

Enterprise Generative AI Adoption

The 3 fundamentals most programmes miss

Who this is for

If you're a CIO, CTO, or CDO responsible for moving generative AI beyond pilots, this guide is for you.

What you'll learn

- How to make GenAI outcomes measurable inside real workflows
- How to design trust using human-in-the-loop as a control system
- How to build AI assistants and agents that actually get adopted
- What to prioritise on Monday to move from experimentation to scale

The real problem with enterprise GenAI adoption

Most enterprise conversations focus on three things:

- Security
- Data quality
- ROI

Those matter. But they are not what determines whether GenAI actually gets used.

The reason most programmes stall is simple: GenAI is treated like a software purchase, not an operating capability.

So while platforms are secured and business cases are approved, adoption quietly fails.

The gap sits in three areas:

- Outcomes are unclear
- Trust is unmanaged
- Tools are too generic

Fix those, and adoption becomes predictable.



1. Make AI outcomes measurable, not mystical

GenAI creates value in ways that are harder to measure:

- Faster thinking
- Better drafts
- Reduced rework

Without clear measurement, ROI becomes belief, not evidence.

What works

Define outcomes *inside workflows*, not in reports.

Example patterns

Scenario-based productivity

- Define a specific task (e.g. architecture design)
- Compare time and quality with and without AI

Embedded workflow metrics

- Track iteration count, cycle time, rework
- Measure before vs after AI introduction

Agent usage as time-shift proxy

- If an agent saves 10–15 minutes per task at scale, that's measurable value

Checklist: AI workflow measurement

For each use case, define:

- Workflow in scope
- Baseline (time, cost, errors)
- Target outcome
- Measurement owner
- Review cadence

About 848

848 helps enterprises turn generative AI from isolated experiments into scalable, governed operating capability.

2. Treat human-in-the-loop as a control system

AI introduces a new dynamic:

- Some users under-trust and disengage
- Others over-trust and stop checking
- Both create risk.

Human-in-the-loop is not optional.

It is a **control mechanism for trust, quality, and accountability.**

What works

Train for iteration, not prompting

Users must challenge and refine outputs

Embed review into workflows

Certain steps require mandatory human sign-off

Limit AI authority externally

AI can inform, not decide, in high-impact scenarios

Define the trade-off explicitly

For each workflow:

- **Improve:** speed and productivity
- **Accept:** higher variance in outputs
- **Control with:** review, sampling, guardrails

Checklist: Trust & Control Design

- Where is human review mandatory?
- Where is sampling sufficient?
- Where must AI not act independently?
- What are the escalation paths?

3. Design AI agents for real operations

Most pilots succeed in controlled conditions. Operations don't.

Giving people a generic AI tool rarely works.

People don't want to be prompt engineers. They want their work to be easier.

What works: Scenario-based agents

Document-specific agents

- Structured outputs for specific document types

Portfolio-aware assistants

- Grounded in your services, products, and assets

Role-specific SME agents

- Designed around real responsibilities and constraints

Design principles

- Defined persona and tone
- Standard output structure
- Clear guardrails
- Specific task focus

Checklist: Agent Design

- What workflow is this agent for?
- What output does it produce?
- What can it NOT do?
- Who owns the final decision?

What I'd do on Monday

If you're moving beyond pilots, start here:

1. Define measures for three live use cases

No more "it feels useful"

2. Map human checkpoints

Make review part of the process

3. Convert one tool into a real agent

Pick a specific workflow and design for it

4. Set external boundaries

Be explicit about what AI can and cannot say

5. Write one trade-off statement

For your most important workflow:

Improve: _____

Accept: _____

Control with: _____

The decision in front of you

GenAI adoption is already happening.

The real choice is:
Fragmented experiments
or
A governed, scalable capability

Move too fast → **risk and rework**
Move too slow → **lost learning**

The answer is not speed or safety.
It's intentional design.