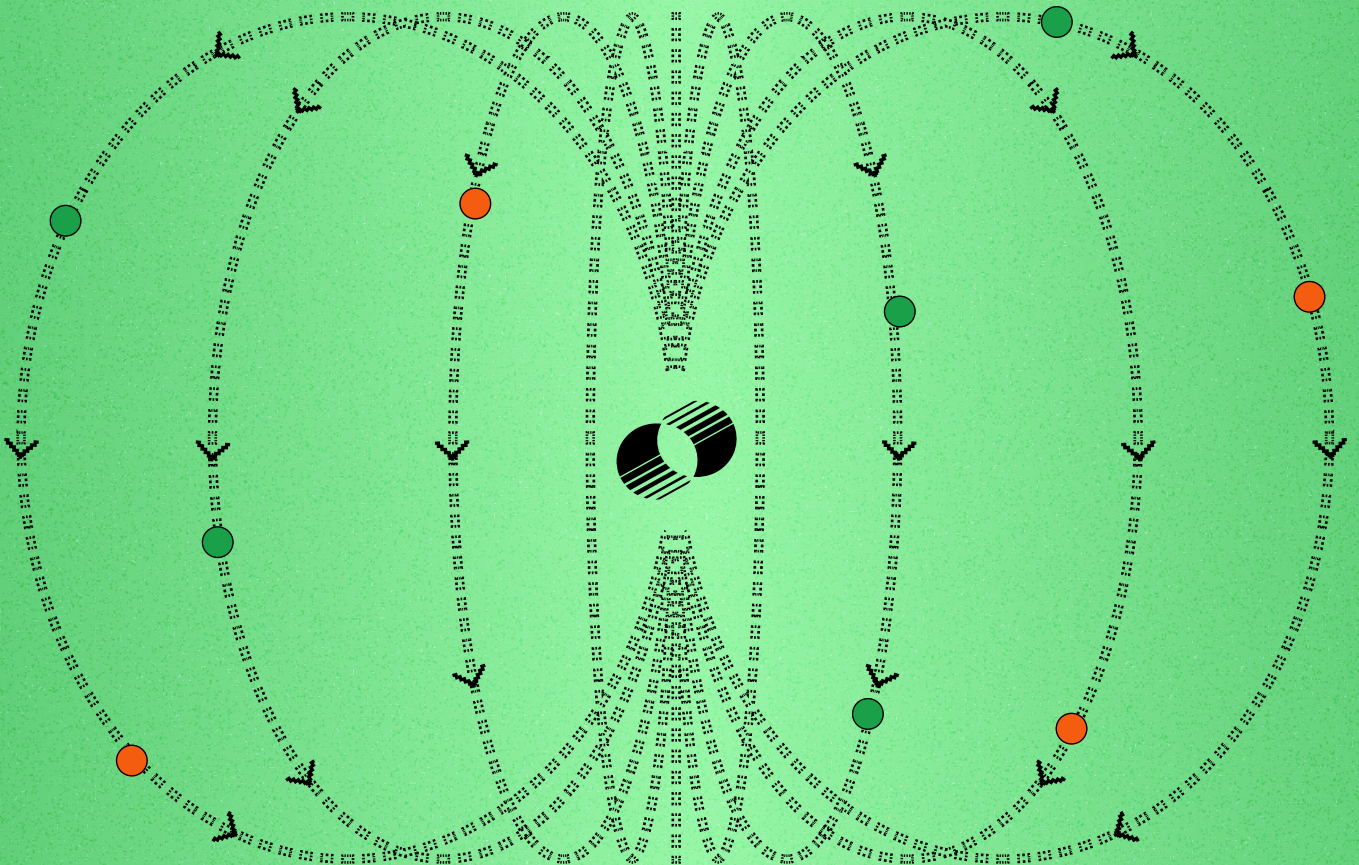


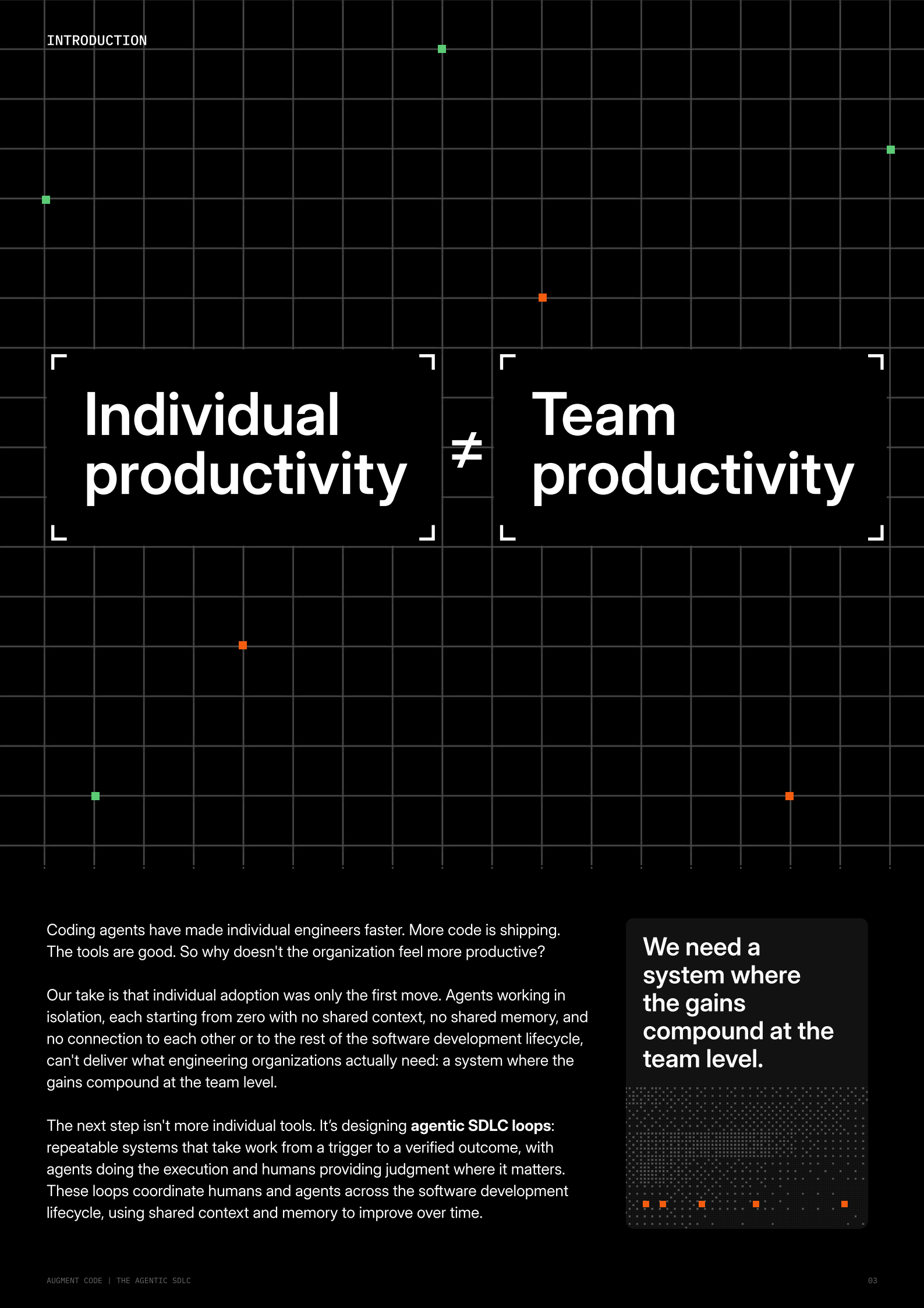
augment code

From Agents to Loops

■ HOW SOFTWARE TEAMS MOVE FROM
INDIVIDUAL CODING AGENTS TO
TEAM-LEVEL SOFTWARE DELIVERY →



■	INTRODUCTION	03
■	WHERE MOST TEAMS ARE TODAY (AND WHY IT ISN'T ENOUGH)	04
■	HOW HUMANS AND AGENTS WORK TOGETHER INSIDE AGENTIC SDLC LOOPS	05
■	THE STACK FOR RUNNING AGENTIC SDLC LOOPS	08
■	BUILDING COMPOUNDING AGENTIC LOOPS	09
■	THE SHIFT FROM INDIVIDUAL AGENTS TO TEAM-LEVEL AGENTIC LOOPS	10



Individual
productivity

≠

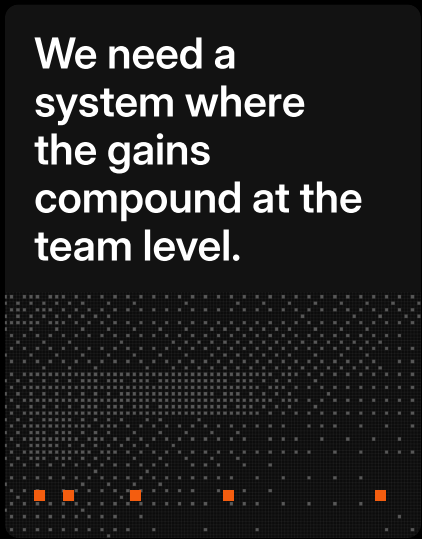
Team
productivity

Coding agents have made individual engineers faster. More code is shipping. The tools are good. So why doesn't the organization feel more productive?

Our take is that individual adoption was only the first move. Agents working in isolation, each starting from zero with no shared context, no shared memory, and no connection to each other or to the rest of the software development lifecycle, can't deliver what engineering organizations actually need: a system where the gains compound at the team level.

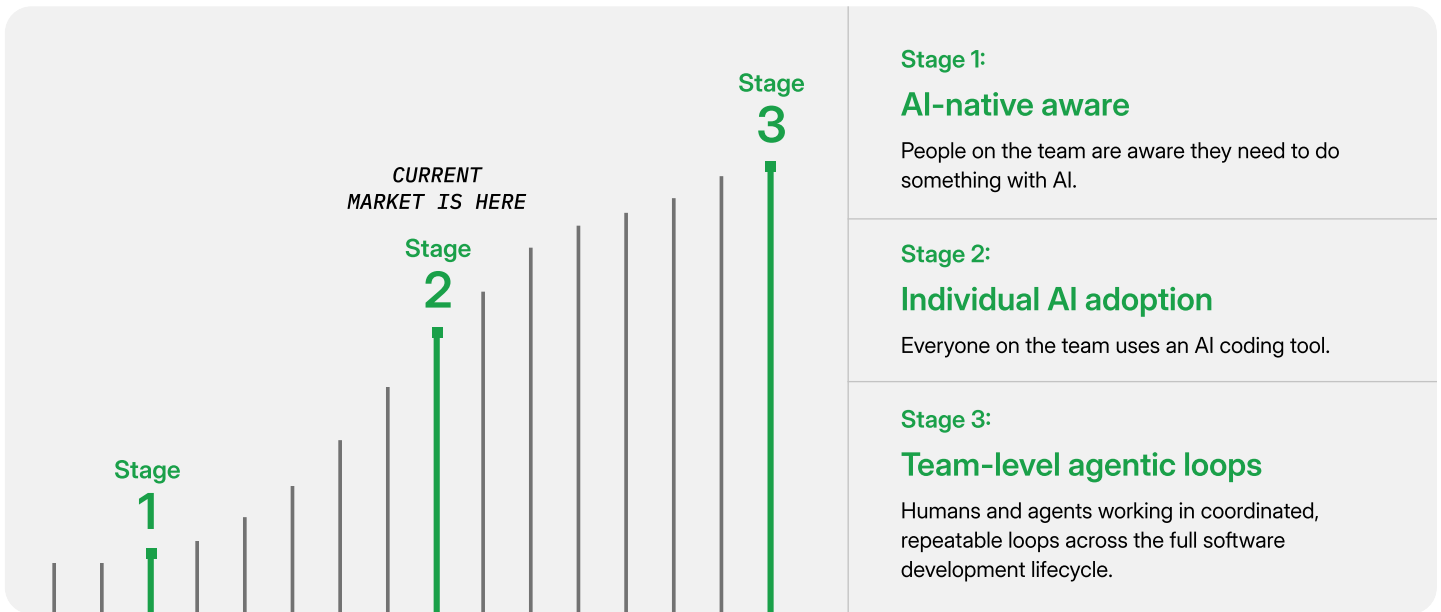
The next step isn't more individual tools. It's designing **agentic SDLC loops**: repeatable systems that take work from a trigger to a verified outcome, with agents doing the execution and humans providing judgment where it matters. These loops coordinate humans and agents across the software development lifecycle, using shared context and memory to improve over time.

We need a
system where
the gains
compound at the
team level.



Where most teams are today (and why it isn't enough)

Most engineering leaders equate AI native with individual adoption: everyone on the team is using a coding agent. That's a real milestone. But it's stage two of three.



The organizations stuck at Stage 2 share a common set of symptoms: engineers buried in agent-generated PRs, quality slipping, 2am alerts up, and a growing gap between what was promised and what's showing up in org-level outcomes. The productivity gains stayed with the individual. Nothing compounded.

What's missing isn't more licenses. It's a system where humans, agents, code, tools, and memory work together at the org level. One system, not a swarm of disconnected workers.

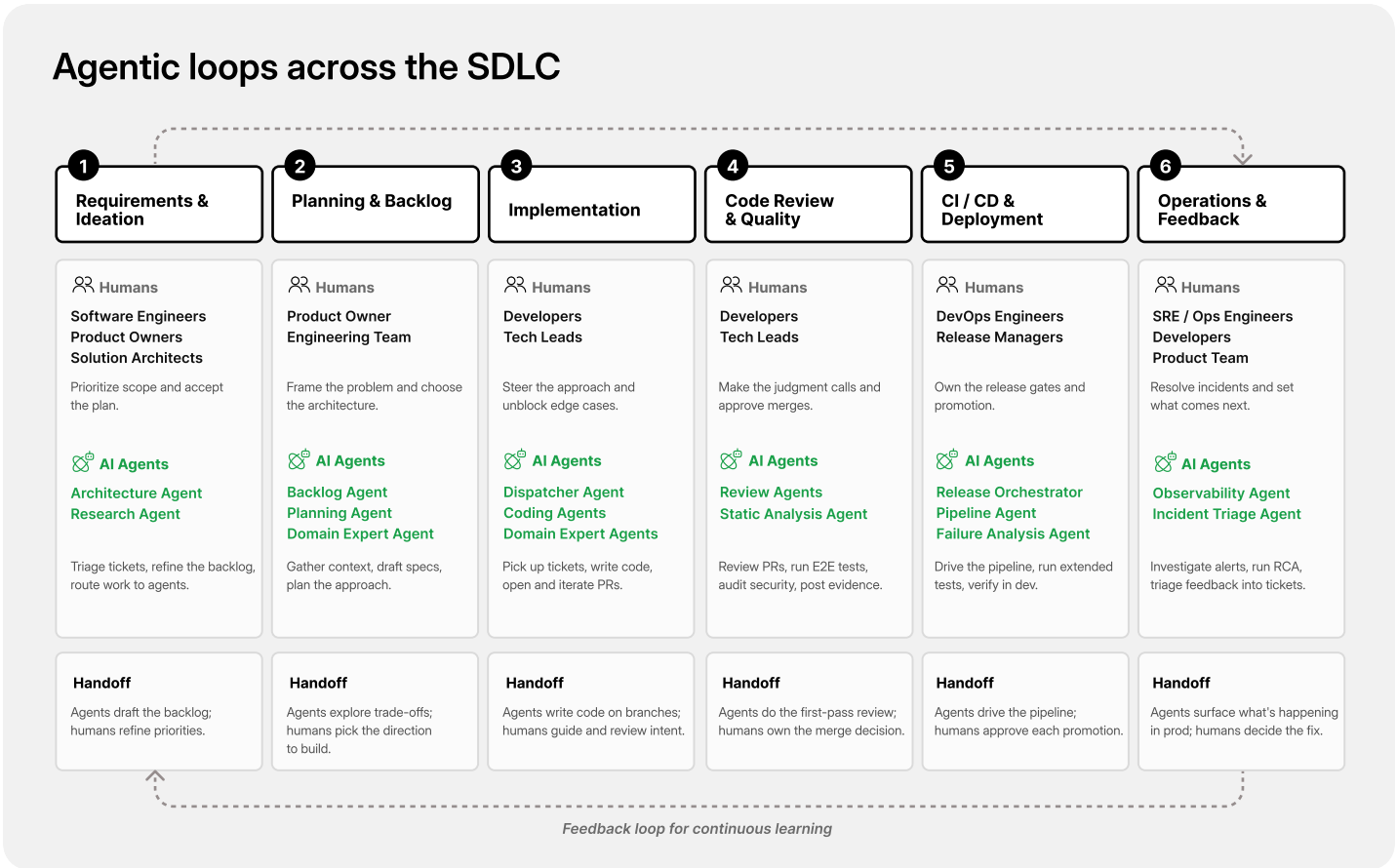
How humans and agents work together inside agentic SDLC loops

Team-level AI native doesn't mean agents do everything. It means designing agentic SDLC loops in which humans and agents each do what they're best at across every stage of software delivery, not just inside the IDE.

An agentic SDLC loop starts with a trigger, gives agents responsibility for the repetitive execution work, brings humans in where judgment matters, and ends in a verified outcome. Some loops operate within one stage of the SDLC; others connect several stages together.

Our philosophy at Augment: humans steer; agents execute. Humans provide direction, make architectural decisions, and own quality. Agents take on the execution work at speed and scale, iterating until the outcome meets its acceptance criteria or human judgment is required.

Humans steer;
agents execute



[0 1] Planning and triage

Planning used to be a human-only activity. Engineers gathered requirements, debated scope, and tried to break work into actionable tasks, often without full visibility into what the codebase would actually need. The result was backlogs full of underspecified tickets and surprises mid-sprint.

Agents can now change what comes into a planning session. They can analyze the existing codebase and architecture to surface dependencies, flag risks, and generate the engineering context that makes work items actionable before a sprint starts. In mature setups, agents observe the backlog continuously and propose routing, matching tasks to the right agent or engineer based on codebase expertise and workload. Humans own priorities and decisions; agents do the preparation work that used to eat planning time.

Humans own priorities and decisions; agents do the prep work.

[0 2] Requirements and design

The traditional gap between what gets documented in a requirements phase and what actually lives in the codebase tends to widen over time. Agents can now help close it. They contribute across multiple lenses simultaneously: security constraints, platform standards, operational requirements, cost implications. Design decisions get stress-tested against what already exists before anyone writes a line of code.

The more important shift is structural. Design discussions stop being isolated, one-on-one conversations with an AI tool and become shared engineering discussions where humans and agents participate together. That visibility matters. Decisions stay legible to the whole team rather than living in one engineer's chat history.

Decisions stay legible to the whole team.

[0 3] Implementation

This is where human-agent collaboration is most intensive, and where the failure mode is most common. Agents can now implement features, generate tests, and execute multi-step workflows. Whether that works in practice comes down to one thing: context. An agent with access to the team's actual codebase, coding standards, and architectural guidance produces reliable output. An agent without it doesn't.

The investment that unlocks this stage isn't more tooling. It's context infrastructure. Structured repository conventions, shared architectural guidance, and standardized patterns make the codebase legible to an agent at enterprise scale. Teams that build this foundation find that some work items become primarily agent-driven; others remain primarily human-authored. The mix shifts over time as context accumulates.

The investment that unlocks isn't more tooling. It's context infrastructure.

[0 4] Code review and verification

Code review is historically a bottleneck. A small number of engineers review everything, catch what they can, and slow down as volume increases. Agents don't just speed up that review; they change the shape of it.

Rather than one generic check, multiple specialized agents can now run review passes in parallel: security, performance, architecture compliance, observability compliance, API governance. Each focused on its domain. They also integrate with the quality tooling teams already run, interpreting CI pipeline failures, and surfacing root-cause analysis rather than raw logs. Engineers aren't triaging; they're deciding. One principle holds across every agentic implementation: verification quality has to exceed generation speed. As agents write more code faster, the review layer has to keep pace.

Engineers aren't triaging; they're deciding.

[0 5] CI/CD and deployment

Deployment has always been the stage where problems that should have been caught earlier become expensive. Agents can now move the catch point upstream. They validate infrastructure definitions before anything ships, flagging security issues, policy violations, cost concerns, and architectural inconsistencies at the point where they're cheapest to fix.

Progressive rollout strategies like feature flags, canary deployments, and blue/green releases also become easier to enforce consistently when agents are part of the deployment workflow rather than relying on engineers to apply them correctly under pressure.

Agents can now move the catch point upstream.

[0 6] Observability and feedback

Operations has historically been isolated from development. Incidents get resolved; the learnings rarely make it back into the backlog in any systematic way. This is the stage where the compounding effect of a team-level system becomes most visible, and where the absence of one is most costly.

Agents can now monitor production continuously, correlate deployment events with anomalies, and surface incidents with context (affected systems, likely cause, estimated blast radius) rather than raw alerts. But the more important shift is what happens next: operational findings flow back into the engineering workflow automatically. Technical debt items, observability gaps, recurring patterns, resilience improvements get created, tracked, and addressed rather than disappearing between sprints. The feedback loop from production to backlog closes. That's where the system starts to compound.

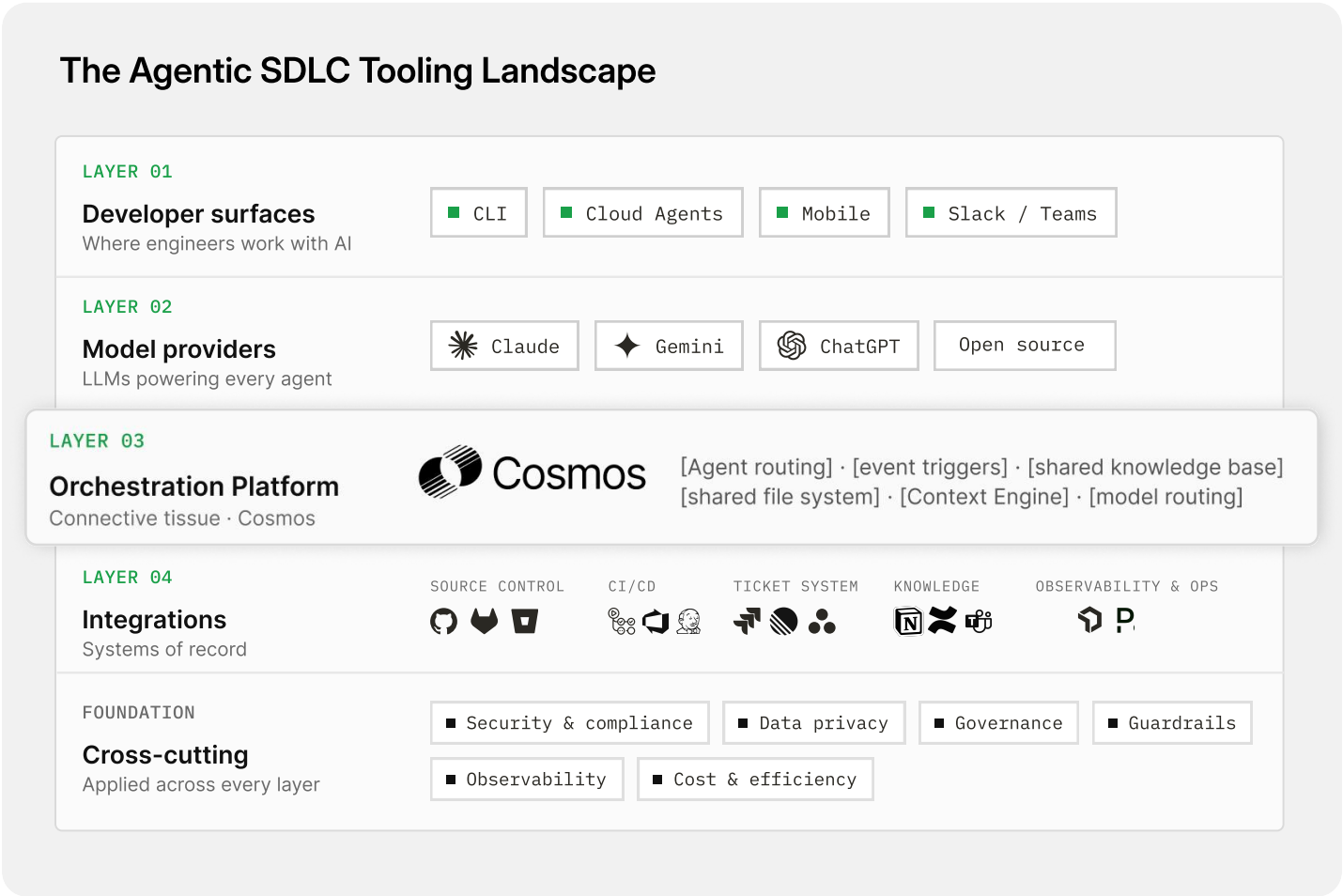
Operational findings flow back automatically.

The stack for running agentic SDLC loops

Running agentic loops across the software development lifecycle requires more than a collection of individual tools. It requires a platform that connects triggers, agent execution, verification, human checkpoints, outcomes, and learning.

Most teams trying to build this today are stitching point solutions together: one tool for review, another for coding, separate tools for triage and testing.

Each has its own auth, its own data model, its own context. Building a workflow that spans more than one stage is a custom integration project every time. That's the problem a unified platform solves.



Developer surfaces

Where engineers already work: CLI, Slack, mobile. Agents don't require a separate console when they can meet the team where it already is.

Agent execution

Agents that handle the repetitive middle of each loop: investigating issues, implementing changes, running tests, reviewing results, and iterating until the acceptance criteria are met or human judgment is required.

This is the layer that turns disconnected agent actions into observable, repeatable loops. Cosmos coordinates the triggers, agents, models, tools, human checkpoints, and outcomes involved in each loop. Its shared context engine, tenant-wide memory, and event-driven runtime allow work to continue across tools and sessions without starting from zero. It routes work to the right agent and model, makes every action observable and auditable, and retains the traces, patterns, and corrections that help future runs improve. Loops built by one engineer can be promoted into shared capabilities that the whole organization can use.

Key capabilities in this layer: workflow orchestration, model routing (BYOK across major providers), context engine, tenant-wide memory, observability.

Integrations

The engineering systems of record your team already runs: source control, issue tracking, CI/CD, artifact repositories, observability tooling. These systems supply the triggers that start loops—PRs, tickets, alerts, and deployment events—and receive the verified outcomes when the loops complete.

Foundation

Governance, security, identity, and compliance that cuts across every layer. Agent permissions are scoped. Every action is auditable. Human-in-the-loop rules are set by the team and enforced by the platform.

Building compounding agentic loops

Not all agentic setups are equal. The loops that get better over time share three properties.

[01]

Model agnostic by design

Betting on a single model lab means betting your entire agentic strategy on one provider's roadmap, pricing, and ecosystem. A system built for the long term brings your own keys across Anthropic, OpenAI, Bedrock, Vertex, and open-source models, routing work to the right model for each task without rebuilding workflows when the market shifts.

[02]

Shared memory that accumulates

Every agent on the market today starts from zero. A compounding system is different: patterns, corrections, and institutional knowledge persist across sessions and across the team. Agents built by one engineer get promoted into shared loops the whole org draws on. Each completed loop produces a trace that can improve future runs across the entire team.

[03]

Governance by default, not by retrofit

The hardest thing to add later is control. Organizations that bake governance in from the start: observable actions, auditable runs, scoped permissions, explicit human-in-the-loop rules are the ones that can scale agent deployment without incident. Teams set the policies; the platform enforces them.

The shift from individual agents to team-level agentic loops

The teams modernizing fastest, organizations like Stripe, Ramp, and Uber, are building this system themselves. They're defining how humans and agents divide the work, building shared memory across the team, and connecting agents to every stage of the software development lifecycle.

Most engineering organizations don't have the platform team to build it from scratch. That's not a small gap, it's a multi-year investment, and even the most sophisticated internal teams are finding it harder than expected. Augment Cosmos is what those teams would have built, productized.

The move from Stage 2 to Stage 3 happens when teams stop managing agents one task at a time and start designing repeatable loops around outcomes. The unit of optimization changes from the productivity of an individual engineer to the performance of the whole system: whether work finishes, whether it meets the team's quality standards, and whether every run makes the next one better.

This isn't a tool swap. It's a systems redesign. The organizations that get there first won't just generate more code; they'll build agentic loops that reliably deliver merged fixes, resolved incidents, remediated vulnerabilities, and other concrete outcomes. Every run, correction, and successful pattern will add to the system instead of disappearing at the end of a session.

That's what we think making a software team truly AI native means.

The path from individual agents to team-level agentic loops doesn't have to be a multi-year internal build. Augment Cosmos is the system for running those loops, productized.

**It's not a tool swap.
It's a systems
redesign.**

**Augment Cosmos
is a platform for
building agentic
SDLC loops.**

[See how it works →](#)

