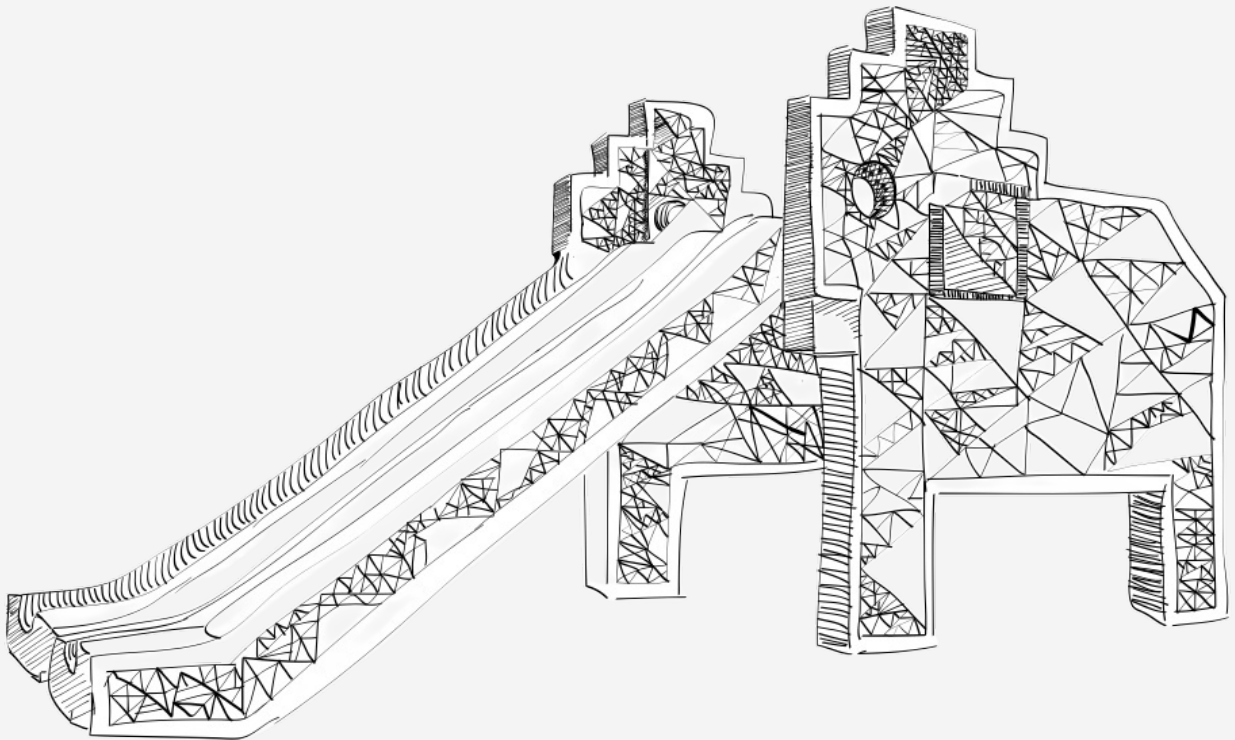


JOB AQ • THE METHOD

The Whole Elephant

How well do you actually work with AI?



JOBAQ.ME

Contents

Why an Elephant?	iii
1 The Problem with “AI Skills”	1
2 The Four Dimensions	2
3 Tailored to Your Job	4
4 Belief versus Behaviour	6
5 Your AI Persona	8
6 The Skills Map	10
7 The Whole Elephant	12
8 For Teams and Leaders	13
9 Where You Are on the Curve	14
10 If-Then Plans: A Short Workbook	16
11 For Organisations: Measure Behaviour, Not Reactions	18
12 The Quarter as a Loop	20
13 Where to Start	22

Why an Elephant?

An old story. Six blind men meet an elephant. One grabs the trunk: “Ah a snake!” One touches the ear: “No, it’s just a fan.” One hugs a leg: “It’s a tree.” The flank is a wall. The tusk a spear. The tail a rope. All very certain of themselves. All totally wrong.

SIX CERTAIN MEN, ONE ELEPHANT

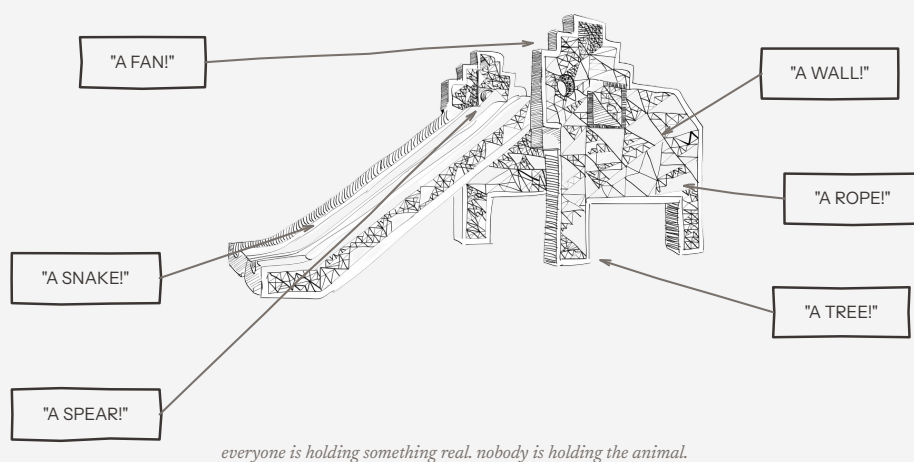


FIGURE 0.1 *Six certain men, one elephant. Everyone is holding something real. Nobody is holding the animal.*

That’s most of us with AI right now. We’ve each touched one part – a few prompt tricks, one good demo, one project that went wrong – and we call our part the whole animal.

Job AQ (jobaq.me) walks you around the whole elephant. It’s a free behavioural assessment of how well you actually work with AI – not what you know about it – built from scenarios drawn from your own job. Think of it as IQ, but for AI.

At the end, you get four things. A score out of 100. Your four dimensions, which show where you’re strong and where your blind spot is. Your AI persona. And a Skills Map showing where AI would pay off in your job.

This book explains what the assessment measures and what to do with the results.

1

The Problem with “AI Skills”

Everyone claims AI skills now. Résumés say “proficient in AI tools.” Profiles collect certificates. Teams report adoption numbers to the board. But almost none of it tells you whether someone can work with AI where and when it matters.

Think about what each of these actually proves. A certificate proves you sat through a course. A tool list proves you opened the tools. A quiz proves you can recall definitions. Each is a part of the animal mistaken for the whole – the certificate is the ear, the tool list is the tail.

Now think about what actually happens at work. AI hands you a confident, polished, partly wrong piece of work. Do you understand why it does what it does? Do you know how to catch the issues? Do you know which part of the task to keep for yourself? And when it goes out under your name and something is off, do you understand what that means for you and your organization?

None of that is prompt engineering. All of it is behaviour.

The gap. Ask people how good they are with AI. They’ll tell you. Then watch what they do. Do the two agree?

Generic tests make it worse. AI doesn’t show up in your work as a general topic. It shows up in your tasks. A generic test never meets your real life and work.

So Job AQ does three things differently. It uses scenarios, not quizzes. It tailors the scenarios to your job. And every scored question focuses on a specific way AI actually works and fails.

And remember – a weak dimension doesn’t mean you’re bad at your job. It means there’s room to grow.

2

The Four Dimensions

Job AQ measures four things. Together they cover the lifecycle of working with AI. Before, during, after, and throughout.

Dimension	What it tells you	When
Understand AI	Do you know where AI helps, and where it makes things worse?	Before
Work With AI	Do you hand AI the right amount of each task?	During
Evaluate AI	Do you catch AI's mistakes in what it produces?	After
Own AI Risk	Do you own the outcome of AI-assisted work?	Throughout

THE WALK AROUND THE LIFECYCLE

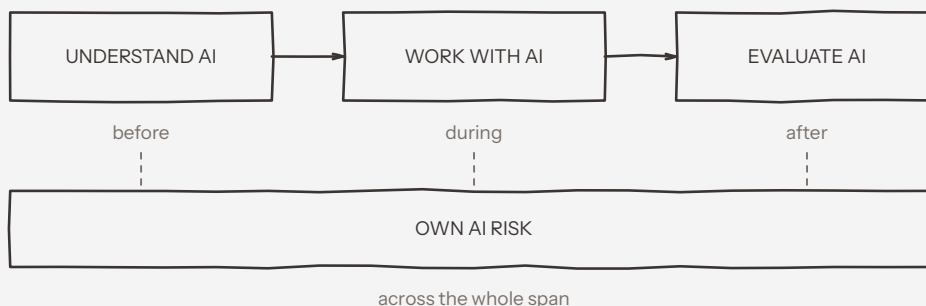


FIGURE 2.1 *The walk around the lifecycle. Understand AI before, Work With AI during, Evaluate AI after, and Own AI Risk running underneath the entire span.*

Someone strong on all four can be handed AI-heavy work without worry. Someone strong on only three will fail in a predictable place – and the assessment tells you where.

Scenarios, not questions. Job AQ doesn't ask you about any of this – it puts you inside the scenario. In one, you decide how much of a real task to hand over, and sometimes the right answer is more AI, not less. In another, you read a polished AI-written deliverable and mark what you'd check before signing off. In another, AI-assisted work has gone out under your name and turned out wrong, and you work through what happens next.

Why not another AI quiz. Most AI assessments measure recognition. Can you define hallucination? Which prompt is better? That's awareness, the lowest level of skill – which is why people who can define hallucination still miss a fabricated benchmark in a real document. Job AQ measures behaviour instead. It also scores both directions of error the same way: handing too much to AI costs you, and so does refusing it for work it does well.

Sixteen signals. Each dimension breaks into four sub-skills. We cross-checked the sixteen against the AI-fluency frameworks already out there – the ones published by the big AI labs, government skills taxonomies, academic instruments. We cover most of what they cover, plus five we didn't find in any of the ones we checked:

- **Manage AI's effort** – knowing when the time AI saves you in drafting gets eaten up by the time you spend checking its work.
- **Spot the irreversible** – recognising decisions that are hard to undo once they spread downstream.
- **Check the substance** – noticing when output sounds confident but says nothing useful.
- **Own the mistakes** – raising AI-assisted errors quickly, and being clear about AI's part in them.
- **Own your voice** – making sure what goes out reflects how confident you actually are, not AI's default certainty.

One caveat. Each run has only about one question per sub-skill, so treat the sixteen as pointers to explore, not precise scores.

3

Tailored to Your Job

Everyone is scored the same way. Nobody sees the same scenarios.

Same scoring for everyone, so scores can be compared. Different scenarios for everyone, so the questions actually look like your work.

TWELVE WAYS OF WORKING

STRATEGIST	BUILDER	ANALYST
STORYTELLER	CLOSER	DESIGNER
OPERATOR	CULTIVATOR	GUARDIAN
STEWARD	EDUCATOR	LEADER

a mode of working, not a job title. pick one, blend two.

FIGURE 3.1 *The twelve job archetypes, three columns by four rows, as they appear at intake.*

Twelve archetypes. Modes of working, not job titles. A “Senior Manager” title tells you seniority and your company’s naming habits. Not whether the week is spent building models, closing deals, or reviewing compliance.

Archetype	The work	For example
Strategist	Plans and sets direction; decides what to build, where to invest	Product manager, strategy consultant, chief of staff
Builder	Designs, writes, and delivers technical work	Developer, data engineer, tech lead
Analyst	Works with data to inform decisions	Data analyst, data scientist, quant
Storyteller	Shapes how a message, brand, or idea is communicated	Marketer, copywriter, content lead, PR
Closer	Drives sales, acquires customers, closes deals	Sales, account manager, partnerships
Designer	Shapes how a product looks, feels, and works	UX, product, or graphic designer
Operator	Keeps day-to-day operations and projects running	Project manager, ops lead, programme manager
Cultivator	Manages how people are hired, developed, supported	HR, recruiter, people lead
Guardian	Protects the org from legal, financial, operational risk	Compliance, internal audit, counsel
Steward	Manages and reports on money and resources	Finance manager, accountant, FP&A
Educator	Helps others build skills and knowledge	Teacher, L&D, trainer, coach
Leader	Sets vision; accountable for outcomes; makes the final calls	CEO, VP, director, founder

The twelve come from large occupational datasets – over a thousand occupations, tens of thousands of task statements – cross-checked against skills taxonomies and AI-exposure research. Each archetype carries its own skills, each rated for how much AI can genuinely help. Those ratings come back later, in your Skills Map.

Most real jobs span more than one archetype. So you can pick a primary and a secondary, with a slider for how your time splits. A product-minded engineer may be a Builder-Strategist. A founder may be closer to a Leader-Closer.

AI also hurts each archetype differently. Take the Analyst. When AI writes a weak strategy memo, people may sense the bullshit and push back on it. But when AI makes up a number in an analysis, people may take it as fact and build decisions on it. And a made-up number looks exactly like a real one, so there is nothing odd on the page to warn anyone. That's why your archetype's scenarios are built around the specific ways AI causes problems in your kind of work, not generic ones. And the archetype changes what you see, never how you're scored. An Analyst's 74 and a Storyteller's 74 mean the same thing.

4

Belief versus Behaviour

Before the scored scenarios, Job AQ asks twelve questions about your habits of working with AI. How you believe you use AI. How much you hand over. How carefully you check. No right (or wrong) answers. And nothing here moves your score.

Then the scenarios watch what you actually do.

Here's why we bother asking. A gap between belief and behavior is more common than we think: people rate their AI fluency higher than their behaviour supports. Self-assessment is unreliable, and confidence about your AI use is a poor guide to your capability. That's the reason the whole assessment is behavioural. What you say shapes your persona. What you do makes your score.

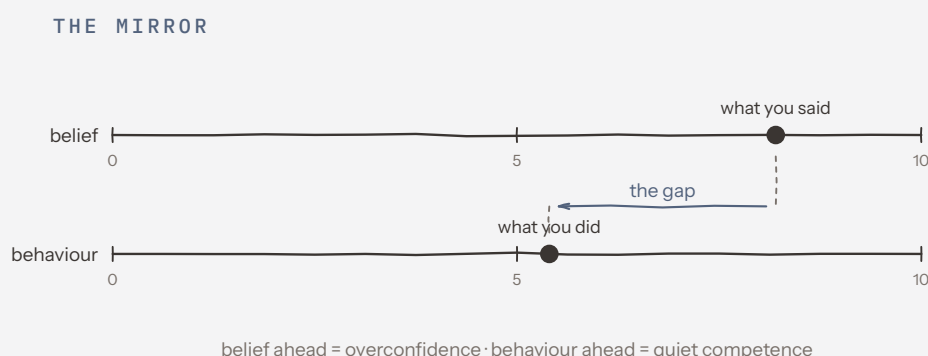


FIGURE 4.1 *Belief and behaviour on the same scale. When belief runs ahead of behaviour the gap is overconfidence; when behaviour runs ahead it is quiet competence.*

Belief ahead of behaviour: overconfidence. You rate yourself a careful verifier. You missed the planted fabrication. Confident people get handed the high-stakes work where blind trust costs most. Which failures did you miss, and what would checking for them look like?

Behaviour ahead of belief: quiet competence. You call yourself a beginner, then delegate, verify, and push back like an expert. People in this quadrant undersell themselves in exactly the meetings where their judgment should count.

Where do you think you are? And where do you want to be? The assessment shows you the gap, and the report gives you a path to close it.

5

Your AI Persona

Your persona comes from a two-by-two on two of your belief axes: how much you hand over to AI, and how much you scrutinise what comes back.

FOUR WAYS OF RELATING TO AI

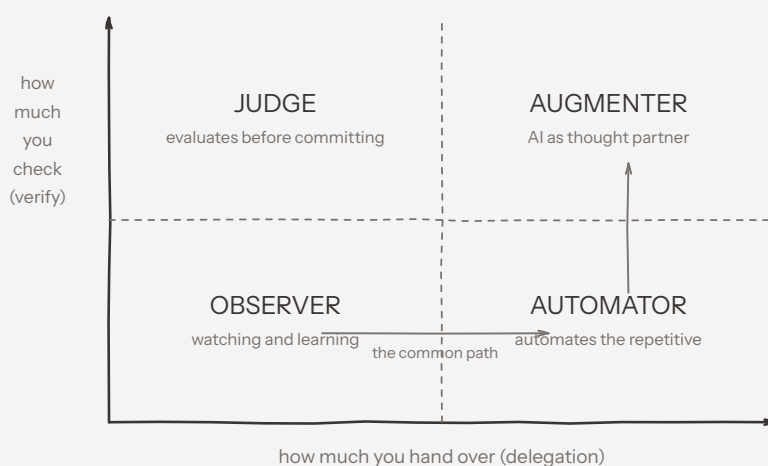


FIGURE 5.1 *The four AI personas on two axes. Delegation across, verification up. Automator, Augmenter, Observer, and Judge each hold one quadrant.*

- **Automator** - high delegation, lighter scrutiny. You automate the repetitive. Your risk is the day the repetitive turns out to matter.
- **Augmenter** - high delegation, high scrutiny. AI as a thought partner. Hand over a lot, check a lot. Sustained high-trust AI work lives here.

- **Judge** – low delegation, high scrutiny. You evaluate before you commit. Deliberate and quality-first, but sometimes too slow to catch the opportunity.
- **Observer** – low delegation, lighter scrutiny so far. Watching and learning. Everyone starts somewhere.

Observer to Automator to Augmenter. Adoption first, judgment next. Or maybe you are the rarer Judge – scrutiny before adoption.

The score indicates how you work with AI. The persona shows how you relate to it. Two people can share a 70 but be different. The Automator should verify more. The Judge should hand more over.

It's a snapshot. It comes from what you said about yourself. And stances change. Retake after a quarter and you may well land in a different quadrant.

6

The Skills Map

The score tells you how well you work with AI. The map tells you where to use it next.

It's built from the tasks you're responsible for, picked at intake, and how much AI is already part of each one.

THE SKILLS MAP

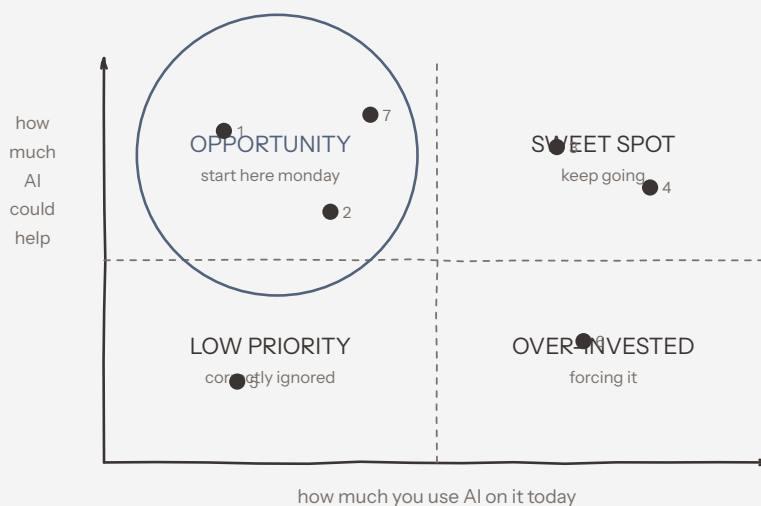


FIGURE 6.1 *The Skills Map. AI relevance up, your adoption across. The upper left quadrant, high relevance but low adoption, is the Opportunity corner.*

The map has two axes: how much AI could genuinely help with each task, and how much you use it on that task today. The relevance rating comes from the same occupational research as the archetypes, computed on the backend. You can't nudge it.

-
- **Sweet Spot** – high relevance, high adoption. Keep going, go deeper.
 - **Opportunity** – high relevance, low adoption. AI could change these tasks and you haven't started.
 - **Low Priority** – low relevance, low adoption. Correctly ignored. Leaving these alone is good judgment.
 - **Over-Invested** – low relevance, high adoption. You're forcing AI into work it doesn't fit. Usually the verification overhead eats the drafting speed-up.

The Opportunity corner is where to start: the highest-relevance task you're not using AI on yet.

7

The Whole Elephant

The results page is the one thing the blind men never got.

Your Job AQ, out of 100. A plain average of the four dimensions. One number to compare with your own retake in three months.

Your identity. Your persona and your job archetype together. An Augmenter Analyst. A Judge Storyteller.

The four dimensions, side by side. Where you're strongest, and where your gap is.

A written summary, not a number dump. A short plain-language read of how you work with AI, generated from what you actually did.

Your Skills Map. Opportunity corner included.

Once you sign-up, you get the full report. The sixteen signals. Your badges. Sign-in is a magic link to your email. No password. Signing up also allows you to save your result, and makes the quarter to quarter tracking possible.

The score is a baseline, not a verdict. The assessment is free and repeatable. Take it. Work differently for a quarter, starting in your Opportunity corner. Take it again. Your history charts each take.

AI fluency is a practice, not a trait.

8

For Teams and Leaders

A team is multiple people holding different parts of the elephant. The engineer has the trunk: AI is a coding tool. The marketer has the ear: a drafting tool. The compliance lead has the tail: a risk register entry. All right about their part. And the team's AI conversations go in circles.

You don't need everyone to know the whole animal. You need a map of who's holding what.

What the map tells you. Where the shared blind spot is. A team weak on Evaluate AI is putting out unverified AI output right now, whatever its adoption numbers say. Who to trust with high-stakes work. The belief-behaviour gap is a list of who sounds fluent versus who is fluent. Who the hidden assets are - the quiet Augmenters whose behaviour is better than their self-rating. And where the AI spend should go - the weak dimension, not another generic workshop.

One quarter.

1. Week 1: everyone takes it. Fifteen minutes each, free, no procurement.
2. Week 2: read the map together. One shared blind spot, one owner, one fix.
3. The quarter: invest where the map points. Everyone starts on their own Opportunity corner.
4. Week 13: everyone retakes. A before-and-after, not a satisfaction survey.

Most teams have never measured this once. Measuring it twice puts you ahead.

The blind men never compared notes. Never walked around the animal. Your team can do both, this quarter, for free.

Take it at jobaq.me. Send it to the other people holding your team's elephant. This book is free too - share it.

9

Where You Are on the Curve

We've explained what Job AQ measures. Now, let's see what one can do with the results. It borrows from a few old frameworks.

The first is the four stages of competence, a model that has been floating around management training since the 1970s. For any skill, you're at one of four stages.

Unconscious incompetence: you don't know what you don't know.

Conscious incompetence: you know your gap, but can't close it yet.

Conscious competence: you can do it, with effort.

Unconscious competence: you do it without thinking.

THE FOUR STAGES OF COMPETENCE

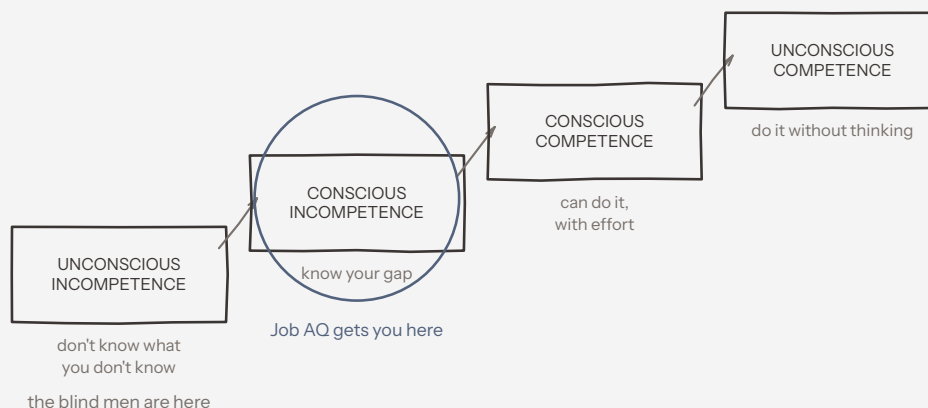


FIGURE 9.1 *The four stages of competence. The blind men sit at stage one. The assessment's job is to get you to stage two.*

The blind men in our story are stage one: certain, wrong, and unaware of it. That's where most of us sit with AI. And it's the dangerous one, because at stage one, confidence tells you nothing.

An assessment has one job: move you from stage one to stage two. Before, you had a feeling about how you work with AI. After, you know which dimension, how far off, and how that compares to what you believed. That's progress, even if the score stings a bit.

10

If-Then Plans: A Short Workbook

Knowing your gap doesn't close it. Good intentions about working better with AI usually fail the way all resolutions fail.

Psychology has a fix for this, known as implementation intentions.

Fancy name, simple idea: plans written as if-then. If this situation comes up, then I do this.

People who write plans this way act on them far more often than people who just set goals, because the plan removes the decision.

One warning that most advice skips: if-then plans only work when you actually care about the goal. A plan written to feel productive changes nothing. So the worksheet starts with the reason.

Step 1, your moment. Think of the last time AI output got past you, or nearly did. The number you almost quoted. The clause you almost missed. If nothing comes to mind, that is worth thinking about too - it may mean you haven't been looking.

Step 2, the cost. Now imagine that same miss or near-miss on the most important thing you'll touch - the board paper, the client deliverable, the hiring decision.

Step 3, your gap. Your miss or near-miss usually points at one dimension. Missed what AI got wrong: Evaluate AI. Handed over the wrong thing: Work With AI. Reached for AI where it couldn't help: Understand AI. Went quiet when it went wrong: Own AI Risk.

Step 4, three plans. The best "if" is your moment from Step 1: *if an AI answer contains a number, then I check it against the source before I use it* · *if an error would be expensive or*

hard to see, then I do the task myself · if someone asks how I checked, then I say what AI did, what I verified, and what changed. Vague plans perform no better than no plan at all.

STEP 1 · MY MOMENT
the last time AI output got past me, or nearly did

STEP 2 · THE COST
the same miss, on the most important thing I'll touch this quarter

STEP 3 · MY GAP
understand AI · work with AI · evaluate AI · own AI risk

STEP 4 · MY THREE PLANS
1 if _____ then I will _____
2 if _____ then I will _____
3 if _____ then I will _____

FIGURE 10.1 *The worksheet. Fill it in here, or copy it somewhere you will see it - the plan only works if the “if” moment finds it.*

At the end of the quarter, question yourself: Did the situation come up? Did you do the thing? Keep the ones that worked, sharpen the ones that didn't, and replace any whose situation never came up. Then retake - your second score tells you if the plan worked.

11

For Organisations: Measure Behaviour, Not Reactions

The classic way to judge whether training works is the Kirkpatrick model, from the 1950s. Four levels. Reaction: did people like it? Learning: can they pass a test on it? Behaviour: do they work differently on the job? Results: did anything the business cares about move?

FOUR LEVELS OF TRAINING EVALUATION

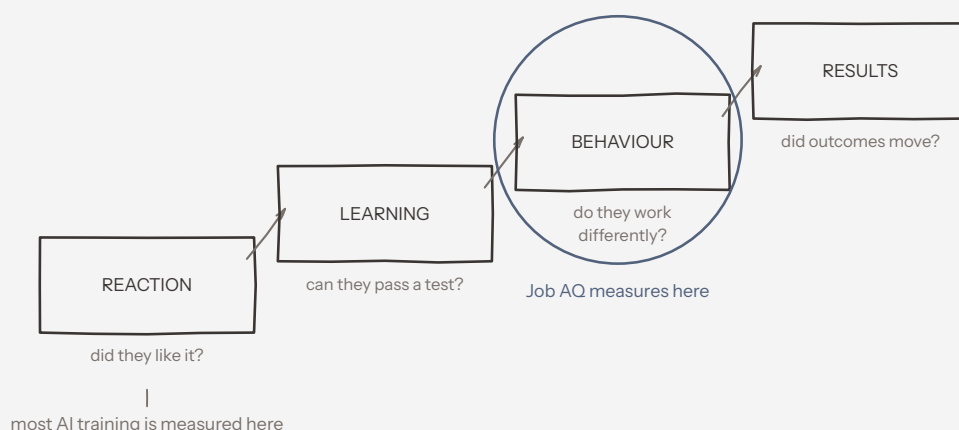


FIGURE 11.1 *Four levels of training evaluation. Reaction, learning, behaviour, results. Most AI training is measured at the first level. Job AQ measures at the third.*

Almost all AI training is evaluated at level one, with a satisfaction survey at the end of the workshop. Some reaches level two, with a quiz. Behaviour – whether anyone actu-

ally works differently on a normal working day – almost never gets measured, because it's hard to measure. Which is a pity, because that's the level where the money and the risk sit.

Job AQ is a level-three measurement you can run in fifteen minutes per person. Before any training, it tells you which behaviour actually needs to change – a team weak on Evaluate AI needs verification practice, not another prompting workshop. After the training, a retake tells you if behaviour moved.

The modern version of the Kirkpatrick model adds one more useful idea: plan backwards. Start from the result you want (fewer unverified AI outputs reaching clients, faster turnaround without more errors), work back to the behaviours that produce it, and only then choose the training. Most organisations run this exactly in reverse – they buy the course first and hope for a result.

The measurement also shows you some things worth acting on. First, some people work with AI better than they rate themselves. They make good reviewers of other people's AI-assisted work, but because they undersell themselves, nobody asks them. Second, teams often share the same weakness. One person weak on Evaluate AI is a gap. A whole team weak on it is a risk, because everyone assumes someone else is doing the checking.

12

The Quarter as a Loop

The oldest improvement trick in management is also the simplest: plan, do, check, act. Make a plan, do the work, measure what happened, adjust, and go around again.

The loop is the point. A one-off push - a course, a policy, a burst of good intentions - fades. And nobody ever finds out if it worked. A loop ends with a measurement. So you always know what to do next.

AI fluency works the same way. And a quarter is about right. Long enough for your if-then situations to come up in real work. Short enough that you actually do the retake.

THE QUARTER AS A LOOP

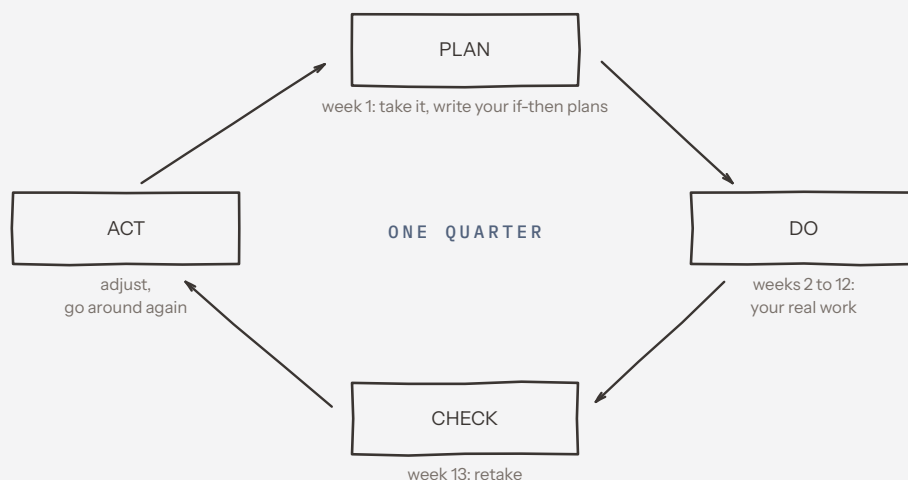


FIGURE 12.1 *The quarter as a loop. Plan in week one, do the work in weeks two to twelve, check with a retake in week thirteen, adjust, and go around again.*

Plan, week one. Take the assessment. Your weakest dimension is the behaviour to work on. The Opportunity corner of your Skills Map is the task to start using AI on. Write your two or three if-then plans from the last chapter. That's the plan – just one behaviour, one task, three sentences.

Do, weeks two to twelve. Work as normal. Nothing new goes into your calendar. Use AI on the Opportunity task, with a light check. And follow your if-then plans whenever the situations come up.

Check, week thirteen. Retake. Fifteen minutes. The difference between the two scores tells you if your behaviour actually changed.

Act. Score moved? Pick the next dimension and go again. Score didn't move? Your plans were probably too vague. Make the "if" more specific, make the "then" smaller, and run the same dimension again.

Same loop for a team, just with more people. Week one, everyone takes it. Week two, look at the results together and pick one shared weakness, one owner, one fix. Spend the quarter on it. Week thirteen, everyone retakes. Now you have a clear before and after.

Nothing here is about trying to be clever.

It's about being simple enough to repeat. And improve.

13

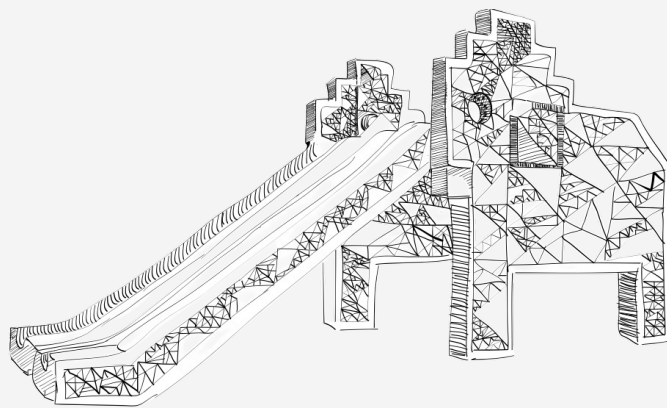
Where to Start

Take the assessment at jobaq.me. Fifteen minutes, free. Everything else in this half builds on it - the if-then plans, the team measurement, the quarterly loop.

Once done, you've walked around the animal once - more than the blind men ever did. Come back in a quarter and walk it again.

And if this book helped, share it. It's free. The other people holding your team's elephant still think it's a snake, a fan, and a tree.

*We have each touched one part
and decided we know the whole animal.*



Job AQ is a free, fifteen-minute behavioural assessment
of how well you actually work with AI.

Not what you know about it. What you do with it.

JOBAQ.ME