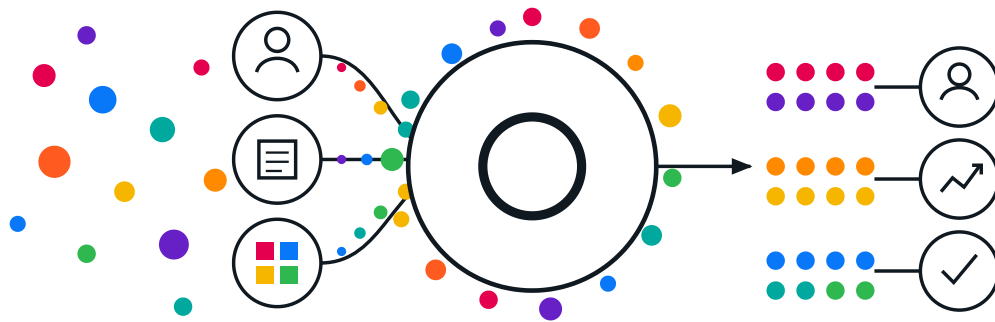


RUDI

RESPONSIBLE USE OF DIGITAL INTELLIGENCE

THE WORKPLACE AI ENABLEMENT PLAYBOOK



A human-centered framework for building literacy, redesigning work, evaluating AI investments, and developing durable organizational capacity.

RUDI

RESPONSIBLE USE OF DIGITAL INTELLIGENCE

THE WORKPLACE AI ENABLEMENT PLAYBOOK

A human-centered operating framework for moving from AI experimentation to durable organizational capacity.

START HERE — GET ORIENTED BEFORE YOU ORGANIZE

Before you begin

Most organizations do not begin AI enablement with a clean slate.

Employees may already be experimenting with AI. Leaders may be receiving vendor proposals. Different departments may be moving at different speeds. Policies, training, approved tools, and clear ownership may still be incomplete — or may not exist at all.

That is a normal place to begin. You do not need to be an AI expert, have a platform selected, secure a large budget, or arrive with a complete AI strategy. You need a willingness to understand what is already happening, examine how work is actually performed, and make deliberate decisions about what should happen next.

This playbook helps an organization move from scattered experimentation to responsible, repeatable capacity. It provides a sequence, but it is not meant to be absorbed all at once. Each phase produces a practical artifact and prepares you for the next decision.

— WHO SHOULD USE THIS PLAYBOOK?

This work usually begins with one person, but it cannot remain one person's responsibility. Your authority and starting point may differ. The method still begins with the same questions: Where are we now? What is already happening? Who needs to participate? What decision must we make next?

EXECUTIVE SPONSORS Read the Executive Quick Start and make the five leadership decisions described there.	PROGRAM AND ENABLEMENT LEADERS Begin by establishing ownership, participation, and decision rights.
DEPARTMENT LEADERS Begin by documenting the work and identifying workflows worth investigating.	INTERNAL CHAMPIONS Begin by creating a fact-based leadership conversation. Do not launch an unofficial program on your own.

— YOU DO NOT NEED TO DO EVERYTHING AT ONCE

The complete method contains ten phases, four decision gates, and fifteen working tools. Those pieces describe the full organizational journey. They are not a checklist you must finish before taking action.

THE SMALLEST WORKABLE COUNCIL COVERS FIVE FUNCTIONS

Executive authority; one accountable enablement lead; IT and security; frontline workflow knowledge; and risk, legal, compliance, or privacy review. These are functions, not necessarily five different people. In a smaller organization, one person may cover more than one function and specialist review may be external or on call. The responsibility cannot be omitted simply because there is no internal legal or compliance department.

— YOUR FIRST MOVE

Schedule a 60-minute orientation conversation with the smallest group capable of moving the work forward. Include, where possible, an executive sponsor, a prospective enablement lead, a people or HR representative, a technology or security representative, someone close to the work, and a risk, legal, compliance, or privacy representative when sensitive or regulated work is involved.

FIRST MEETING Four questions, one next decision

ASK

1. Where is AI already being used — formally or informally?

2. Where does the organization believe it is on the enablement journey?

3. Who will own the next 30 days of work?

4. What leadership decision is needed before proceeding?

PRODUCE

An honest current-position statement, a named executive sponsor, a named owner for the next step, and a date for the next working session.

The goal of your first week is not to launch an AI initiative. It is to establish enough clarity and ownership to make the next responsible decision.



CONTENTS — ●●●●●●●

—	Before you begin	
—	The RUDI Workplace AI Enablement Method	
—	How to use this playbook	
—	Executive Quick Start	
I	Understand AI enablement	§1–2
II	Organize leadership and accountability	§3–5 · GATE 1
III	Establish the baseline	§6–7
IV	Create shared understanding	§8 · GATE 2
V	Understand departmental work	§9–10
VI	Redesign work at the task level	§11–12
VII	Apply human-centered principles	§13
VIII	Prioritize and implement	§14–16 · GATE 3
IX	Measure and institutionalize	§17–18 · GATE 4
—	The complete methodology at a glance	



THE SYSTEM — ONE METHOD, FIVE LAYERS

The RUDI Workplace AI Enablement Method

The method is one system with five layers. Each layer answers a different leadership question, and none of them works alone.

LAYER	WHAT IT IS	LEADERSHIP QUESTION IT ANSWERS
Maturity progression	Literacy → Enablement → Adoption → Capacity → Transformation	Where are we taking the organization?
Operating methodology	Organize → Assess → Orient → Diagnose → Decompose → Classify → Evaluate → Prioritize → Implement → Institutionalize	What do we do, in what order?
Task decisions	Automate → Augment → Avoid → Add	How should each piece of work change?
Human-centered lens	RESPECT	How do we judge whether a change is responsible?
Governance structure	Executive sponsor + AI Enablement Council + implementation teams + champions network	Who decides, who builds, who spreads it?

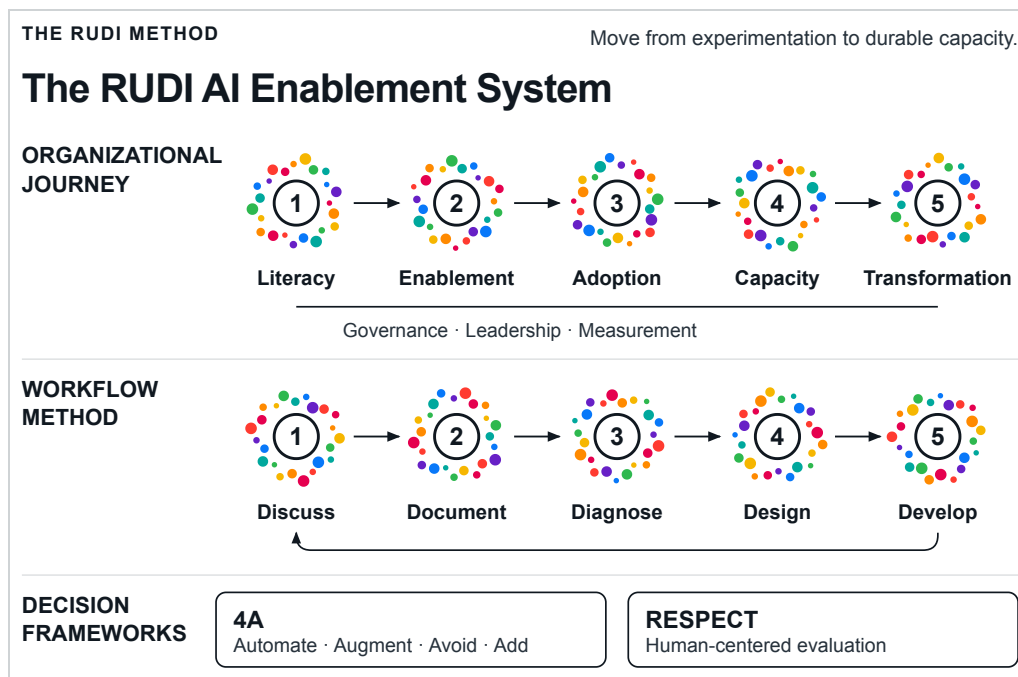


FIGURE 1 The RUDI AI Enablement System — the organizational journey, the workflow method, and the decision frameworks, held together by governance, leadership, and measurement.

The operating methodology describes what the organization does. The maturity progression describes what changes as a result. RESPECT describes how every decision along the way is evaluated. Governance is who holds the pen.

WORKING DOCUMENT — NOT A READING ASSIGNMENT

How to use this playbook

This is a working document, not a reading assignment. Every numbered section ends with a **Field Work** block in a consistent format:

FIELD WORK The format	
DO	The exercise to run before moving on.
PRODUCE	The named artifact this section generates — each maps to a toolkit appendix.
YOU SHOULD NOW HAVE	A checklist for whether the section is actually done.
GATE	Where it appears, the explicit leadership approval required before proceeding.

If you complete every Field Work block in order, you finish the playbook holding a complete, signed, evidence-backed AI enablement program: charter, baseline, opportunity maps, task disposition maps, RESPECT assessments, pilot charters, and a measurement dashboard.

● **MERIDIAN — THE RUNNING EXAMPLE**

One running example threads through the entire method: **Meridian Community Services**, a composite 140-person nonprofit that administers early-childhood, family-support, and workforce programs across a three-county region. Meridian is fictional, but every detail is drawn from patterns we see repeatedly in real organizations of this size. Watch for the **Meridian** callouts to see each abstract step become concrete.



TWO PAGES — FOR THE LEADER WHO DECIDES

Executive Quick Start

Two pages for the leader who has to decide whether this happens, before deciding how.

— **THE THESIS**

Your organization is already using AI. The only question is whether that use is designed or accidental. Right now, employees are experimenting with unapproved tools because approved paths don't exist, departments are buying subscriptions IT has never reviewed, and nobody can tell you what any of it produces. The fix is not a policy memo or an enterprise license. It is an operating system: named ownership, an honest baseline, shared language, task-level work redesign, human-centered evaluation criteria, disciplined pilots, and measurement against baselines. This playbook is that operating system.

The core reframe: stop asking which jobs AI can replace. Work is a collection of tasks, and the productive question is which tasks to automate, which to augment, which to deliberately avoid, and what new work to add. Analyze at the task level. Redesign at the job level.

— **THE SEQUENCE AND WHAT EACH STEP PRODUCES**

#	PHASE	OUTPUT
1	Organize	Signed council charter; participation structure; interim acceptable-use guidance
2	Assess	Organizational baseline report, segmented by department and role
3	Orient	Shared-language training delivered; workforce ready for workflow conversations
4	Diagnose	Departmental opportunity maps and workflow inventory (5D interviews)
5	Decompose	Task decomposition worksheets for priority workflows

#	PHASE	OUTPUT
6	Classify	Task disposition map — every task marked Automate / Augment / Avoid / Add
7	Evaluate	RESPECT assessment per use case, each ending in go / revise / pilot / reject
8	Prioritize	Ranked use-case portfolio; signed pilot charters
9	Implement	Running pilots with baselines, human-review plans, and stop conditions
10	Institutionalize	Adoption dashboard; decision register; scale / stop decisions; capacity assessment

— THE FOUR DECISION GATES

Leadership approval is required at four points. Everything between gates is delegated.



GATE 1 — CHARTER · END OF ORGANIZE

Sponsor approves the council charter, the named enablement lead, a budget envelope, and interim acceptable-use guidance.



GATE 2 — BASELINE · END OF ORIENT

Council accepts the baseline findings and approves the training plan and priority departments for interviews.



GATE 3 — PORTFOLIO · END OF PRIORITIZE

Sponsor approves the pilot portfolio and budget; each pilot charter is signed by its sponsor and workflow owner.



GATE 4 — SCALE · END OF EACH PILOT REVIEW

Council decides, in writing, whether each pilot scales, changes, or stops.

— THE FIRST 90 DAYS

- **Weeks 1–2:** Name the executive sponsor and enablement lead. Draft the charter. Seat the council. **Gate 1.**
- **Weeks 3–6:** Inventory current and shadow AI use. Issue interim acceptable-use guidance. Launch the baseline survey.
- **Weeks 5–8:** Deliver foundational training. Analyze the baseline; build department profiles and learning cohorts. **Gate 2.**
- **Weeks 7–10:** Run 5D interviews in two or three priority departments. Build the workflow inventory.
- **Weeks 9–12:** Decompose and classify the strongest candidate workflows. Run RESPECT. Select two or three pilots and sign charters. **Gate 3.**
- **Day 90:** Pilots launch with documented baselines. First council review scheduled.

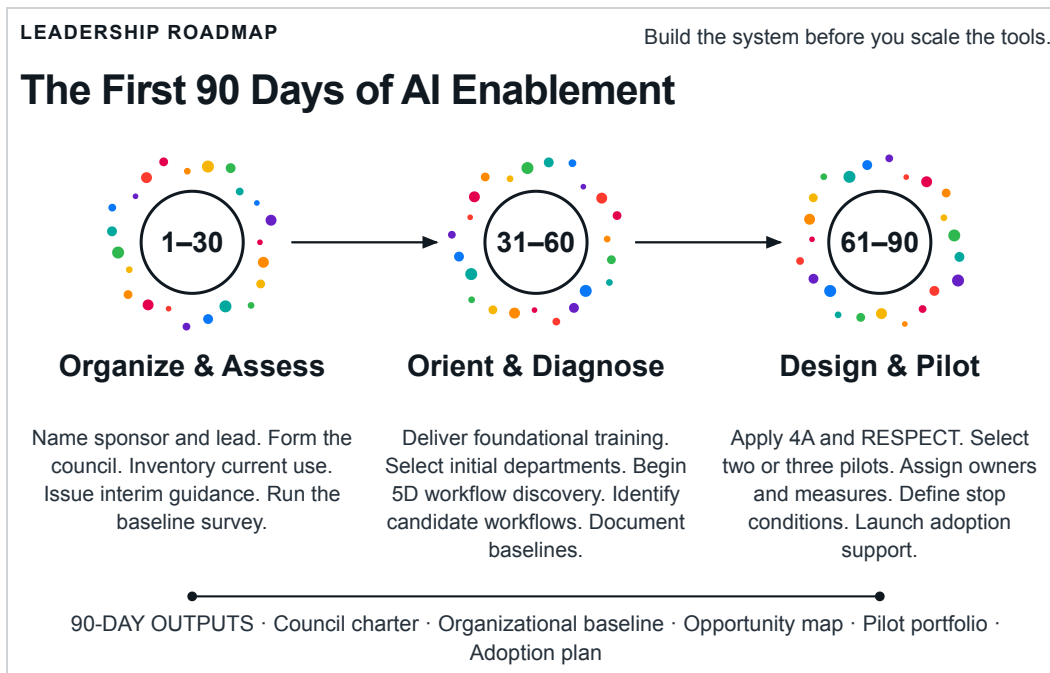
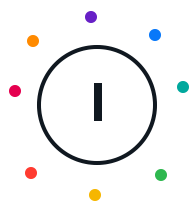


FIGURE 2 The first 90 days — build the system before you scale the tools.

— THE DECISIONS ONLY YOU CAN MAKE

An executive reading this owes the program five decisions, and only five: who sponsors it, who runs it, what the budget envelope is, what the interim rules are, and which pilots get funded. Everything else in this playbook is designed to be delegated — and to give you the evidence to make those five decisions well.



PART I — SECTIONS 1-2

Understand AI enablement

Shared definitions and the task lens — the two ideas every later step depends on.

1 Define the terms

Leadership conversations about AI stall when the same words mean different things to different people in the room. Establish these distinctions early and use them consistently:

AI literacy is understanding: what the technology does, where it fails, what responsible use requires. Literacy is measured by what people know and how they judge situations, not by whether they use tools.

AI enablement is equipping: approved tools, data access, role-relevant training, policies people can actually follow, and somewhere to go with questions. An organization can have high literacy and poor enablement — people who understand AI but have no sanctioned way to use it. That gap is where shadow AI grows.

AI adoption is behavior change: people using AI in real work, repeatedly, in ways that show up in how the work gets done. Adoption is measured in workflows, not license counts.

AI capacity is durability: the organization can identify opportunities, evaluate tools, redesign workflows, and support employees on its own, repeatedly, after the consultants leave.

AI transformation is organizational change: new services become possible, roles are redesigned around new task mixes, and strategy accounts for capabilities that did not exist before.

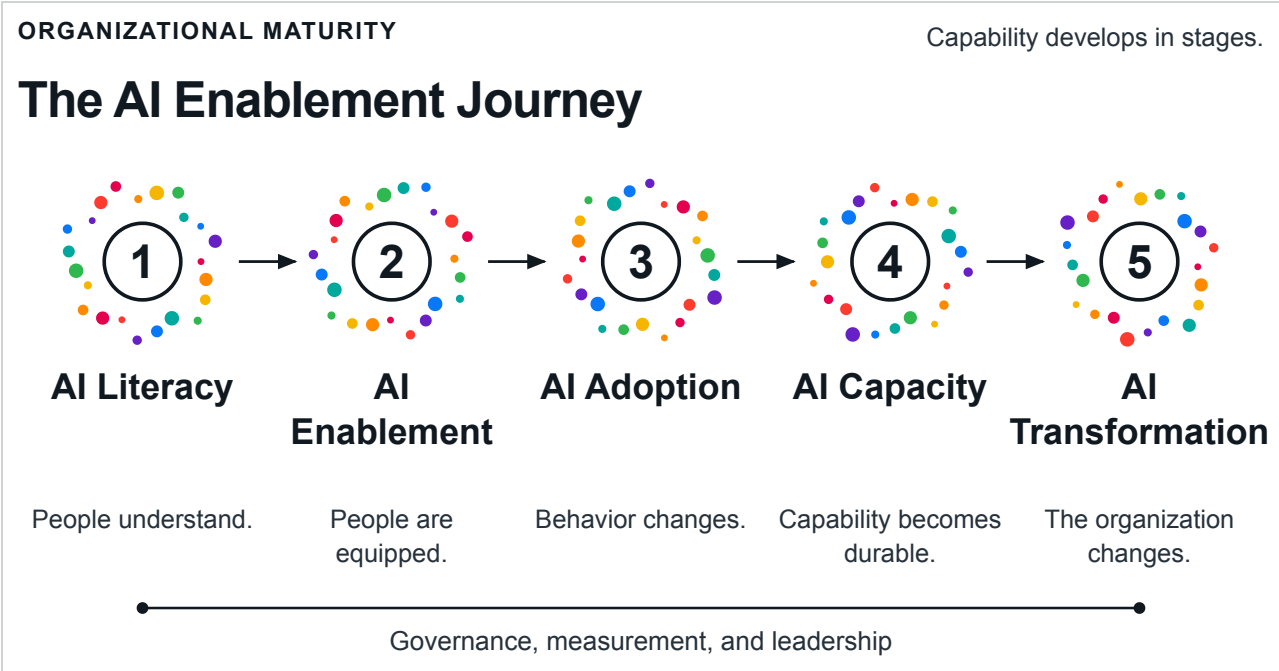


FIGURE 3 The AI Enablement Journey — capability develops in stages, and governance runs beneath every one of them.

Skipping stages is the most common failure pattern. Training without enablement produces awareness with no outlet. Tool purchases without literacy produce shelfware. Pilots without governance produce risk.

FIELD WORK Definitions alignment	
DO	In your next leadership meeting, ask each leader to write down, privately, where they believe the organization sits on the five-stage progression and why. Compare answers aloud.
PRODUCE	A one-paragraph statement of current stage, with the evidence for it, adopted by the leadership team.
HAVE	☐ Shared definitions in use ☐ An honest, agreed answer to "where are we now" ☐ Disagreements surfaced rather than papered over

Shift from occupations to tasks

Work is a collection of tasks. A job title bundles those tasks together with things a task list does not capture: accountability for outcomes, relationships with colleagues and clients, institutional knowledge, and judgment built over years. Both halves of that sentence matter.

The task view is what makes AI strategy tractable. "Can AI do a case manager's job" is unanswerable and threatening. "Which of the forty tasks in this case management workflow involve summarizing documents, and which involve earning a family's trust" is answerable and useful. When you decompose work into tasks, decisions, handoffs, and human interactions, most workflows turn out to contain some tasks where AI helps enormously, some where it helps a little, and some where it does not belong.

The human-elements view is what keeps the redesign honest. If task analysis strips out accountability and relationships, the redesigned workflow will look efficient on paper and fail in practice, because the parts that made it work were never on the task list.

Analyze at the task level. Redesign at the job level.

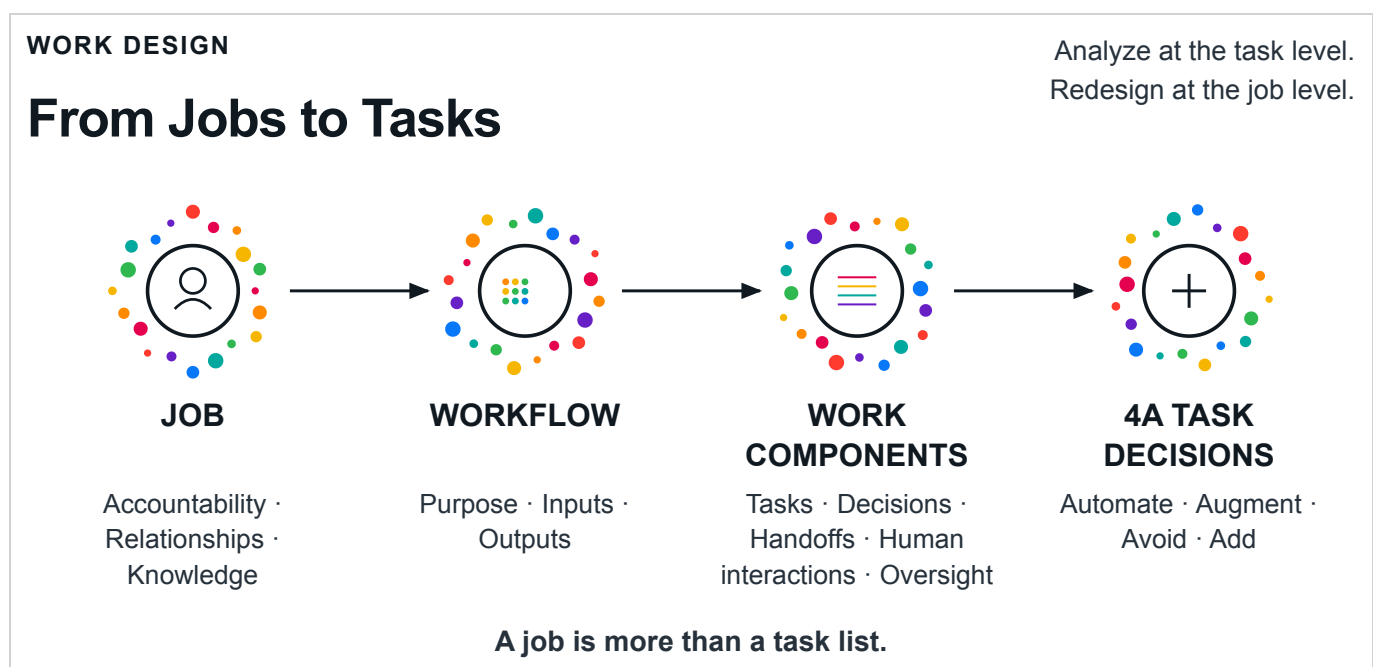


FIGURE 4 From jobs to tasks — a job is more than a task list, and the 4A decisions happen at the component level.

• MERIDIAN

Meridian's executive director had been asked by her board whether AI would "replace the family advocates." She reframed the question for the board using the task lens: a family advocate's month contains roughly sixty distinct tasks, from drafting case notes to sitting with a parent during an eligibility appeal. The board discussion changed immediately — from an unanswerable question about a job to a productive one about which of the sixty tasks deserved help.

DO

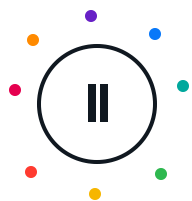
Pick one role in your organization that people worry about. List fifteen to twenty tasks that actually make up that person's week. Mark which involve judgment, relationships, or accountability.

PRODUCE

A one-page task list for one real role — your first proof that the task lens works here.

HAVE

- ☐ A concrete task list
- ☐ At least three tasks nobody would want automated
- ☐ Language ready for "is this about cutting jobs"



PART II — SECTIONS 3–5 · GATE 1

Organize leadership and accountability

Named ownership, a working council, clear participation boundaries, and the machinery for evaluating tools.

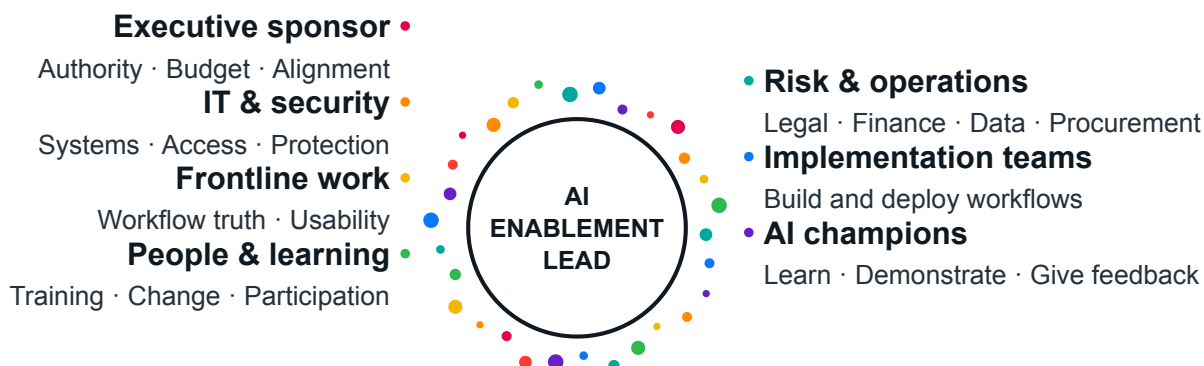
3

Establish the AI Enablement Council

For an organization of roughly fifty people or more, AI enablement cannot be delegated entirely to IT or left to individual enthusiasts. IT can provision a tool but cannot decide which workflows matter. Enthusiasts can demonstrate what is possible but cannot commit budget or accept risk on the organization's behalf. The gap between those two groups is where a council belongs.

One refinement before forming the council: name the executive sponsor and one accountable enablement lead first. A committee can make decisions, but an individual must own execution. Without a named lead, every council member can support the work while nobody owns the outcome, and the initiative dies between meetings.

The AI Enablement Council



Council decides. Specialists assess. Teams implement. Employees participate.

FIGURE 5 The AI Enablement Council — one accountable lead, cross-functional decisions.

— MINIMUM VIABLE COUNCIL

ROLE	WHAT THEY BRING	WHAT THEY OWN
Executive sponsor	Budget authority or strong influence over it (CEO, COO, CFO, CIO)	Resolving conflicts, connecting the work to organizational priorities, funding decisions
AI enablement lead	Program management capacity and standing to coordinate across functions	Running the program, managing the use-case portfolio, accountability between meetings
IT / security representative	Knowledge of systems, identity, data protection, and vendor technical evaluation	Provisioning, access, integrations, technical support
Frontline workflow representative	Deep knowledge of how the work actually happens	Reality-testing whether proposed solutions are useful, usable, and realistic
Risk / legal / compliance / privacy representative	Knowledge of regulatory obligations, contracts, privacy, intellectual property, records, and escalation; may be internal or external	Risk triage, identifying when specialist review is required, and maintaining the review and escalation path

These are five functions, not necessarily five different people. In a smaller organization, one person may cover more than one function and legal or compliance expertise may be external or on call. The risk function may not attend every working session, but it must have a named owner and a defined review path before Gate 1. Do not let the absence of internal counsel become the absence of legal or compliance review.

— EXPANDED COUNCIL FOR LARGER OR REGULATED ORGANIZATIONS

When the organization is regulated or routinely handles sensitive data, make the risk, legal, compliance, or privacy function a standing council seat rather than an on-call reviewer. As the organization grows, also add: **HR / learning and development** (role changes, training, communication, adoption); **finance or procurement** (business cases, vendor contracts, budget discipline, benefits validation); **data or analytics** (data availability, quality, permissions, measurement); and **operational leaders** from the functions that own priority workflows.

— THE COUNCIL CHARTER

Form the council with a short written charter, one to two pages, that names: the executive sponsor; the accountable program lead; members and the functions they represent; scope and the business outcomes the work serves; decision rights — what the council decides, recommends, and delegates; budget authority; the use-case intake process; risk and approval pathways; meeting cadence; initial 90-day deliverables; and measures of success.

During the first 90 days, the core council meets every two weeks; pilot teams may meet weekly. Once the operating model stabilizes, the council can move to monthly.

— THE COUNCIL'S FIRST ASSIGNMENTS

The council should not begin by writing a comprehensive AI policy or buying licenses for everyone. Its first assignments are narrower and generate the information every later decision depends on:

1. Inventory current AI use, including unofficial or "shadow" use
2. Establish interim acceptable-use guidance — enough to make experimentation safe, not a final policy
3. Identify important workflows and organizational pain points
4. Define a simple risk-tiering process
5. Select two or three pilot workflows
6. Assign each pilot an owner, users, a baseline, and success measures
7. Determine what tools, data, training, and approvals each pilot requires
8. Establish a feedback and support mechanism
9. Review results and decide what should scale

• MERIDIAN

Meridian's council: the executive director as sponsor, the director of programs as enablement lead (four hours a week, protected on her calendar), the two-person IT team's senior member, a family advocate with eleven years in the field, and the HR manager, who also served as the named privacy and risk owner with an external counsel escalation path. Five people covering six functions. Their first meeting produced two artifacts: a charter the sponsor signed, and a one-page interim guidance memo — which tools were provisionally fine, which data could never leave Meridian's systems, and one named person to ask when unsure.

FIELD WORK

Council formation

DO

Name the sponsor and enablement lead. Recruit the minimum viable council. Draft the charter using the toolkit template. Schedule the first 90 days of meetings before the first meeting occurs.

PRODUCE

A signed AI Enablement Council charter and interim acceptable-use guidance.

HAVE

☐ Named sponsor with budget influence

☐ One accountable lead with protected time

☐ IT and frontline voices seated

☐ Risk / legal / compliance / privacy review path named

☐ Charter signed

☐ Interim guidance published to all staff

G1

GATE 1 — CHARTER

The executive sponsor approves the charter, the named lead, a budget envelope for the first 90 days, and the interim guidance. Nothing else in this playbook proceeds without this signature.

Templates: Appendix 1 — Council charter · Appendix 2 — Council membership worksheet

4

Establish the participation structure

One committee should not be asked to do everything. Separate the structure into distinct groups with distinct jobs:

GROUP	PURPOSE
Executive sponsor	Authority, budget, organizational alignment
AI Enablement Council	Prioritization, governance, investment, measurement
Implementation teams	Build and deploy specific workflows
AI champions network	Peer learning, feedback, demonstrations, adoption

The employee-facing group deserves care in naming. Call it an **AI champions network** or **AI community of practice** rather than an employee resource group — ERGs conventionally refer to employee-led affinity and identity groups, and borrowing the term creates confusion for HR.

The champions network should include people from different departments who test approved tools and workflows, surface useful applications, identify friction and risks, demonstrate practices to colleagues, hold office hours, provide feedback to the council, and help translate training into everyday work.

The boundary matters as much as the role: champions inform adoption. They do not approve tools, set policy, or accept organizational risk. Those remain council decisions.

DO

Fill in the four-group table with real names. Invite champions by asking department heads for their most curious people, not their most senior. Write one sentence per group describing what it may and may not decide.

PRODUCE

A participation structure map with named people and decision boundaries.

HAVE

- ☐ Champions from at least three departments ☐ Written decision boundaries
- ☐ No group assigned work belonging to another

Template: Appendix 3 — Decision-rights matrix

5

Evaluate partners and internal tools

The council owns the process, portfolio, standards, and decisions for AI vendors and internal tools. It should delegate detailed technical, security, legal, and workflow evaluations to the appropriate specialists rather than performing every review itself — the council decides; specialists assess.

Council responsibilities: evaluating external AI vendors and implementation partners; approving internal AI tools and use cases; maintaining an inventory of approved, experimental, and prohibited tools; establishing evaluation and risk criteria; selecting and funding pilots and assigning implementation owners; monitoring adoption, outcomes, and risks; deciding whether pilots scale, change, or stop; and coordinating literacy, training, and capacity-building investments.

— VENDOR AND TOOL EVALUATION RUBRIC

Every evaluation should consider, at minimum:

1. **Business and workflow fit** — does it address a workflow the organization has actually prioritized?
2. **Security and privacy** — architecture, certifications, breach history, data handling
3. **Data retention and model-training terms** — does the vendor retain inputs? Train on them? Can that be turned off contractually?
4. **Identity, access, and administration** — SSO, role-based access, admin visibility
5. **Integration requirements** — what it must connect to, and what that costs
6. **Accuracy and human-review requirements** — error rates for this task, and the review burden they create
7. **Accessibility and employee usability** — can people across roles and abilities actually use it?
8. **Training and support** — what the vendor provides versus what the organization must build
9. **Total cost and expected value** — licenses plus integration, training, oversight, and administration
10. **Vendor stability and exit options** — what happens to your data and workflows if the vendor folds or the contract ends

— THE COUNCIL'S OPERATING ARTIFACTS

The council runs on documents, not meetings. Its core artifact set: the vendor evaluation rubric; a tool and use-case intake form; the approved-tool registry; a risk-tiering and review process; the pilot charter template; the organizational baseline; an adoption and outcome dashboard; and an AI risk and decision register.

FIELD WORK Evaluation machinery	
DO	Adopt the rubric above, adding any criteria your regulatory context requires. Stand up the approved-tool registry with three columns: approved, experimental, prohibited. Publish the intake form so any employee can nominate a tool or use case.
PRODUCE	A working evaluation rubric, a live tool registry, and an open intake channel.
HAVE	☐ Rubric adopted by the council ☐ Registry published where staff can find it ☐ Intake form live ☐ Every shadow-AI tool placed in a registry column
Template: Appendix 11 — Vendor evaluation rubric	



PART III — SECTIONS 6–7

Establish the baseline

Assess before you prescribe — measure literacy, adoption, and readiness as three different conditions.

6 Conduct the literacy and readiness survey

Before prescribing training, buying more tools, or selecting pilots, establish where the organization actually is. Use a short organization-wide survey, supplemented by departmental interviews or focus groups.

Comfort and familiarity are useful signals, but they are self-reported and they conflate three different things. The assessment should distinguish:

- **Literacy** — what people understand about AI, its limitations, and responsible use
- **Adoption** — how people are currently using AI in real work
- **Readiness** — whether tools, policies, leadership, workflows, and support make productive use possible

An organization can score high on one and low on the others, and the intervention differs in each case. Low literacy calls for education. Low adoption despite high literacy usually signals an enablement gap — no approved tools, unclear permission, or managers who discourage experimentation. Low readiness is a leadership and infrastructure problem no amount of employee training will fix.

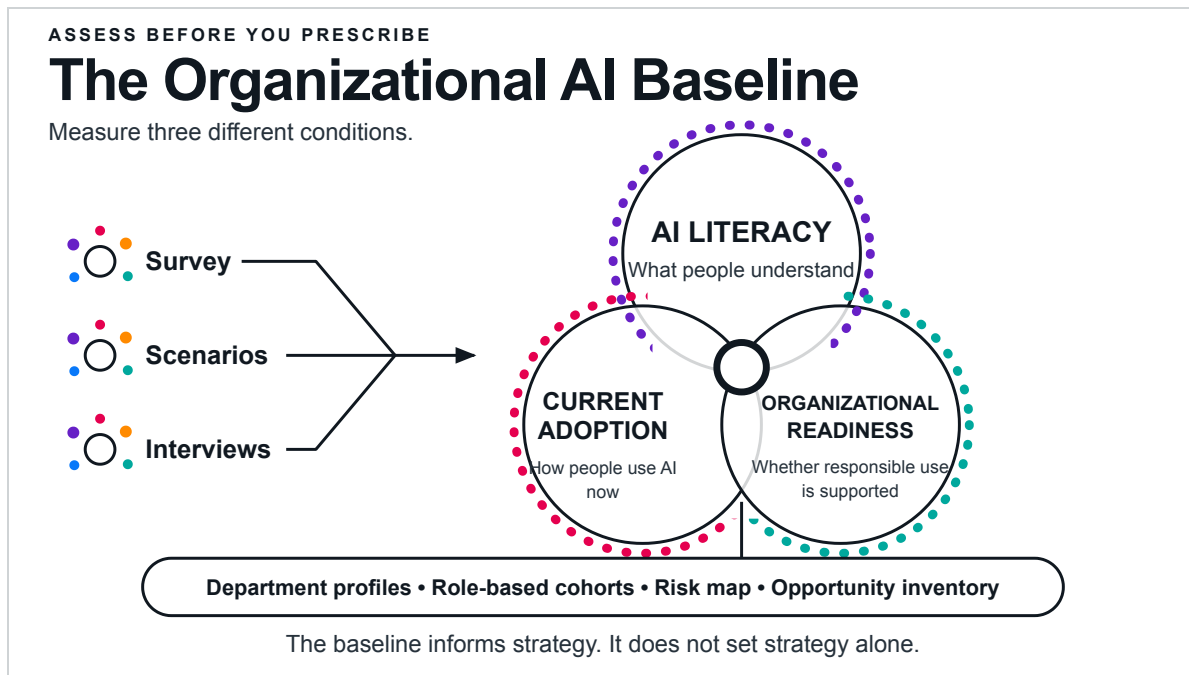


FIGURE 6 The organizational AI baseline — measure three different conditions; the baseline informs strategy, it does not set it alone.

— DIMENSIONS TO MEASURE

1. **Awareness** — understanding of available AI capabilities and common limitations
2. **Current use** — tools being used, frequency, tasks, and whether usage is approved or informal
3. **Practical capability** — ability to apply AI to actual workflows, evaluate outputs, and maintain human oversight
4. **Risk judgment** — understanding of privacy, confidential data, intellectual property, bias, accuracy, and escalation
5. **Confidence and willingness** — comfort experimenting, perceived usefulness, concerns, fear of displacement or surveillance
6. **Access and support** — approved tools, data access, technical support, policies, examples, manager encouragement
7. **Workflow opportunity** — repetitive tasks, bottlenecks, knowledge gaps, administrative burdens, and quality problems employees believe AI might improve

— SAMPLE SCALED ITEMS (FIVE-POINT AGREEMENT SCALE)

- I understand what generative AI can and cannot do.
- I know which AI tools I am permitted to use at work.
- I know what organizational information I may enter into an AI system.
- I can identify tasks in my role where AI may be useful.
- I feel capable of reviewing AI output for errors or unsupported claims.
- I currently use AI for work at least once per week.
- My manager encourages responsible experimentation with AI.
- I know where to go when I have an AI-related question.

- I am concerned that using AI may negatively affect how my work is evaluated.
- I have repetitive or time-consuming tasks that may benefit from AI assistance.

— OPEN-ENDED ITEMS

Which tasks consume disproportionate time? Where do handoffs or approvals create delays? What AI tools are you already using? What concerns prevent you from using AI? What would help you feel capable and confident? What is one workflow you would nominate for exploration?

— SCENARIO ITEMS

Include a few scenario-based questions so the baseline measures judgment, not only confidence — self-assessed confidence and actual judgment routinely diverge. For example:

An employee wants to paste a confidential client document into a publicly available AI assistant to summarize it faster. What should they do?

— PROTECT EMPLOYEE TRUST

The survey instrument should state explicitly: this is not a performance evaluation; individual responses will not be shared with managers; do not submit confidential prompts, data, or client information in your answers; the purpose is to design appropriate tools, training, policies, and support; findings will be reported in aggregate.

Without these protections the data is worthless: employees conceal unofficial AI use, and answers drift toward what people believe leadership wants to hear. The survey's honesty depends entirely on whether employees believe answering honestly is safe.

• MERIDIAN

Meridian's survey went to all 140 staff and got 112 responses — because the executive director recorded a two-minute video promising, on camera, that no individual answers would reach supervisors. The scenario question caught something the confidence items missed: 71% felt "capable of reviewing AI output," but only 34% correctly answered what to do with a confidential client document. Eighteen staff disclosed regular use of free consumer chatbots for work, including two who had pasted case-note fragments into one.

DO

Build the instrument from the dimensions above (or adapt your existing readiness assessment to them). Have the sponsor personally deliver the trust commitments. Field for two weeks with one reminder. Target a response rate above 60% — below that, your segments will be too thin to act on.

PRODUCE

A fielded survey and raw response set.

HAVE

- ☐ All seven dimensions covered
- ☐ At least two scenario items
- ☐ Trust commitments delivered by the sponsor, not by email signature
- ☐ Response rate above 60%

Template: Appendix 4 — AI literacy and readiness survey instrument

7

Interpret the organizational baseline

Do not reduce the results to one organization-wide score. A single number hides exactly the variation the council needs to act on. Produce instead:

- An organization-wide literacy baseline
- Department-level readiness profiles
- A current tool and usage inventory
- A shadow-AI risk map
- Role-based learning cohorts
- An employee concern and confidence profile
- An initial workflow opportunity inventory
- A list of functions requiring additional governance or support

This segmentation is what lets you avoid generic training. One group may need basic literacy; another may need workflow-design workshops; managers may need change-leadership preparation; technical teams may need evaluation and governance practices.

The survey informs strategy; it does not set it alone. The baseline gives the council its strategic hypotheses — which departments need foundational support, where unofficial adoption is already happening, where the greatest risks sit, which functions are ready for pilots, what concerns leadership must address, and where tools, policy, or access create barriers. The departmental interviews in Part V then test those hypotheses against how work actually happens.

● MERIDIAN

The segmentation surfaced three cohorts: program staff needed foundational literacy plus clear data rules; the finance and development teams were already experimenting and needed workflow-design support, not basics; and mid-level managers scored lowest on "I encourage experimentation" — several admitted in focus

groups they didn't know whether they were allowed to. The council's hypothesis list put grant reporting and case documentation at the top: both appeared repeatedly in the open-ended friction answers.

FIELD WORK		Baseline report
DO	Analyze responses into the eight outputs above. Present to the council as department profiles, not one score. Write three to five strategic hypotheses the interviews will test. Report findings back to all staff in aggregate — closing the loop is part of keeping the trust you promised.	
PRODUCE	The organizational baseline report and hypothesis list.	
HAVE	☐ Department-level profiles ☐ Shadow-AI risk map ☐ Named learning cohorts ☐ Written hypotheses ☐ All-staff report-back delivered	



PART IV — SECTION 8 · GATE 2

Create shared understanding

A common vocabulary before departments are asked to evaluate their own work.

8 Deliver foundational training

The first training should not be a tool tutorial. Its purpose is to establish a common organizational vocabulary before departments are asked to evaluate their own work. Skip this step and the departmental interviews produce requests shaped by fear, hype, or a narrow view of what is possible — people cannot nominate good workflows for a technology they misunderstand.

The foundational training should cover:

1. What AI is and is not
2. What generative AI can and cannot reliably do
3. The distinction between automation, assistance, and decision-making
4. Approved tools and data boundaries — what may go into which systems
5. Human responsibility and oversight — the person remains accountable for the output
6. The RESPECT principles (Part VII)
7. Concrete examples of low-, medium-, and high-risk uses drawn from this organization's work
8. How employees can nominate a workflow for exploration
9. Where to go for support or escalation

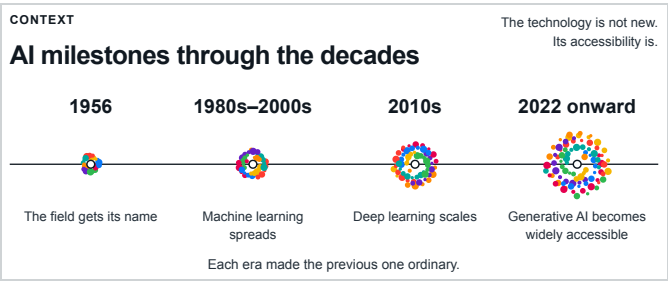


FIGURE 7 AI milestones through the decades — context for "what AI is and is not."

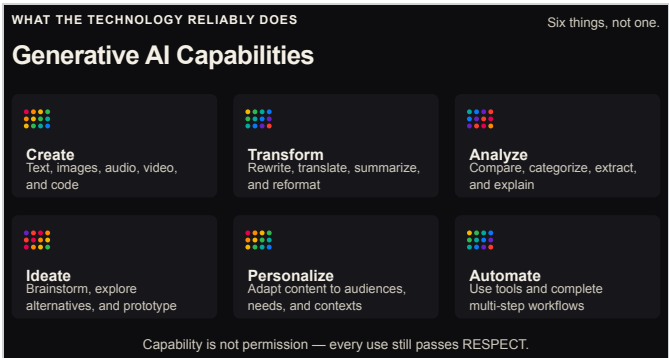


FIGURE 8 Six generative AI capabilities — the vocabulary for what the technology can reliably do.

The measure of success for this training is not completion rates. It is whether the departmental interviews that follow are productive — whether employees arrive able to talk about their tasks, their friction, and realistic possibilities in a shared language.

● **MERIDIAN**

Meridian ran three 90-minute sessions segmented by the baseline cohorts rather than one all-staff webinar. The risk examples were Meridian's own: a low-risk use (drafting a volunteer-appreciation email), a medium-risk use (summarizing anonymized survey feedback), and a high-risk use (anything touching a family's case file). The session ended with each person writing one workflow nomination on a card. Ninety-one cards came in; grant reporting appeared on nineteen of them.

FIELD WORK Shared language

DO	Deliver cohort-based sessions built on the nine topics. Use your organization's real work in every example. End every session with a workflow nomination exercise. Check comprehension with three scenario questions, not a satisfaction survey.
PRODUCE	Training delivered to all cohorts; a stack of workflow nominations; scenario-based comprehension results.
HAVE	☐ All cohorts trained ☐ Scenario comprehension materially above baseline ☐ Workflow nominations collected ☐ Support and escalation paths named in the session
G2	GATE 2 — BASELINE The council formally accepts the baseline report, confirms the learning-cohort plan, and approves the priority departments for 5D interviews. This is where the council commits to where it will look first.
Template: Appendix 5 — Foundational training agenda	



Understand departmental work

From organization-wide data to departmental depth — examine work, don't collect wish lists.

9 Conduct 5D departmental interviews

With shared language established, move from organization-wide data to departmental depth. The interviews should examine work — not simply ask "where would you like to use AI?" That question invites wish lists. Examining work surfaces the real material:

- Repetitive tasks and high-volume information processing
- Delayed decisions and manual handoffs
- Difficult knowledge retrieval and documentation burdens
- Quality inconsistencies
- Work dependent on a few experienced employees
- Customer or employee pain points
- High-stakes activities where AI may be inappropriate
- The outcomes the department is accountable for

Structure the interviews with the 5D framework — the RUDI workflow method. Humans define and validate the work; AI helps structure, analyze, prototype, and test it:

1	Discuss What actually happens?	Listen to the people doing the work.	OUTPUT Conversation + problem brief
2	Document Is the workflow accurate?	Make actors, decisions, exceptions, inputs, and outputs explicit.	OUTPUT Validated current-state workflow
3	Diagnose Where are friction and risk?	Separate repetition from judgment, agency, privacy, and safety.	OUTPUT Opportunity + risk map
4	Design What should change?	Assign human and AI responsibilities.	OUTPUT Future-state workflow
5	Develop Does it work in practice?	Prototype, test with real cases, observe, and refine.	OUTPUT Tested prototype + evidence

Humans validate before the work advances. New evidence returns the team to Discuss.

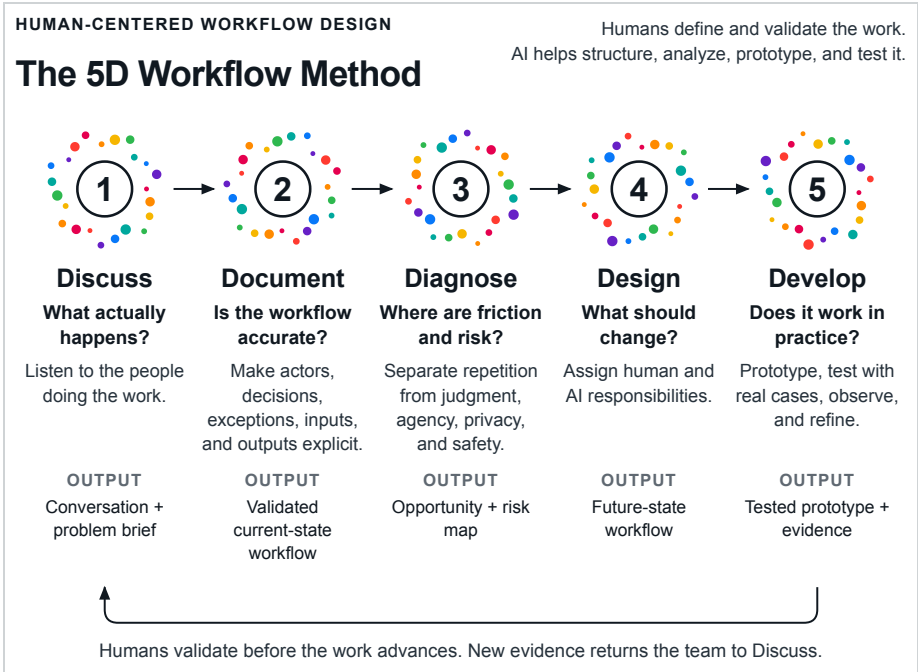


FIGURE 9 The 5D Workflow Method — humans define and validate the work; AI helps structure, analyze, prototype, and test it.

The output of each interview is a structured departmental opportunity map, not an unfiltered wish list: candidate workflows, the friction inside them, the constraints around them, and the department's own judgment about where AI does not belong.

● **MERIDIAN**

The council interviewed three departments: family services, finance/grants, and development. The grants interview ran two hours with four staff. The signal finding: quarterly grant reporting across eleven funders consumed roughly 26 staff-days per quarter, most of it re-assembling the same program data into eleven different formats — and two reports had gone out late in the past year, one triggering a funder inquiry. The team also drew a hard line unprompted: nothing AI-generated goes to a funder without the grants manager's review, ever. That line went into the opportunity map as a design constraint, not an obstacle.

DO Schedule 90–120 minute interviews with each priority department. Bring the 5D interview guide and the department's own baseline profile. Walk actual work products — reports, forms, queues — not abstractions. Record where the department says AI does not belong, verbatim.

PRODUCE A departmental opportunity map per interviewed department.

HAVE

- ☐ Each priority department interviewed
- ☐ Candidate workflows with friction quantified
- ☐ Explicit "not here" boundaries recorded
- ☐ Department's accountable outcomes documented

Template: Appendix 6 — 5D departmental interview guide

10 Build the workflow inventory

For each workflow surfaced in the interviews, document: purpose; owner; participants; inputs and outputs; systems and data involved; decisions made within it; exceptions and their frequency; handoffs; baseline performance (time, volume, quality, cost — whatever the department already tracks); and current problems.

The baseline row is the one leaders skip and later regret. Without a documented starting point, no pilot can demonstrate improvement, and the council ends up debating anecdotes.

• MERIDIAN — WORKFLOW INVENTORY ENTRY (ABRIDGED)

Workflow: Quarterly grant reporting. **Owner:** Grants manager. **Participants:** Grants manager, two program directors, finance associate. **Volume:** 11 funder reports per quarter. **Baseline:** ~26 staff-days per quarter; 2 late submissions in the past 4 quarters; each report averages 4 revision cycles. **Systems:** Program database, spreadsheet exports, funder portals. **Known problems:** Manual data re-assembly, inconsistent narrative quality across authors, single point of failure in the grants manager.

DO Complete an inventory entry for every candidate workflow from the opportunity maps. Chase down real baseline numbers — timesheets, submission logs, queue counts. Where no number exists, have the team estimate and mark it as an estimate.

PRODUCE The workflow inventory, with a baseline for every entry.

HAVE

- ☐ Every candidate workflow documented
- ☐ A baseline number (or flagged estimate) on each
- ☐ Owners named
- ☐ Systems and data mapped

Template: Appendix 7 — Workflow inventory template



Redesign work at the task level

Decompose workflows into their components, then make a deliberate 4A decision about every task.

11 Decompose workflows

Break each prioritized workflow into its components:

- **Tasks** — the units of activity
- **Decisions** — points where someone exercises judgment
- **Handoffs** — transfers between people, teams, or systems
- **Human interactions** — moments where relationship, empathy, or legitimacy is part of the work
- **Oversight responsibilities** — existing review, approval, and quality-control activity

For each task, document: purpose, owner, frequency, time required, inputs and outputs, systems and data involved, required judgment, variability and exceptions, consequences of error, human interaction requirements, current pain points, and baseline performance.

• MERIDIAN

Decomposing grant reporting produced fourteen tasks. Among them: pull program data from the database (2 days/quarter), reconcile financials (3 days), draft narrative sections (8 days across three authors), internal review cycles (5 days), format per funder template (4 days), certify compliance and submit (grants manager, half a day), and respond to funder questions (variable). Two tasks were decisions: which outcomes to feature, and whether a variance requires proactive explanation. One was a pure handoff — program directors passing drafts to the grants manager — that accounted for most of the calendar delay.

FIELD WORK Task decomposition

DO

With the workflow owner and one frontline participant in the room, decompose the top two or three workflows using the decomposition worksheet. Time-box it: one workflow per 90-minute session. Tag every item as task, decision, handoff, or human interaction.

PRODUCE

A completed task decomposition worksheet per priority workflow.

HAVE

- ☐ Every task tagged by type
- ☐ Time and frequency on each
- ☐ Consequences of error rated
- ☐ Owner agrees the decomposition matches reality

Template: Appendix 8 — Task-decomposition worksheet

12 Apply the 4A Task Framework

Classify every task into one of four dispositions. The verbs are deliberate: each names a decision the organization makes, not a property the task has.

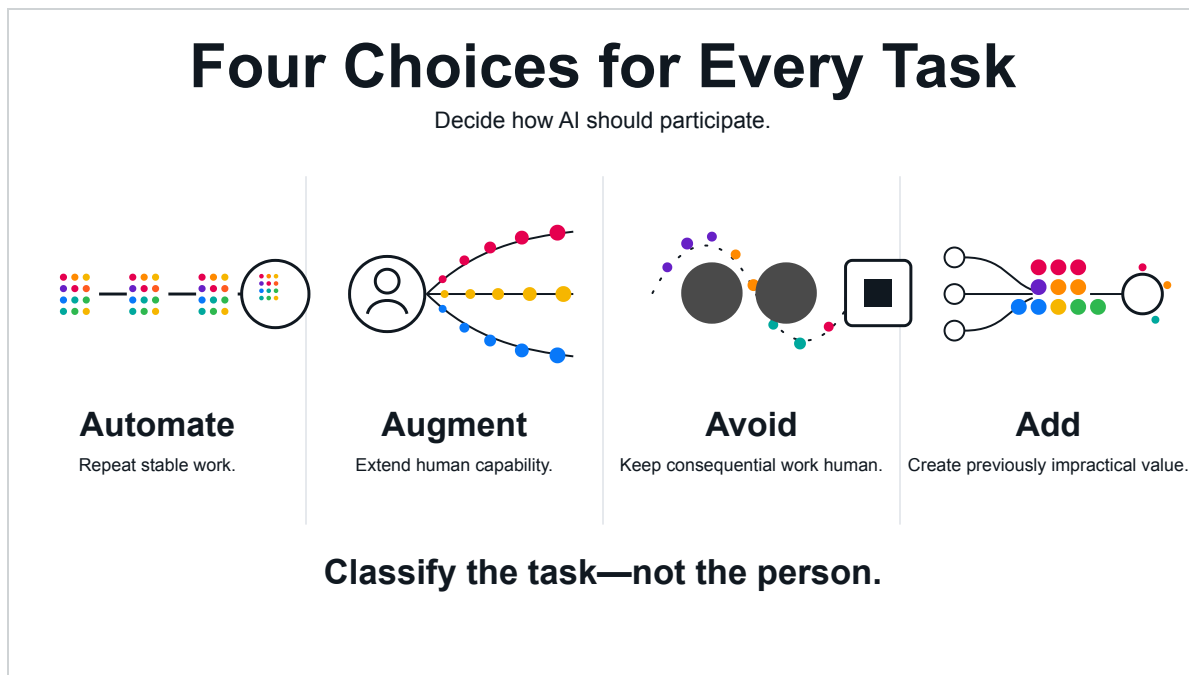


FIGURE 10 Four choices for every task — decide how AI should participate.

— AUTOMATE

AI or automation performs most or all of the task within established controls. The best candidates share a profile: repeated steps, clear inputs and outputs, stable decision rules, sufficient volume to justify the setup cost, detectable errors, manageable consequences, and limited need for empathy or contextual judgment. Automation may still require review, exception handling, and escalation.

Example: Extracting standard fields from invoices and routing exceptions to an employee.

— AUGMENT

AI assists a person; the person retains judgment, responsibility, and control. Strong candidates involve analysis, drafting, research, preparation, comparison, summarization, scenario development, and complex or contextual decisions. Most early organizational use cases will fall here, and that is appropriate — augmentation delivers value while the organization is still building the oversight muscle that automation requires.

Example: AI prepares a summary of employee survey responses; a leader interprets the findings and decides how to respond.

— AVOID

The organization deliberately decides not to use AI for a task, or restricts it to tightly controlled supporting functions. Grounds for avoidance include: high consequences of error, unreliable or unverifiable outputs, sensitive personal or organizational data, legal or regulatory restriction, essential human relationships, decisions requiring empathy or moral accountability, absence of meaningful human review, or insufficient value relative to risk.

Avoid is a design decision, not resistance to innovation. An organization that cannot articulate what it will not use AI for has not finished thinking. Documented avoidance decisions also protect the organization later — they are the evidence that risk was considered, not overlooked.

Example: AI does not independently make a final hiring, termination, medical, disciplinary, or benefits-eligibility decision.

— ADD

AI makes a valuable new task, service, control, or capability possible — or creates new work required to use AI responsibly. This includes new personalized services, new forms of analysis, continuous quality evaluation, AI-output review, use-case inventory management, model and vendor monitoring, prompt and agent maintenance, incident escalation, AI coaching and office hours, and documentation and audit activity.

Add matters because it corrects the framing everyone brings into the room. AI transformation is not only subtraction. Some of the highest-value outcomes are analyses and services that were previously too expensive to produce at all, and some of the new work — review, monitoring, coaching — is the price of doing the rest responsibly.

Example: A department begins a monthly analysis of customer-support patterns that was previously too time-consuming to produce.

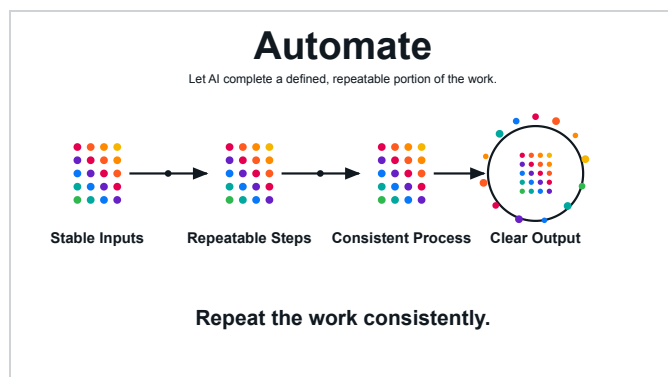


FIG 11A Automate — repeat the work consistently.

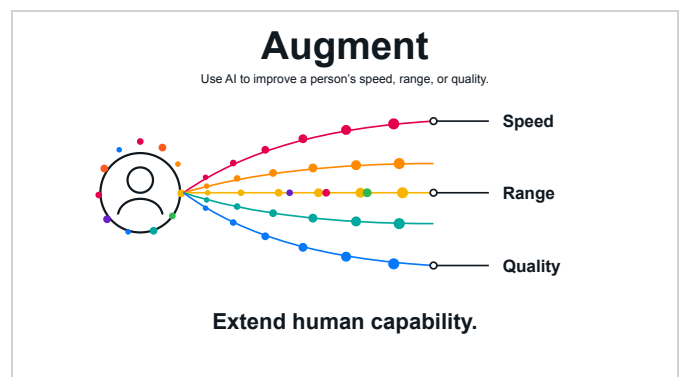


FIG 11B Augment — extend human capability.

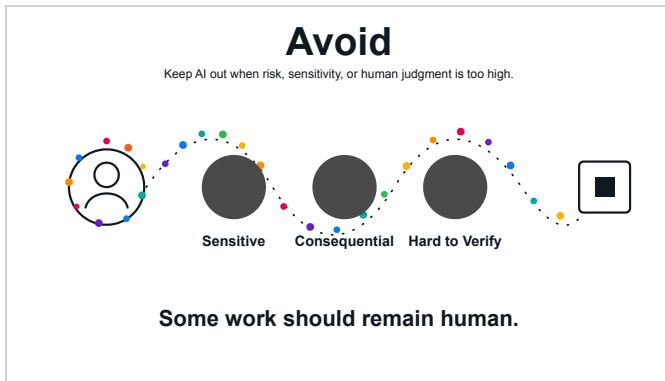


FIG 11C Avoid — some work should remain human.

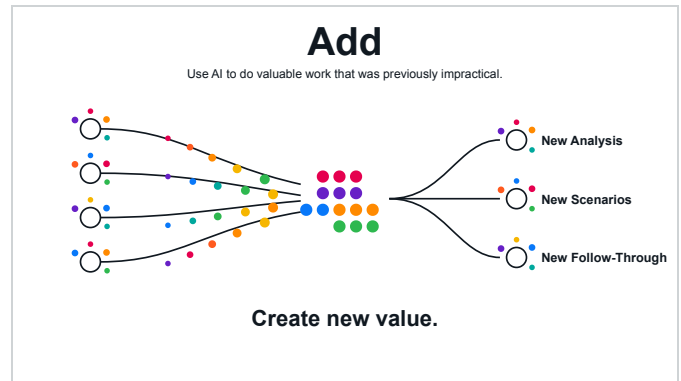


FIG 11D Add — create new value.

— WORKED EXAMPLE: A HIRING WORKFLOW

TASK	4A DISPOSITION	REASON
Schedule interviews	Automate	Repetitive, rules-based, reversible
Draft a job description	Augment	AI drafts; humans define the actual role
Summarize interview notes	Augment	Useful assistance; context and verification remain human
Make the final hiring decision	Avoid	Requires accountability, judgment, and bias safeguards
Audit AI-assisted hiring activity	Add	New control required for responsible use
Create role-specific interview guides	Augment	AI assists; hiring leaders retain control

For each classification, record the reasoning, supporting evidence, identified risks, and required human oversight. The classification is a claim; the documentation is what makes it reviewable when circumstances change.

● MERIDIAN — GRANT REPORTING DISPOSITION MAP (ABRIDGED)

TASK	DISPOSITION	REASON
Pull program data from database	AUTOMATE	Repetitive, rules-based, errors detectable against source
Reconcile financials	AUGMENT	AI flags variances; finance associate resolves them
Draft narrative sections	AUGMENT	AI drafts from data + prior reports; program directors revise and own
Decide which outcomes to feature	AVOID (AI)	Strategic judgment tied to funder relationships
Format per funder template	AUTOMATE	Mechanical transformation, verifiable output
Certify compliance and submit	AVOID	Legal accountability stays with the grants manager — the team's own hard line
Quarterly cross-funder outcomes analysis	ADD	Previously too time-consuming; now feasible and valuable to the board

FIELD WORK Task disposition map

DO In a working session with the same people who did the decomposition, classify every task. Argue about the hard ones — the argument is the analysis. Write the reason next to every disposition, including every Avoid.

PRODUCE A task disposition map per priority workflow, with documented reasoning.

HAVE

- ☐ Every task classified ☐ A written reason per task
- ☐ At least one Avoid you can defend to a skeptic
- ☐ At least one Add — if there are none, look again
- ☐ Oversight noted for every Automate and Augment

Template: Appendix 9 — 4A classification worksheet



Apply human-centered principles

RESPECT — the cross-cutting lens that judges whether a change is responsible, not merely possible.

13 Use the RESPECT framework

RESPECT is the cross-cutting decision lens applied to vendors, internal tools, workflows, task dispositions, pilots, and scaled deployments. It answers a question the 4A analysis does not: a task may be technically automatable and still fail as a human-centered design.

RESPECT is a human-centered framework for determining whether an AI-enabled workflow is not merely possible, but reliable, accessible, safe, privacy-preserving, worthwhile, controllable, and transparent.

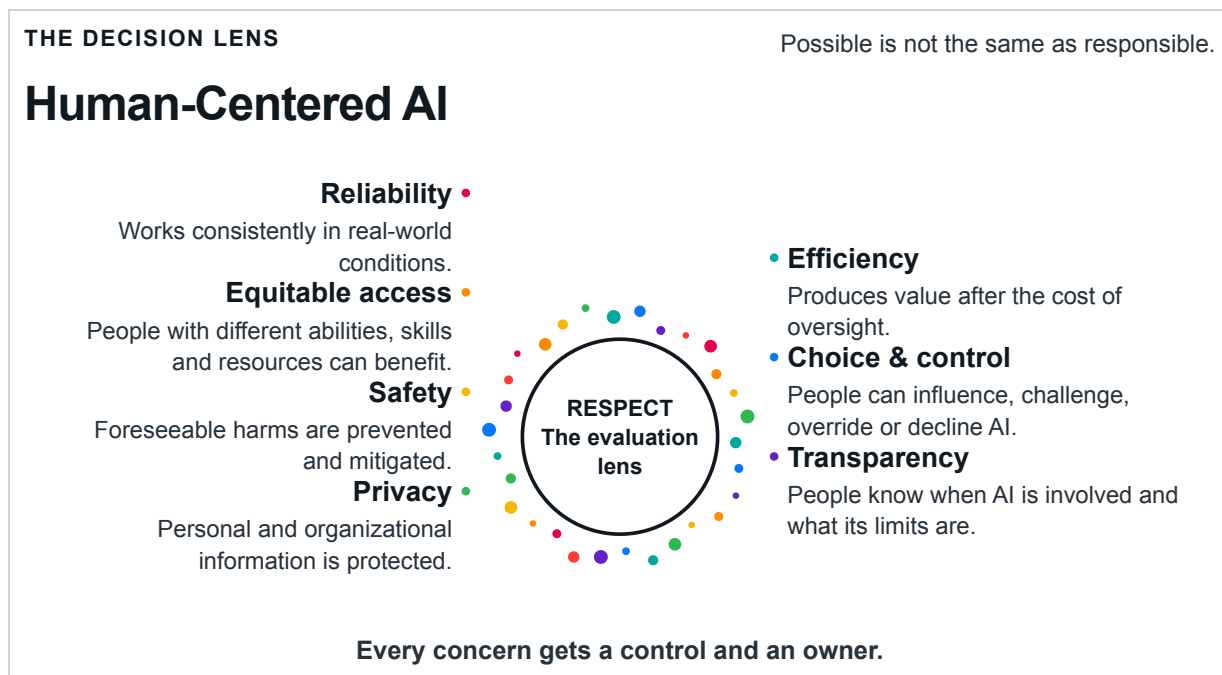


FIGURE 12 The seven human-centered principles behind RESPECT.

PRINCIPLE	ORGANIZATIONAL QUESTION
R — Reliability	Does it perform consistently enough for this particular task, and how will errors be detected?
E — Equitable access	Can people across roles, locations, abilities, and levels of technical confidence benefit appropriately?

PRINCIPLE	ORGANIZATIONAL QUESTION
S — Safety	What foreseeable harm could result, and what controls or escalation paths are required?
P — Privacy	What data is collected, entered, retained, shared, or used for model training?
E — Efficiency	Does it meaningfully improve time, cost, quality, capacity, or service after accounting for oversight?
C — Choice & control	Can people review, challenge, override, escalate, or opt out where appropriate?
T — Transparency	Do users and affected people understand when AI is involved, what it does, and what its limitations are?

RESPECT should not collapse into one vague score. A single score lets a strong efficiency rating paper over a privacy failure. Each proposed use case should instead receive: supporting evidence for each principle; identified concerns; required controls; an accountable owner; and a decision — **go, revise, pilot, or reject**.

Note the efficiency principle's phrasing: *after accounting for oversight*. A workflow that saves an hour of drafting and adds ninety minutes of review has negative efficiency, and RESPECT is designed to catch exactly that arithmetic before the pilot does.

• MERIDIAN

RESPECT on the grant-reporting redesign surfaced two issues 4A had not. Privacy: narrative drafting would work best with program examples, but program examples contain family details — so the control became a de-identification step before anything reaches the drafting tool, owned by the program directors. Equitable access: one of the three narrative authors was far less comfortable with the tools than the others, which without support would have concentrated the work rather than distributing it — so the pilot added paired sessions with a champion for the first two cycles. Decision: **pilot**, with those two controls written into the charter.

FIELD WORK RESPECT assessment	
DO	Run every proposed pilot through the RESPECT template with the council. For each principle, write evidence and concerns — "no concerns" requires a sentence explaining why. Assign a control and an owner to every concern. End with one of four words: go, revise, pilot, or reject.
PRODUCE	A completed RESPECT assessment per use case, each with an explicit decision.
HAVE	☐ Seven principles addressed with evidence ☐ Every concern paired with a control and owner ☐ Oversight cost counted in the efficiency judgment ☐ A one-word decision on record
Template: Appendix 10 — RESPECT assessment template	



Prioritize and implement

Turn the inventory into a portfolio, the portfolio into chartered pilots, and the pilots into supported adoption.

14 Build the use-case portfolio

By this point the organization has an opportunity inventory with 4A dispositions and RESPECT evaluations. Prioritization turns that inventory into a portfolio. Score candidates against:

1. **Organizational value** — impact on outcomes the organization is accountable for
2. **Employee value** — does it remove friction employees actually feel?
3. **Feasibility** — can it be built or configured with available capability?
4. **Data readiness** — does the required data exist, in usable form, with permission to use it?
5. **Integration requirements** — what it must connect to, and the cost of connecting
6. **Adoption readiness** — is the affected team willing and prepared?
7. **Risk** — from the RESPECT evaluation and risk tiering
8. **Cost** — total, including oversight and support
9. **Measurability** — can improvement be demonstrated against a baseline?
10. **Reusability** — does solving this teach the organization something it can apply elsewhere?

The first pilots should score well on employee value, feasibility, and measurability even if they are not the highest-value opportunities on paper. Early wins that employees feel build the trust the harder projects will need.

The Use-Case and Pilot Lifecycle



FIGURE 13 The use-case and pilot lifecycle — every pilot ends with a decision.

FIELD WORK Portfolio ranking	
DO	Score every RESPECT-cleared use case on the ten criteria in a council session. Rank. Select two or three pilots — no more, however tempting the list looks. Record why the non-selected candidates wait, so the decision survives staff turnover.
PRODUCE	A ranked use-case portfolio with pilot selections and deferral reasons.
HAVE	☐ Every cleared use case scored ☐ Two or three pilots selected ☐ Early pilots strong on employee value and measurability ☐ Deferred candidates documented, not discarded
Template: Appendix 12 — Use-case prioritization matrix	

15 Select initial pilots

Every pilot gets a one-page charter before any work begins. No charter, no pilot. Each charter names: executive sponsor; workflow owner; participating users; baseline (documented before the pilot starts); intended outcome; approved tool; RESPECT evaluation summary; human-review plan; success measures; stop conditions; and review date.

Stop conditions deserve emphasis. A pilot without stop conditions cannot fail, and a pilot that cannot fail cannot produce evidence. Deciding in advance what result would end the pilot is what separates experimentation from commitment-by-drift.

● MERIDIAN — PILOT CHARTER (ABRIDGED)

Pilot: AI-assisted grant reporting, two funders for two quarters. **Owner:** Grants manager. **Users:** Two program directors, finance associate. **Baseline:** 26 staff-days/quarter across all funders; 4 revision cycles per report; 2 late submissions in 4 quarters. **Intended outcome:** 40% time reduction on the two pilot reports with equal or better funder acceptance. **Human-review plan:** De-identification before drafting; grants manager reviews and certifies everything; nothing reaches a funder without her sign-off. **Success measures:** Staff-days, revision cycles, on-time submission, author experience survey. **Stop conditions:** Any confidentiality breach stops the pilot immediately; if review burden exceeds 50% of time saved after the first quarter, the council reconsiders the design. **Review date:** End of quarter two.

🛑 **Gate 3 at Meridian:** the executive director approved the portfolio and a \$9,400 first-year budget (tool licenses, champion time, training hours); the charter was signed by sponsor and owner in the same meeting.

FIELD WORK Pilot charters

DO Draft a one-page charter per selected pilot using the template. Document the baseline before anyone touches a tool. Write stop conditions that name numbers and events, not sentiments. Get physical or digital signatures.

PRODUCE Signed pilot charters.

HAVE

- ☐ Charter per pilot, one page
- ☐ Baseline captured pre-launch
- ☐ Stop conditions with numbers
- ☐ Review date on the council calendar
- ☐ Signatures from sponsor and workflow owner

G3 🛑 **GATE 3 — PORTFOLIO**

The executive sponsor approves the pilot portfolio and budget. Each pilot charter is signed by its sponsor and workflow owner. Unsigned pilots do not launch.

Template: Appendix 13 — Pilot charter template

16 Support adoption

Tools do not adopt themselves, and training alone does not change behavior. Build the surrounding support system:

- **Manager enablement** — managers set the local norms; an employee whose manager is skeptical will not use the tool no matter what the policy says
- **Role-based learning** — matched to the cohorts identified in the baseline
- **Office hours** — a standing, low-stakes place to bring questions
- **AI champions** — peer demonstration travels farther than central communication
- **Help and escalation paths** — known, named, and responsive

- **Reusable playbooks** — each solved workflow documented so the next team starts from the last team's work
- **Internal demonstrations** — real colleagues showing real work, not vendor demos
- **Feedback loops** — a route from user friction back to the council

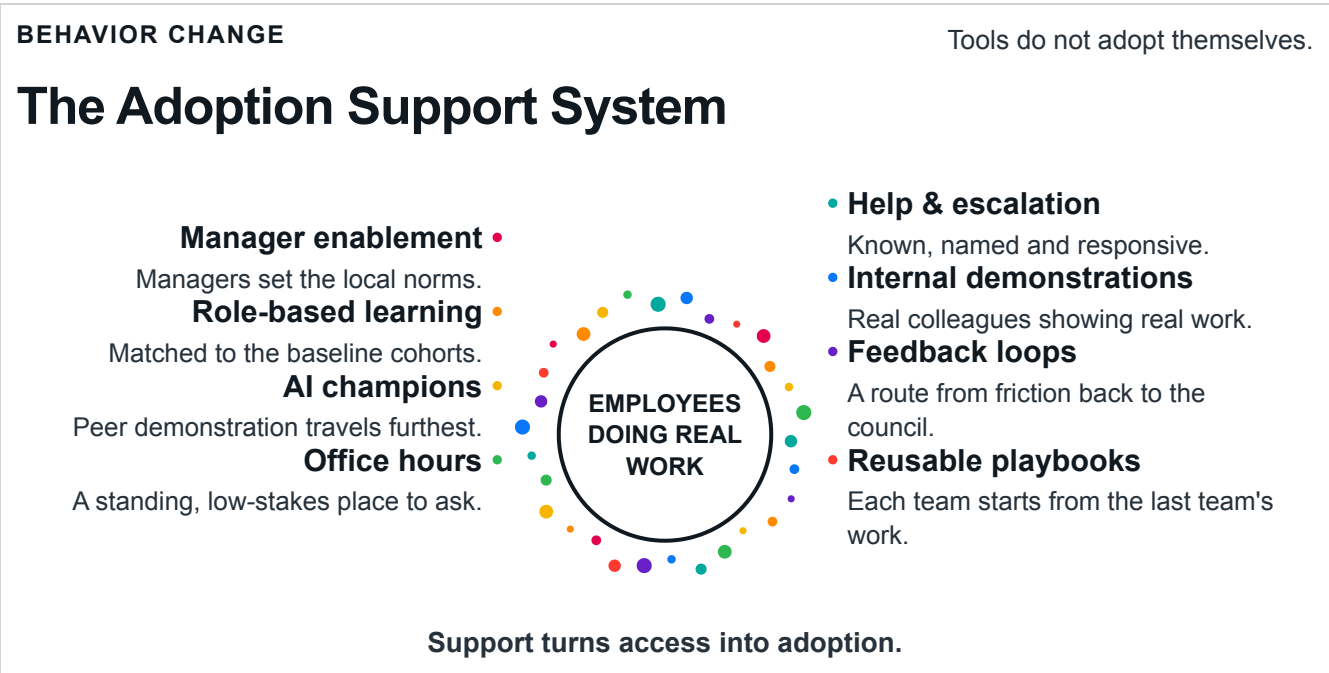


FIGURE 14 The adoption support system — tools do not adopt themselves; support turns access into adoption.

FIELD WORK Adoption support system	
DO	Stand up weekly office hours before the pilots launch, not after problems appear. Brief every manager whose team touches a pilot on what to encourage and what to escalate. Assign a champion to each pilot team. Open a visible feedback channel to the enablement lead.
PRODUCE	A running support system: office-hours schedule, manager briefings delivered, champions assigned, feedback channel live.
HAVE	☐ Office hours on the calendar ☐ Every affected manager briefed ☐ A champion per pilot ☐ A feedback route employees can name when asked



Measure and institutionalize

Measure the full chain from access to outcome, and build the capacity to scale, revise, and — hardest of all — stop.

17 Measure more than training completion

Training completion and license counts measure spending, not change. The measurement system should track the full chain from access to outcome:

LAYER	MEASURES
Access	Approved tools provisioned; employees with access
Usage	Active use; repeat use; use embedded in named workflows
Performance	Time and cycle-time changes against baseline; quality; error and rework rates
Experience	Employee experience of the changed work; customer or constituent outcomes
Risk	Risk events; near misses; human-review burden and whether it is sustainable
Capacity	Internal ownership; reusable organizational knowledge; reduced dependence on outside help

Report against the baselines documented in Parts III and V. The measurement question is always the same: compared to what?

The AI Adoption Measurement Chain

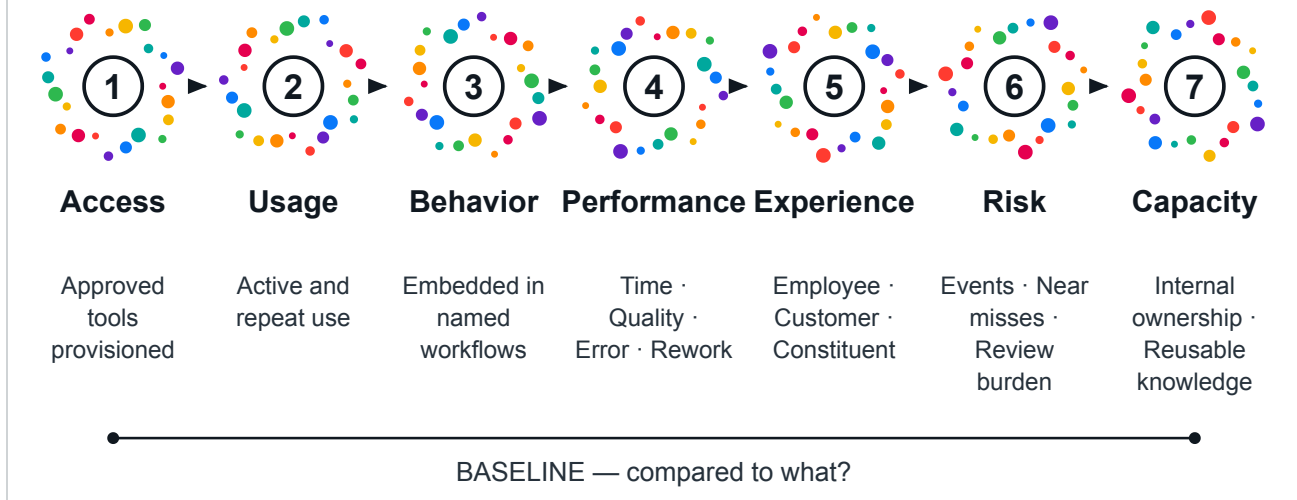


FIGURE 15 The AI adoption measurement chain — training completion and license counts are inputs, not outcomes.

● MERIDIAN

After two quarters: the two pilot reports took 11 staff-days against a 16-day baseline for those funders — a 31% reduction, short of the 40% target but real. Revision cycles fell from four to two. Both reports went out early. The review burden ran at 20% of time saved, well under the stop threshold. One near miss went into the risk register: a draft briefly included an insufficiently de-identified example, caught in review — which is what the review was for. The author experience survey split: two authors wanted to expand immediately; the third still found the tool more burden than help.

FIELD WORK Adoption dashboard

DO

Build a one-page dashboard on the six layers, fed by the pilot measures. Review at every council meeting. Log every risk event and near miss in the register, including the ones the controls caught — those are evidence the system works, not embarrassments to bury.

PRODUCE

A live adoption and outcome dashboard; a maintained risk and decision register.

HAVE

- ☐ All six layers on one page
- ☐ Every number paired with its baseline
- ☐ Risk register in use
- ☐ Dashboard reviewed on a standing agenda

Template: Appendix 14 — Adoption dashboard template

18 Build durable capacity

The organization has developed capacity — the fourth stage of the maturity progression — when it can repeatedly, without outside help: identify worthwhile opportunities; evaluate tools and partners; redesign workflows; manage risk; implement pilots; support employees; measure results; scale what works; and stop what does not.

That last item is the one most organizations cannot do. Sunsetting a tool or ending a pilot feels like admitting error, so failed experiments linger, consuming licenses and attention. A council that has stopped something and survived has demonstrated more institutional maturity than one that has only launched things.

Transformation — the fifth stage — is not a project with an end date. It is what becomes possible once capacity exists: services redesigned around new capabilities, roles rebuilt around new task mixes, strategy that assumes the organization can absorb the next wave of tools because it has a working system for absorbing this one.

• MERIDIAN

At the Gate 4 review, the council made three decisions in writing: scale the grant-reporting workflow to five more funders next quarter, with the third author paired with a champion for one more cycle; hold the case-documentation candidate until the de-identification process proves itself at larger volume; and stop a small scheduling-tool experiment the development team had started — usage data showed two logins in six weeks. Stopping it took ten minutes and no one was blamed. That meeting, more than any pilot metric, was the evidence Meridian was building capacity.

FIELD WORK Capacity review

DO At each pilot review date, run Gate 4: scale, revise, or stop, decided in writing with reasons. Twice a year, run the capacity maturity assessment against the nine capabilities above and report the result to the sponsor and, where relevant, the board.

PRODUCE Written scale/revise/stop decisions; a semi-annual capacity maturity assessment.

HAVE

- ☐ Every pilot reviewed on its charter date
- ☐ At least one documented stop decision in year one
- ☐ Reusable playbooks from scaled pilots
- ☐ A capacity score trending, honestly, toward durable



🚫 GATE 4 — SCALE

The council decides each pilot's fate in writing: scale, revise, or stop. No pilot continues by default.

Template: Appendix 15 — Capacity maturity assessment



SUMMARY — TEN PHASES · FOUR GATES

The complete methodology at a glance

#	PHASE	FIELD WORK ARTIFACT	GATE
1	Organize	Signed council charter; participation map; interim guidance; tool registry	Gate 1 — Charter
2	Assess	Baseline survey; baseline report and hypotheses	
3	Orient	Cohort training; workflow nominations; comprehension results	Gate 2 — Baseline
4	Diagnose	5D opportunity maps; workflow inventory	
5	Decompose	Task decomposition worksheets	
6	Classify	4A task disposition maps	
7	Evaluate	RESPECT assessments with go/revise/pilot/reject decisions	
8	Prioritize	Ranked portfolio; signed pilot charters	Gate 3 — Portfolio
9	Implement	Running pilots; adoption support system	
10	Institutionalize	Dashboard; risk and decision register; capacity assessment	Gate 4 — Scale

Running beneath the sequence: the maturity progression (Literacy → Enablement → Adoption → Capacity → Transformation) describes what changes in the organization; the operating sequence describes what the organization does; RESPECT describes how decisions are evaluated at every step; governance holds the pen throughout.

The central proposition:

AI transformation does not begin with eliminating occupations. It begins by understanding the tasks that constitute work and deliberately deciding what should be automated, augmented, avoided, or added.



Toolkit appendices

Each appendix is the template behind a Field Work artifact. Open a worksheet in the browser, click into any response area to type, and print with background graphics enabled.

1	AI Enablement Council charter template	§3
2	Council membership worksheet	§3
3	Decision-rights matrix	§4
4	AI literacy and readiness survey instrument	§6
5	Foundational training agenda	§8
6	5D departmental interview guide	§9
7	Workflow inventory template	§10
8	Task-decomposition worksheet	§11
9	4A classification worksheet	§12
10	RESPECT assessment template	§13
11	Vendor evaluation rubric	§5
12	Use-case prioritization matrix	§14
13	Pilot charter template	§15
14	Adoption dashboard template	§17
15	Capacity maturity assessment	§18

PRODUCTION NOTES ON THIS EDITION

5D framework (§9): the source draft carried a placeholder for the five "D" labels. This edition fills it verbatim from the established RUDI 5D Workflow Method (Discuss · Document · Diagnose · Design · Develop), as published in the method diagram reproduced as Figure 9.

Survey instrument (Appendix 4): buildable directly from the existing Tally AI Readiness Assessment once its questions are mapped to the literacy / adoption / readiness / risk / opportunity dimensions.

Meridian: currently a nonprofit composite. If the flagship version should read sector-neutral, a parallel private-sector vignette (or a swap to one) is a contained edit — the method steps don't change.