

AI's Moat Moves Into Workflows, Inference, and Proof

VC Tech Radar

2026-07-31

AI's Moat Moves Into Workflows, Inference, and Proof

By VC Tech Radar • July 31, 2026

Current signals point beyond raw model access: domain-specific agents are executing labor, inference is being routed and repriced, and enterprise buyers are asking for auditable boundaries.

1. Funding & Deals

Revoy

Revoy announced a \$27M Series A led by Standard Capital. The company says it is launching the United States' first hybrid-electric, cost-competitive long-haul freight network; its powered converter dolly cuts diesel use by 95%, while letting shippers use the network without converting legacy fleets or buying new infrastructure. The initial lanes are in the Pacific Northwest. [1]

The team is a meaningful part of the signal: Dalton Caldwell identifies Ian Rust (YC W22) as founder/CTO and former employee #1 at Cruise, while Peter previously founded Segment (YC S11, acquired by Twilio for \$3.2B) and Charm Industrial. Revoy says it is already operational and accepting freight customers. [2] The underwriting thesis is a retrofit-and-network wedge: adoption can begin through freight capacity rather than a wholesale fleet replacement.

WorkWeave

Dalton Caldwell announced a \$13.5M Series A for WorkWeave, led by Standard Capital. The accompanying founder interview centers on the model-router market and argues that "tokenmaxxing" is misguided. [3] The round is an early signal that infrastructure for choosing and controlling model usage is attracting capital alongside the models themselves.

2. Emerging Teams

Lassie

Lassie is an unusually concrete example of AI selling labor execution to small businesses rather than another copilot. The team says it serves hundreds of dental practices in a market of roughly 160,000 US practices, charges five figures for a first agent that handles about 30 hours of work per month, and has reached roughly 98% automation; much of its growth is word of mouth among dentists. [4] The founders' backgrounds include Superhuman and Robinhood, and they say the product began with the founders doing the office work themselves before automating their own process. [4]

The potential moat is operational knowledge: the team says general models do not encode the workflows of small businesses, while Lassie can learn from ERP history and staff feedback. [4] That is a stronger diligence story than generic "agent" positioning, but it depends on maintaining high automation and reliable integrations as the company expands beyond dentistry.

Infisical

Infisical is turning an open-source agent-security wedge into a commercial product. Its Agent Vault reportedly passed 2,000 GitHub stars, reached tens of thousands of installations and millions of agent runs, and was used by agents including Claude Code and OpenClaw. The commercial Agent Proxy is stateless, fetches secrets from Infisical rather than exposing them to agents, supports more than 30 service presets, and inherits rotation, RBAC, audit-log, and policy features. [5]

The underlying product thesis is that agents should not hold credentials directly: prompt injection can originate in user input or ingested data, so a proxy should attach credentials at the network boundary. [5] This is a credible infrastructure category to watch as agents move from isolated demos into production systems.

Camber

Camber's AI-native revenue-cycle platform shows where vertical AI can create value without displacing the core clinical workflow. Its article reports that only about 3% of claims are denied outright, but roughly 28% require rework; manual resolution can cost up to 6.3 times the automated alternative, with an estimated \$21B annual savings opportunity. [6] Camber says the share of claims requiring manual intervention fell by more than half within two months, versus a clinic's roughly two-year learning curve to improve first-pass billing. [6] The investable pattern is measurable recovered revenue and labor savings in a complex domain, not another general-purpose seat license.

3. AI & Tech Breakthroughs

Jeff Dean’s next bet is automated experimentation

Jeff Dean says current models are roughly at the level of a junior engineer for agent-based, long-running coding tasks, and that progress on complex tasks and non-coding domains has been faster than he expected. His 2027 prediction is that ML systems will improve themselves by decomposing problems, running automated experiments, and recombining the results; he expects the same loop to extend into science and engineering wherever objectives are measurable. [7]

The infrastructure consequence is equally important: Dean expects specialized, low-energy, low-latency inference hardware, noting that a 50x latency improvement would change what people build. He also says capable agents can run for days or weeks on difficult tasks, while the TPU precedent delivered 30–80x better energy efficiency and 20–30x lower latency than CPUs and GPUs of its time. [7]

Kimi K3 exposes the systems layer behind open-weight progress

A current-period walkthrough of Moonshot’s technical report says Kimi K3 is an open-weight frontier model ranked fourth of 580 by Artificial Analysis. The new detail worth tracking is systems efficiency: Kimi Delta Attention reduces the reported memory requirement for a one-million-token context from 104.6 GiB to 27.2 GiB; Quantile Balancing addresses expert-load imbalance across 896 experts; and the AgentENV Firecracker runtime reportedly created 51 million training sandboxes with 133 ms checkpoints and 49 ms resumes. [8] The broader investment signal is co-design across memory, routing, and training infrastructure—not simply scaling parameter count.

Parsing is becoming a routing problem

LlamaIndex’s open-source Parse Gateway estimates document complexity page by page, keeps simple pages in a free in-process path, and sends scans, dense tables, garbled text, and image-heavy pages to more capable parsing tiers. It is also exposed as an MCP server so agents can choose the parsing tier themselves. [9, 10] This is a concrete example of the same economic logic behind model routers: use expensive intelligence only where the task requires it.

4. Market Signals

Inference pricing is resetting

OpenAI announced an 80% price cut for GPT-5.6 Luna, a 20% cut for Terra, and a Fast mode for Sol offering up to 2.5x the speed for twice the price at the same intelligence. Sam Altman framed the strategy as finding the best price–intelligence tradeoff at every level. [11, 12] The implication for infrastructure

startups is that raw token access is becoming a weaker moat; routing, latency, workflow integration, and reliability matter more.

Adoption may be expensive before it is visibly productive

Exponential View notes that Barclays has not yet observed broad AI adoption lifting productivity, while half of CEOs in a BCG survey say their jobs depend on getting AI strategy right; public disclosures of net AI returns remain limited. [13] Its model argues that companies pay the learning bill first—new processes, skills, and organizational change—so a successful rollout can look expensive or irrational before it becomes productive. [13] For diligence, pilot count is less useful than evidence that an organization is carrying learning from one deployment into the next; the essay distinguishes bounded adopters from companies that accumulate projects without learning. [13]

Software is easier to ship, harder to distribute and retain

A SaaS founder who says he sold a prior company for tens of millions reports that roughly 15,000 subscription apps now launch each month, versus about 2,000 three years ago. The same post cites falling search clicks, Google Ads costs up roughly 25% since 2023, 12-month retention of 6.1% for monthly AI subscriptions versus 9.5% for non-AI subscriptions, and median net revenue retention of 48% for AI-native companies versus 82% for traditional B2B SaaS. It concludes that software is cheaper to build but more expensive to reach customers and easier to lose them, while Claude can now handle much of the software market. [14] Treat the numbers as founder-reported directional data, but the diligence consequence is clear: distribution, trust, and durable workflow adoption deserve more weight than feature velocity.

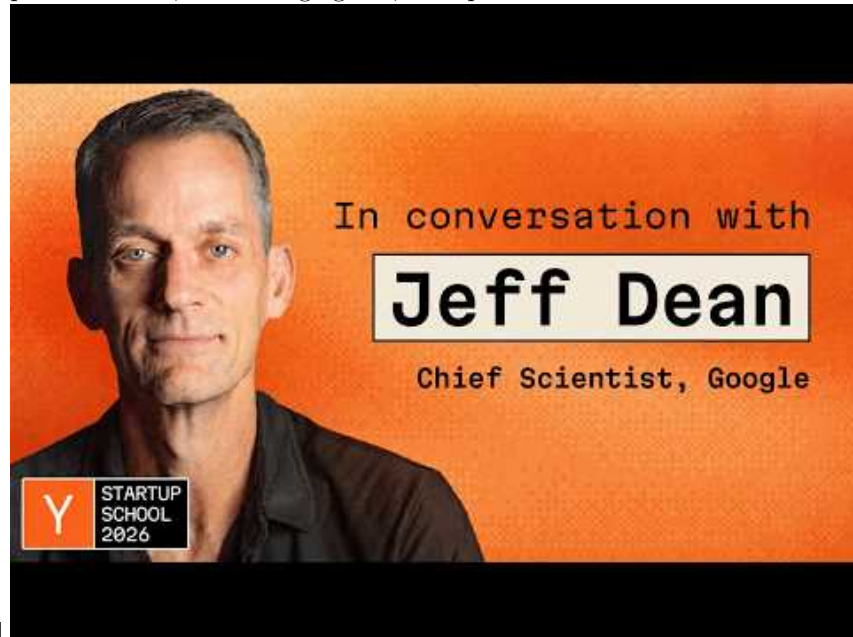
AI adoption is acquiring an operating owner—and an audit surface

A Reddit-posted analysis of 17M public job postings reports that “Head of AI” hiring tripled in nine months, with 1,142 companies currently advertising an AI-leadership role and 69% of hiring companies outside tech. It says 95% of those companies had not posted an AI-leadership role before 2026; titles skew toward Enablement and Transformation, and companies hiring AI leaders adopt agent frameworks at four to five times the base rate. [15]

At the deployment level, the trust question is becoming more specific. A practitioner says a client did not want exported traces; it wanted the rules governing the agent and proof that it stayed within them. The thread’s proposed “operating receipt” or evidence pack records permitted scope, policies and hard stops, inputs, actions, diffs, exceptions, and human checkpoints, backed by tool calls, deployment IDs, policy checks, and timestamps in an append-only manifest. [16, 17, 18] That missing product surface may become a procurement requirement for agent vendors, especially in regulated workflows.

5. Worth Your Time

- **Jeff Dean on self-improving ML systems and long-running agents.** The conversation connects the junior-engineer threshold to automated experimentation, weeks-long agents, and specialized inference hard-



ware. [7]

Jeff Dean: The 1% Rule for Building in AI (0:37)

- **How Lassie built an agent that does dental-office labor.** The interview is useful for the human-in-the-loop-to-automation transition, the 98% target, and the product work required to onboard small businesses without asking them to replace their systems. [4]



How Lassie Is Automating Healthcare Administration (17:57)

- **Exponential View’s AI adoption J-curve.** A compact framework for separating expensive learning from genuine failure in enterprise AI rollouts. [13]
- **The Kimi K3 technical walkthrough.** Read it for the concrete memory, expert-routing, and sandbox-runtime choices behind an open-weight model’s frontier performance. [8]

Sources

1. X post by @reinp
2. X post by @daltonc
3. X post by @daltonc
4. How Lassie Is Automating Healthcare Administration
5. X article by @dangtony98
6. X article by @camberhristophe
7. Jeff Dean: The 1% Rule for Building in AI
8. r/MachineLearning post by u/noninertialframe96
9. X post by @llama_index
10. X post by @jerryjliu0
11. X post by @sama
12. X post by @sama
13. For AI adopters, success and failure look identical — at first

14. r/SaaS post by u/yannis_ps
15. r/artificial post by u/vilnitskiy
16. r/SaaS post by u/Entire-Atmosphere-65
17. r/SaaS comment by u/Calm-Dimension3422
18. r/SaaS comment by u/Plane-Marionberry380