

AI Leaders

Expert Exchange

Executive summary
May 20, 2026

kyndryl.



Overview

A cross-industry roundtable of AI and technology leaders at the Q1 Kyndryl AI Leaders Expert Exchange explored key challenges in enterprise adoption. Four primary themes emerged: the struggle to accurately measure the financial value of new software investments, the challenge of governing unauthorized employee tool usage, the urgent need to accelerate security updates against imminent threats, and the importance of transforming workforce processes with technology.

Host

Chinh Vo
Vice President, Data and AI
US Consult

SME

Justin Haney
Vice President, Security & Resiliency
Kyndryl US

Key topics

PAGE

- 04 Proving the Financial Value of Artificial Intelligence
- 05 Navigating Governance and Unauthorized Technology Use
- 06 Accelerating Security Responses for Emerging Threats
- 07 Transforming Workforce Processes and Employee Education

Kyndryl AI Leaders Expert Exchange: Executive brief



“Organizations frequently fail because leaders focus entirely on the technology, completely neglecting the fundamental changes required in people and operational processes.”

— AI Leaders Expert Exchange Member

Key takeaways

- Leaders are struggling to prove the financial value of AI, especially beyond efficiency gains into revenue impact
- Governance remains a major challenge as employees continue to use unauthorized tools, creating risk exposure
- Organizations must accelerate security response cycles to keep pace with emerging AI-driven threats
- Workforce and operating model transformation is critical—technology alone does not drive success

Implications

- Traditional metrics are insufficient—new frameworks are needed to measure AI value (speed, quality, outcomes)
- Risk management models must evolve to balance control with employee-driven innovation
- Security risk is increasing due to faster, more automated threat development and response gaps
- Organizations must invest in continuous education and workflow redesign to fully realize AI benefits

Key topics

Financial value of AI

- Traditional financial frameworks fail to capture productivity, risk avoidance, and customer impact
- Leaders are shifting toward localized, outcome-based and use-case-driven measurement models
- Industry-specific metrics (e.g., healthcare outcomes) complicate standard ROI models

Governance and risk

- Employees frequently circumvent approved tools, exposing enterprise data
- Debate: strict enforcement vs. flexible, risk-based governance models
- Effective governance depends heavily on organizational culture and monitoring capabilities

Security and resilience

- Security teams face pressure to rapidly accelerate patching and remediation cycles
- Emerging threats may chain vulnerabilities faster than legacy systems can respond
- Shift toward protecting critical data paths and isolating risk, not just patching systems

Workforce transformation

- AI success requires redesigning workflows—not layering technology onto existing processes
- Empowering employees with tools increases productivity but requires clear education and guardrails
- Organizations face tension between efficiency gains and workforce restructuring decisions

Proving the Financial Value of Artificial Intelligence

- Participants widely agreed that traditional financial frameworks fail to capture the true value of artificial intelligence, prompting a need for updated measurement strategies that account for speed, quality, and industry-specific outcomes.
 - A manufacturing executive expressed an ongoing internal struggle to find initiatives that drive market growth rather than just delivering basic operational efficiency. Leaders widely agreed that while they observe obvious productivity gains, tracking direct revenue expansion requires entirely new measurement frameworks.
- Highlighting a diverging perspective on metrics, a healthcare leader passionately argued that their primary measurements revolve strictly around patient care outcomes rather than traditional corporate accounting. They shared a subjective view that preventing negative medical incidents delivers absolute value, but translating that value into financial justifications remains exceptionally difficult.
 - A financial services participant pointed out a massive gap between outdated accounting models and the actual speed benefits modern software delivers. They argued that traditional business cases consistently fail to capture the combined value of productivity, risk avoidance, and customer satisfaction..
- Proposing an alternative approach, one member advocated for empowering employees with individual computing budgets based on personal justification rather than strict macro-level metrics. By leaning into localized time savings and driving individual adoption, their organization successfully achieved rapid, measurable increases in output quality..

“We still experience an internal struggle around finding use cases that deliver significant financial return, especially regarding driving market growth rather than just internal efficiency.”

— AI Leaders Expert Exchange Member

Unlock your Inside Edge with Kyndryl Agentic AI Digital Trust

[Learn more](#)

Navigating Governance and Unauthorized Technology Use

- | | | |
|---|---|---|
| <ul style="list-style-type: none">- The conversation revealed a shared challenge across industries regarding employees exposing corporate data to unapproved tools, leading to intense debates between strict enforcement and flexible risk management.- Executives from multiple regulated sectors voiced identical concerns regarding employees continuously adopting unauthorized consumer applications to bypass slower corporate systems. One leader detailed a disturbing incident where an employee actively transferred an entire customer database to a personal account strictly to utilize an unapproved external software model. | <ul style="list-style-type: none">- Participants widely agreed that stopping unauthorized use requires significant effort, but individuals diverged on the best response. A technology leader noted that resolving unauthorized practices required a forceful, top-down mandate and automated network scans to actively police corporate devices.- Conversely, a software executive offered a counterargument, explaining how their organization entirely abandoned rigid tool lists for a flexible, risk-based policy. This alternative approach categorizes projects by risk level, allowing general employee research while mandating executive approval for tasks involving confidential data. | <ul style="list-style-type: none">- The Kyndryl host highlighted that internal policing heavily depends on an organization's specific culture and investigation capabilities. They pointed out that while some agencies strictly monitor and block actions, smaller businesses lack this ability, meaning leaders cannot universally apply rigid risk management. |
|---|---|---|

“Organizations can implement governance policies and lists of approved tools, but you can never completely stop someone who believes they have a better way to solve a problem.”

— AI Leaders Expert Exchange Member

Accelerating Security Responses for Emerging Threats

- Participants expressed concerns about the emergence of Claude Mythos, agreeing that organizations must urgently accelerate legacy software maintenance and establish stronger operational defenses.
- The discussion uncovered a collective fear that upcoming advanced threat systems will automatically chain together minor software vulnerabilities to rapidly overwhelm standard corporate defenses. Multiple executives agreed that current maintenance cycles for older computer code operate too slowly to handle simultaneous critical fixes.
- Demonstrating a proactive approach, an infrastructure leader shared that their organization formed a dedicated task force to aggressively shrink foundational software update cycles from forty days down to thirty days. They actively partner with major vendors to test automated incident remediation, showing how advanced systems effectively fight external threats.
- Providing a strategic counterargument, Kyndryl advised moving away from traditional vulnerability scoring to focus purely on protecting critical data paths. He recommended establishing defensive controls, such as network isolation, to heavily protect older systems that technology teams cannot update quickly enough.
- A government technology leader introduced a diverging concern, arguing that the true obstacle involves a severe lack of available staff to manage the flood of automated security alerts. He questioned how organizations can successfully address these threats when daily operational maintenance entirely consumes the current workforce.

“Security teams drive with extreme urgency to accelerate patching and remediation processes before developers release highly advanced threat models.”

—AI Leaders Expert Exchange Member

Transforming Workforce Processes and Employee Education

- | | | |
|---|--|--|
| <ul style="list-style-type: none">- The executives agreed that deploying advanced technology remains fundamentally insufficient without actively restructuring legacy business operations and heavily investing in continuous employee education.- Kyndryl shared that roughly three-quarters of technology projects fail strictly due to neglected human elements rather than technical limitations. They argued that companies must analyze how entire business workflows change, warning that merely inserting technology into existing operations yields only marginal improvements. | <ul style="list-style-type: none">- Highlighting a point of strong agreement, multiple leaders praised the concept of empowering non-technical staff by placing powerful development tools directly into their hands. A member enthusiastically noted that removing bureaucratic obstacles allowed teams to independently build customized workflows, significantly supercharging their daily output.- An executive introduced a counterargument regarding this widespread democratization, emphasizing that open access requires rigorous education to prevent massive resource waste. They stressed that leaders must explicitly teach employees when to use simple tools versus highly complex computing models to avoid unnecessary processing costs. | <ul style="list-style-type: none">- A healthcare leader voiced deep frustration over the difficulty of reorganizing operational personnel even after team realizes efficiency gains. They noted an ongoing battle between finance teams and operational partners, where business units strongly resist eliminating positions despite obvious technology-driven time savings. |
|---|--|--|

“Organizations frequently fail because leaders focus entirely on the technology, completely neglecting the fundamental changes required in people and operational processes.”

—AI Leaders Expert Exchange Member

Agentic AI at scale: The CD&AI scorecard

Agentic AI delivers CD&AI success only when speed, trust and governance are engineered together.

Time-to-value and scale	Reliability of outcomes	Governance and auditability by design	Sustainable operating model
(From pilots to enterprise workflows)	(Sustaining trust in AI decisions)	(Scaling AI without losing trust)	(Autonomy with control, humans with agents)
What CD&AI leaders care about: Speed from concept to production; ability to scale beyond isolated use cases	What CD&AI leaders care about: Consistency of decisions; confidence in AI-driven outcomes over time	What CD&AI leaders care about: Explainability, audit readiness, regulatory alignment	What CD&AI leaders care about: Adoption at scale; clear decision rights; manageable operating cost
What the research shows: - AI-native enterprises move beyond disconnected pilots by embedding AI at the core of workflows, enabling faster sensing, decision-making and action - Agentic AI enables continuous execution across systems — but only when designed for enterprise-wide orchestration, not one-off copilots	What the research shows: - Agentic AI introduces new risk surfaces — including adversarial prompts, poisoned data and opaque decision paths that traditional application security does not address - Static, perimeter-based controls fail when AI systems are adaptive, autonomous and continuously learning	What the research shows: - Policy as code translates business rules and regulatory requirements into machine-enforceable guardrails that govern how AI agents act - This creates traceable, reviewable, and explainable agent decisions — critical for regulated and mission-critical environments	What the research shows: - AI-native organizations are built around human-agent collaboration, with clear boundaries on what agents can do independently and when humans intervene - Treating security and governance as a final checkpoint increases friction and cost; agentic AI requires security and policy embedded from design through runtime
Implication: Speed comes from designing agentic AI as a repeatable operating model, not as experiments.	Implication: Outcome reliability requires continuous controls embedded across the AI lifecycle, not post-hoc validation.	Implication: Governance must move from documentation to execution-level enforcement.	Implication: The winning model balances autonomy and oversight — enabling scale without chaos.



Additional Resources:

Security

- [Redefining security for the agentic AI era](#)

AI and policy as code

- [Journey to AI Native](#)
- [Policy As Code - Article](#)
- [Policy As Code Homepage](#)

The AI Expert Exchange is hosted by Kyndryl. Please contact Chinh Vo (chinh.vo@kyndryl.com) or Justin Haney (justin.haney@kyndryl.com) with any questions about Kyndryl or this Exchange.

© Copyright Kyndryl, Inc. 2026

Kyndryl is a trademark or registered trademark of Kyndryl Inc. in the United States and/or other countries. Other product and service names may be trademarks of Kyndryl Inc. or other companies.

This document is current as of the initial date of publication and may be changed by Kyndryl at any time without notice. Not all offerings are available in every country in which Kyndryl operates. Kyndryl products and services are warranted according to the terms and conditions of the agreements under which they are provided.