

Asking the Questions

Sponsoring, governing, and leading people through a technology nobody has mastered

Parts One and Two gave you the vocabulary and the judgement. This final guide is about the three things you will actually be asked to do: sponsor an AI initiative, govern AI in your function, and lead people through a change that is making them simultaneously curious and quietly anxious.

The reframe for this territory, and it is different from every other in this library: there is no expert who owns the answers. Vendors are selling, consultancies are generalising, and technologists own capability rather than consequence. Which means the questions are unusually valuable and unusually available — the director who asks them well is not keeping up with this field, they are one of the few people bringing discipline to it.

Everything here is usable this quarter, and none of it requires technical skill.

Sponsoring an AI initiative that survives production

Promotion criteria, agreed in advance

The idea: Before any pilot begins: what measured result graduates it to production, by when, and what result stops it. Written down at approval, not negotiated afterwards.

How to use it: This single discipline is the antidote to the proof-of-concept graveyard, and it is entirely within a sponsor's gift. It also builds the funding credibility the finance series described: a sponsor who has demonstrably stopped something is trusted with larger things.

The question to ask: *"What result promotes this, what result stops it, and who decides on which date?"*

Benefit ownership — with an address

The idea: The finance-series discipline applied to AI's hardest case: which budget line changes, by how much, when — and who has agreed to be accountable for it. Time saved is not a benefit until it is converted into capacity redeployed, cost removed, or volume absorbed without hiring.

How to use it: Making this explicit at approval prevents the productivity-claim-nobody-can-find pattern, and it forces the honest version of the case. Sometimes the honest version is 'this improves quality and speed rather than cost' — which is a perfectly good case, stated truthfully.

The question to ask: *"If this works, what specifically will be different that someone can point to — and who owns making that happen?"*

Measure on your own work, briefly and early

The idea: Before the large business case, run a two-week measurement on real tasks with real people: time before, time after, correction rate. Rough numbers are sufficient and vastly more persuasive than vendor benchmarks.

How to use it: This is the cheapest credibility available in the entire AI conversation. Bringing measured evidence into a room running on assertion changes who the room listens to — and it protects you from sponsoring something that measures badly.

Design the control before the capability

The idea: Decide the checking model — who reviews what, at what volume, with what information — as part of the design, not as a compliance afterthought.

How to use it: The oversight-that-isn't pattern is created at design time, by assuming review capacity that nobody costed. Sponsors who ask 'how many items per reviewer per hour, and is that realistic?' at approval are the reason some initiatives survive audit and others don't.

Governing AI in your function

Whatever your organisation's central position, your function needs four things settled locally — and if the centre hasn't provided them, sponsoring them yourself is a fast, visible contribution:

- **The one-page rules.** Approved tools, data boundary, what always gets checked, who is accountable for output. One page, communicated once properly, refreshed when it changes. If this doesn't exist, your team is deciding it individually today.
- **A fast sanctioned path.** Shadow AI is a friction signal (Part One). If someone in your function wants to try something, there must be a route that takes days rather than quarters — otherwise governance is theoretical and practice is invisible.
- **A use-case triage.** Three buckets: proceed (internal, augmentation, cheap errors), review first (customer-facing, or touching personal data), and stop — anything affecting hiring, promotion,

credit, or eligibility, which carries real regulatory weight and belongs with legal before anyone builds a prototype.

- **A view on context, not just tools.** Approved tools are the easy half. The harder half is what they are allowed to see and whether it is current, because that is what determines whether the answers are any good. If your function's AI is grounded in documents nobody has updated in two years, the tool is not the problem. Ask who owns the material, not just the licence.
- **A record of what you tried.** What was piloted, what was measured, what was adopted or stopped and why. Cheap to keep, and it turns scattered experimentation into institutional learning — the difference between a function that compounds and one that repeats itself.

03

Agentic AI — governing action, not just output

Agents change the governance question from 'is the output right?' to 'what did it do, and can we undo it?' The framework:

- **Scope of action, written down.** Precisely which actions, in which systems, within what limits. Start shorter than the vendor proposes; expand on evidence.
- **Establish it is an agent at all.** Before applying agentic governance, ask what decides the next action, the model or a rule somebody wrote. Plenty of what arrives labelled agentic is scripted automation, and knowing which you have determines whether you need scope limits and reversibility rules or simply ordinary change control.
- **Reversibility as the selection criterion.** Automate first where mistakes are visible and undoable. Anything that touches money, contracts, employment, or a customer relationship keeps a human checkpoint until the evidence is overwhelming — and even then, be slow.
- **Detection of silent failure.** The dangerous agent error is not dramatic; it is quiet and repeated — the escalation never raised, the record wrongly updated, the case closed that shouldn't have been. Ask specifically: how would we know if it did nothing, or did the wrong thing consistently, for a month?
- **Accountability, named.** When an agent acts, a person owns the outcome — by name, before deployment, not discovered during an incident review. 'The system did it' has never once survived a regulator, a customer, or a board.

04

Leading people through it

The part of this territory that is unambiguously yours, and the part most senior leaders handle badly by saying nothing.

- **Name the anxiety plainly.** Your people are wondering whether this replaces them. Silence is heard as confirmation. You do not need certainty to be honest: what you know, what you don't, what you'll tell them when you know it. Directors who address it directly find their teams experiment openly rather than secretly — which is also better governance.
- **Be explicit about judgement.** AI drafts; humans decide and are accountable. Said clearly and repeatedly, this both protects quality and tells people where their value sits.
- **Teach the agreement trap.** These tools tend to side with whoever is asking, which means an AI that endorses a business case is not corroboration. Make it a stated expectation in your function that anyone using AI in analysis asks it for the counter-case, not for approval. It costs nothing and it prevents a whole class of confidently wrong work reaching your desk.
- **Watch the capability question.** AI raises the floor and can quietly erode the practice that builds judgement — juniors skipping the productive struggle, experienced people out of practice. Decide which work stays human because that is where skill is built, and say why. Almost nobody at your level is asking this question, which is precisely why asking it marks you out.

- **Build fluency, not just access.** Licences without capability produce the dabbler's plateau at enterprise scale. A short internal session where people show each other what actually worked beats any procurement decision for return, and costs an hour.

05

In the room — the twelve questions

The distillation — director's edition. Asked plainly in an AI conversation, these read as the most senior person present:

1. "Show it working on our data, with our messiness, end to end."
2. "Where did that productivity number come from — our work, or a vendor benchmark?"
3. "What's the error rate on production-like data — and what does checking cost at full volume?"
4. "If this works, which budget line changes — and who owns making that happen?"
5. "What result promotes this pilot, what result stops it, and who decides when?"
6. "How many items per reviewer per hour — and is that oversight or decoration?"
7. "What can the agent do without a human, and how would we know if it failed quietly?"
8. "Does this touch hiring, credit, or eligibility — and has legal seen it?"
9. "What does this cost at full adoption, and what are we switching off to fund it?"
10. "What is this doing to what our people can do without it?"
11. "Is this genuinely agentic, or an automation at agentic pricing — and what decides its next action?"
12. "Is the analysis independent, or is the tool agreeing with whoever prompted it?"

The end of nodding through the AI item

Three guides ago, the AI agenda item was something to survive. Now: you have the vocabulary (Part One), you read business cases, vendor claims and failure patterns the way executives read forecasts (Part Two), and you own the sponsor's disciplines, the governance minimum, and the people conversation (Part Three).

You are not an AI expert. In this territory, uniquely, almost nobody is — which is exactly why the questions matter more here than anywhere else in this library. The vendors own the claims, the technologists own the capability, and the consequence is unowned until a leader picks it up.

And the closing observation, which by now will be familiar: what looks like an AI problem is almost always a structure problem in new clothing. The pilot graveyard is an absence of promotion criteria. The unfindable productivity is unowned benefit. The oversight that isn't is a control designed without capacity. Shadow AI is friction your people couldn't get resolved. Organisations that adopt AI well will not be the ones that bought earliest — they will be the ones that were already clear about who decides, who owns outcomes, and how they learn. Which means the most valuable thing you can bring to this conversation is not technical knowledge at all. It is the clarity you bring to everything else.

John Obidipe is a business transformation specialist and Gallup-Certified CliftonStrengths coach. He works with founders and leadership teams on the structural side of growth — the part the numbers can only hint at. · coach@catalystgrowthcoach.co.uk · catalystgrowthcoach.co.uk