
THE STANDARD · v1.0

The Seven-Persona AI Readiness Standard

Version 1.0 — an open standard for measuring the human side of AI adoption.

PRELIMINARY · PRE-RELEASE · CC BY (INTENDED)

Publisher The AI Readiness Institute

Published June 30, 2026 · **Status** Preliminary — population figures indicative, validation forthcoming

Suggested citation: The AI Readiness Institute (2026). The Seven-Persona AI Readiness Standard, v1.0 (preliminary).

Abstract

The Seven-Persona AI Readiness Standard is an open framework for measuring how individual employees adopt artificial intelligence at work. It places each person on two empirically convergent axes — **attitude** (sentiment and trust toward AI) and **skill/usage** (capability and frequency of use) — modified by **enablement** (whether the organization equips the person to succeed). From these it derives seven personas, from Champion to Resistor, and five readiness dimensions that roll up from the individual to the team and organization. The Standard is designed for the small- and mid-sized organizations underserved by enterprise AI-maturity tools, and it is published openly so that researchers, employers, and practitioners can cite, apply, and challenge it. This version is **preliminary**: the conceptual model and persona definitions are stable and research-grounded; population estimates are indicative pending the Institute's own benchmark, and a pre-registered validation study is forthcoming.

1. Purpose & scope

What this is. A common, citable language for the *human* layer of AI adoption — the individual employee — and a method for aggregating it to a team and organizational picture.

What it is not. It is not a technical-skills test (it does not certify prompt-engineering proficiency), not an organizational-maturity audit (it does not score infrastructure, data, or strategy), and not a national-readiness index. Those instruments exist and are useful; this Standard measures the layer they omit — where adoption actually happens or stalls: the person.

Who it is for. HR, L&D, and Chief AI Officer leaders who must move a workforce from experimentation to impact; researchers studying technology adoption; and any organization — particularly in the **SMB and mid-market** — that needs an evidence-based read on its people without an enterprise budget.

Design principles. 1. **Open.** Publicly documented, versioned, free to cite and adopt. 2.

Grounded. Built on established adoption science and the convergent findings of the major

2024–2026 workforce studies. 3. **Dimensional, not typological.** Personas are *positions on continuous axes*, not fixed personality types — a person can move as enablement changes (see §8). 4. **Actionable.** Every persona resolves to a recommended intervention; the Standard exists to change what a leader does Monday morning.

2. The conceptual model

2.1 Two axes and a modifier

Every serious model of workplace AI adoption converges on the same structure: **attitude × skill/usage, with the organization's enablement as a modifier**. The Standard adopts this explicitly.

- **Attitude** — the person's sentiment toward and trust in AI: enthusiasm vs. indifference vs. opposition; confidence in AI's reliability; concern about its effect on their role.
- **Skill / Usage** — demonstrated capability and frequency of use: from non-use through occasional use to fluent, high-frequency, value-generating use.
- **Enablement (modifier)** — whether the organization gives the person the tools, permission, training, and norms to use AI well. Enablement does not change *who a person is*; it changes whether their latent readiness becomes real adoption. The same Eager Novice flourishes under good enablement and stalls under bad.

2.2 Theoretical lineage

The Standard inherits from established, validated work, which is both intellectually honest and a source of credibility:

- **Rogers, *Diffusion of Innovations*** — the adoption-curve backbone (innovators → laggards) that orders the personas from Champion to Resistor.
- **The Technology Readiness Index (TRI)** (Parasuraman) — the validated precedent for measuring an *individual's* propensity to adopt technology along attitudinal contributors (optimism, innovativeness) and inhibitors (discomfort, insecurity). The TRI is, to date, the only psychometrically validated anchor the AI-readiness field can point to; the Standard is designed to extend that lineage to AI specifically.
- **UTAUT** (Unified Theory of Acceptance and Use of Technology) — the constructs (performance expectancy, effort expectancy, social influence, facilitating conditions) that inform the enablement modifier.

2.3 Empirical convergence

The two-axis structure is not a theoretical preference; it is where the evidence independently lands. BCG's adoption personas, Slack Workforce Lab's workplace-AI personas, Microsoft's two-axis readiness index, and McKinsey's attitude segments all resolve to attitude × capability with organizational support as the swing factor (see §7). The Standard's contribution is to consolidate that convergence into one open, complete, citable taxonomy.

3. The five readiness dimensions

Personas describe *who* a person is; dimensions describe *how ready* they are, and they are what roll up to a team score. Each is scored 0–100 at the individual level and aggregated (above a minimum-respondent floor) to the team.

DIMENSION	WHAT IT MEASURES	LOW END → HIGH END
Skill	Demonstrated capability to use AI tools effectively	Cannot use AI productively → fluent, evaluates output critically
Adoption / Usage	Frequency and breadth of real use in the role	Never uses AI → uses it daily across core tasks
Trust & Attitude	Sentiment toward AI and confidence in its reliability	Hostile / distrustful → confident and constructive
Enablement Need	How much the person depends on the org to progress (<i>high = needs more support</i>)	Self-sufficient → blocked without tools/permission/training
Governance Awareness	Understanding of safe, appropriate, policy-compliant use	Unaware of risks/policy → consistently uses AI responsibly

Reading note. *Enablement Need* is inverted relative to the others: a *high* score is a flag, not a strength — it marks people whose readiness is gated by the organization. It is the dimension most directly actionable by leadership.

4. The seven personas

Each persona is a region on the attitude × skill plane. Definitions are stable in v1.0. **Indicative prevalence** is mapped from adjacent published research (primarily Slack Workforce Lab's 2024 persona study and BCG's adoption research) and is explicitly provisional — the Institute's own benchmark will replace these figures.

1. Champion

One-liner: All-in, fluent, and openly evangelizing AI to peers. **Attitude:** Very positive · **Skill/Usage:** High and visible **Defining signals:** Frequent daily use across tasks; high self-rated skill; high enthusiasm; openly shares their AI practice; high trust in (and realistic understanding of) AI. **Behavioral markers:** Builds workflows others copy; mentors peers; pushes the org to adopt; experiments with new tools first. **What they need:** Scope to lead — a visible champion role, advanced tooling, a channel to spread practice. Under-using Champions is wasted leverage.

Risk if ignored: Disengagement or attrition; the org loses its most effective internal multiplier and a natural change agent. **Leadership action:** Formalize them as enablers — a champion network, internal teaching, governance input. They are your fastest path to diffusion. **Indicative prevalence:** ~25–30% (cf. Slack “Maximalist” 30%). *Provisional.*

2. Quiet Power-User

One-liner: Skilled, heavy AI user who keeps it low-profile — often on unsanctioned tools. **Attitude:** Positive · **Skill/Usage:** High, but hidden **Defining signals:** High frequency and skill *paired with low visibility* — does not share usage; may use personal/unauthorized tools; the distinguishing axis from Champion is **visibility**, not capability. **Behavioral markers:** “Shadow AI” — real value generated off the books; reluctant to admit AI use for important work; routes around policy gaps. **What they need:** Permission and safety to surface their practice; clear, non-punitive policy; sanctioned equivalents of the tools they already use. **Risk if ignored:** This is the richest signal *and* the clearest governance exposure. Their hidden practice is the organization’s fastest latent adoption path — and its largest unmanaged-risk surface. Suppressing them destroys value; ignoring them invites data/compliance incidents. **Leadership action:** Treat shadow AI as a demand signal. Amnesty + sanctioned tools + light governance converts hidden value into shared, safe practice. **Indicative prevalence:** ~18–22% (cf. Slack “Underground” 20%; aligns with BCG’s ~54% who would use AI even if unauthorized). *Provisional.*

3. Eager Novice

One-liner: Enthusiastic and willing, but barely using AI yet. **Attitude:** Positive · **Skill/Usage:** Low **Defining signals:** High sentiment with low actual usage and low self-rated skill; wants to engage; blocked by lack of time, training, or access rather than by resistance. **Behavioral markers:** Asks how to start; tries tools sporadically; cheers AI but hasn’t integrated it. **What they need:** The highest-ROI enablement in the workforce — structured training, time, starter use-cases, and access. Willingness is already present; only capability is missing. **Risk if ignored:** Enthusiasm decays into the Cautious Observer’s indifference. The single biggest *missed* opportunity, because the activation cost is lowest. **Leadership action:** Prioritize them for enablement. This is where training budget produces the fastest measurable adoption gain. **Indicative prevalence:** ~15–18% (cf. Slack “Superfan” 16%). *Provisional.*

4. Structured Adopter

One-liner: Pragmatic; adopts when given tools, permission, and a clear path. **Attitude:** Neutral → positive · **Skill/Usage:** Moderate **Defining signals:** Moderate usage that tracks organizational support; not an early experimenter, not a resister; responds to structure, defaults, and norms. **Behavioral markers:** Uses sanctioned tools once rolled out; follows established workflows; adopts when it’s the path of least resistance. **What they need:** Defaults and scaffolding — integrated tools, role-specific playbooks, manager modeling, clear expectations. **Risk if ignored:** Stays at low-impact usage indefinitely; represents the bulk of the “stuck in the middle” majority where most adoption stalls. **Leadership action:** Make good AI use the default path. This is the

largest segment to *move*, and it moves on structure, not inspiration. **Indicative prevalence:** ~15–20% (the “early/late majority” core; interpolated). *Provisional*.

5. Cautious Observer

One-liner: Aware and watching, but indifferent and not yet engaged. **Attitude:** Indifferent · **Skill/Usage:** Low **Defining signals:** Low usage and low-to-neutral sentiment — not hostile, not enthusiastic; aware AI exists but sees no personal relevance yet. **Behavioral markers:** Waits to see if AI “sticks”; minimal experimentation; neither advocates nor opposes. **What they need:** Relevance and a nudge — concrete, role-specific demonstrations of value; social proof from peers; low-stakes first wins. **Risk if ignored:** Drifts toward Skeptic or Resistor if the organization’s rollout is clumsy, or simply remains inert ballast on the adoption curve. **Leadership action:** Show, don’t tell — peer demonstrations and relevance beat mandates. A small nudge moves them; pressure backfires. **Indicative prevalence:** ~15% (cf. Slack “Observer” 16%). *Provisional*.

6. Skeptic

One-liner: Capable but unconvinced — distrusts reliability, worries about deskilling. **Attitude:** Guarded / critical · **Skill/Usage:** Moderate → high **Defining signals:** The crucial distinction from a Resistor — a Skeptic can be a *frequent, skilled* user who nonetheless distrusts AI. Attitude must be probed **separately** from usage. Concern centers on reliability, accuracy, ethics, and erosion of craft. **Behavioral markers:** Uses AI but double-checks everything; raises failure modes; resists over-reliance; often a thoughtful critic rather than an opponent. **What they need:** Proof and guardrails — evidence of reliability, transparency about limits, quality controls, and a voice in governance. Skeptics are frequently *right*, and make excellent governance partners. **Risk if ignored:** Dismissed as laggards, they become credible internal opponents; engaged well, they become the organization’s quality conscience. **Leadership action:** Win them with evidence, not enthusiasm. Convert their scrutiny into a governance asset. **Indicative prevalence:** ~10–15% (split from Slack “Rebel” 19% and the broader skeptical cohort; interpolated). *Provisional*.

7. Resistor

One-liner: Actively opposed — sees AI as a threat, and as unfair when peers use it. **Attitude:** Strongly negative · **Skill/Usage:** Low / avoidant **Defining signals:** Low usage paired with strong negative sentiment and, often, job-security concern; may view colleagues’ AI use as cheating or as a personal threat. **Behavioral markers:** Avoids AI; objects to its adoption; can dampen team norms; concern is frequently rooted in fear (of displacement, of deskilling, of unfairness). **What they need:** Respect, security, and agency — honest communication about role impact, reskilling pathways, and a stake in how AI is introduced. The need is psychological safety first, capability second. **Risk if ignored:** Becomes a vocal blocker and a drag on team adoption; left unaddressed, fear hardens into entrenched opposition. **Leadership action:** Engage with respect, never pressure. Address the fear directly; a Resistor met with security and agency can move —

one met with mandates digs in. **Indicative prevalence:** ~12–19% (cf. Slack “Rebel” 19%, inclusive of some Skeptics). *Provisional.*

A core principle: even the Resistor is treated with respect. Most AI-readiness messaging traffics in fear; this Standard does the opposite — it meets people where they are. That stance is both ethically correct and operationally superior: fear suppresses the honest answers that make the measurement worth anything.

5. The level-set mechanic (codebook)

Personas are assigned from a short level-set of six signals. Publishing the logic is what makes this a *standard* rather than a black-box quiz; it is intended to be replicable.

The six level-set signals: 1. **Usage frequency** — how often the person uses AI for work (never → daily). 2. **Self-rated skill** — confidence in their ability to use AI effectively. 3. **Sentiment** — overall positive/neutral/negative feeling toward AI. 4. **Visibility** — whether they share their AI use openly or keep it hidden. 5. **Trust / reliability** — confidence in AI’s accuracy and appropriateness. 6. **Job-security concern** — degree of worry about AI’s effect on their role.

Classification logic (the codebook): - **Frequency × Sentiment** separates the high-usage personas: Champion, Quiet Power-User, and Skeptic can *all* be frequent users — sentiment and the signals below distinguish them. - **Visibility** splits **Champion** (shares openly) from **Quiet Power-User** (hides usage). This is the defining cut between the two. - **Sentiment × low Usage** separates the low-usage personas: positive → **Eager Novice**; indifferent → **Cautious Observer**; negative → **Resistor**. - **Skill × Trust** isolates the **Skeptic**: capable and frequent but low-trust — which is why attitude must be measured independently of usage. - **Enablement signals** (access, training, time, permission) and **job-security concern** refine **Structured Adopter** vs. **Eager Novice** vs. **Resistor**, and set the *Enablement Need* dimension.

Implementation note. In the Readigence application, these thresholds and weights live in configuration (not hard-coded), so the classification can be recalibrated against real response data without changing the published Standard’s logic. The Standard defines *what* is measured and *how personas are derived*; calibration tunes the boundaries.

6. Reading a team: persona mix → action

The Standard's payoff is at the team level, where the **persona mix** reveals patterns a single score hides. Common, named patterns:

- **The eager-but-untrained majority** — a large Eager Novice + Structured Adopter share with low Skill scores. *Highest-ROI intervention: structured enablement*. The org has willingness it isn't converting.
- **Hidden power-users** — an outsized Quiet Power-User share. *Latent value + governance exposure*. Surface and sanction.
- **A forming skeptic/resistor cluster** — concentration of Skeptics or Resistors in a function (often where AI threatens established craft). *Engage with proof and security before it hardens*.
- **Champion-thin** — too few Champions to seed diffusion. *Cultivate and empower visible leaders*.

Team and organizational results are only ever shown **above a minimum-respondent floor** (default N=5), so no individual is identifiable from an aggregate — a requirement, not an option (see the Institute's data covenant).

7. Relationship to other models

The Standard is built on, and credits, the best prior work — and is the more complete superset.

INSTITUTE (7)	BCG (5)	SLACK (5)	WEF "5 FACES" (5)	MCKINSEY (4)
Champion	AI Champion	The Maximalist	AI Enthusiast	(Bloomer/ Zoomer)
Quiet Power-User	Independent Explorer	The Underground	—	—
Eager Novice	—	The Superfan	AI Curious	—
Structured Adopter	Organizational Adopter	—	—	—
Cautious Observer	Passive Observer	The Observer	AI Cautious	(Gloomer)
Skeptic	Cautious Skeptic	(part of Rebel)	AI Sceptic	(Doomer)
Resistor	—	The Rebel	AI Opposed	(Doomer)

INSTITUTE (7)	BCG (5)	SLACK (5)	WEF "5 FACES" (5)	MCKINSEY (4)
Axes	attitude x confidence	attitude (usage- flavored)	attitude only	attitude only
<i>Open standard?</i>	No (report)	No (quiz)	No (article)	No (segments)

What the Standard adds: the **Eager Novice** (high willingness / low skill — the highest-ROI segment, absent from BCG and McKinsey) and a clean separation of **Skeptic** (capable, distrustful) from **Resistor** (avoidant, opposed) that the others blur. It grounds all seven explicitly in **attitude x skill x enablement** — not attitude alone — and, uniquely, publishes as an **open, versioned standard** rather than a one-off report, quiz, or proprietary segmentation.

Sources: BCG, "The AI Adoption Puzzle" (2025); Slack Workforce Lab (2024); World Economic Forum, "The 5 Faces of Human Readiness for AI Adoption" (2026); McKinsey, "Superagency in the Workplace" (2025); Microsoft, 2026 Work Trend Index. Full citations in the Institute's evidence dossier.

8. Validation status & roadmap

Status: preliminary. The conceptual model, the seven personas, and the five dimensions are stable and research-grounded. What remains to be established empirically, and is openly flagged as such:

1. **Population estimates** are indicative (mapped from adjacent studies), pending the Institute's own benchmark.
2. **Classification thresholds** will be calibrated against real response data.
3. **Predictive validity** — whether persona membership predicts useful outcomes (adoption velocity, training uptake, risk behavior) — is the subject of a **pre-registered validation study**, beginning with a design-partner cohort and conducted with an academic partner. Findings, including limitations, will be published openly.

On the standing critique of personas. Critics rightly note that persona models can be unfalsifiable and temporally unstable. The Standard answers this by design: its personas are **empirically derived positions on continuous, measured axes**, not fixed personality types, and are **re-benchmarked over time** — a person can and should move (an Eager Novice becomes a Champion as enablement lands). Readiness is a dynamic state, and the Standard is built to track movement, not to file people permanently.

9. Versioning & changelog

The Standard uses semantic versioning. Material changes to persona definitions, axes, or the codebook increment the major version; calibration and clarifications increment the minor version. All versions remain citable.

- **v1.0 (2026-06-30) — preliminary pre-release.** Initial publication of the conceptual model, seven personas, five dimensions, and level-set codebook. Population estimates indicative; validation forthcoming.

References

Selected primary sources (full, verified citations with URLs in INSTITUTE-evidence-dossier-260630.md):

- BCG (2025). *The AI Adoption Puzzle: Why Usage Is Up but Impact Is Not*.
- BCG (2025). *AI at Work 2025: Momentum Builds, but Gaps Remain*.
- McKinsey (2025). *Superagency in the Workplace; The State of AI in 2025*.
- Microsoft (2026). *Work Trend Index Annual Report*.
- Slack Workforce Lab (2024). *The Workforce Index; AI personas study*.
- World Economic Forum (2026). *The 5 Faces of Human Readiness for AI Adoption; (2025) Future of Jobs Report*.
- Parasuraman, A. *Technology Readiness Index (TRI)*.
- Rogers, E. M. (2003). *Diffusion of Innovations* (5th ed.).
- Venkatesh et al. *UTAUT*.

This is a preliminary pre-release of an open standard. Definitions are research-grounded; population figures and validation results are provisional and will be updated. © 2026 The AI Readiness Institute. Intended for release under CC BY 4.0.