

Audit your own AI SDLC.

Six questions matter more than the framework label.

01

Can it preserve intent?

02

Can it block unsafe mutation?

03

Can it prove validation?

04

Can it retain human authority?

05

Can it roll back production?

06

Can incidents improve evaluations?

If one answer is unclear, you have found the next design problem.

AI-native delivery needs both.

A trustworthy control plane

for intent, evidence, provenance and authority



An operational learning loop

for deployment, rollback, telemetry and incidents

Save the six questions. Use them on your own workflow, regardless of vendor or framework.

A framework that cannot learn is not ready to govern learning agents.

How two AI-native SDLCs arrived at the same conclusion

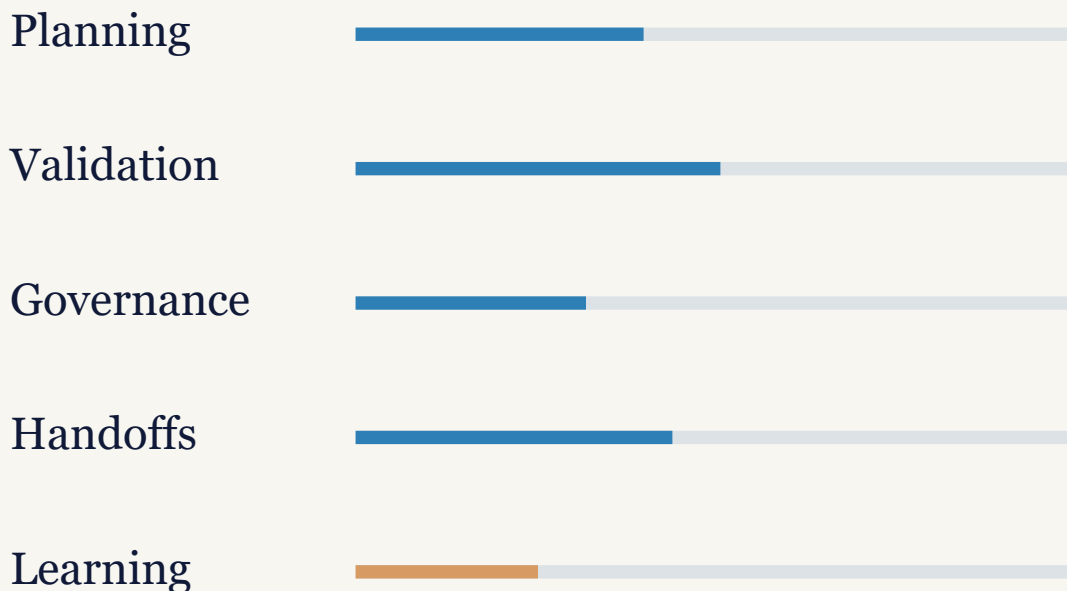
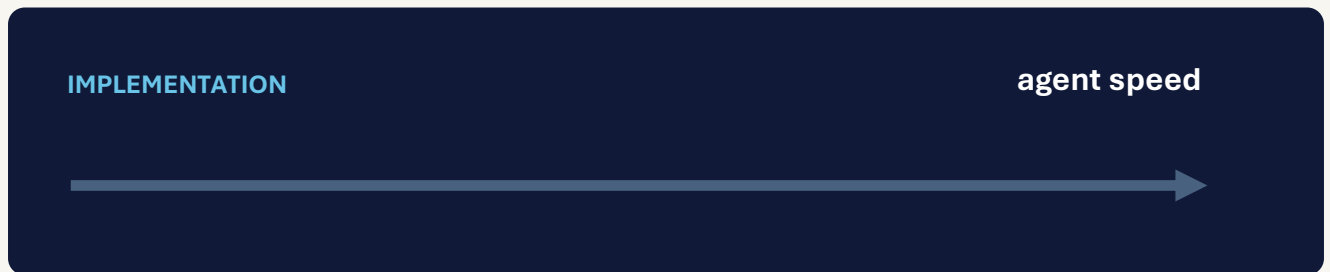
One portable control plane. One Claude-native playbook.



The comparison, without a winner

The plot twist: they agree.

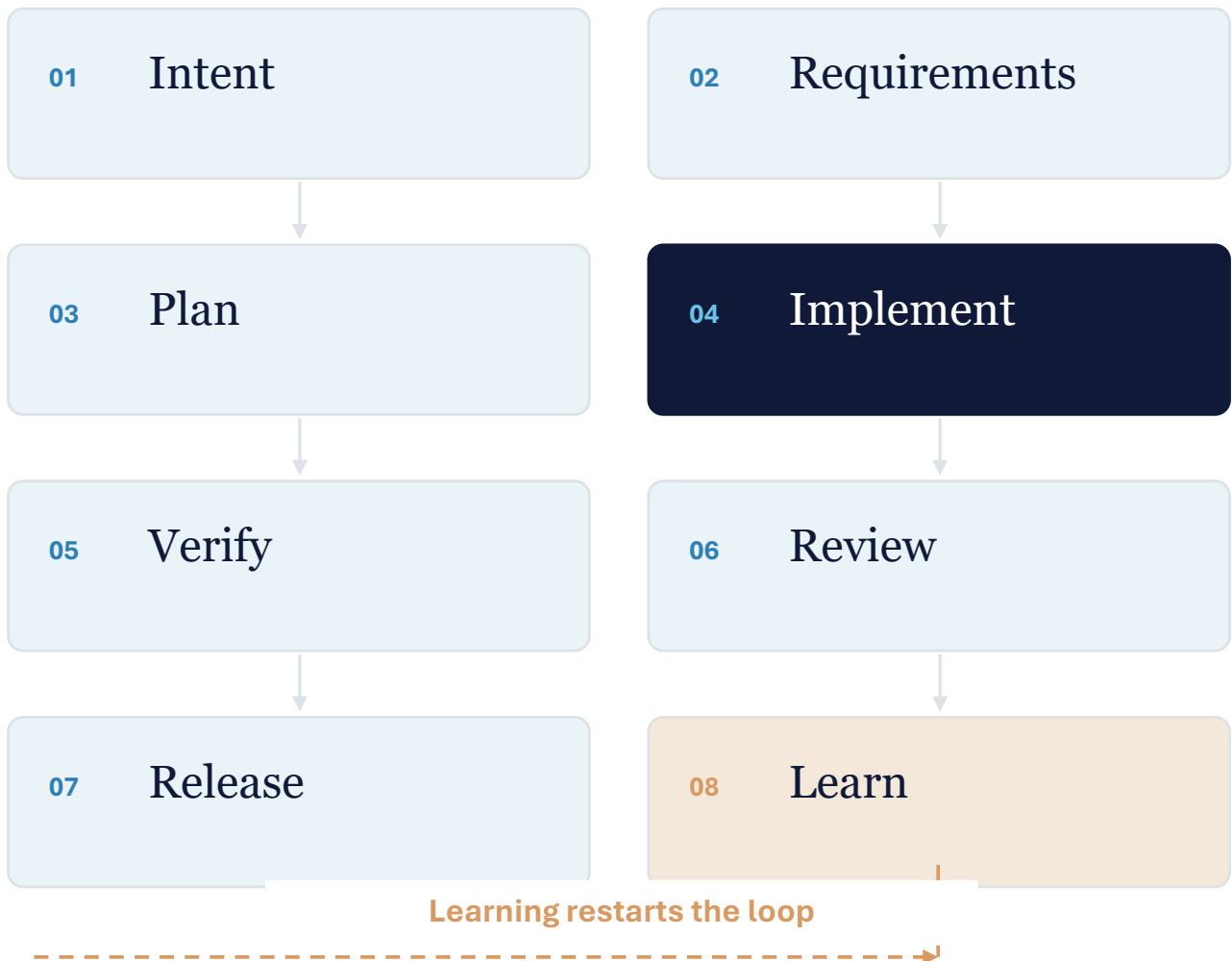
AI made implementation faster. The surrounding delivery system now has to catch up.



The bottleneck moved from producing code to coordinating trustworthy change.

The lifecycle still matters.

Both approaches preserve the work around implementation, not only implementation itself.



Chat memory is not governance.

Both turn transient agent reasoning into reviewable, versioned evidence.

AgentFlow

Issues

Role handoffs

Commits and checks

PR evidence

Anthropic

Intent

Spec and plan

Delivery history

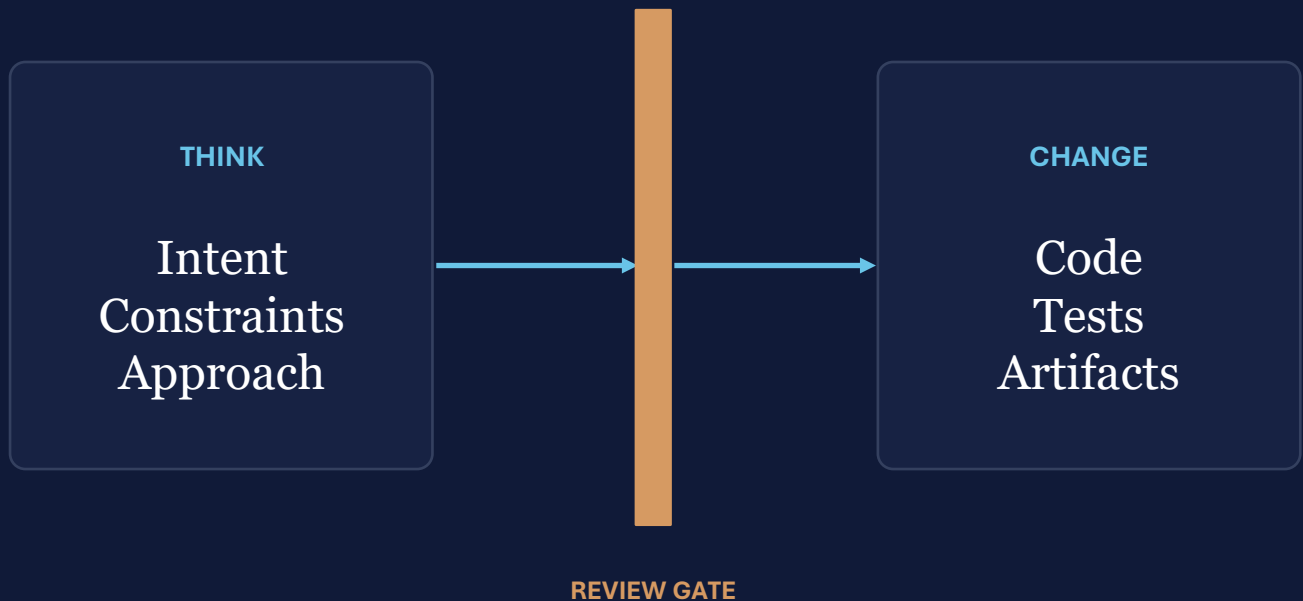
Incident record

Different artifact grammar

Same principle: preserve enough context to resume,
review and learn.

Plan before mutation.

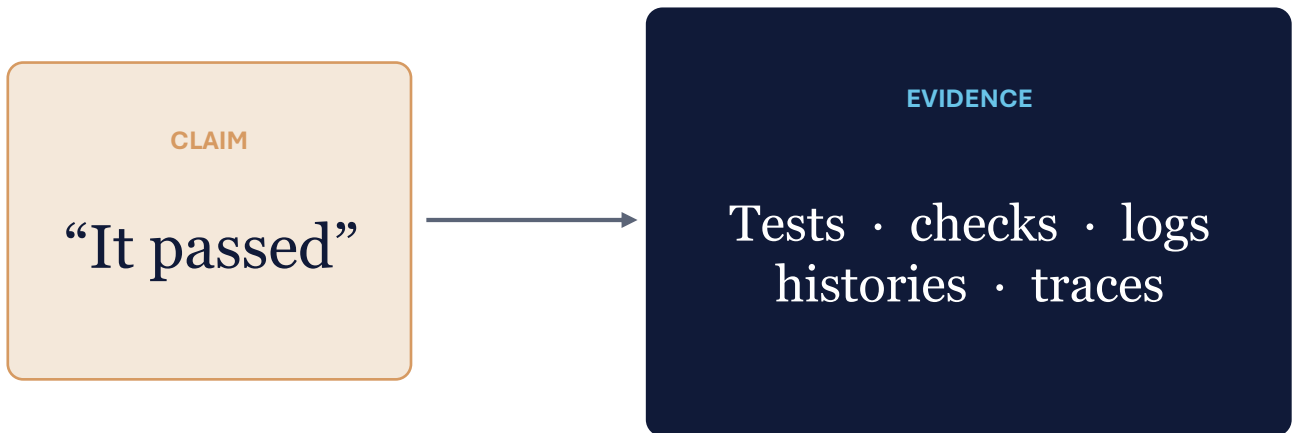
A plan is a reviewable control point, not the agent's private preamble.



The gate makes disagreement cheap before mutation makes it expensive.

Evidence before confidence.

The agent's conclusion is a claim. The result lives in the evidence behind it.



- What ran?

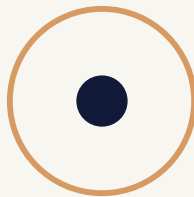
- Against what?

- What failed?

- Who reviewed?

Autonomy is not abdication.

Automation can expand. Accountability remains visible.



Consequential judgment stays outside the agent's self-approval loop.

Same problem. Different centre.

The distinction is useful because the approaches optimise different layers.

AgentFlow

- Portable governance
- Normalized provenance
- Defensible claims
- Cross-tool control

Anthropic

- Runtime controls
- Deployment and rollback
- Telemetry
- Incident learning

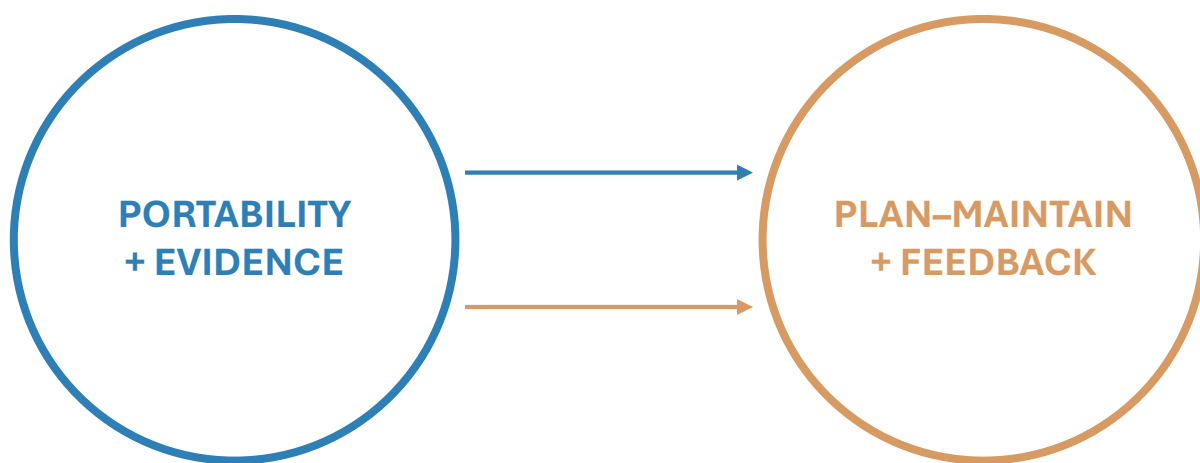
CONTROL PLANE

OPERATIONAL PLAYBOOK

Not a winner. A layering choice.

Each approach sharpens the other.

The interesting outcome is not convergence. It is what each makes more visible.



BEST SYNTHESIS

A trustworthy control plane with an operational learning loop.

Use the six questions. Keep what each approach makes visible.