

You Prepare the People. We Build the *Foundation*.

ProjectBits Consulting — for change-management professionals · 2 pages + fit checklist

PAGE 1 — TWO DISCIPLINES, ONE READINESS

The lane, stated plainly: we don't run your change program, facilitate your workshops, or own the people side of an AI adoption. Readiness, communication, stakeholder alignment, and adoption are your seat — the seat most AI projects skip. What we do is build and govern the thing your readiness work assumes is already there: a **financial and data foundation** an AI can actually reason over — governed books, integrated operational data, and a policy layer that makes its outputs trustworthy and defensible.

You coach the organization to be ready. We make sure there's something real to be ready **for**. The integration point is your engagement: everything our system produces gives your readiness assessment a baseline to score against and your adoption plan a system that won't embarrass it at go-live.

The pattern you've named a hundred times

You know this client. Leadership wants AI; a vendor demoed something impressive; a tool got bought before the problem was defined, the problem got defined before the return category was named, and nobody asked whether the organization could sustain the governance the tool required. The deployment underperforms, and the post-mortem blames adoption resistance — **your** domain — or pins it on the operator with a "you don't understand AI well enough" label.

That framing is wrong, and it costs you credibility you didn't earn the blame for. The failure is rarely the people and rarely the technology — it's a sequence of decisions made in the wrong order, on top of books and data the AI was never able to read. Our position paper, [The Third Perspective: People, Preparation, and Readiness](#), makes the full argument and names **your** discipline as the third perspective the two technical ones can't deliver without. It's written to be forwarded to a client about to make exactly this mistake.

How the foundation gets built

AI projects fail for the same reasons ERP, CRM, and RPA projects failed — undocumented processes, missing baselines, removed-but-not-replaced human review. Our pipeline is the structural fix, automated where practical and professionally reviewed where it isn't:

QuickBooks Online, connected and governed. Reads are scoped and logged; writes pass a policy gate built on the Open Policy Agent (OPA) engine with rules written as code in its Rego language — deny-by-default, human approval for anything material, every decision logged with the rule that fired. That governed path is the harness the AI runs inside.

Augmented in our PostgreSQL layer. Our analytics database adds the dimensions QBO can't carry — cost behavior, per-client and per-job profitability, forecast models — the integrated, entity-resolved data that lets an AI answer a real question instead of untangling the business's history.

Governed by a declared harness position. This is where your work and ours meet. Every AI use case gets a declared human-oversight position — **Above, In, On, or Under the Loop** — documented and authorized before deployment. That declaration is at once a governance control **and** a change intervention: it names who reviews, who monitors, who can stop the system, and what each role must be trained for. It's the artifact your adoption plan has been missing.

The harness framework, in one table

HARNESS POSITION	THE HUMAN'S JOB	WHAT YOUR CHANGE PLAN MUST PREPARE
Above the Loop	Sets policy, reviews aggregate outcomes	Governor literacy — reading performance signals, knowing what triggers policy review
In the Loop	Reviews each AI output before action	Funded review time, clear criteria, workflow discipline
On the Loop	Monitors a running process, intervenes	Anomaly interpretation, intervention authority
Under the Loop	AI acts within approved policy, no per-decision review	Demonstrated maturity at every prior position first

An undeclared harness position defaults to **Under the Loop by omission** — the highest-risk position, arrived at by accident. The single most common AI governance failure is a workflow that says "human review" on paper while providing no time, no criteria, and no measurement for it. That gap is exactly what your readiness assessment is built to catch — once there's a system concrete enough to assess.

PAGE 2 — THE READINESS INSIGHT, AND HOW A REFERRAL WORKS

Readiness is scoreable — and most AI projects never score it

Here's the piece of our methodology built for your readiness conversations. An organization's preparedness to deploy a **specific** AI system at a **specific** harness position targeting a **specific** return category can be scored across five dimensions — and if the client can't score themselves before deployment, they aren't ready:

- **Leadership alignment** — documented agreement on AI tier, harness position, return category, and control ownership.
- **Workforce capability** — role-specific skills for the declared harness position (reviewer, monitor, or governor — they're different curricula).
- **Process documentation maturity** — processes understood and documented at the standard the harness position requires.
- **Governance maturity** — policies explicit, approved, and enforced; control documentation current.
- **Measurement infrastructure** — baselines established and attribution designed **before** a single user touches the system.

When readiness is low, it's often not a discipline problem you can coach away — for two of those dimensions it's a structural absence in the books and data: the baseline doesn't exist because the books can't produce it. We build the foundation; you run the readiness program on top of it. **The assessment itself is the first change intervention, and the score is the receipt that makes the eventual return claim credible.**

What a referral looks like — and why it makes your engagement stickier

1. **You spot the signal:** the adoption with no defined return category, the "human review" step nobody funded, the owner who can't answer "how do you know it's working?"
2. **You hand off with one sentence:** "Their AI readiness depends on a financial and data foundation they don't have yet — ProjectBits builds it and declares the harness positions your adoption plan needs." (Or forward [The Third Perspective](#) — Section 7 names your seat.)
3. **We start with the Tax Ready Assessment** — fixed scope, plain-language readout, honest timeline. Low commitment for your client; bounded exposure for your name.
4. **Your engagement gets stronger:** the readiness score has a real system to measure, the harness positions become documented controls instead of hopes, the return materializes because it was baselined, and the owner credits the program **you** led. Adoptions that work renew the coach who led them. That's the loop.

We reciprocate deliberately: our clients who get their foundation to decision-grade and then ask "now how do we get the **organization** to actually adopt this?" are precisely your prospects.

Mechanics and boundaries, in writing

- Engagement letters name the lane: financial function, data foundation, and AI governance architecture; change-program facilitation, communication, and adoption coaching out of scope.
- With client permission, you get the monthly narrative the leadership team sees, plus a standing 90-day check-in on any referred client.
- We serve businesses roughly \$750K–\$10M (professional services and MSPs, trades, financial services and real estate especially) in Northern Virginia and remote.

Start with one

Think of the client whose AI adoption you watched stall for reasons that weren't really about the people. Send them the paper, or call us first and pressure-test the fit — we'll tell you honestly if it isn't one.

THE FIT CHECKLIST (FORWARDABLE — "IS THIS CLIENT A REFERRAL?")

A client likely needs the foundation built before the adoption can succeed if **three or more** apply:

- ☐ An AI tool was **bought before the problem was defined** or the return category named
- ☐ A workflow says **"human review"** but no one funded the time, criteria, or measurement for it
- ☐ No **baseline** was established before the AI went live, so the return can't be attributed
- ☐ The books **can't produce** the number the AI is supposed to improve — it was never recorded
- ☐ No one can name the **harness position** (Above/In/On/Under the Loop) for the AI in production
- ☐ The owner can't answer **"how do you know this is working?"** with anything but an assertion
- ☐ Revenue roughly **\$750K–\$10M**

[Talk to ProjectBits →](#)