

We talked to 100+ C-suite executives about AI. Here are their 10 biggest questions.

Not the ones that make headlines. The ones that come up in the room, once the demo is over and someone has to actually own the outcome.

1. How do we prove an agent is good enough to trust?

Traditional QA assumes deterministic output. Our test suites weren't built for software that answers differently twice

2. An agent just did eight hours of work. How does a person verify it in ten minutes?

Without rubber-stamping it.

3. Our data lives in forty systems and none of them agree.

What actually has to get fixed before AI works, and what can we leave broken?

4. What is safe to put in, and who decides?

Especially in legal, finance, HR, and anything a regulator can subpoena.

5. How do we let the front line build without letting them near the systems that can't break?

Everyone wants citizen development until someone automates against the general ledger.

6. Who owns an agent after it ships?

Whose budget, whose headcount, whose phone rings at 2am when it drifts.

7. How do we control spend without turning every request into an approval queue?

And how do we attach a dollar of value to a dollar of consumption.

8. How do we move a multi-thousand person org from AI literate to AI first?

Adoption is easy in the top decile. The other 90% is the actual problem.

9. How do we bet big on one lab without being locked in two years from now?

Model agnostic sounds obvious until you price what it costs to stay that way.

10. How do we capture what our best people know?

The judgment that never made it into a single SOP, so that agents inherit it instead of relearning it.

“

Almost none of these are model questions. They're operating model questions. The companies pulling ahead worked that out early.

”