

# The AI Autonomy Decision Guide

Generative AI creates. Agentic AI acts. The useful question is not which is better, but how much autonomy a given task actually deserves. Run a task through this before you hand it over.

## Step 1. Name the task

### What is the task, in one sentence?

Be concrete. 'Draft product descriptions' not 'improve content'.

## Step 2. Score it on four axes

Higher scores mean the task is safer to hand to an agent that acts on its own.

### Reversibility

If it goes wrong, we can undo it easily and cheaply.

1

2

3

4

5

### Blast radius

A mistake affects little, and nothing customer-facing at scale.

1

2

3

4

5

### How well understood

The workflow is repetitive and we already know the right answer looks like.

1

2

3

4

5

### Oversight available

Someone will actually see the output before it matters.

1

2

3

4

5

**Reading your score:** 16 or above, safe to give an agent real autonomy. 10 to 15, let it act but approve each step. Below 10, keep it generative: let the model draft, and you decide.

## Step 3. Set the dial

| Level           | What the AI does                                        | Use for this task? |
|-----------------|---------------------------------------------------------|--------------------|
| Generative      | Produces a draft and stops. You review everything.      |                    |
| Tool-assisted   | Uses a tool or two, you approve each step.              |                    |
| Semi-autonomous | Strings actions together, checks in at decision points. |                    |
| Autonomous      | Runs end to end, surfaces exceptions only.              |                    |

## Step 4. Before you switch it on

### What permissions does it need?

List only what the task requires. Anything more is unnecessary risk.

### What is the worst thing it could do?

Write it down. If you cannot live with it, narrow the permissions.

### How will you know it failed?

An agent that is wrong about its own success is the dangerous case.

**Start small.** Give agents low-stakes, reversible tasks first and widen their reach as you learn to trust them. The reward is bigger with autonomy, but so is the blast radius.