

Five questions your board should be asking about AI oversight

Why each one matters, and what a good answer sounds like

Julie Hendry FBCS · juliehendry.com · August 2026

In May 2026 the European Union agreed to delay the AI Act's high-risk obligations, including the Article 14 requirement for effective human oversight. Those obligations now apply from 2 December 2027 for standalone high-risk systems, and from August 2028 for AI embedded in regulated products. Regulators do not defer enforcement because the rules stopped mattering. They defer because the organisations being regulated could not have complied. The delay is a readiness window, and what fills it is not more technology governance, which the market already measures exhaustively, but the layer almost nobody measures: whether the people expected to oversee AI systems can still do it. The five questions below each probe one dimension of that human layer. They are confrontations, not rating scales, and the most honest answer to most of them is "no" or "I don't know". That is the point. Ask them now, while there is still time to change the answer.

1. Has anyone in your organisation been formally assessed for their ability to critically evaluate AI-generated output? Not trained. Assessed.

Why it matters. Article 14 requires that a person overseeing a high-risk system can detect when something is going wrong, interpret the output correctly, and decide to override it. Training records show attendance. They do not show whether that capability exists, or whether it has survived sustained AI use. A Harvard Business School and Boston Consulting Group study of 244 consultants found that roughly 27 per cent had become what the researchers call Self-Automators: they delegated entire workflows to the AI and developed neither real AI skill nor the domain judgment their role assumes. Separately, Vicente and Matute showed in *Scientific Reports* that people who worked with a biased AI reproduced that bias in their own later decisions, even after the AI was removed. The ability to evaluate AI output is measurable, and it degrades. If nobody has been assessed, your oversight regime rests on an assumption.

WHAT A GOOD ANSWER SOUNDS LIKE

A specific answer names who was assessed, how, and what changed as a result. "We assessed the reviewers on our highest-risk system this spring, found two who could no longer catch a planted error, and redesigned the role." Anything that begins "everyone completed the training" is answering a different question.

2. Can you name the last time a human overrode an AI recommendation in your organisation? And what happened next?

Why it matters. The override is where oversight is either real or decorative. You can build the cleanest override button in the industry, and the obligation still fails if the person holding it has stopped being able to

tell when to press it. A control that has never fired is either governing a perfect system or not governing at all, and no organisation running AI in hiring, credit, healthcare triage or fraud detection is running a perfect system. The second half of the question carries as much weight as the first. If the person who overrode was supported and the case examined, human judgment is genuinely in the loop. If the override was quietly reversed, or the person learned not to do it again, the policy describes something that does not happen. Testing one documented override tells you more about your Article 14 readiness than the policy that describes it.

WHAT A GOOD ANSWER SOUNDS LIKE

A named instance with a date and an outcome. "In March one of our analysts rejected the model's recommendation, her manager backed her, and the case fed into the next model review." If the honest answer is "I would have to check whether that has ever happened", that is your finding.

3. If your AI governance self-assessment scored you 4 out of 5, would you trust that score?

Why it matters. The gap that decides Article 14 readiness is the gap between documented governance and observed practice: between the oversight that is written down and the oversight that happens. Self-assessments live on the documented side of that gap. Meanwhile the people accountable are telling a different story. An IBM Institute for Business Value survey of 2,000 CIOs and CTOs, published in June 2026, found 67 per cent held accountable for AI systems they do not fully control, and 77 per cent saying AI adoption is already outpacing their organisation's governance capability. Executives report losing control while self-assessments come back reassuring. Both cannot be true. A high self-score is evidence of who wrote the assessment, not of what a regulator would find, and an organisation that cannot see the difference has a governance problem the score will never surface.

WHAT A GOOD ANSWER SOUNDS LIKE

"No, not on its own", followed by what you do about it: testing the risk register against operational reality, comparing what leadership believes with what frontline AI users experience, or bringing in someone with no stake in the number. Trusting the score outright is itself an answer.

4. Do you measure what AI has done to your team's confidence, autonomy, and judgment? Or only what it's done to their productivity?

Why it matters. Adoption metrics measure the wrong layer. In February 2026, Acemoglu, Kong and Ozdaglar described the underlying dynamic as knowledge collapse: as AI substitutes for human cognitive effort, people rationally stop investing in the knowledge and judgment that the effort used to build. The longer the exposure, the thinner the capability, and none of it appears in a productivity dashboard, which will show gains for exactly as long as the erosion continues. There is also a clock on this one that no regulator can extend. Every month of unmeasured deployment lowers the baseline you will eventually measure against. An organisation that waits until 2027 to look will be measuring an already diminished workforce and calling it normal. The honest

baseline of what your people could do before AI reshaped their work has a short shelf life, and it is shortening now.

WHAT A GOOD ANSWER SOUNDS LIKE

At least one measure of human capability tracked alongside the output metrics, dated and repeated. "We baselined independent judgment on one high-risk workflow in June and we re-measure quarterly." A list of productivity gains, however impressive, is a "no".

5. If an AI-assisted decision caused serious harm tomorrow, could you name the accountable individual within 30 seconds?

Why it matters. The 30 seconds is not theatre. Hesitation is diffusion, and diffusion is what accountability looks like just before it fails. The question a regulator or a board will eventually ask is not who appears on the org chart but who would answer for the decision under pressure, and whether that person would agree they are answerable. The IBM finding is stark here: 67 per cent of the executives surveyed are held accountable for AI systems they do not fully control. Accountability without the capability to exercise it is not oversight. It is exposure, for the individual and for the organisation. Committees, shared ownership and distributed responsibility all feel safer, and all produce the same result when harm arrives: a room full of people explaining why it was not quite their decision.

WHAT A GOOD ANSWER SOUNDS LIKE

A name, said quickly, plus confidence that the named person knows it, agrees with it, and has the authority and access to act on it. "It depends on the system" is acceptable only if you can name someone for each system that matters.

What to do with your answers

This is a mirror, not the assessment. There are no scores here, no maturity levels and no rankings, because scoring an organisation on questions like these takes evidence, not self-report. What the mirror shows is enough: if most of your answers were "no" or "I don't know", you are in the majority, and the window before December 2027 is what you have to change that.

These questions come from CHART, the openly published framework for assessing the human layer of AI governance (juliehendry.com/CHART/framework.html). If you want to go one step further, the free self-assessment at juliehendry.com/CHART/self-assessment.html takes a few minutes and gives you a directional read across all five dimensions. And if you want a conversation about what your answers mean for your organisation, the contact page at juliehendry.com is the place to start it.

SOURCES

Digital Omnibus and AI Act high-risk timeline: European Parliament approval 16 June 2026; Council formal adoption 29 June 2026; deferral to 2 December 2027 (Annex III) and August 2028 (Annex I).

AI use archetypes (Cyborgs, Centaurs, Self-Automators): Randazzo et al., "Cyborgs, Centaurs and Self-Automators", Harvard Business School Working Paper 26-036, December 2025.

Executive accountability and governance lag: IBM Institute for Business Value with Oxford Economics, survey of 2,000 CIOs and CTOs across 33 countries, June 2026.

Bias transfer and persistence: Vicente, L. and Matute, H. (2023), "Humans inherit artificial intelligence biases", *Scientific Reports*, DOI 10.1038/s41598-023-42384-8.

Knowledge collapse: Acemoglu, Kong and Ozdaglar (2026), NBER Working Paper w34910, DOI 10.3386/w34910.

Julie Hendry's research on the cognitive impact of AI on organisational decision-making: SSRN, DOI 10.2139/ssrn.6786758.

© 2026 Julie Hendry. CHART (Cognitive Health Assessment for Responsible Technology) is the work of Julie Hendry. juliehendry.com