

Agentic AI on one page.

What an agent is, when you need one and what keeps it in check. Six blocks to save and pass on, each with a path into the full glossary.

01 What an agent is

03 The building blocks

05 Autonomy in levels

02 Script, RPA or agent?

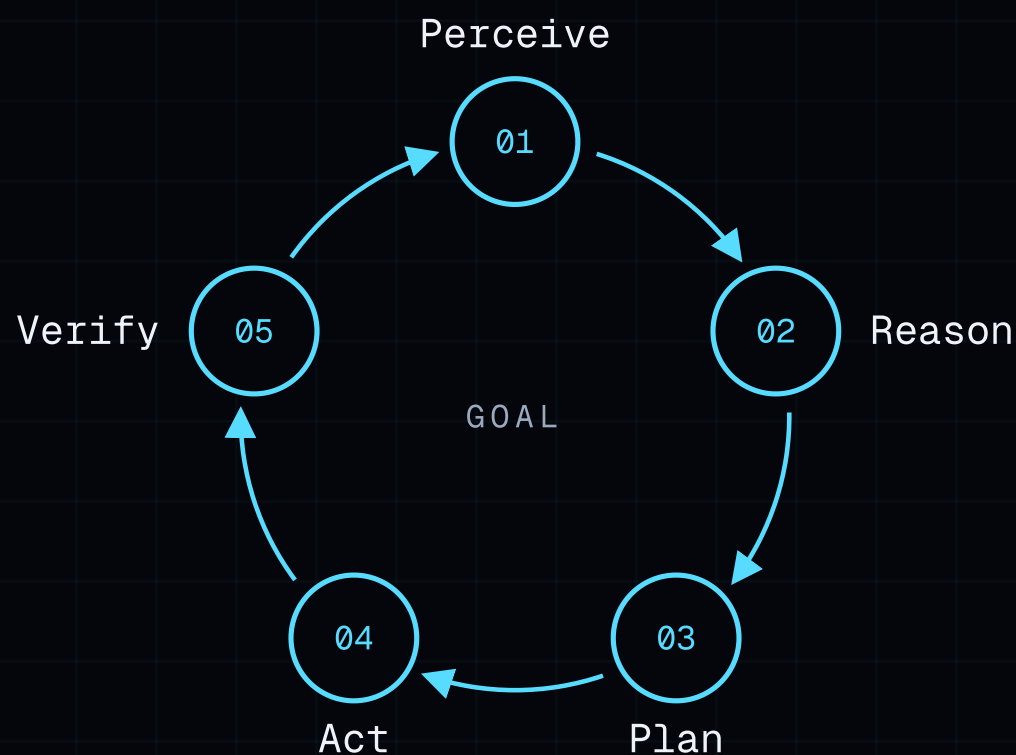
04 The governance ring

06 EU AI Act: where agents sit

01 / 06

What an agent is

- It is given a **goal** instead of a sequence of steps.
- It works in a loop until the goal is reached or a rule stops it.
- Chatbot answers · copilot suggests · **agent executes**.
- The test: who sets the steps? If the system chooses its path and tools itself, it is an agent – what it may execute without approval is set by the levels of autonomy.



In the glossary: AI agent · Agent vs. chatbot vs. copilot

02 / 06

Script, RPA or agent?

Stable interface, fixed rules

Script or API

No interface, the same screens every time, rare exceptions

RPA

Many variants, unstructured documents, case-by-case judgement

Agent

No **agent** as long as there are no real cases to test against or nobody owns the exceptions. In practice, one process often combines all three – and people.

In the glossary: RPA · Process mining

03 / 06

The building blocks

MODEL

Understands language, but guarantees nothing.

CONTEXT & RAG

Your documents as verifiable evidence instead of guesswork.

TOOLS & MCP

The agent's hands. Narrowly scoped, they are the key control – screen or code tools open everything its account may do.

EVALS

Quality as a measurement on real cases – outcome and path, tested repeatedly.

In the glossary: [LLM](#) · [RAG](#) · [Tools & function calling](#) · [Evals](#)

04 / 06

The governance ring

IDENTITY

Its own agent identity with an accountable owner, short-lived tokens, never a shared service account.

AUDIT

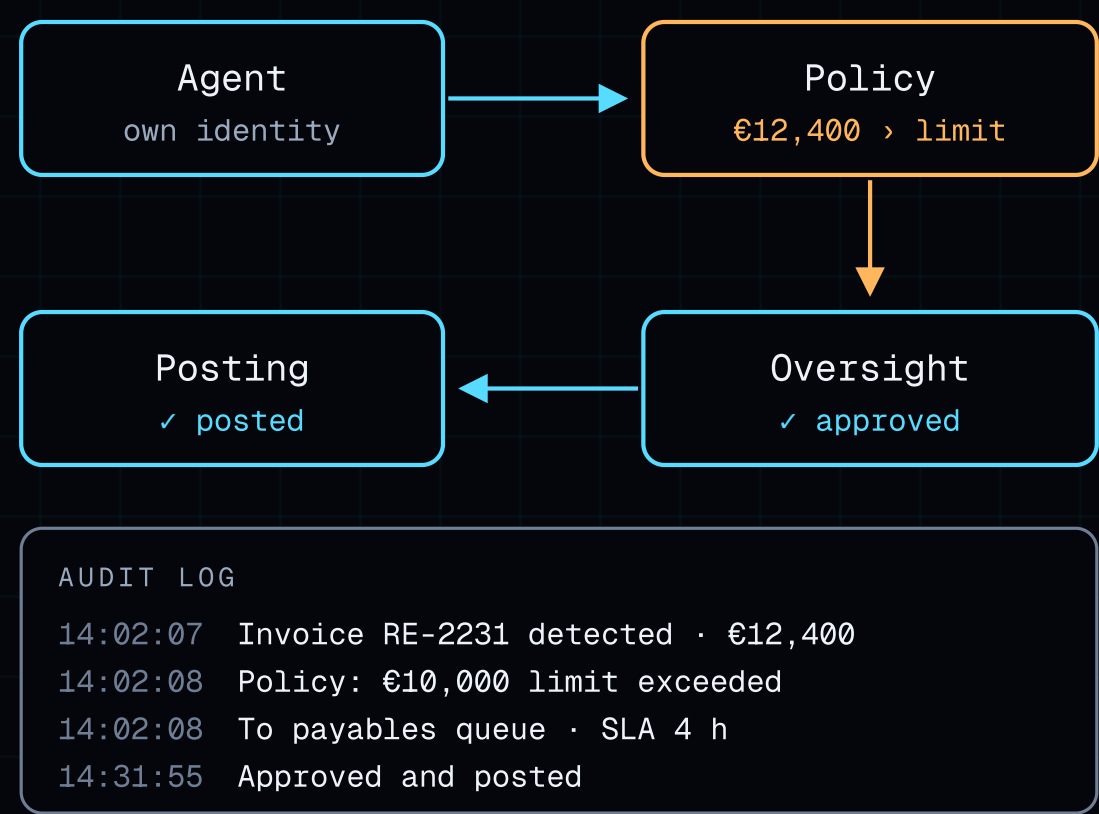
Every step logged tamper-evidently: inputs, tool calls, policy decision, reasoning.

POLICY

Limits live in code and the tool layer, not in the prompt – any email or PDF can carry hidden instructions.

OVERSIGHT

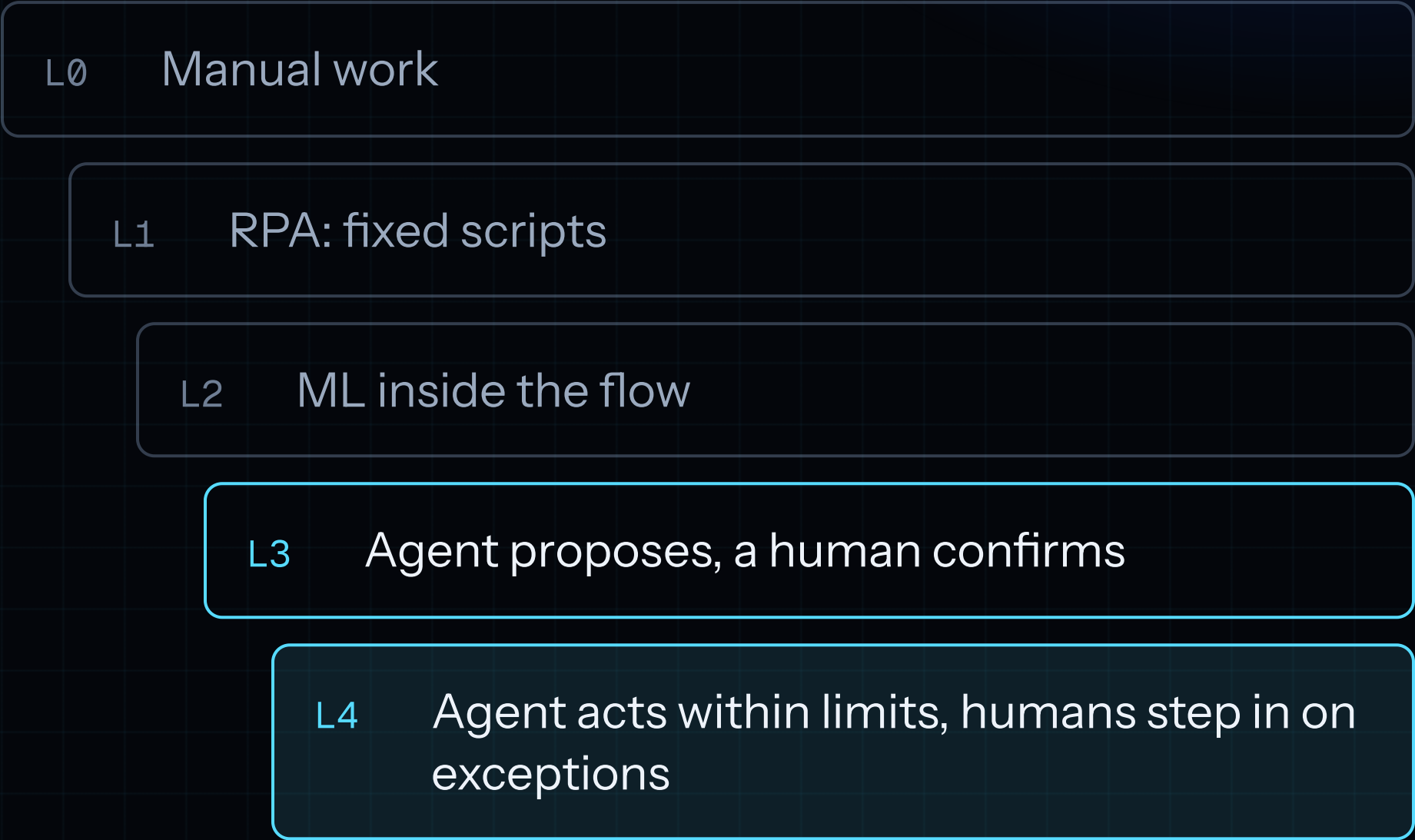
Exceptions land in a queue with an SLA; the stop switch works outside the model.



In the glossary: Agent identity · Guardrails & policy · Audit trail · Human-in-the-loop · Prompt injection

05 / 06

Autonomy in levels



A maturity model per group of cases, not a standard: move up a level only when evals, the audit trail and the exception rate justify it.

In the glossary: Levels of autonomy L0-L4

06 / 06

EU AI Act: where agents sit



Most back-office agents sit in the lower two tiers. They move up as soon as their decisions directly affect people: pre-screening job applications is high-risk.

since 2 Aug 2026

Article 50 transparency obligations

2 Dec 2026

Deadline for machine-readable marking of existing systems

2 Dec 2027

High-risk under Annex III

2 Aug 2028

High-risk under Annex I

As of September 2026, reflecting the Digital Omnibus (Reg. (EU) 2026/1744). Not legal advice.

In the glossary: Risk classes · The timeline

Every term explained in the glossary.

[botforce-technology.com
/en/glossary](https://botforce-technology.com/en/glossary)

Four chapters with real-world examples:
automation, AI, agentic AI and the EU AI Act.

As of September 2026 · EU AI Act details
reflect the Digital Omnibus; not legal
advice.