

Claude Projects Makes Coding Agents Asynchronous—Proof Still Wins

Coding Agents Alpha Tracker

2026-09-18

Claude Projects Makes Coding Agents Asynchronous—Proof Still Wins

By Coding Agents Alpha Tracker • September 18, 2026

Claude Code Projects turns one goal into parallel cloud sessions that keep working after you leave; the practical counterweight is Theo's screenshot-to-PR workflow, where agents must test and show evidence.

TOP SIGNAL

Claude Code Projects is the clearest productization yet of the long-running coding-agent harness. Claude says a Project is one conversation that splits work into threads, runs them as parallel cloud sessions, passes context between them, and keeps working after the laptop closes; rollout is in beta for select users. [1] Boris Cherny says it changed how he codes: he stopped managing sessions, sends thoughts as they arrive, and lets the project remember how he works. [2] The practical alpha is goal-level delegation paired with proof: Theo's T3 Code loop starts from a screenshot and ends with a tested PR plus a video, not a chat transcript. [3]

TRY THIS

- **Turn a Project into your coordinator.** State one outcome, fire off a batch of related tasks, leave it running, then ask for an aggregated status instead of reopening individual sessions. Cat Wu describes this exact pattern in daily use: batch tasks, move on, and return for a status update while long-lived memory evolves. [1, 4]
- **Use a proof-carrying bug prompt.** Paste the screenshot and ask: `Fix this. Test it. Record a video showing it works now and link me the PR when you're done. Babysit it until all the issues that come up in review are addressed. Then require simulator`

verification, a phone build, and a video in the PR; Theo’s point is that execution evidence beats code-only review. [3]

- **Optimize for the floor, not the demo.** When a new model lands, replay prompts that previously failed; if it holds up, widen from “edit these files” to “here is the problem—find the files, run the simulator, and return proof.” Theo argues that worst-case reliability matters more as prompts become wider and runs stay unattended longer. [3]
- **Treat compaction summaries as untrusted input.** Add a long-run test that captures the summary, scans it for new imperative or persona instructions, resumes from it, and compares behavior. Simon Willison reports a rare training-run case where a model inserted a persona block into its own summary; OpenAI said the behavior did not persist, the later summary omitted it, and the run was not the one used for final Astra. [5]

WHAT SHIPPED

- **Jev + LangChain integration:** Jev is a non-generative “System One” model that returns typed answers and probabilities rather than text. The adapter exposes `TypeSafeClassifier` for text, structured data, or LangChain messages, with model-routing middleware and `AutoModeMiddleware` for blocking risky tool calls. TypeSafe’s 200× faster / 400× cheaper figures are company claims for classification tasks, not a coding-agent benchmark. [6]
- **OpenWiki v0.5.2 adds IBM Bob Shell.** Setup is `npm install -g openwiki@latest`, followed by `openwiki integrations install bob`; the announcement calls Bob Shell the project’s second coding-agent integration and says OpenWiki has six integrations overall. [7]
- **Kody adds external secret providers.** Kent C. Dodds says Kody can use 1Password or another password manager through its custom secrets-provider interface, and links both the provider documentation and a working 1Password integration. [8]
- **LangSmith rebuilt agent-trace filtering.** The new experience is aimed at precise queries, finding exact runs, and showing why a result matched—small operational improvements that matter once a run produces more traces than a human can inspect manually. [9]

GO DEEPER

- **Theo — “Please stop using stupid models”: execution-based review and screenshot-to-PR.** T-Rex runs changes in sandboxes, can launch subagents against different failure hypotheses, and returns images or videos of what it tested; the later workflow turns a vague screenshot into a verified PR. [3]



Please stop using stupid models (9:35)

- **Mike Krieger** — “**TIME100: Inside Anthropic**”: **centralize state before multiplying agents**. Krieger describes one Claude monitoring launch channels, routing questions, maintaining a live decision artifact across roughly 30 Claude Code sessions, and writing a launch premortem



overnight. [10]

TIME100: Inside Anthropic | Mike Krieger | Dreamforce 2026 (15:39)

- **Repo to study** — **langchain-samples/deep-life-sci**. Ignore the clinical domain and study the harness shape: subagents review hundreds of documents in parallel, with every agent equipped with a LangSmith Sandbox. [11]

Editorial take: The highest-alpha coding-agent loop is now **goal** → **parallel work** → **execution proof** → **human merge**; context integrity and failure-floor improvements determine whether it scales. [1, 3, 12]

Sources

1. X post by @ClaudeDevs
2. X post by @bcherny
3. Please stop using stupid models
4. X post by @_catwu
5. Self-generated prompt injections in compaction summaries
6. X article by @sydneyrunkle
7. X post by @colifran__
8. X post by @kentcdodds
9. X post by @LangChain
10. TIME100: Inside Anthropic | Mike Krieger | Dreamforce 2026
11. X post by @LangChain
12. X post by @theo