

AI organisation readiness

AI is an amplifier, not a fixer.

It takes whatever your organisation already does and does more of it, faster.
Most companies asking which AI tools to buy should ask these six questions first.

1 Decide

- 01 Prioritisation clarity
- 02 Ownership and review capacity

2 Learn

- 03 Feedback loops
- 04 Documentation and knowledge

3 Scale

- 05 Technical and data foundations
- 06 Change tolerance

01 Prioritisation clarity

The team knows what matters and can say no to everything else.

✓ What good looks like

- The top three priorities can be stated without hesitation, and the answers agree across the team.
- There is a single backlog, not competing lists spread across tools and teams.
- New requests meet a working mechanism for saying no, and people can recall it being used.

⊗ Red flags

- Everything is labelled P1, so the label carries no information.
- The roadmap shifts weekly in response to the loudest voice.
- Quick wins keep jumping the queue and displacing planned work.

Field test

The three-answer test. Ask three people, separately, to name the top three priorities this quarter.

Reveals: If you get three different answers, you have your answer.

If this is weak

AI speeds up delivery. If the priorities are wrong, you just build the wrong things faster.

02 Ownership and review capacity

One named owner per outcome, with decision-making authority.

✓ What good looks like

- A person, not a team, owns each significant outcome, and everyone can name them.
- Owners decide within their remit without assembling stakeholders.
- Senior people have time to review work properly, not just wave it through.

⊗ Red flags

- Committees own things, so no one is really responsible.
- Decisions require five stakeholders and two weeks of alignment.
- Everything waits for one person, because no one else feels able to decide.

Field test

The failed-project test. Pick a recent project that went wrong and ask who owned it. **Reveals:** A team name means ownership is diffuse. Multiple names mean it is contested.

If this is weak

AI makes producing work easy. Someone still has to review it. If no one owns it, it piles up unread.

03 Feedback loops

Retrospectives happen, produce actions, and the actions change behaviour.

✓ What good looks like

- Retros run on schedule, including when things are busy, which is when they matter most.
- Actions have owners and get completed, and the next retro checks that they were.
- Problems surface without blame, and customer feedback reaches those building the product.

⊗ Red flags

- Retros are the first thing skipped under pressure.
- The same issues appear quarter after quarter.
- Post-mortems focus on who, not why, so people stop surfacing problems.

Field test

The action audit. Pull the actions from the last three retros and count how many were done.

Reveals: A retro that produces uncompleted actions is a meeting, not a team that learns.

If this is weak

Getting value from AI needs evaluations and learning. AI makes mistakes. If no one looks at what went wrong, the same mistakes keep happening and there is no progress.

04 Documentation and knowledge

How the organisation works is written down, not held in one person's head.

✓ What good looks like

- A new joiner can learn core workflows without asking five people.
- Decisions are recorded with their rationale, not just their outcome.
- Architecture is understood by more than one person, and written down.

⊗ Red flags

- Critical knowledge is held by one person.
- Documentation is six months out of date and everyone knows it.
- Onboarding is “shadow someone for a few weeks”.

Field test

The bus-factor count. For each critical system, ask who else could run it tomorrow. **Reveals:** Every answer that comes back ‘nobody’ is a dependency AI cannot help with.

If this is weak

AI is only as good as the context you give it. If how things work lives in people's heads, there is nothing to point it at.

05 Technical and data foundations

Trusted data, trusted deployments, and clear rules on what AI is allowed to do.

✓ What good looks like

- One source of truth for the metrics that leadership reviews.
- Automated deployments with rollback, and a test suite engineers trust.
- An explicit rule for what AI may touch unattended and what needs sign-off.

⊗ Red flags

- Two teams give two different numbers for the same metric.
- Production errors are found by customers, not by monitoring.
- “We will fix the tech debt after this release”, for two years running.

Field test

The metric trace. Take one number leadership reviews and trace it back to raw data. Count the manual steps and name who would notice if one broke. **Reveals:** Fragile manual pipelines that any AI application will inherit.

If this is weak

AI answers always look right. Built on bad data, they are wrong at scale, and they still look right.

06 Change tolerance

The organisation has adopted new ways of working before, and made them stick.

✓ What good looks like

- Recent tool and process adoptions have stuck, and the organisation can say why.
- Failed experiments are shared openly and examined, not buried.
- Leaders admit being wrong, which makes it safe for everyone to learn in public.

⊗ Red flags

- The last three rollouts were quietly abandoned.
- “We tried that and it did not work” ends conversations.
- Waiting out new initiatives is a viable strategy, and everyone knows it.

Field test

The adoption ledger. List the last five changes you attempted and count how many stuck past six months. **Reveals:** Fewer than three, and AI is just another change that will not stick.

If this is weak

AI adoption is a change problem, not a tooling problem. Push it on a team that is not ready, and ‘we tried AI and it did not work’ joins the graveyard of past initiatives.

Score each 1 to 4

Your AI ambition is set by your **lowest** score.

1 **Blocking**
Visible in a single afternoon of interviews.

2 **Inconsistent**
Works on paper. Breaks down under pressure.

3 **Functional**
Reliable with minor gaps, generally recovered.

4 **Embedded**
Strong practice the organisation would defend.

Below 2 on any dimension, **close that gap first.**

At that level AI amplifies the dysfunction rather than routes around it. A contained pilot is a fine way to surface these gaps. Betting core operations on it before closing them is not.

Try running the full diagnostic.
Let us know if we can help.