

The Board Questions on AI

Twelve questions, and the answer that should worry you

2026 edition · Published 26 August 2026 · Next edition August 2027

Free to copy, attach to a board pack and use. No registration, no licence, attribution appreciated.

CAPABILITY

1 · What can we no longer do without these systems?

Worrying answer: "Nothing, we could always go back to how we did it before." Nobody has checked. Capability does not disappear on a schedule anyone tracks, so capability debt is only visible when something fails or a novel question arrives.

2 · Who could tell if this system were wrong?

Worrying answer: a function rather than a person, or a person who could not have produced the work themselves. Fluency carries no signal: a plausible wrong answer looks exactly like a correct one to anyone who cannot independently evaluate the content. See who owns verification.

3 · How are we keeping people good at the work we have automated?

Worrying answer: "That is the point of automating it." Reasonable, until you need someone to supervise it, override it, or do it when it fails. Nobody budgets for maintaining capability in automated work, because it looks like paying people to do something a machine does faster.

EVIDENCE

4 · Does this arrangement beat the better of the human alone or the system alone?

Worrying answer: "We have not measured that." A meta-analysis of 370 effect sizes across 106 experiments found human-AI combinations performing worse on average than the stronger party alone, with the losses concentrated in exactly the arrangement most organisations have installed. A pairing that has never been tested against that baseline may be subtracting.

5 · Are we measuring usage or capability?

Worrying answer: seat counts, licences, prompt volumes, adoption percentages. Those measure activity. Nothing in them says anyone got better at anything. See usage theatre.

6 · Whose evidence are we relying on, and what does it not show?

Worrying answer: a vendor case study, or a consultancy survey from a firm selling the remedy. Ask what the study does not support. Ask whether an average is concealing opposite effects on different people: in the radiology evidence, the effect of AI assistance ran from strongly positive to strongly negative between individual readers and was not predicted by experience.

ACCOUNTABILITY

7 • Who is accountable, by name, for each stage of this work?

Worrying answer: the team, the function, or a name that appears only after something goes wrong. Accountability assigned retrospectively is attribution, and it behaves very differently under pressure. The Delegation Boundary Map makes the gaps visible in about ninety minutes.

8 • How many times has anyone overridden the system this quarter?

Worrying answer: zero. That evidences an untested right rather than a good system, and possibly a workload that makes real scrutiny impossible. Article 14 of the EU AI Act now requires that overseers of high-risk systems can decide not to use them or disregard their output, which is a capacity, not a permission.

9 • Which decisions are we now approving rather than making?

Worrying answer: a defensive one. This is the personal version of the question and the most uncomfortable, because it applies to the board as much as to anyone below it. Executives reading machine summaries instead of source material are making a different kind of decision than they think they are.

THE PIPELINE

10 • Where are our next senior people coming from?

Worrying answer: "We will hire them." Everyone is planning to hire them, from a pool that is being drained by the same automation. The work that built senior judgement is the work most easily automated. See the missing rungs and synthetic seniority.

11 • What is our AI literacy programme actually producing?

Worrying answer: a completion rate. Completion is not capability. Article 4 of the EU AI Act has required a sufficient level of AI literacy since February 2025, at every risk tier, and enforcement began in August 2026. It asks about understanding relative to context and to the people affected, not about tool familiarity. See what AI literacy means for leaders.

THE ONE THAT MATTERS MOST

12 • What evidence would make us reverse this?

Worrying answer: silence, or "we would look at it if something went wrong." A deployment with no stated reversal condition is a commitment rather than a decision. Asking this before rollout is cheap. Asking it afterwards is a crisis meeting.

Every claim behind these questions is graded at thesuperskills.com/research/evidence, with what it does not prove stated alongside what it does. The next edition is due August 2027. If a question in it changes, the change will be shown rather than made quietly.