

BLACKRIDGE AI & LEADERSHIP · 2026

What Is AI Leadership Readiness?

Why being ready to adopt AI is not the same as being ready to lead in an AI-shaped role.

6,047 establishments

OECD cross-national algorithmic-management survey [1]

758 consultants

BCG capability-frontier experiment [4]

5,179 support agents

workplace GenAI deployment across 3M+ chats [5]

28% reported

unclear accountability when algorithmic recommendations were wrong [1]

WHAT THIS PAPER EXAMINES

AI readiness and AI leadership readiness are related, but they answer different questions. Organizational readiness asks whether the system around AI is prepared: governance, data, infrastructure, talent, policy, risk controls, investment, and the operating conditions for trustworthy use. Leadership readiness asks whether a particular person is prepared for what AI changes about a particular role.

BlackRidge uses AI leadership readiness as a working definition, not as a claim that a new psychometric construct has already been validated. The central question is practical: when AI can accelerate information, draft analysis, generate recommendations, and reshape coordination, can the leader still judge what matters, challenge what is wrong, lead people through the consequence, and personally own the decision?

This paper brings together experimental evidence, workplace deployments, algorithmic-management research, public-sector AI workforce guidance, and leadership-readiness theory. The evidence is strongest on AI-assisted task performance and decision risk. Direct studies linking AI exposure to transition-specific leadership readiness and subsequent leader performance remain limited. Where evidence is indirect, BlackRidge labels the implication rather than presenting it as causal proof.

BLACKRIDGE DISTINCTION

AI readiness asks whether the organization is prepared to adopt AI. AI leadership readiness asks whether the leader is prepared for what AI changes about the role.

CONTENTS

The Readiness Map

Nine passages move from definition to evidence, then from evidence to assessment, development, and organizational responsibility.

00	Executive Summary	The distinction in one sentence: system readiness and leader readiness are not the same level of analysis.
01	Two Readiness Questions	What organizational AI readiness covers - and why it does not answer whether a leader is prepared.
02	A Working Definition	What BlackRidge means by AI leadership readiness, term by term.
03	Tool Fluency Is Not Readiness	Why fast, polished AI-assisted output can hide dependency or untested judgment.
04	The Jagged Frontier	Why leaders need calibrated reliance rather than blanket trust or blanket skepticism.
05	Accountability Is More Than Approval	What human ownership requires when recommendations are automated or AI-assisted.
06	Readiness Evidence	Which behaviors become especially revealing in AI-shaped leadership roles.
07	Development + Application	How to practice transfer, challenge, voice, and consequence instead of software usage alone.
08	Research + Boundaries	What the evidence supports, what remains open, and what BlackRidge deliberately does not claim.

EXECUTIVE SUMMARY

An organization can be ready to deploy AI while a leader is unready to carry what AI changes.

AI can improve access to information, distribute expert practices, accelerate routine work, and create new ways to coordinate. None of those conditions by themselves establish that a leader is ready. Readiness still has to be demonstrated in judgment, relationships, decision ownership, adaptation, and the ability to create results through other people when technology is part of the operating environment.



01 System readiness is context

Governance, data, infrastructure, policy, skills, and risk controls create the conditions in which AI can be used responsibly.

02 Leader readiness is person-role fit

The question is whether a named leader can carry the judgment, people, and accountability demands of the role ahead.

03 AI can hide weak evidence

Assisted output can rise faster than independent judgment. Readiness therefore requires process, challenge, transfer, and ownership.

DECISION POINT

Do not ask only, "Can this leader use AI?" Ask, "Can this leader judge, lead, and own the outcome when AI changes the information, speed, and coordination around the role?"

AI leadership readiness is not a technology score. It is a time-stamped judgment about preparation for a specific leadership role in an environment where AI changes what information is available, how work is performed, and where human responsibility becomes most consequential.

01 • TWO READINESS QUESTIONS

The organization and the leader can be ready at different times.

Organizational AI readiness concerns the system around adoption. OECD and NIST guidance emphasizes governance, data, infrastructure, skills, risk management, oversight, and other enabling conditions for trustworthy use. Leadership readiness sits inside that environment, but it asks a different unit-of-analysis question: is the person prepared for the responsibilities the role will carry? [2,3]

Organizational AI readiness Typical focus: • governance and risk controls • data and digital infrastructure • skills and talent • policies, procurement, investment • model/system evaluation • change management and adoption Question: Can the organization deploy and govern AI responsibly?	AI leadership readiness BlackRidge focus: • judgment and calibrated reliance • accountability for outcomes • people leadership and legitimacy • pattern awareness and context • resource choice and escalation • transfer when the tool changes or fails Question: Can this leader carry the role well in an AI-shaped environment?
--	--

A BLACKRIDGE INTERPRETATION

System readiness and leader readiness should be read together, not substituted for one another.

	LEADER DEVELOPING	LEADER READY
SYSTEM READY	Capable adoption The system supports trustworthy AI use and the leader can exercise judgment, voice, and accountability.	Governed tools, weak decisions Technology and policy may be sound while the leader over-relies, avoids ownership, or cannot carry people consequences.
SYSTEM DEVELOPING	Ready person, weak system A capable leader may still be constrained by poor data, unclear decision rights, missing appeal paths, or unsafe deployment.	Compounded readiness risk Neither the surrounding system nor the person has enough support to carry consequential AI-enabled decisions reliably.

KEY DISTINCTION

Organizational readiness is an enabling condition. It is not evidence that a particular leader is prepared for a particular role.

02 · WORKING DEFINITION

AI leadership readiness is a person-in-role judgment, not an AI skill score.

BlackRidge working definition: AI leadership readiness is the degree to which a leader is prepared to exercise judgment, lead people, make accountable decisions, and carry the responsibilities of a role whose work and operating environment are increasingly shaped by artificial intelligence.



Degree Readiness is not binary. It can strengthen, weaken, or change as the role and environment change.	Prepared The comparison is to future-role demands, not merely current performance or enthusiasm for AI.	Exercise judgment The leader can verify, challenge, compare, escalate, and decide when evidence is incomplete or conflicting.
Lead people The leader creates clarity, trust, standards, development, and legitimate process - not just efficient output.	Carry responsibility The leader can explain the decision, hear challenge, correct error, and personally own the consequence.	AI-shaped role AI changes the information and coordination around the job; it does not erase the human obligations of leadership.

WHAT IT IS NOT

AI leadership readiness is not prompt skill, AI optimism, technology adoption, or an organization-wide maturity score. Those may matter. They do not answer whether a leader is prepared to carry the role.

03 • TOOL FLUENCY IS NOT READINESS

AI can make performance look mature before judgment is transferable.

A real customer-support deployment followed 5,179 agents across more than three million chats. AI assistance raised issues resolved per hour by 15%, with the largest gains among less-experienced workers. But the research could not determine whether workers had independently learned the underlying reasoning or were following high-quality recommendations. [5]

15% more issues per hour

Workplace GenAI deployment; peer-reviewed estimate. Strong evidence of task-performance gain, not leadership readiness.

5,179 agents

The sample is unusually large and workplace-based, but it involved support agents - not newly promoted managers.

Transfer not measured

The study could not separate durable learning from successful reliance on recommendations.

OBSERVED BEHAVIOR	WHAT IT CAN SHOW	WHAT IT CANNOT ESTABLISH
Prompts well	Resource fluency	Independent judgment
Produces a fast draft	Execution speed	Standards or decision quality
Uses AI frequently	Adoption behavior	Calibrated reliance
Accepts a recommendation	Compliance with advice	Ownership or fairness
Explains the output	Communication	Whether the reasoning is correct
Gets a strong result	Performance on this task	Transfer to a changed problem

A CHI 2025 survey of 319 knowledge workers points to the same developmental risk from another angle: greater confidence in GenAI was associated with less self-reported critical-thinking effort, while greater task-specific self-confidence was associated with more. The study is correlational, but it reinforces why readiness evidence should include how a person checks, integrates, and stewards AI-assisted work - not only whether they can produce it. [7]

READINESS BOUNDARY

Assisted performance shows what a person can do with a resource. It does not by itself prove independent judgment, transfer, or ownership.

04 • CALIBRATED JUDGMENT

The best leader is not the person who always trusts AI - or always distrusts it.

In a preregistered experiment with 758 BCG consultants, AI improved speed and output on tasks inside the model's capability frontier. On a deliberately different managerial problem outside that frontier, AI users were 19 percentage points less likely to reach the correct answer. The risk was not obvious: the AI could still produce persuasive explanations. [4]

Inside the frontier

AI users completed 12.2% more tasks and worked 25.1% faster on realistic tasks the model handled well.

Readiness implication: use the resource.

Boundary is uncertain

The hard part is recognizing when a seemingly similar task has moved outside the model's reliable capability.

Readiness implication: test task fit.

Outside the frontier

Correctness fell by 19 percentage points on the outside-frontier task despite fluent AI output.

Readiness implication: challenge and override.

WHAT CALIBRATED RELIANCE LOOKS LIKE

1. Task fit

What kind of problem is this, and is AI likely to be reliable here?

2. Stakes

What happens if the recommendation is wrong, and is the choice reversible?

3. Evidence

Can I reconstruct the inputs, assumptions, and missing context?

4. Challenge

What would change my answer? What independent check is appropriate?

5. Decision

Can I explain the reasoning and own the consequence without hiding behind the tool?

LEADERSHIP IMPLICATION

Confidence should be strong enough to act, but conditional on task fit, evidence quality, stakes, and reversibility. Fluency is not a substitute for calibration.

05 • ACCOUNTABILITY, VOICE, LEGITIMACY

A human in the loop is not the same as human accountability.

In the OECD survey, 28% of respondents reported unclear accountability when an algorithmic recommendation was wrong, and 64% reported at least one trustworthiness concern. Separate personnel-selection experiments involving 1,403 participants found that incorrect AI advice reduced decision quality and that explanations did not reliably eliminate overreliance. [1,6]

28%

reported unclear accountability when
algorithmic recommendations were
wrong

Approval can be ceremonial. Accountability is deeper. The leader must be able to understand what evidence drove the recommendation, identify missing context, hear challenge from the affected person, detect repeated or group-level error, and explain why the final decision is justified. If the answer cannot survive those tests, clicking “approve” does not make the process genuinely human-governed.

FIVE QUESTIONS THAT MAKE ACCOUNTABILITY REAL

01	What evidence drove the recommendation?	Can I reconstruct the inputs, assumptions, and data sources?
02	What context is missing?	What facts, history, constraints, or employee perspective are invisible to the system?
03	Is this one error - or a pattern?	Does the output behave differently across cases, groups, or time?
04	Who is affected and who can challenge it?	Is there a genuine path for voice, appeal, correction, or review?
05	Can I explain and own the decision?	If challenged by the employee, HR, my manager, or a regulator, can I defend the reasoning rather than cite the tool?

DECISION POINT

Accountability becomes real when the leader can trace, verify, challenge, hear voice, correct, explain, and personally carry the consequence - not merely approve the output.

06 • WHAT READINESS LOOKS LIKE

Six existing readiness behaviors become especially revealing when AI is in the workflow.

BlackRidge does not treat these as a new AI competency model. They are evidence-informed emphases within the broader readiness system. The strongest direct evidence concerns ownership, calibrated decision confidence, pattern detection, self-monitoring, interpersonal interpretation, and the choice of resources. Direct validation across AI-shaped leadership roles remains a research agenda.

READINESS SIGNAL	WHAT IT LOOKS LIKE IN AN AI-SHAPED ROLE	WHAT IT IS NOT
Owning the Outcome	Traces, challenges, corrects, and personally owns AI-assisted decisions.	"The system recommended it."
Decision Confidence	Acts with enough confidence to decide, but adjusts reliance to evidence, stakes, and reversibility.	Blanket trust or blanket skepticism.
Pattern Awareness	Detects repeated error modes, group disparities, missing context, and changed conditions.	Spotting one obvious mistake.
Inner Awareness	Notices deference, reactance, confirmation seeking, and cognitive offloading in one's own use.	Assuming one is unbiased because a human is present.
Reading the Room	Sees silence, fear, surveillance concerns, resistance, or perceived unfairness that dashboards miss.	Treating sentiment as noise.
Resource Confidence	Chooses appropriately among AI, data, peers, HR/legal, escalation, and direct conversation.	Treating AI as the universal expert.

EVIDENCE BOUNDARY

These signals are priorities to observe and validate, not a claim that BlackRidge has already established a separate validated "AI leadership readiness" scale.

READINESS IS ROLE-SPECIFIC

AI does not create one universal leadership profile.

The relevant readiness question changes with the leadership transition. AI may alter the information, pace, coordination, and decision environment at every level, but the responsibility that must be carried is different. The same person can be ready for one role and underprepared for another.

SHIFT

Individual contributor
→ first-line manager

Can the leader stop being the answer source and become accountable for people, standards, decisions, and outcomes?

From expertise to responsibility

ARENA

Manager → leading
through managers

Can the leader scale judgment through other leaders, define decision rights, and resist becoming the AI-enabled bypass?

From direct control to distributed authority

ASCENT

Functional leader →
broader integration

Can the leader integrate technology, functions, people, and priorities without letting speed replace alignment?

From functional excellence to enterprise integration

EDGE

Enterprise leadership

Can the leader govern institutional consequence, ethics, risk, legitimacy, and accountability when AI affects the whole system?

From decision quality to institutional stewardship

BLACKRIDGE POSITION

If the role changes, the evidence of readiness must change with it. AI should be added to the context of the transition - not treated as a generic skill that replaces transition-specific assessment.

GATHER BETTER EVIDENCE

Do not score the output alone. Observe how the leader gets there.

A polished AI-assisted answer can conceal weak reasoning just as easily as it can reveal good judgment. Assessment should therefore vary support conditions and make the leader's process visible. The purpose is not to test whether someone can "beat" AI. It is to see whether judgment remains calibrated when the technology is useful, uncertain, absent, or wrong.

01 Assisted output

What can the person accomplish with AI available?

02 Reasoning trace

Can the person explain the problem framing, evidence, assumptions, and standard?

03 Challenge + verification

Does reliance change when the task, evidence, or model fit changes?

04 Transfer

Can the reasoning survive when AI is unavailable, weaker, or the problem shifts?

05 Human consequence

Can the leader hear voice, read context, address fairness, and preserve legitimacy?

06 Ownership after outcome

Can the leader correct course and own what happened rather than outsource responsibility?

USEFUL ASSESSMENT DESIGNS

- paired scenarios inside and outside the AI capability frontier
- one task with AI, one after AI feedback, and one without AI
- a series of individually plausible outputs containing a systemic pattern
- a recommendation that creates an employee-voice or fairness consequence

- a resource-allocation choice among AI, data, peer expertise, HR/legal, and escalation
- a requirement to explain what evidence would change the decision
- a follow-up outcome that forces the leader to correct, communicate, and own the consequence

EVIDENCE PRINCIPLE

Readiness is stronger when the behavior transfers across support conditions. A single high-quality AI-assisted answer should never carry the whole judgment.

SCENARIO • HUMAN CONSEQUENCE

The AI recommends denying the promotion. Now what?

Imagine a leader receives an AI-assisted recommendation to deny an employee's promotion. The model cites recent performance indicators and a pattern of missed targets. The output is concise, confident, and consistent with the dashboard. Then the employee raises context the system does not contain: a changed territory, two long-term absences on the team, and a prior manager who assigned stretch work without updating the formal goals.



01 Trace the evidence Which data, time periods, proxies, and assumptions shaped the recommendation?	02 Add missing context Which facts are decision-relevant but not represented in the system?	03 Check the pattern Is the recommendation part of a repeated error or disparity across similar cases?	04 Hear voice Can the employee challenge the evidence and be heard before consequence is imposed?	05 Decide + own What should happen, why, and who will explain and carry the result?
--	--	---	--	--

The point is not to reject AI because a human has additional context. The point is to govern the decision. A ready leader can use the recommendation as evidence, test its fit, hear the person affected, identify what the system cannot see, and still make a timely decision that can be explained and defended.

WHAT THIS SCENARIO REVEALS

Owning the Outcome, Decision Confidence, Pattern Awareness, Reading the Room, Inner Awareness, and Resource Confidence all become visible in one consequential decision.

DEVELOP TRANSFER, NOT DEPENDENCE

Leadership development should vary the level of AI support - not maximize it.

If development always occurs with the tool available and performing well, organizations may strengthen workflow fluency without learning whether judgment transfers. Better practice deliberately changes the support condition while keeping the leadership responsibility constant.



DEVELOPMENT METHOD	WHAT PRACTICE SHOULD INCLUDE	WHAT IT REVEALS
Variable AI availability	Practice with AI available, partially available, and unavailable.	Does reasoning survive support changes?
Frontier shifts	Use two similar cases where AI is reliable in one and misleading in the other.	Does reliance change with task fit?
Error-pattern practice	Present several plausible outputs containing a repeated group-level or temporal pattern.	Can the leader detect systemic rather than isolated error?
Human consequence	Attach employee voice, conflict, fairness, or trust to an otherwise efficient recommendation.	Can the leader preserve legitimacy while deciding?
Reasoning debrief	Score framing, evidence, challenge, alternatives, thresholds, and ownership - not just the final answer.	Can the leader explain the process?

DEVELOPMENT PRINCIPLE

The aim is not independence from AI. The aim is independence of judgment: a leader who can use the resource, challenge it, learn from it, and still carry the responsibility when conditions change.

THE ORGANIZATION IS PART OF READINESS

A capable leader can still enter an unready AI environment.

Leadership readiness is not solely an individual burden. NIST and OECD guidance make clear that governance, data, infrastructure, skills, oversight, and organizational capacity shape whether AI can be used responsibly. A leader cannot demonstrate sound accountability inside a system that withholds the evidence, blurs decision rights, or gives affected people no meaningful path to challenge. [2,3]

ORGANIZATIONAL OBLIGATION	WHAT GOOD SUPPORT LOOKS LIKE	WHAT FAILURE CREATES
Clarify decision rights	Define where AI may inform, recommend, automate, or require human decision.	Leaders inherit responsibility without knowing what authority they actually hold.
Create traceability	Make relevant inputs, data provenance, assumptions, and system limitations accessible.	The leader is asked to explain a decision that cannot be reconstructed.
Protect voice + appeal	Provide a real channel for employees or stakeholders to challenge consequential outcomes.	Human review becomes ceremonial because no one affected can contest the process.
Provide escalation	Make HR, legal, risk, technical, and senior-leader support available for high-stakes cases.	The manager either over-trusts the tool or escalates everything.
Observe behavior	Review real AI-assisted decisions, conversations, corrections, and transfer - not adoption metrics alone.	Tool usage is mistaken for readiness or maturity.
Develop the role ahead	Connect development to the actual transition, decision environment, and support conditions.	Generic AI training is expected to solve a leadership problem it was not designed to address.

BLACKRIDGE POSITION

Do not ask leaders to prove AI leadership readiness inside conditions that make accountable leadership impossible. The surrounding system is part of the evidence.

RESEARCH BOUNDARIES

What this paper is - and is not.

This paper synthesizes research about AI-assisted work, algorithmic management, decision quality, human oversight, workforce readiness, and leadership development. It uses that evidence to define a BlackRidge working concept and to identify assessment and development implications. It does not claim that a new construct has already been psychometrically validated.

It is A research-grounded working definition for thinking about person-role readiness when AI changes the work environment.	It is not A peer-reviewed validation study, organization-wide AI maturity score, or certification of “AI-ready leadership.”	Strongest evidence Task performance, calibration risk, algorithmic-management use, accountability concerns, and decision overreliance.	Important unknowns Which readiness signals predict successful leadership across different AI-shaped roles, industries, cultures, and technologies.
--	--	---	---

PUBLICATION GUARDRAILS

- 01 Use “AI-assisted” unless a study demonstrates autonomous execution.
- 02 Separate faster output from durable learning, judgment, or leadership maturation.
- 03 Treat organizational flattening and promotion changes as separate from AI causation unless jointly measured.
- 04 Do not infer readiness from tool usage, current-role performance, or polished output alone.
- 05 Treat BlackRidge assessment implications as evidence-informed until predictive validity is established.

SELECTED RESEARCH

- [1] OECD (2025)
Algorithmic Management in the Workplace. Employer survey of 6,047 establishments across six countries.
- [2] NIST (2023/2024)
AI Risk Management Framework 1.0 and Generative AI Profile.
- [3] OECD (2026)
Building an AI-ready Public Workforce.
- [4] Dell’Acqua et al.
Navigating the Jagged Technological Frontier. 758 BCG consultants.
- [5] Brynjolfsson, Li & Raymond (2025)
Generative AI at Work. Quarterly Journal of Economics.
- [6] Cecil et al. (2024)
Incorrect AI advice in personnel selection. Scientific Reports.
- [7] Lee et al. (2025)
The Impact of Generative AI on Critical Thinking. CHI.
- [8] Dell’Acqua et al. (2025)
The Cybernetic Teammate, P&G field experiment.

ABOUT THE WORKING DEFINITION

“AI leadership readiness” is BlackRidge organizing language for an emerging leadership problem. The evidence supports the problem and many of its behavioral implications; validation of the construct itself remains future research.

BLACKRIDGE CLOSING POSITION



**AI readiness prepares the organization to use the technology.
AI leadership readiness prepares the person to carry what the
technology changes.**

AI changes what information is available, how quickly options appear, and how work is coordinated. It can improve performance and still create new risks of overreliance, borrowed execution, blurred accountability, and weakened human voice. That makes readiness more - not less - dependent on judgment, context, development of others, and ownership of consequence.

THE BLACKRIDGE QUESTION

Not "Is the leader good at AI?" but "Is the leader prepared for the role ahead when AI is part of the environment?"

BUILD READINESS FOR THE ROLE AHEAD - NOT JUST ADOPTION OF THE TOOL.

blackridgeinstitute.com/insights/ai-leadership-readiness/