
WORKING PAPER No. 1

The Measurement Gap

Why AI adoption is rising, impact isn't, and the missing layer is human.

PRELIMINARY · PRE-RELEASE · LIMITED CIRCULATION

Authors Deven Spear; David Wilson; The AI Readiness Institute (advisory review pending)

Published June 30, 2026 · **Status** Pre-release for limited circulation — not for public distribution

Suggested citation: The AI Readiness Institute (2026). The Measurement Gap. Working Paper No. 1 (preliminary).

Abstract

Organizations have adopted artificial intelligence almost universally, yet most report little measurable return. In 2025, roughly **88% of organizations** said they used AI regularly in at least one function, but only **39%** could attribute any earnings impact to it — and most of those at under 5% of EBIT. This paper argues that the gap between adoption and impact is not primarily a problem of infrastructure, strategy, or model quality. It is a **measurement problem**: organizations measure AI readiness at the wrong altitude. The dominant instruments assess national policy environments or organizational maturity — strategy, data, governance, tooling — while leaving unmeasured the layer where adoption actually happens or stalls: the individual employee, and specifically their **attitude, skill, and enablement**. Drawing on a structured synthesis of the major 2024–2026 workforce studies, we show that the human layer is both decisive and systematically unmeasured, that the field is already converging on persona-based segmentation of it, and that no existing model of that layer is independent, open, or validated. We introduce the **Seven-Persona AI Readiness Standard** as an instrument for the missing layer, and argue that such a standard must be independent, openly published, and empirically validated — precisely because the proprietary alternatives are none of these. We close with the Institute's commitments and an invitation to cite, adopt, and challenge the Standard.

1. The paradox: usage is up, impact isn't

The defining fact of enterprise AI in 2025 is a contradiction. Adoption is nearly universal and still climbing; measurable impact is not.

On the adoption side, the numbers are unambiguous. McKinsey's global survey found that **78% of organizations used AI in at least one business function in 2024, rising to roughly 88% in 2025** — the figure had been 55% as recently as 2023 [McKinsey, 2025; Stanford HAI AI Index, 2025]. Generative AI specifically more than doubled in a single year. At the individual level, **about 75% of knowledge workers** reported using generative AI at work by mid-2024 [Microsoft & LinkedIn, Work Trend Index, 2024], and Gallup found everyday AI use among U.S. employees **nearly doubled from 21% to 40%** between 2023 and 2025 [Gallup, 2025]. By any historical standard, this is one of the fastest workplace-technology diffusions ever recorded.

On the impact side, the picture inverts. In McKinsey's 2025 survey, **only 39% of organizations attributed any earnings (EBIT) impact to their AI use — and most of those reported the impact at less than 5% of EBIT** [McKinsey, 2025]. BCG reached the same destination from a different road: **about 60% of companies were generating no material value from AI** [BCG, 2025]. The Boston Consulting Group's developer-focused study of adoption *quality* — how deeply AI is

woven into work, not merely whether it is touched — found that **more than 85% of workers remained at shallow “assistance” stages, while fewer than 10% had reached genuine AI-augmented collaboration** [BCG, “The AI Adoption Puzzle,” 2025]. (*This last figure is drawn from surveys of software developers, treated as a leading indicator; we flag the sample explicitly.*)

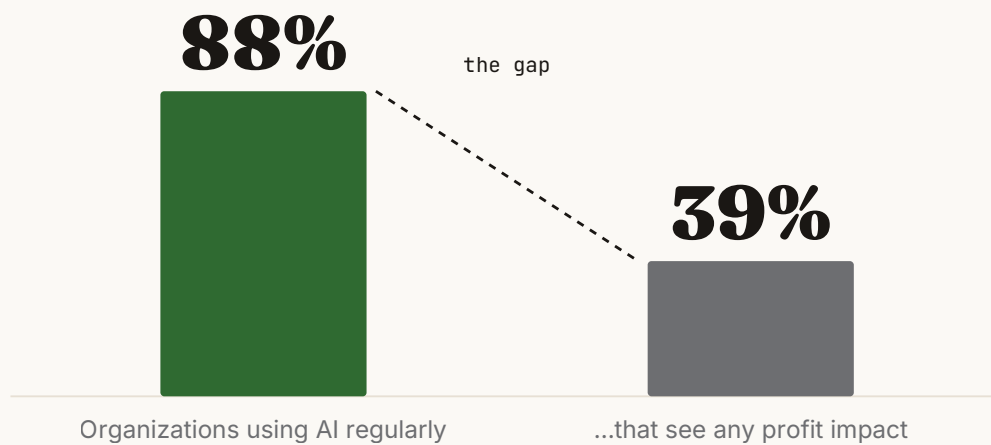


FIGURE 1 — The adoption-impact gap. Usage has scaled; value has not. (McKinsey, 2025.)

This is the measurement gap’s surface symptom. Enormous sums are flowing into AI tools, licenses, and training, and the spending is not translating into outcomes at anything like the same rate. The natural managerial response is to ask *which technology, which use case, which platform* — to treat the shortfall as a tooling or strategy problem. We argue this is the wrong question, and that asking it is itself a symptom of measuring at the wrong altitude.

2. Why the gap persists: the wrong altitude

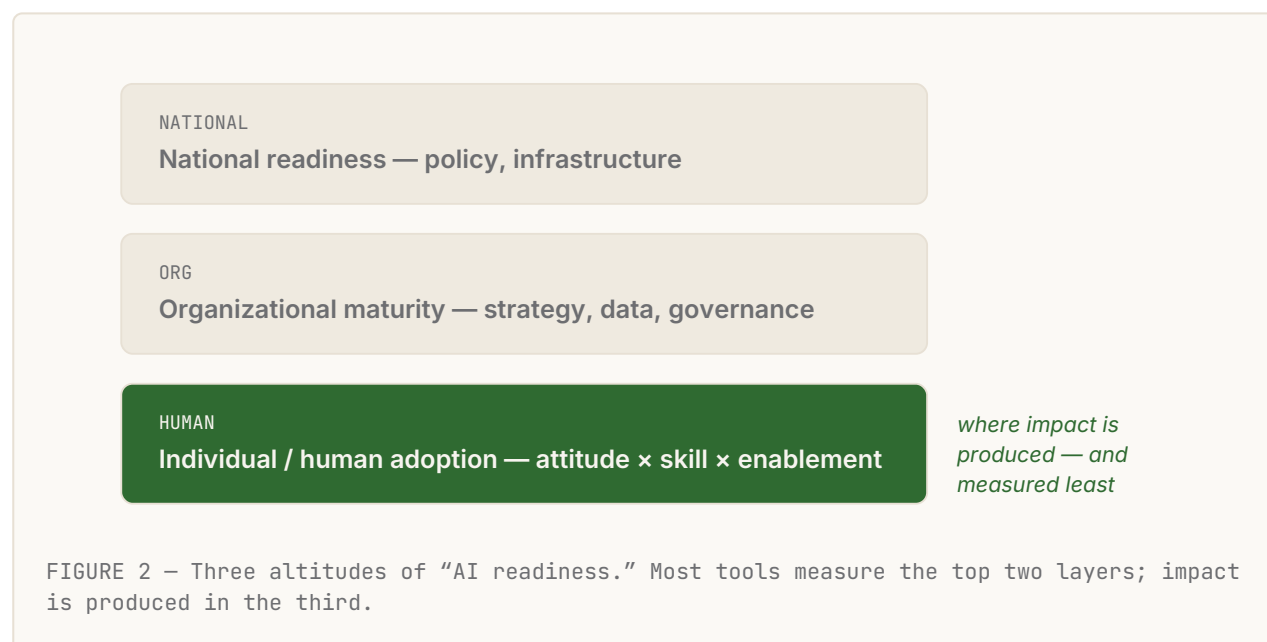
If you want to know why AI investment isn’t producing impact, you would expect to consult an “AI readiness” assessment. A revealing thing happens when you do: almost every available instrument measures something other than the people doing the work.

Consider the landscape of what “AI readiness” currently means in practice:

- **National readiness** — instruments such as Oxford Insights’ Government AI Readiness Index, UNESCO’s Readiness Assessment Methodology, and the OECD’s AI policy observatory rank *countries* on policy, infrastructure, and talent supply. The unit of analysis is a nation.
- **Organizational maturity** — Cisco’s AI Readiness Index, Microsoft’s and ServiceNow’s maturity indices, Gartner’s AI Maturity Model, and the major consultancies’ frameworks score

a *company* on strategy, data, infrastructure, governance, and operating model. The unit of analysis is an organization's plumbing.

Both layers matter. But notice what neither measures: the individual employee's readiness to actually change how they work. A company can score well on data infrastructure and AI strategy and still see no impact, because impact is produced when a marketing coordinator, a claims adjuster, or an account manager changes their daily practice — and whether they do is a function of their skill, their attitude, and whether the organization has enabled them. That layer sits below the altitude at which the dominant instruments operate.



The consequence is that organizations are, in effect, flying with the wrong instruments. They can see their strategy and their stack; they cannot see their people. And the evidence — to which we now turn — indicates that the people are exactly where the impact gap lives.

It is worth being precise about the unit-of-analysis confusion that pervades this field, because it obscures the gap. When a vendor reports “78% AI adoption,” that is *organizations* using AI somewhere. When Pew reports that **only 16% of U.S. workers say AI does even some of their work**, that is *individuals* [Pew Research Center, 2025]. Both are true; they describe different altitudes. The headline “adoption is everywhere” is an organizational statistic. At the human altitude, adoption is shallow, uneven, and — critically — unmeasured.

3. The human layer, in the evidence

If the human layer is where impact is won or lost, what does the evidence say is happening there? Four findings recur across independent studies, and together they describe a workforce whose readiness is real but unrealized.

First, adoption is wide but shallow. Usage has spread far faster than capability or integration. BCG's adoption-quality work (developer sample) places the overwhelming majority of workers at early "assistance" stages rather than deep, value-generating collaboration [BCG, 2025]. Gallup found that even among users, only **16% strongly agree the AI tools their organization provides are useful** [Gallup, 2025]. Breadth of access has outrun depth of use.

Second, the enablement deficit is severe. This is the single most consistent theme in the literature. Slack's Workforce Lab found that **only 15% of desk workers strongly agree they have the training they need to use AI effectively** [Slack Workforce Lab, June 2024, n=10,045]. BCG's 2026 edition found **just 36% of workers feel they have received adequate upskilling**, even as 72% say expectations for their skills have shifted [BCG, 2026]. When workers are asked what blocks them, the top answer is not fear or distrust — it is **lack of time, cited by 41%** [Study.com, 2026], with roughly a third reporting *no* AI training of any kind. The pattern is not a workforce that refuses AI; it is a workforce that has not been equipped to use it.

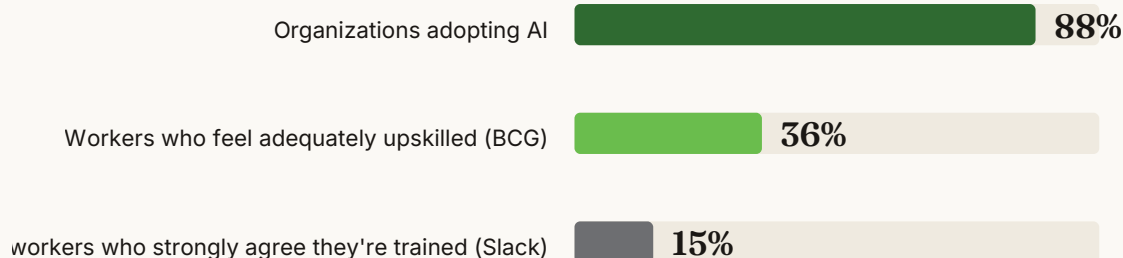


FIGURE 3 — The enablement deficit. Adoption has scaled; enablement has not. (Slack, 2024; BCG, 2026.)

Third, latent readiness exceeds what leaders perceive. McKinsey's *Superagency in the Workplace* found that **employees are markedly more ready for AI than their leaders assume** — executives estimated that a small fraction of employees use generative AI for a meaningful share of their work, when the true figure was about **three times higher** — and that **47% of C-suite leaders say their own organizations are developing and deploying AI too slowly** [McKinsey, 2025]. The bottleneck the data describes is not employee reluctance; it is organizational pace

and enablement. Notably, **just 1% of companies described themselves as “mature” in AI deployment** [McKinsey, 2025].

Fourth, workers route around the vacuum — and that is a signal, not just a risk. When an organization fails to enable AI use, motivated employees do not stop; they go underground. **54% of employees say they would use AI tools even if their employer did not authorize them** [BCG, 2025], and **78% of AI users report bringing their own tools to work** [Microsoft & LinkedIn, 2024]. “Shadow AI” is typically framed as a compliance threat. It is also the clearest possible demand signal: it marks exactly where employees have found value faster than their organization could sanction it. A measurement instrument that surfaces this layer turns a governance liability into an adoption map.

Underlying all four findings is the trust dimension. Globally, **only 46% of people say they are willing to trust AI systems**, even as two-thirds use them regularly [KPMG & University of Melbourne, 2025]. In the U.S. workforce, **52% are more worried than hopeful** about AI’s role at work [Pew, 2025]. Trust is not uniform, and it is not the same as usage — a person can use AI daily and distrust it deeply. That distinction turns out to matter enormously for what an organization should do, and it is invisible to any instrument that measures usage alone.

The composite picture is consistent across sources: a workforce that is more willing than its leaders believe, substantially unequipped, shallowly engaged, quietly working around the gaps, and unevenly trusting. None of that is measured by a maturity audit of the organization’s data strategy. All of it is the human layer — and all of it is actionable, *if it can be seen*.

4. The field is already converging on personas

The encouraging news is that the most serious researchers have independently arrived at the same conclusion — that the human layer must be measured, and that the right way to measure it is by **segmenting people, not scoring organizations**. The convergence is striking, and it is the strongest evidence that the persona approach is correct.

Slack’s Workforce Lab built five workplace-AI personas — **the Maximalist (30%), the Underground (20%), the Rebel (19%), the Superfan (16%), and the Observer (16%)** — explicitly along attitude and openness [Slack Workforce Lab, 2024]. BCG’s “AI Adoption Puzzle” defined five adoption personas — **AI Champion, Independent Explorer, Organizational Adopter, Passive Observer, and Cautious Skeptic** — along motivation and confidence [BCG, 2025]. McKinsey segmented workers by attitude into memorable categories; the World Economic Forum, in mid-2026, published “**The 5 Faces of Human Readiness for AI Adoption**” — AI Enthusiasts, Curious, Cautious, Sceptics, and Opposed — framed explicitly as *human* readiness [WEF, 2026]. Microsoft’s 2026 Work Trend Index introduced an “**Organization vs. Employee AI Readiness Index**,” placing each respondent on two axes: individual readiness and organizational readiness [Microsoft, 2026].

Strip away the branding and these models resolve to the same underlying structure: **attitude × skill/usage, with organizational enablement as the swing factor**. This is the convergence thesis stated empirically — when independent teams measuring the same phenomenon from different starting points arrive at the same axes, the axes are probably real.

But convergence on the *idea* has not produced a usable *standard*, and here the field's work falls short in four consistent ways:

1. **It is proprietary and one-off.** Each model is a single report, quiz, or segmentation tied to a vendor's marketing cycle — not a maintained, citable standard others can adopt.
2. **It is unvalidated.** Remarkably, across all of these models, none has published peer-reviewed validation of its persona taxonomy. The one psychometrically validated anchor the field can point to is the academic **Technology Readiness Index** (Parasuraman), which predates the generative-AI era. The category has built an edifice of archetypes on an unexamined foundation.
3. **It is enterprise-skewed.** The instruments are built for, and sold to, large enterprises. The small- and mid-sized organizations that employ much of the workforce are left with static templates.
4. **It is incomplete.** The taxonomies blur or omit important segments. Few cleanly separate the **Skeptic** (capable but distrustful) from the **Resistor** (avoidant and opposed) — yet these require opposite interventions. And several omit the **Eager Novice** (high willingness, low skill) — which our analysis suggests is the *highest-return* segment to enable.

The field, in short, has correctly identified the human layer and the persona method, and has not yet produced an open, validated, complete instrument for it. That is the gap the next section addresses.

5. A standard for the human layer

We propose to close the gap with a published instrument: the **Seven-Persona AI Readiness Standard** (described in full in the Institute's companion Standard document; summarized here).

The Standard places each employee on two empirically convergent axes — **attitude** (sentiment and trust toward AI) and **skill/usage** (capability and frequency) — modified by **enablement** (whether the organization equips them to succeed). From a short level-set of six signals — usage frequency, self-rated skill, sentiment, visibility, trust, and job-security concern — it derives seven personas:

- **Champion** — all-in, fluent, openly evangelizing; your force multipliers.
- **Quiet Power-User** — skilled and heavy, but hidden, often on unsanctioned tools; the shadow-AI signal made visible.
- **Eager Novice** — willing and ready to learn, but barely using AI yet; the highest-ROI enablement target.
- **Structured Adopter** — pragmatic; adopts when given tools, permission, and a path.

- **Cautious Observer** — aware and watching, but indifferent; a nudge moves them.
- **Skeptic** — capable but unconvinced; win them with proof and guardrails.
- **Resistor** — actively opposed; engage with respect, not pressure.

Two design choices distinguish the Standard from the prior art. First, it **measures attitude independently of usage**, which allows it to separate the Skeptic (a frequent, skilled user who distrusts AI) from the Resistor (an avoider) — a distinction the other models blur and one that determines whether a leader should respond with evidence or with reassurance. Second, it includes the **Eager Novice**, the willing-but-unskilled segment that the enablement evidence (Section 3) identifies as the fastest available win and that BCG's and McKinsey's taxonomies omit.

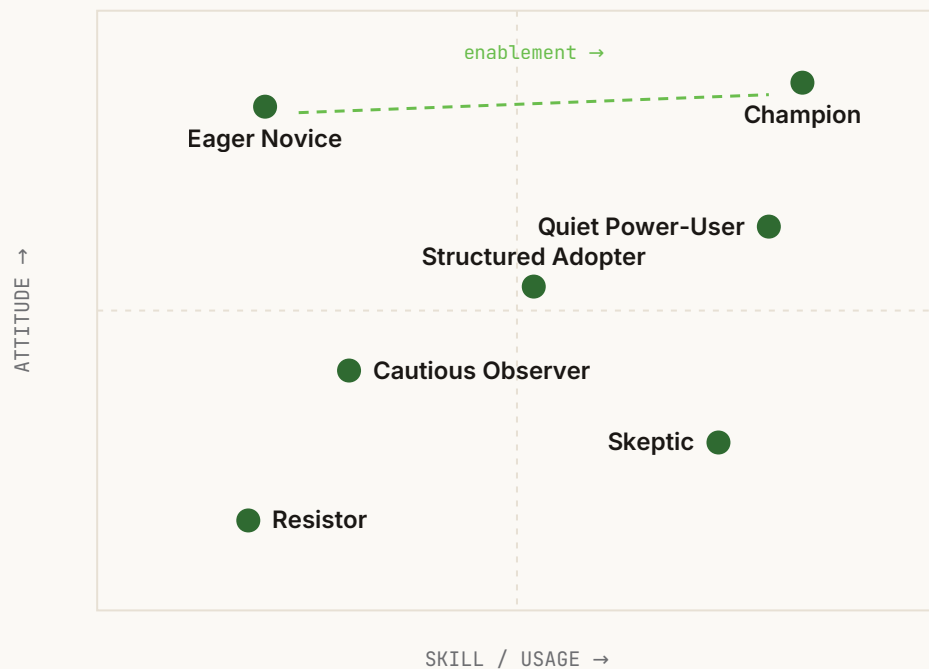


FIGURE 4 — The persona plane. Readiness is a position you can move, not a type you are.

The personas roll up, through five readiness dimensions — **Skill, Adoption/Usage, Trust & Attitude, Enablement Need, and Governance Awareness** — into a team and organizational picture: a persona mix, a dimension heatmap, and a set of named risk patterns. Crucially, the Standard is **published openly and versioned**, with its classification logic documented (the codebook), so that it can be cited, replicated, and adopted — the property that turns a proprietary segmentation into a standard.

A note on what the Standard deliberately is *not*. It is not a skills certification (it does not test prompt-engineering proficiency); that territory is well served by skills platforms. It is not an organizational-maturity audit; that is the middle altitude. It measures the layer those instruments omit, and it is designed to sit alongside them, not to replace them.

6. From measurement to impact

Measurement is only worth the action it enables. The reason to read the human layer by persona — rather than as a single readiness score — is that the persona mix tells a leader *what to do*, and the actions differ sharply by segment.

Consider the patterns the enablement evidence predicts most organizations will find:

- **An eager-but-untrained majority.** A large share of Eager Novices and Structured Adopters with low Skill scores. The evidence (15% adequately trained; lack of time the top barrier) suggests this is the modal organization. The intervention is not persuasion — these people are already willing — but **structured enablement**: time, training, role-specific starter use-cases. This is where training budget produces the fastest measurable adoption gain, and it is invisible to an instrument that reports only an aggregate score.
- **Hidden power-users.** An outsized Quiet Power-User share signals both latent value and governance exposure (recall: 54% would use AI unsanctioned; 78% bring their own tools). The intervention is **amnesty plus sanctioned tooling plus light governance** — converting shadow practice into shared, safe practice. Treating it purely as a compliance violation destroys the organization's fastest adoption path.
- **A forming skeptic or resistor cluster.** Concentrations of Skeptics or Resistors — often in functions where AI threatens established craft — call for **proof and security**, not mandates. Skeptics engaged well become an organization's quality conscience; Resistors met with respect and reskilling pathways can move, while Resistors met with pressure entrench.
- **A champion shortage.** Too few Champions to seed diffusion calls for **cultivating and empowering visible internal leaders** — the multiplier the diffusion-of-innovations literature predicts.

The throughline is that the impact gap (Section 1) closes not by buying more tooling but by acting on the human layer — and acting well requires seeing it. An organization that knows it has, say, “31% Eager Novices starved of training, 20% Quiet Power-Users on unsanctioned tools, and a Skeptic cluster in finance” has an action plan. An organization that knows only that its “AI readiness score is 63” has a number.

This also reframes readiness as **dynamic**. Personas are positions, not fates: the entire purpose of enablement is to move an Eager Novice toward Champion, and a well-run organization re-

measures to track that movement. Readiness is not a one-time audit but a managed trajectory — which is also why a point-in-time score, however precise, is the wrong instrument for it.

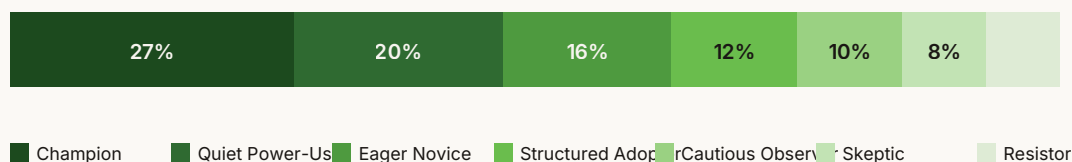


FIGURE 5 — Illustrative team persona mix. ILLUSTRATIVE • PRELIMINARY — to be populated with anonymized benchmark data.

7. Independence, openness, and validation

An instrument for the human layer is only as useful as it is trusted, and trust is exactly where the existing models are weakest. Every persona model in the market today is published by an organization that sells the thing it measures readiness for — a tooling vendor, a consultancy, a training company. That does not make their findings wrong, but it does mean the field has no neutral, validated reference standard. We believe it needs one, and that the standard must satisfy three conditions the proprietary models do not.

Independence. The Standard should be stewarded by a body whose conclusions are not subject to a commercial sponsor's approval, with its funding and relationships disclosed. The AI Readiness Institute is structured for this: it publishes openly, names its funders (including Readigence, the assessment platform that applies the Standard and contributes anonymized data), and maintains editorial independence from them. Disclosed relationships, transparently governed, are what make a standard citable by parties who compete with its sponsor.

Openness. A standard that cannot be freely cited and adopted is not a standard; it is a product. The Seven-Persona Standard is published with its methodology and classification logic documented and is intended for release under an open license, so that researchers can test it, employers can apply it, and even competitors can build on it. Adoption is the point.

Validation. This is the condition the category has conspicuously failed, and therefore the Institute's central commitment. We are establishing a **pre-registered validation program**: stating in advance what the personas should predict (adoption velocity, training uptake, risk behavior), testing the instrument against real workforce cohorts beginning with a design partner, grounding the construct in the validated lineage of the Technology Readiness Index and UTAUT, and

publishing the results — including limitations — openly, with an academic partner. A persona standard that is independently and psychometrically validated would be the first of its kind.

We are aware of the standing critique that persona models can be unfalsifiable and unstable over time. We accept it, and answer it by design: the Standard's personas are empirically derived positions on continuous, measured axes, re-benchmarked over time, explicitly built to track movement rather than to file people permanently. Turning that critique into a design principle is itself part of the rigor we are committing to.

8. Limitations and what's next

This is a preliminary paper, and intellectual honesty about its limits is part of the standard we intend to hold ourselves to.

It is, at this stage, a **synthesis** — an argument built on public evidence and an established theoretical lineage, not yet on the Institute's own benchmark data. The population estimates for the personas are indicative, mapped from adjacent studies (principally Slack's and BCG's), and will be replaced by the Institute's measured distribution. The validation of predictive accuracy described in Section 7 is **forthcoming, not complete**. And much of the underlying public evidence carries its own caveats, which we have tried to surface rather than hide: several key studies are vendor-commissioned; some headline figures (notably BCG's adoption-quality stages) rest on software-developer samples used as leading indicators; and adoption statistics must always be read with their unit of analysis (organization, knowledge worker, or individual) in view. We have led, wherever possible, with the most independent anchors — the World Economic Forum, Stanford HAI, Pew, Gallup, and the KPMG–University of Melbourne study.

What does not depend on forthcoming data is the core argument, which the weight of independent evidence strongly supports: **adoption has outrun impact; the gap lives in the human layer; that layer is systematically unmeasured; the field is converging on personas as the way to measure it; and no open, validated standard for it yet exists**. Those claims stand on the public record today.

What comes next is the work of closing the gap we have named: completing the validation, publishing the inaugural *State of AI Readiness* benchmark, and maintaining the Seven-Persona Standard as a living, openly-governed reference. We invite researchers, practitioners, and even competitors to cite it, apply it, and challenge it. A standard improves by being used and tested — and the measurement gap closes only if the instrument for the missing layer is one the whole field can trust.

References

Selected primary sources. Full, verified citations with URLs, sample sizes, and methodological notes are maintained in the Institute's evidence dossier (INSTITUTE-evidence-dossier-260630.md). All figures in this paper reflect that verified base.

1. McKinsey & Company (2025). *The State of AI in 2025: Agents, Innovation, and Transformation*.
2. McKinsey & Company (2025). *Superagency in the Workplace: Empowering People to Unlock AI's Full Potential*.
3. Boston Consulting Group (2025). *The AI Adoption Puzzle: Why Usage Is Up but Impact Is Not*.
4. Boston Consulting Group (2025). *AI at Work 2025: Momentum Builds, but Gaps Remain*.
5. Boston Consulting Group (2026). *AI at Work: Why Strategy Matters More Than Tools*.
6. Stanford HAI (2025). *Artificial Intelligence Index Report 2025*, Ch. 4 (Economy).
7. Microsoft & LinkedIn (2024). *Work Trend Index: AI at Work Is Here. Now Comes the Hard Part*.
8. Microsoft (2026). *Work Trend Index Annual Report*.
9. Slack Workforce Lab (2024). *The Workforce Index* (June; Fall) and the AI personas study.
10. Gallup (2025). *AI Use at Work Has Nearly Doubled in Two Years*; (2026) *Rising AI Adoption Spurs Workforce Changes*.
11. Pew Research Center (2025). *U.S. Workers Are More Worried Than Hopeful About Future AI Use in the Workplace*.
12. World Economic Forum (2025). *Future of Jobs Report 2025*; (2026) *The 5 Faces of Human Readiness for AI Adoption*.
13. DataCamp & YouGov (2026). *The State of Data & AI Literacy Report 2026*.
14. KPMG & University of Melbourne (2025). *Trust, Attitudes and Use of AI: A Global Study 2025*.
15. Study.com (2026). *State of AI Jobs and Skills Report 2026*.
16. ISACA (2026). *AI Pulse Poll*.
17. Parasuraman, A. *Technology Readiness Index (TRI)*.
18. Venkatesh, V. et al. *Unified Theory of Acceptance and Use of Technology (UTAUT)*.
19. Rogers, E. M. (2003). *Diffusion of Innovations* (5th ed.). Free Press.

Preliminary pre-release draft. The argument and public-source evidence are stable; the Institute's proprietary benchmark and the persona validation are forthcoming and are flagged as such. © 2026 The AI Readiness Institute. Prepared for limited circulation.