20et%20al.%20--%20Failing%20to%20Learn%20from%20Others%20%28wit.pdf>

[8] [9] [15] <https://store.hbr.org/product/teaching-smart-people-how-to-learn/91301?srsId=AfmBOoqshcK0lHOqlBYJCh38BVrNrZmJE13qi7glHyz0eTW9xm7OUtnr>

<https://store.hbr.org/product/teaching-smart-people-how-to-learn/91301?srsId=AfmBOoqshcK0lHOqlBYJCh38BVrNrZmJE13qi7glHyz0eTW9xm7OUtnr>

[11] [65] <https://exp-platform.com/Documents/controlledExperimentDMKD.pdf>

<https://exp-platform.com/Documents/controlledExperimentDMKD.pdf>

[12] [19] [21] [23] [39] <https://ideas.repec.org/a/inm/ororsc/v2y1991i1p88-115.html>

<https://ideas.repec.org/a/inm/ororsc/v2y1991i1p88-115.html>

[13] <https://sjbae.pbworks.com/f/march%2B1991.pdf>

<https://sjbae.pbworks.com/f/march%2B1991.pdf>

[16] [25] [30] [57] <https://mitsloan.mit.edu/centers-initiatives/institute-work-and-employment-research/ai-work-redesign-playbook>

<https://mitsloan.mit.edu/centers-initiatives/institute-work-and-employment-research/ai-work-redesign-playbook>

[20] [24] [58] <https://mitsloan.mit.edu/ideas-made-to-matter/when-humans-and-ai-work-best-together-and-when-each-better-alone>

<https://mitsloan.mit.edu/ideas-made-to-matter/when-humans-and-ai-work-best-together-and-when-each-better-alone>

[22] [35] [53]

https://cistr.mit.edu/publication/2024_1201_EnterpriseAIMaturityModel_WeillWoernerSebastian

https://cistr.mit.edu/publication/2024_1201_EnterpriseAIMaturityModel_WeillWoernerSebastian

[26] [28] [29] [50] https://www.hbs.edu/ris/Publication%20Files/10-074_0bf3c151-f82b-4592-b885-cdde7f5d97a6.pdf

https://www.hbs.edu/ris/Publication%20Files/10-074_0bf3c151-f82b-4592-b885-cdde7f5d97a6.pdf

[27] [34] <https://exp-platform.com/Documents/2007-10EmetricsExperimentation.pdf>

<https://exp-platform.com/Documents/2007-10EmetricsExperimentation.pdf>

[31] [71] <https://web.mit.edu/manso/www/mi.pdf>

<https://web.mit.edu/manso/www/mi.pdf>

[33] [66] <https://ai.stanford.edu/~ronnyk/2007GuideControlledExperiments.pdf>

<https://ai.stanford.edu/~ronnyk/2007GuideControlledExperiments.pdf>

[36] [72] <https://onejustice.org/wp-content/uploads/2020/02/HBR-Article-The-Hard-Side-of-Change-Management.pdf>

<https://onejustice.org/wp-content/uploads/2020/02/HBR-Article-The-Hard-Side-of-Change-Management.pdf>

[37] [41] <https://nvlpubs.nist.gov/nistpubs/ai/nist.ai.100-1.pdf>

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

[40] <https://hbr.org/1993/07/building-a-learning-organization>

<https://hbr.org/1993/07/building-a-learning-organization>

[42] <https://www.library.hbs.edu/working-knowledge/learning-in-action>

<https://www.library.hbs.edu/working-knowledge/learning-in-action>

[43] <https://www.library.hbs.edu/working-knowledge/managing-to-learn-how-companies-can-turn-knowledge-into-action>

<https://www.library.hbs.edu/working-knowledge/managing-to-learn-how-companies-can-turn-knowledge-into-action>

[44] <https://hbr.org/1991/05/teaching-smart-people-how-to-learn>

<https://hbr.org/1991/05/teaching-smart-people-how-to-learn>

[46] <https://www.library.hbs.edu/working-knowledge/the-hard-work-of-failure-analysis>

<https://www.library.hbs.edu/working-knowledge/the-hard-work-of-failure-analysis>

[48] https://www.hbs.edu/ris/Publication%20Files/20-081_9d2608f9-5a49-4bc3-a902-307d3427477b.pdf

https://www.hbs.edu/ris/Publication%20Files/20-081_9d2608f9-5a49-4bc3-a902-307d3427477b.pdf

[49] https://www.hbs.edu/ris/Publication%20Files/AB_Testing_R_R_08b97538-ed3f-413e-bc38-c239b175d868.pdf

https://www.hbs.edu/ris/Publication%20Files/AB_Testing_R_R_08b97538-ed3f-413e-bc38-c239b175d868.pdf

[51] <https://www.library.hbs.edu/working-knowledge/managing-alignment-as-a-process>

<https://www.library.hbs.edu/working-knowledge/managing-alignment-as-a-process>

[52] <https://www.hbs.edu/faculty/Shared%20Documents/conferences/2019-imo/Sandino%20paper.pdf>

<https://www.hbs.edu/faculty/Shared%20Documents/conferences/2019-imo/Sandino%20paper.pdf>

[54] [63] https://cistr.mit.edu/publication/2024_0501_AIEverybodysBusiness_WixomBeath

https://cistr.mit.edu/publication/2024_0501_AIEverybodysBusiness_WixomBeath

[56] <https://mitsloan.mit.edu/ideas-made-to-matter/whats-your-companys-ai-maturity-level>

<https://mitsloan.mit.edu/ideas-made-to-matter/whats-your-companys-ai-maturity-level>

[59] <https://mitsloan.mit.edu/ideas-made-to-matter/how-ai-reshaping-workflows-and-redefining-jobs>

<https://mitsloan.mit.edu/ideas-made-to-matter/how-ai-reshaping-workflows-and-redefining-jobs>

[60] <https://mitmgmtfaculty.mit.edu/jsterman/business-dynamics/>

<https://mitmgmtfaculty.mit.edu/jsterman/business-dynamics/>

[61] https://web.mit.edu/nelsonr/www/Repenning%3DSterman_CMRSu01_.pdf

https://web.mit.edu/nelsonr/www/Repenning%3DSterman_CMRSu01_.pdf

[64] <https://www.nist.gov/itl/ai-risk-management-framework>

<https://www.nist.gov/itl/ai-risk-management-framework>

[67] <https://exp-platform.com/Documents/2017-08%20KDDMetricInterpretationPitfalls.pdf>

<https://exp-platform.com/Documents/2017-08%20KDDMetricInterpretationPitfalls.pdf>

[68] <https://pubsonline.informs.org/doi/10.1287/opre.2021.2135>

<https://pubsonline.informs.org/doi/10.1287/opre.2021.2135>

[69] <https://academic.oup.com/biomet/article/25/3-4/285/200862>

<https://academic.oup.com/biomet/article/25/3-4/285/200862>