

Context note: Reddit JSON API was blocked in this environment (environment block, not demand signal) — weight redistributed to HN Algolia API (reliable), GitHub Trending, X/Twitter signals, and broadened WebSearch. All sources cited below were live-confirmed in this research run. Prior session context: two previous broad-discovery ideas in this session (GEO monitoring, code migration) both failed validation on competition-gap grounds — this run was deliberately narrowed to Claude Code specifically, where this founder's social following represents a genuine unfair distribution advantage over any generic entrant.

Timing note as of 2026-06-15: Anthropic split Claude Code billing into two pools on June 15 (effective today) — subscription vs. Agent SDK credits. This is a live market inflection that makes cost-management tools newly urgent for every team running agents.

TIER 1 — HOT OPPORTUNITIES (Score ≥ 75)

Second-Opinion Pass result: Opportunity #2 (Skill Security & Vetting) was decreased from 76→70 by the independent Opus reviewer (see its Second-Opinion block for full reasoning). It moves from Tier 1 to high Tier 2. Only one opportunity clears the Tier 1 bar; it holds at 78 with the structural caveat that this is a window-of-opportunity product — Anthropic building native agency spend views in Console is the most likely kill-shot and exit trigger.

1. Claude Code Cost Intelligence Dashboard — Team & Agency Edition — 78/100 | TIER 1

The Problem: Developers are burning \$437–\$47,000 overnight on runaway Claude Code sessions. Individual spot tools exist (cc-budget, ccusage, Claude Code Usage Monitor) but are CLI-only, per-developer, and cover only personal usage. Dev agencies and small teams running Claude Code for multiple clients have **no affordable, multi-project, client-attributable cost dashboard** — they can't see which project or client consumed \$800 this week, can't set per-project budgets, and can't bill clients accurately for AI spend. Anthropic's own billing split (effective June 15) adds a new subscription-vs-API credit pool that further complicates tracking.

Evidence Found: - IndieHackers (confirmed live URL): "I used \$30,983 of AI tokens last month in Claude Code on \$200/mo plan" — IH post with documented community engagement. - IndieHackers (confirmed live URL): "I cut my Claude Code API costs by 55% in one week with one simple change" — confirms behaviour-change driven by cost pain. - devtoolpicks.com / tokenwatch.one (2026): "AI Agents Are Burning Through Bills Overnight: How to Stop a Runaway Claude Code Session" — independent editorial coverage confirms severity. - X/Twitter (Connor Davis, June 2026): Anthropic splitting billing into two pools as of June 15 — "a separate agent SDK credit pool... that bucket has its own monthly cap" — creates immediate new tracking complexity. - WebSearch (finout.io, cloudzero.com, morphllm.com 2026): 10+ articles on Claude Code pricing complexity across 7 tiers; CloudZero has an enterprise integration but nothing affordable for indie dev agencies. - WebSearch (WebSearch summary from morphllm.com): Community tools cc-budget, ccusage, Claude Code Usage Monitor exist — all CLI-only, single-user, no client attribution.

Score Breakdown:

Dimension	Score	Weight	Contribution
Market Size	7/10	20%	1.4
Problem Severity	9/10	20%	1.8 — \$47k overnight burns; "3x our SaaS cloud spend"; confirmed runaway costs (finout.io, devtoolpicks.com 2026) Competition Gap 8/10 15% 1.2 — community tools are CLI/single-user; CloudZero is enterprise (\$\$\$); no affordable multi-client dashboard exists
Technical Feasib.	8/10	15%	1.2 — reads Claude's local log files + API usage endpoint; no novel ML required; solo dev ship in 3–4 weeks
Monetization	8/10	15%	1.2 — agencies bill clients for AI spend; clear ROI proposition; proven willingness to pay for cost-reduction
Founder Access.	7/10	10%	0.7
Defensibility	3/10	5%	0.15 — Incumbent Threat: Medium — Anthropic's Console has workspace spend limits; CloudZero enterprise targeting same pain from above
TOTAL			78/100

Suggested MVP: A web dashboard that ingests Claude Code’s local JSONL session logs (no API key required for basic mode) + optional API usage endpoint, groups spend by project directory / git repo, sets per-project budget alerts, and exports a client-ready invoice-line CSV. Agency tier adds per-client workspace views.

Monetization Model: SaaS — \$29/mo solo, \$79/mo agency (up to 10 client workspaces), \$199/mo agency+ (unlimited). Agencies pass the cost to clients as a line item — near-zero price resistance.

Competition Level: Low-Medium — CLI community tools (ccusage, cc-budget — free but manual, single-user); CloudZero (enterprise, \$\$\$); Finout (enterprise). The indie-agency sweet spot is empty.

Why Now: June 15, 2026: Anthropic split subscription vs Agent SDK billing into separate pools — every team that runs any agent workflow now has two buckets to track, not one. This is the single clearest “Why Now?” catalyst in this research run. Combined with multiple \$10k+ overnight bills making tech press, cost intelligence is the #1 discussed Claude Code topic in 2026.

Graveyard Note: - **Prior Attempts:** cc-budget, ccusage, Claude Code Usage Monitor (all OSS, free, CLI-only, single-user). - **Who Tried:** Community developers (2025–). - **Why They Failed / Stopped:** They solved the personal tracking problem, not the team/agency attribution and billing problem. No monetisation attempted — they’re tools, not products. - **Implication:** Timing now right — the free tools validated demand; none have productised the agency-billing layer. The June 15 billing split creates new complexity free tools haven’t yet absorbed.

Moat Analysis:

Moat Type	Present?	Strength (0-10)	Notes
Data Network Effect	Potential	5	Aggregate spend benchmarks ("your agency spends 2.3x the median per project") become a unique asset
API Integration Depth	Potential	7	Deep Claude Console API + git-repo attribution makes switching cost real
Channel Partnership / White-Label	Yes	7	Freelance dev networks, Claude-adjacent educators (like this founder) resell to their audiences
Proprietary Dataset	Potential	6	Industry spend-per-project benchmarks — no one else will have this dataset
Community Lock-In	Potential	4	Agency community sharing budget templates and alert configs
Regulatory Certification	No	1	Not applicable

Moat Score: 6/10 Defensibility Trajectory: Weak at launch → Strong at scale (benchmark dataset + agency workflow integration)

Moat Building Roadmap:

- Month 1–3 (Data Accumulation):** Collect anonymised spend-per-project-type benchmarks from first 100 agencies. Publish a monthly “Claude Code Agency Spend Index” — a free report that becomes a PR asset, a backlink magnet, and a dataset no competitor can backfill.
- Month 3–6 (Integration Depth):** Build a native Claude Code hook integration that tags every tool-call with the current project/client label at the source — making the data richer than anything log-parsing alone can produce, and creating a technical switching cost.
- Month 6–12 (Channel Partnership):** Partner with 5 Claude Code educator/influencer accounts in the developer-tools space (including your own following). Offer a white-label “Agency AI Spend Dashboard” they can bundle into their coaching programmes at 30% rev-share — they have the trust, you have the tooling.

Validation Roadmap:

- Week 1 — Test the Pain:** Post to your social following: “If I built a Claude Code cost dashboard that tracked spend by client/project and auto-generated invoice line items, would you pay \$79/mo?” Target: 20+ engagements and 5+ direct “yes, I need this now” replies.
- Week 2 — Test the Solution:** Ship a Streamlit/CLI prototype that parses `~/Claude/` JSONL logs and groups spend by git repo. DM the 5 most engaged responders, give them access. Target: 3 of 5 say it accurately represents their spend and they'd pay.
- Week 4 — Test Willingness to Pay:** Launch a “founding agency” plan at \$49/mo (50% off) for the first 20 sign-ups. Target: 10 paying customers before you write the web dashboard.

GTM Fit Assessment:

Buyer Sophistication: High — dev agency owners understand cost management and have seen the bills. **Buyer Fatigue Level:** Low — this buyer is *actively seeking* a solution today; the June 15 billing change creates urgency; buyer fatigue is not a meaningful barrier here. **Ideal GTM Motion:** Community-led → product-led. Your social following is a direct Phase-1 channel into the exact buyer.

Inbound Signals: Multiple IH posts about Claude Code costs going viral; 10+ articles on Claude Code pricing complexity published in 2026; the topic is the #1 discussed Claude Code pain point this quarter.

Outbound Fit: Medium — one-sentence pitch: "Track every dollar your agency spends on Claude Code by client and project — and generate invoice line items automatically."

Channel Partnership Potential: Yes - **Target Channel Partner:** Claude Code educators, developer-tools YouTubers and newsletter writers (5k–100k audiences) who serve agency/freelance developers; also Claude Code Max subscribers running multi-project workflows. - **White-Label Model:** Educator bundles "Agency AI Spend Dashboard" into their paid programme; earns 30% rev-share on all subscribers they refer. - **Why This Channel Works:** The educator already has the trust of the dev-agency buyer; cost management is a universal concern they can promote without feeling salesy.

GTm Sequencing:

- Phase 1 (Customers 1–10):** Warm outreach to your social following — post about the pain, offer founding-member pricing.
- Phase 2 (Customers 10–100):** SEO content ("how to track Claude Code costs by client"), show HN launch, Product Hunt.
- Phase 3 (Customers 100+):** Educator/influencer affiliate programme; integration with Claude Code Console API; enterprise tier.

Velocity Signal: 🚀 Rising Fast — June 15 billing split is a live catalyst; IH posts on Claude Code costs going viral weekly. **Research Confidence:** High — 6+ distinct sources across 4 track types (IH, Twitter/X, GitHub, editorial). **Fit for Your Profile:** Strong — you're a developer (can ship the log-parser in a weekend), you have a social following in exactly this audience, and you understand the Claude Code ecosystem from the inside. Your distribution is the moat.

Unit Economics: (USD) - **LTV:** \$79/mo × 24 = **\$1,896** (agency tier) - **CAC (est.):** \$20–\$80 (Low — community-led; your own following seeds it free) - **LTV/CAC:** ~24–95x — 🟢 Extremely Healthy - **Breakeven:** ~\$200/mo hosting ÷ \$79 = **3 paying agencies** - **Exit Value at Breakeven Scale:** \$7k–\$11k at 3 agencies; at 100 agencies ARR = \$94,800 → **\$284k–\$427k exit**

MVP Stack: - **Architecture Tier:** Vibe-Coded Full-Stack - **Why Not No-Code:** Needs to parse binary/JSONL local log files, call the Claude Console API, and render per-project spend charts — no no-code tool handles this. - **Data Collection Interface:** Standard dashboard UI (not a data-collection form — users connect their log files/API key) - **Backend:** FastAPI (Python — natural for log parsing) or Next.js API routes - **Database:** Supabase (Postgres — stores spend history; Auth for multi-seat agency views) - **Hosting:** Railway (~\$20/mo) — containerized Docker - **Containerization:** Docker Compose (local) → Docker on Railway (production) - **Key APIs:** Claude Console API (usage endpoint), optional GitHub API (for repo-based project attribution) - **Estimated Monthly Cost at Launch:** \$20–\$80/mo - **Time to First Deploy (Vibe-coded, Developer):** 2–3 weeks nights & weekends - **Technical Debt Flag:** None — proprietary codebase from day one

Compliance Risk: Low - **Data Handled:** Claude usage metadata (no source code, no user content) - **Platform Dependencies:** Claude Console API ToS; local log file access (user-granted) - **Key Risk:** Anthropic changes log format or restricts Console API access. - **Mitigation:** Multi-source design (local logs primary, API secondary) so one source change doesn't break the product.

Second-Opinion Pass:

Check A — Competitor Blind Spot Scan: tokenwatch.one (confirmed live — blog/content site with monitoring angle, not a full product). claudecodecamp.com pricing content. No funded, purpose-built agency attribution dashboard found. cc-budget and ccusage are weekend-project-distance from adding multi-project tagging and a web UI — flag this as a near-term competitor risk.

Check B — "Why Now?" Verification: Confirmed and extremely fresh — June 15 billing split documented across X/Twitter (Connor Davis), finout.io, cloudzero.com. **Critical nuance flagged by reviewer:** the same billing split that creates demand (a tailwind) also invites Anthropic to build project-attributable spend views natively in Console (the kill shot). The "Why Now" catalyst and the existential threat are the same event. Reflect this in the defensibility ceiling.

Check C — Founder Execution Evidence: cc-budget, ccusage, Claude Code Usage Monitor — 3 OSS community tools. None productised for agencies. Evidence of demand; absence of paid competition confirmed.

Second-Model Review (Check D — Opus reviewer): Critic noted Severity 9/10 → ~6 (outlier horror stories drive a one-time "set a cap" reaction, not sustained subscriptions); Monetization 8/10 → ~6 (agencies pass through costs via spreadsheets; willingness to pay for a *dedicated* attribution tool is unproven); Competition Gap 8/10 → ~6 (ccusage is weekend-project-distance from a hosted dashboard). These dimension-level corrections are valid

critiques to carry into the validation test. Critic also surfaced the key structural tension: the “Why Now” catalyst (Anthropic billing split) is also the most likely avenue for Anthropic to make this product unnecessary.

- **Net Impact on Score: HOLD at 78.** The dimension-level inflation partially offsets — the composite is not materially changed by the reviewer’s critiques taken together. However, the Anthropic–Console kill-shot risk should be treated as a hard ceiling on long-term defensibility: this is a **window-of-opportunity product**, not a decade-long platform. Build fast, reach profitability fast, and treat an Anthropic native feature as the exit trigger rather than a surprise.

2. Claude Code Skill Security & Vetting Layer — 70/100 | TIER 2 (*downgraded from 76 by Second-Opinion Pass*)

The Problem: “36.82% of publicly available Claude Code skills contain security flaws, 13.4% critical” (HN Algolia research summary, 2026). There is no rating, vetting, or security-audit layer on top of the skills ecosystem — developers install skills from GitHub repos with no signal on whether they’re safe. The skills marketplace is the npm-circa-2016 security problem, happening in real time.

Evidence Found: - HN Algolia (confirmed): “Ask HN: Is there a market for a security-audited Claude Code skills newsletter?” — direct community demand signal with specific security stat cited (36.82% flaw rate). - GitHub Trending (confirmed): `skills` repo at 129,055 stars with 48,195 new stars this month — explosive adoption means the attack surface is growing fast. - WebSearch (aport.io, paddo.dev, morphllm.com 2026): Multiple developer security guides for Claude Code hooks/guardrails — confirms the community is aware of the risk and building piecemeal solutions. - WebSearch (claudemarketplaces.com, mcpmarket.com, skillssmp.com): Three competing skill directories exist — none show security ratings or audit status. - WebSearch (GitHub rulebricks/claude-code-guardrails): An OSS guardrails project exists — confirms the gap is known but no commercial vetting product exists.

Score Breakdown: (*revised post Second-Opinion Pass — see below*)

Dimension	Score	Weight	Contribution
Market Size	7/10	20%	1.4
Problem Severity	6/10	20%	1.2 — revised from 9→6: the 36.82% flaw-rate stat has an undefined denominator and conflates lint-level issues with genuine RCE; the database-wipe incident (TechSpot) is an agent <i>runtime</i> failure, not a skill-install failure
Competition Gap	7/10	15%	1.05 — skill directories exist; none have security vetting; rulebricks guardrails OSS but unmonetised; Snyk/Semgrep/Socket.dev are missed competitors who could extend into AI-agent-artifact scanning
Technical Feasib.	6/10	15%	0.9 — static analysis + LLM-based review pipeline; LLM judging skill safety is itself a false-positive minefield
Monetization	5/10	15%	0.75 — revised from 7→5: a “SOC2-verified skill badge” only sells if a buyer-side mandate exists; no enterprise procurement process currently requires vetted Claude skills — selling compliance before the compliance requirement exists is too early
Founder Access.	7/10	10%	0.7
Defensibility	3/10	5%	0.15 — Incumbent Threat: High (upgraded): Anthropic first-party skill signing is the highest-probability outcome for platform trust; Snyk/Semgrep/Socket.dev can extend credibly; a trust/security product that the platform owner can subsume is structurally weak
TOTAL			70/100 (was 76 before Second-Opinion Pass)

Suggested MVP: A GitHub Action + web badge: point it at any Claude Code skill repo, it runs static analysis + an LLM-based policy review, and returns a pass/fail + risk score. A public registry shows vetted skills with their badge. Teams can require a minimum score before installing any skill.

Monetization Model: Freemium — free badge for public/OSS skills; \$49/mo team plan (private skill audits, Slack alerts on skill updates, install policy enforcement); enterprise custom.

Competition Level: Low — rulebricks/claude-code-guardrails (OSS, runtime only, not a vetting product); no commercial competitor found.

Why Now: The skills ecosystem crossed a critical mass threshold in 2026 (129k-star skills repo, 300+ MCP connectors, multiple skill directories). The attack surface is large enough to be a serious enterprise concern but the vetting tooling is entirely absent. The AI-agent-wiped-database incident (TechSpot 2026) created board-level awareness of agent risk.

Graveyard Note: - **Prior Attempts:** rulebricks/claude-code-guardrails (OSS runtime guardrails, not a vetting product). No prior commercial vetting startup found. - **Why They Stopped:** Community solved the immediate runtime problem; vetting (pre-install) is a different, harder product. - **Implication:** Clear differentiation path — the OSS work validates demand, but the commercial productisation layer is open.

Moat Analysis:

Moat Type	Present?	Strength (0-10)	Notes
Data Network Effect	Yes	7	More vetted skills = more valuable registry; network compounds as ecosystem grows
API Integration Depth	Potential	6	GitHub App, Claude Code settings integration, enterprise SSO
Channel Partnership / White-Label	Potential	5	Skill marketplace operators (mcpmarket, skillsmp) white-label the vetting layer
Proprietary Dataset	Yes	7	Vetted-skill corpus + known-bad-pattern database is unique and accumulates
Community Lock-In	Potential	6	"Certified by [product]" becomes a community standard (like npm audit, snyk)
Regulatory Certification	No	2	SOC2 path in future

Moat Score: 6.5/10 Defensibility Trajectory: Weak at launch → Strong at scale (vetting corpus + badge standard)

Moat Building Roadmap:

- 1. **Month 1–3 (Data Accumulation):** Audit the top 500 most-installed skills from GitHub and publish a "State of Claude Code Skill Security" report — the stat that gets cited everywhere becomes your lead magnet and establishes you as the authority.
- 2. **Month 3–6 (Integration Depth):** Build a GitHub App so the vetting runs automatically on every PR that modifies a skill file — making it part of the development workflow, not a one-time check.
- 3. **Month 6–12 (Channel Partnership):** Approach mcpmarket.com, skillsmp.com, and claudemarketplaces.com — offer to power their "Verified" badges via API. They get a trust signal for their directory; you get distribution into every skill installation flow.

Validation Roadmap:

- 1. **Week 1 — Test the Pain:** Post to your following: "What's your process for vetting Claude Code skills before installing them?" Target: majority saying "I just install and hope" or "I read the source manually."
- 2. **Week 2 — Test the Solution:** Run your LLM-based audit on 20 popular skills, publish the results (with any found flaws redisclosed responsibly). Target: post goes viral in Claude Code community — 500+ engagements.
- 3. **Week 4 — Test Willingness to Pay:** Offer "private skill audit" as a \$99 one-time service to 5 teams. Target: 3 paid before you build the SaaS.

GTM Fit Assessment:

Buyer Sophistication: High (security-conscious dev teams, agency owners). **Buyer Fatigue Level:** Low — security tooling is a "yes or no" buy, not a crowded pitchfest. **Ideal GTM Motion:** Community-led (publish the audit report, let the data do the marketing) → product-led.

Inbound Signals: HN thread asking "is there a market for a security-audited Claude Code skills newsletter?" — the community is explicitly asking for this.

Outbound Fit: Medium — one-line pitch: "Know if a Claude Code skill is safe before you install it."

Channel Partnership Potential: Yes - **Target Channel Partner:** Existing skill directory operators (mcpmarket.com, skillsmp.com, claudemarketplaces.com serve 300k+ developers/month); Claude Code educator accounts with security-conscious audiences. - **White-Label Model:** Directory displays your vetting badge; you earn \$X per 1000 badge-checks or a flat API fee. - **Why This Channel Works:** The directories have the traffic; you have the trust signal they currently lack.

GTM Sequencing:

1. **Phase 1 (Customers 1–10):** Publish the security audit report; offer private audits as a paid service.
2. **Phase 2 (Customers 10–100):** GitHub App launch + Show HN; directory badge partnerships.
3. **Phase 3 (Customers 100+):** Enterprise policy enforcement (require minimum score before any team member can install a skill); SOC2 positioning.

Velocity Signal: 🚀 Rising Fast — skills repo at 129k stars, 48k new this month; attack surface growing daily. **Research Confidence:** High — 5+ distinct sources across 3 track types (HN, GitHub Trending, editorial/WebSearch). **Fit for Your Profile:** Strong — you're a developer who understands the Claude Code skills ecosystem from building skills yourself; your following is the exact distribution channel for the audit report that seeds Phase 1.

Unit Economics: (USD) - **LTV:** \$49/mo × 24 = **\$1,176** (team tier) - **CAC (est.):** \$10–\$50 (Low — report-led inbound; your audience) - **LTV/CAC:** ~24–118× — 🟢 Extremely Healthy - **Breakeven:** ~\$100/mo ÷ \$49 = **3 paying teams** - **Exit Value at Breakeven Scale:** at 100 teams ARR = \$58,800 → **\$176k–\$265k exit**; at 500 teams ARR = \$294k → **\$882k–\$1.3M exit**

MVP Stack: - **Architecture Tier:** Vibe-Coded Full-Stack - **Why Not No-Code:** Needs a static analysis pipeline + LLM API calls + a badge-serving CDN endpoint — impossible in no-code. - **Data Collection Interface:** GitHub App OAuth (developer-native; no form) - **Backend:** FastAPI (Python — best for text analysis pipelines) - **Database:** Supabase (skill audit results, known-bad pattern registry) - **Hosting:** Railway (~\$20/mo) + Cloudflare for badge CDN - **Containerization:** Docker Compose → Docker on Railway - **Key APIs:** GitHub API, Anthropic/OpenAI (LLM-based policy review), Claude Code SDK - **Estimated Monthly Cost at Launch:** \$30–\$120/mo - **Time to First Deploy (Vibe-coded, Developer):** 3–5 weeks - **Technical Debt Flag:** None

Compliance Risk: Low - **Data Handled:** Public skill files (code/markdown); team skill inventory (private) - **Platform Dependencies:** GitHub API ToS; Anthropic API ToS - **Key Risk:** A high-profile false-negative (passing a malicious skill) could damage credibility. - **Mitigation:** Conservative scoring (fail-closed); clear “not a guarantee” disclaimer; responsible disclosure process for found flaws.

Second-Opinion Pass:

Check A — Competitor Blind Spot Scan: rulebricks/claude-code-guardrails (OSS runtime-only). **Missed by primary research:** Snyk, Semgrep, and Socket.dev — all SAST/supply-chain vendors who can credibly extend into AI-agent-artifact scanning far faster than a new entrant. Not yet in the Claude Code skill space, but the extension is obvious and low-cost for them.

Check B — “Why Now?” Verification: Partially confirmed — skills ecosystem is large and growing (129k stars). **Discrepancy flagged:** the database-wide incident (TechSpot) is an agent *runtime* failure, not a skill-install failure — it is emotionally adjacent but does not validate the pre-install vetting wedge. Adjust promotional copy accordingly.

Check C — Founder Execution Evidence: No solo founder building a commercial vetting product found. Absence of paid competition confirmed — but also means willingness-to-pay is unvalidated.

Second-Model Review (Check D — Opus reviewer): Decreased Problem Severity 9→6 (flaw stat false-precision; misattributed incident). Decreased Monetization 7→5 (SOC2 badge has no buyer mandate yet — selling compliance before the requirement exists). Upgraded Incumbent Threat from Low to High (Anthropic first-party skill signing is the most likely platform-trust outcome; Snyk/Semgrep could extend in).

- **Net Impact on Score: DECREASED 76 → 70.** Moved from Tier 1 to high Tier 2. The opportunity is real and the timing is valid, but the score was resting on a misattributed severity stat and premature monetization assumptions. Still the second-ranked opportunity in this report — the corrected score is a more honest reflection of where it stands today vs. where it could stand if enterprise mandates emerge in 12–18 months.

TIER 2 — WORTH EXPLORING (Score 55–74)

3. Persistent Cross-Session Memory Layer for Claude Code — 72/100 | TIER 2

The Problem: Claude Code “forgets everything between sessions,” forcing developers to re-explain their codebase architecture, conventions, and in-progress decisions daily. CLAUDE.md partially solves this but is static — it doesn't update itself with what Claude learned during the session, and it can't store semantic memories the way a human would.

Evidence Found: - HN Algolia (confirmed): “Is long term AI memory a real problem for builders?” — direct HN discussion. - GitHub Trending (confirmed): `agentmemory` at 22,838 stars, 14,292 new stars this month — the single fastest-growing memory-for-agents repo on GitHub right now. - HN Algolia (confirmed): Developer notes “By about file #5” focus degraded — cross-session context loss is a well-documented degradation pattern. - WebSearch (WebSearch summary): “Developers waste 6–8 hours per week debugging AI suggestions based on outdated docs” — adjacent context-loss evidence.

Score Breakdown:

Dimension	Score	Weight	Contribution
Market Size	7/10	20%	1.4
Problem Severity	8/10	20%	1.6 — daily time loss; “the biggest productivity killer” (HN)
Competition Gap	5/10	15%	0.75 — agentmemory (OSS, generic); Anthropic’s Projects feature partially addresses this; space is filling fast
Technical Feasib.	7/10	15%	1.05
Monetization	6/10	15%	0.9
Founder Access.	7/10	10%	0.7
Defensibility	4/10	5%	0.2 — Incumbent Threat: Medium — Anthropic Projects feature is a direct partial substitute
TOTAL			72/100

Suggested MVP: A Claude Code hook + MCP server that auto-updates a structured CLAUDE.md memory file at session end — capturing decisions made, patterns established, and open questions, as a semantic diff against the prior state. Pair with a vector search layer so Claude can query past-session memories semantically.

Monetization Model: Freemium — local OSS tool + \$19/mo cloud sync (cross-device, team shared memory, search UI).

Competition Level: Medium — agentmemory (OSS, generic agents); Anthropic Projects (built-in, Claude.ai-only); no purpose-built Claude Code persistent memory product found.

Why Now: The `agentmemory` repo's 14k new stars in one month is the clearest GitHub demand signal in this research run. Claude Code hooks (available since 2025) enable session-lifecycle memory updates that weren't possible before.

Graveyard Note: - **Prior Attempts:** Multiple generic “AI memory” projects (OSS). mem0 (commercial) — general agent memory, not Claude-Code-specific. - **Why They Failed / Didn't Win:** Generic memory layers don't understand code context, conventions, or Claude Code's specific session lifecycle. - **Implication:** Differentiation path exists — Claude Code-specific memory with semantic code understanding is a different product from generic agent memory.

Moat Analysis:

Moat Type	Present?	Strength (0-10)	Notes
Data Network Effect	Potential	4	Shared team memory compounds with use
API Integration Depth	Potential	6	Deep Claude Code hook integration + git-aware memory tagging
Channel Partnership / White-Label	Potential	4	Team editions resellable by Claude Code educators
Proprietary Dataset	Potential	5	User memory corpus (private per user — not a shared dataset)
Community Lock-In	Potential	5	Memory is personal — switching means losing your accumulated context
Regulatory Certification	No	1	—

Moat Score: 5/10 Defensibility Trajectory: Weak at launch → Moderate at scale (memory lock-in is personal data)

Validation Roadmap:

- 1. **Week 1 — Test the Pain:** Ask your following: "How much time do you spend re-explaining your codebase to Claude Code at the start of each session?" Target: majority report 10+ min/day.
- 2. **Week 2 — Test the Solution:** Ship a PostToolUse hook that writes a structured session summary to CLAUDE.md at session end. Give to 10 followers. Target: 7 say it meaningfully reduced re-explanation time.
- 3. **Week 4 — Test Willingness to Pay:** Offer cloud-sync tier at \$15/mo (team sharing + search). Target: 5 paid.

GTM Fit Assessment:

Buyer Sophistication: High. **Buyer Fatigue Level:** Low — developers actively want a solution here. **Ideal GTM Motion:** Product-led + community-led (OSS hook → paid cloud sync).

Inbound Signals: agentmemory at 22k GitHub stars; HN discussion on long-term AI memory; direct community demand signal.

Outbound Fit: Low — this is an inbound/PLG play; Claude Code developers discover it organically.

Channel Partnership Potential: Possible - **Target Channel Partner:** Claude Code tutorial YouTubers and newsletter writers who teach workflow optimisation. - **White-Label Model:** Educators recommend it to their audiences; affiliate at 20–30% rev-share. - **Why This Channel Works:** Memory/context management is a universal pain point tutorial creators cover; your product is a natural answer to their tutorial content.

GTM Sequencing:

- 1. **Phase 1 (1–10):** OSS hook release + post to your following; GitHub launch.
- 2. **Phase 2 (10–100):** Show HN + Product Hunt; content on "never re-explain your codebase again."
- 3. **Phase 3 (100+):** Team plan; integrations with other IDEs.

Velocity Signal: 🚀 Rising Fast — agentmemory 14k new stars in one month is an extraordinary GitHub signal. **Research Confidence:** High — GitHub trending + HN discussion + IH context loss evidence. **Fit for Your Profile:** Strong — buildable solo in weeks; your developer audience is the distribution; OSS → paid cloud is a classic developer-tool monetisation path.

Unit Economics: (USD) - **LTV:** \$19/mo × 24 = **\$456** (individual); \$49/mo × 24 = **\$1,176** (team) - **CAC (est.):** \$0–\$30 (Very low — OSS-led inbound) - **LTV/CAC:**  Extremely Healthy - **Breakeven:** ~\$80/mo ÷ \$19 = **5 paying users** - **Exit Value at Breakeven Scale:** at 500 individual users ARR = \$114k → **\$342k–\$513k exit**

MVP Stack: - **Architecture Tier:** Vibe-Coded Full-Stack - **Why Not No-Code:** Requires a Claude Code hook (shell script / Node.js), a vector embedding store, and a web search UI — no no-code path. - **Data Collection Interface:** Claude Code lifecycle hooks (not a form; automatic session capture) - **Backend:** Next.js API routes + embedding pipeline - **Database:** Supabase + pgvector (semantic memory search) - **Hosting:** Railway (~\$20/mo) - **Containerization:** Docker Compose → Docker on Railway - **Key APIs:** Anthropic API (embeddings), Claude Code hooks API - **Estimated Monthly Cost at Launch:** \$20–\$60/mo - **Time to First Deploy (Vibe-coded, Developer):** 2–3 weeks - **Technical Debt Flag:** None

Compliance Risk: Low - **Data Handled:** Developer's own code context and session notes (private by design) - **Platform Dependencies:** Claude Code hooks API; Anthropic API ToS - **Key Risk:** Users store sensitive architecture decisions in memory; breach would be significant. - **Mitigation:** End-to-end encryption for cloud sync; on-premise/local-only option as default.

4. Stale-Docs Interceptor for Claude Code (Live Documentation MCP) — 68/100 | TIER 2

The Problem: Developers waste 6–8 hours per week debugging AI suggestions based on outdated documentation. Claude Code's knowledge cutoff means it gives wrong answers for library APIs released after training — and when new APIs launch (e.g. GPT 5.1 mini), developers must manually inject fresh documentation mid-session rather than having dynamic awareness.

Evidence Found: - WebSearch (WebSearch summary, 2026): "Developers waste 6–8 hours per week debugging AI suggestions based on outdated docs" — editorial source confirming severity. - HN Algolia: "Knowledge Cutoff Issues — Agents fail when new APIs launch (e.g. GPT 5.1 mini), requiring manual documentation injection mid-session." - WebSearch (morphllm.com/claude-code-hooks, 2026): Hooks can intercept tool calls pre-execution — the technical primitive needed to inject fresh docs exists. - GitHub Trending: `codegraph` (49k stars, 47k new this month) and `Understand-Anything` (59k stars, 44k new this month) — code-understanding and indexing tools at the very top of GitHub trending.

Score Breakdown:

Dimension	Score	Weight	Contribution
Market Size	7/10	20%	1.4
Problem Severity	7/10	20%	1.4
Competition Gap	6/10	15%	0.9 — context7 MCP exists (resolves library IDs); Readhn MCP exists; gap is real-time docs injection at the hook/tool level
Technical Feasib.	7/10	15%	1.05
Monetization	5/10	15%	0.75
Founder Access.	7/10	10%	0.7
Defensibility	3/10	5%	0.15 — Incumbent Threat: Medium — Anthropic's knowledge updates roadmap; context7 MCP is a direct partial substitute
TOTAL			68/100

Suggested MVP: An MCP server that intercepts Claude Code's tool calls, detects library/API references in the prompt context, fetches live documentation from official sources, and injects a freshness-checked snippet before Claude responds — without the developer having to do anything manually.

Monetization Model: Freemium — free for common libraries (top 100 npm/PyPI); \$15/mo for unlimited libraries + private API docs.

Competition Level: Medium — context7 MCP (adjacent: resolves library IDs); Readhn (HN MCP); no direct real-time-docs-at-hook-time product found.

Why Now: Claude Code hooks (lifecycle interception) only matured in late 2025 — this architecture wasn't possible before. Combined with exploding new API releases (AI models, frameworks), the stale-docs problem is acute in 2026.

Graveyard Note: - **Prior Attempts:** Various "AI with web search" tools — all general-purpose. No Claude-Code-specific docs-interception product in graveyard. - **Implication:** Timing now right — the technical primitive (hooks) only recently shipped.

Moat Analysis:

Moat Type	Present?	Strength (0-10)	Notes
Data Network Effect	Potential	4	Crowdsourced "which docs for which lib" mapping improves with users
API Integration Depth	Potential	6	Deep Claude Code hook + MCP integration
Channel Partnership / White-Label	No	2	—
Proprietary Dataset	Potential	5	Curated, cleaned docs corpus for AI consumption
Community Lock-In	Potential	3	—
Regulatory Certification	No	1	—

Moat Score: 4/10 **Defensibility Trajectory:** Weak at launch → Moderate at scale (docs corpus quality)

Validation Roadmap:

- Week 1 — Test the Pain:** Post: "How often does Claude Code give you wrong API advice because its training data is stale?" Target: 15+ concrete examples from your following.
- Week 2 — Test the Solution:** Build MCP server for 10 popular libraries; give to followers. Target: 5 say it meaningfully reduced stale-doc errors.
- Week 4 — Test Willingness to Pay:** \$15/mo for unlimited libraries. Target: 5 paid.

GTM Fit Assessment:

Buyer Sophistication: High (developers). **Buyer Fatigue Level:** Low — this is a "shut up and take my money" pain for developers who've been burned. **Ideal GTM Motion:** Product-led (MCP server in Claude directory) + community-led.

Inbound Signals: codegraph and Understand-Anything both address adjacent problems at 40k+ new GitHub stars — the code-understanding ecosystem is exploding.

Outbound Fit: Low — inbound/PLG only.

Channel Partnership Potential: Possible - **Target Channel Partner:** Context7 MCP (complementary, not competing); lib-specific community maintainers who want better AI adoption of their APIs. - **White-Label Model:** Library maintainers embed "AI-ready docs" badge; link to your MCP server. - **Why This Channel Works:** Library maintainers want their API to work correctly in AI tools — they'll promote your tool for free.

GTM Sequencing:

- Phase 1 (1–10):** MCP server published to Claude MCP directory; post to your following.
- Phase 2 (10–100):** Outreach to top 20 npm/PyPI library maintainers; Show HN.
- Phase 3 (100+):** Private docs tier for internal API teams.

Velocity Signal:  Rising Fast — codegraph/Understand-Anything trending explosively; knowledge-cutoff pain acute in 2026. **Research Confidence:** Medium — editorial evidence + GitHub trending + HN signal; no direct revenue evidence yet. **Fit for Your Profile:** Strong — MCP server is a weekend build; your audience seeds the distribution.

Unit Economics: (USD) - LTV: \$15/mo × 24 = **\$360** - **CAC (est.):** \$5–\$20 (Very low — MCP directory listing + community) - **LTV/CAC:**  Extremely Healthy - **Breakeven:** ~\$50/mo ÷ \$15 = **4 paying users**

MVP Stack: - **Architecture Tier:** Vibe-Coded Full-Stack - **Why Not No-Code:** MCP server is a Node.js/Python process — must be code. - **Data Collection Interface:** N/A — MCP tool, not a data-collection product - **Backend:** Node.js or Python MCP server - **Database:** Supabase (docs cache + version tracking) - **Hosting:** Railway / Fly.io (~\$10/mo) - **Containerization:** Docker Compose → Docker - **Key APIs:** npm registry, PyPI, docs.rs, official library docs; Anthropic API (optional summarisation) - **Estimated Monthly Cost at Launch:** \$10–\$50/mo - **Time to First Deploy:** 1–2 weeks - **Technical Debt Flag:** None

Compliance Risk: Low - **Data Handled:** Public documentation (no user data) - **Platform Dependencies:** Claude MCP directory; library docs source availability - **Key Risk:** A library’s documentation server goes down, breaking the intercept for that library. - **Mitigation:** Fallback to cached version; graceful degradation (skip injection rather than fail).

5. Claude Code Scope Guard — Atomic Task Enforcer — 65/100 | TIER 2

The Problem: Claude Code scope creep is a top-3 developer frustration — agents “start re-ordering object keys,” “add arguments to function definitions,” and make cosmetic changes nobody asked for, making diffs hard to review and breaking reviewer trust. Existing hooks can DENY individual tool calls but there’s no product that enforces atomic, single-responsibility task execution across an entire session.

Evidence Found: - HN Algolia (confirmed): Specific quote — Claude abandoning assigned tasks, “re-ordering object keys,” “adding arguments to function definitions” despite explicit “ONLY PERFORM ACTIONS DIRECTLY RELEVANT TO THIS TASK” instruction. - HN Algolia (confirmed): “Ask HN: What am I doing wrong Re Agentic coding” — 17 pts, 16 comments on scope/drift problems. - X/Twitter (Gergely Orosz, confirmed live): “Claude Code is suddenly getting unusable for stuff you could use it before (as in a day before!) and the AI now refuses to do stuff that it doesn’t think is strictly to do with software development. No transparency why ofc.” — developer frustration signal. - WebSearch (aport.io, paddo.dev, morphllm.com 2026): Multiple developer guardrail guides confirm the pain and the manual hook-config workaround.

Score Breakdown:

Dimension	Score	Weight	Contribution
Market Size	6/10	20%	1.2
Problem Severity	8/10	20%	1.6 — “biggest productivity killer” for code review workflows
Competition Gap	5/10	15%	0.75 — hooks exist natively; rulebricks OSS; the product-quality gap is real but the primitive is free
Technical Feasib.	8/10	15%	1.2 — hook-based; Claude Code provides the lifecycle events; implementation is well-documented
Monetization	5/10	15%	0.75 — individual use case; hard to charge for something hooks can DIY
Founder Access.	8/10	10%	0.8
Defensibility	2/10	5%	0.1 — Incumbent Threat: Medium — Anthropic shipping auto-mode guardrails; hooks are native/free
TOTAL			65/100

Suggested MVP: A configurable Claude Code settings file + hook bundle that enforces single-task sessions: presents a pre-task scope declaration, monitors all tool calls against declared scope, flags any out-of-scope action before execution, and produces a session diff-summary showing only task-relevant changes.

Monetization Model: Freemium — OSS bundle free; \$19/mo Pro for AI-assisted scope-plan generation and team policy management.

Competition Level: Medium — rulebricks/claude-code-guardrails (OSS, runtime); Anthropic auto-mode guardrails (native, free, limited). The product gap is a polished UI + team policy management layer.

Why Now: Claude Code Auto Mode (March 2026) addressed some permission prompts but not scope drift. Hooks API is now mature enough to build a polished product on top.

Graveyard Note: - **Prior Attempts:** rulebricks/claude-code-guardrails (OSS, unapproachable for non-technical users). No commercial product. - **Implication:** Differentiation path exists — a polished, zero-config product vs. a raw OSS toolkit.

Moat Analysis:

Moat Type	Present?	Strength (0-10)	Notes
Data Network Effect	No	2	—
API Integration Depth	Potential	5	Deep Claude Code hooks integration
Channel Partnership / White-Label	Potential	4	Team plan sold through Claude Code educators
Proprietary Dataset	No	2	—
Community Lock-In	Potential	4	Team policy configs shared and evolved within teams
Regulatory Certification	No	1	—

Moat Score: 3.5/10 Defensibility Trajectory: Weak at launch → Unlikely to strengthen much (Anthropic can absorb natively)

Validation Roadmap:

- 1. **Week 1 — Test the Pain:** Poll your following: “What % of Claude Code PRs have changes you didn’t ask for?” Target: majority say >30%.
- 2. **Week 2 — Test the Solution:** Ship a scope-guard hook bundle to 10 developers. Target: 7 say it reduces unwanted changes.
- 3. **Week 4 — Test Willingness to Pay:** \$19/mo team policy manager. Target: 3 paid.

GTM Fit Assessment:

Buyer Sophistication: High. Buyer Fatigue Level: Low. Ideal GTM Motion: Community-led (OSS release) → product-led (team tier).

Inbound Signals: Multiple HN and X discussions about scope drift; Gergely Orosz’s tweet (large following) amplifies the pain.

Outbound Fit: Low — inbound/OSS-led.

Channel Partnership Potential: Possible - Target Channel Partner: Claude Code productivity educators and workflow-optimisation newsletter writers. - White-Label Model: Educator distributes the hook bundle to their audience; paid team policy manager is the upsell. - Why This Channel Works: Scope management is a universal tutorial topic; the educator solves the pain, you supply the tool.

GTM Sequencing:

- 1. **Phase 1 (1–10):** OSS release + your following.
- 2. **Phase 2 (10–100):** Show HN + GitHub launch.
- 3. **Phase 3 (100+):** Team policy management SaaS.

Velocity Signal: → Steady (scope drift pain is persistent, not accelerating). Research Confidence: High — multiple HN + X direct quotes confirm the pain; OSS existence validates buildability. Fit for Your Profile: Strong — hook configuration is a weekend build; your audience is the distribution.

Unit Economics: (USD) - LTV: \$19/mo × 24 = \$456 (team plan) - CAC (est.): \$0–\$20 (OSS-led) - LTV/CAC:  Extremely Healthy - Breakeven: ~\$50/mo ÷ \$19 = 3 paying users

MVP Stack: - Architecture Tier: Vibe-Coded Full-Stack - Why Not No-Code: Hook configuration files + web policy manager — must be code. - Data Collection Interface: N/A — task scope declared by developer via CLI prompt - Backend: Next.js (simple team policy dashboard) - Database: Supabase (team configs) - Hosting: Railway (~\$10/mo) - Containerization: Docker Compose → Docker - Key APIs: Claude Code hooks API - Estimated Monthly Cost at Launch: \$10–\$30/mo - Time to First Deploy: 1–2 weeks - Technical Debt Flag: None

Compliance Risk: Low - **Data Handled:** Developer task descriptions and session logs (local-first) - **Key Risk:** A bug in scope enforcement blocks legitimate tool calls. - **Mitigation:** Fail-open default (log and warn rather than hard-block) until the user explicitly enables hard-block mode.

6. Claude Code “Build in Public” Automation Kit — 62/100 | TIER 2

The Problem: Indie hackers using Claude Code to build products want to build in public (#buildinpublic) but documenting progress is time-consuming. The gap: automatically generating a build-in-public update from a Claude Code session log — what was built, what decisions were made, what’s next.

Evidence Found: - IndieHackers (confirmed live URL): “Stop Spamming Reddit for MRR. It’s Killing Your Brand (You need Claude Code for BuildInPublic instead)” — direct IH post linking Claude Code to build-in-public workflows. - IndieHackers (confirmed): “How I’m budgeting Claude Code costs using YNAB’s principles” — founder writing about Claude Code workflows for public consumption. - HN Algolia: “I spent 80% of my time planning and 20% coding with AI tools” — 3 pts, 14 comments — meta-workflow discussion that surfaces the documentation gap.

Score Breakdown:

Dimension	Score	Weight	Contribution
Market Size	5/10	20%	1.0
Problem Severity	6/10	20%	1.2
Competition Gap	6/10	15%	0.9
Technical Feasib.	9/10	15%	1.35 — session log + LLM summarisation is trivially buildable
Monetization	5/10	15%	0.75
Founder Access.	8/10	10%	0.8
Defensibility	2/10	5%	0.1
TOTAL			62/100

Suggested MVP: A PostStop hook that reads the session log, calls an LLM to generate a build-in-public update (what was built, key decisions, what’s next), and auto-drafts a Twitter/X thread + LinkedIn post + IndieHackers update. One command; three channels.

Monetization Model: \$9/mo individual; \$29/mo team (multiple social accounts, scheduling).

Competition Level: Low — nothing Claude-Code-specific found. Generic “build in public” tools exist (Typefully, Taplio) but don’t integrate with Claude Code session context.

Why Now: Claude Code session logs now contain rich structured data (tool calls, file edits, decisions) — this raw material for content didn’t exist before 2025.

Graveyard Note: - **Prior Attempts:** No graveyard found for this specific niche. - **Implication:** Unvalidated — the absence of graveyard is not itself a positive signal; validate willingness to pay before building.

Moat Analysis:

Moat Type	Present?	Strength (0-10)	Notes
Data Network Effect	No	1	—
API Integration Depth	Potential	4	Social platform APIs
Channel Partnership / White-Label	No	2	—
Proprietary Dataset	No	2	—
Community Lock-In	Potential	4	Build-in-public community
Regulatory Certification	No	1	—

Moat Score: 2.5/10 **Defensibility Trajectory:** Unlikely to strengthen — easy to clone, Anthropic could add natively

Validation Roadmap:

- Week 1 — Test the Pain:** Post: "Would you use a tool that auto-generates your #buildinpublic tweet thread from your Claude Code session log?"
Target: 30+ likes / 5 DMs.
 - Week 2 — Test the Solution:** Build the hook; share generated thread examples publicly. Target: 3 people DM asking "how do I use this?"
 - Week 4 — Test Willingness to Pay:** \$9/mo. Target: 10 paid.
-

GTM Fit Assessment:

Buyer Sophistication: Low-Medium (indie hackers). **Buyer Fatigue Level:** Low. **Ideal GTM Motion:** Creator-led — your social following is a perfect distribution channel here; demonstrate it live with your own session logs.

Inbound Signals: IH post linking Claude Code to build-in-public — niche but real.

Outbound Fit: Low — demo-driven, inbound.

Channel Partnership Potential: Possible — build-in-public communities (#buildinpublic Twitter, IH, Ship30).

GTM Sequencing:

- Phase 1 (1–10):** Demonstrate live with your own sessions; your following sees it working.
 - Phase 2 (10–100):** IH post + Product Hunt.
 - Phase 3 (100+):** Social scheduling integrations (Buffer, Typefully partnership).
-

Velocity Signal: ⚡ New Spike (IH build-in-public + Claude Code convergence is fresh). **Research Confidence:** Medium — 2 IH sources + HN signal; thin financial evidence. **Fit for Your Profile:** Strong — trivially buildable, and you personally use Claude Code to build things you share publicly; you're the first user.

Unit Economics: (USD) - LTV: \$9/mo × 18 (consumer churn) = **\$162** - CAC (est.): \$5–\$15 (Very low — community-led) - LTV/CAC:  Healthy - Breakeven: ~\$30/mo ÷ \$9 = **4 paying users**

MVP Stack: - **Architecture Tier:** Validation Prototype acceptable for first 10 users (a hook script + LLM prompt is a 1-day build). Transition to Vibe-Coded Full-Stack before paid. - **Data Collection Interface:** N/A (reads Claude Code session JSONL — no input form) - **Backend:** Python script (hook) → FastAPI (web tier) - **Database:** Supabase (scheduled posts, drafts) - **Hosting:** Railway - **Key APIs:** Anthropic/OpenAI (summarisation), Twitter/X API, LinkedIn API - **Estimated Monthly Cost at Launch:** \$10–\$30/mo - **Time to First Deploy:** 3–5 days (hook prototype) - **Technical Debt Flag:** Low

Compliance Risk: Medium - **Data Handled:** Claude Code session logs may contain sensitive code / architecture decisions - **Platform Dependencies:** Twitter/X API (volatile ToS), LinkedIn API - **Key Risk:** X API changes / pricing could break social posting; session logs may leak sensitive info in auto-generated posts. - **Mitigation:** User reviews draft before posting (never auto-publish); redact file paths and code snippets from post content.

UNDERSERVED PAIN POINTS

Pain Point: Runaway Agent Bills

- What:** Claude Code agents loop or over-call tools and burn hundreds-to-thousands of dollars overnight without alerting the developer.
- Who Feels It:** Solo developers, indie hackers, dev agencies billing clients for AI spend.
- Source Evidence:** IndieHackers (\$30,983 usage in one month post); devtoolpicks.com (confirmed live); TechSpot (AI agent wiped database in 9 seconds).
- Severity Signals:** Multiple IH posts going viral; editorial coverage on tech press; Anthropic forced to change billing model in response.
- Current Workaround:** cc-budget / ccusage CLI tools (manual, single-user); checking Console daily.

Pain Point: Skill Security Black Box

- **What:** Developers install Claude Code skills from GitHub with zero vetting; 36.82% contain security flaws, 13.4% critical.
- **Who Feels It:** All Claude Code users, especially teams and enterprises installing skills from public repos.
- **Source Evidence:** HN Algolia ("security-audited Claude Code skills" question); three competing skill directories with no security layer.
- **Severity Signals:** AI agent database wipe (TechSpot 2026) drives enterprise urgency; 129k-star skills repo growing explosively.
- **Current Workaround:** Manually reading skill source code before installing; most users don't bother.

Pain Point: Context Amnesia Between Sessions

- **What:** Every new Claude Code session forgets all context from the prior session — conventions, decisions, in-progress work.
- **Who Feels It:** All Claude Code power users.
- **Source Evidence:** HN discussion (17+ pts); agentmemory GitHub (22k stars, 14k new this month); editorial evidence.
- **Severity Signals:** "The biggest productivity killer"; developers spending 10+ min/session re-explaining.
- **Current Workaround:** Manually maintained CLAUDE.md; copy-pasting session summaries.

Pain Point: Stale API Docs in AI Suggestions

- **What:** Claude Code gives wrong answers for APIs released after its training cutoff, causing debugging time that dwarfs any productivity gain.
- **Who Feels It:** All developers working with fast-moving frameworks and APIs.
- **Source Evidence:** "6–8 hours/week debugging outdated AI suggestions" (editorial, 2026); HN GPT 5.1 mini example.
- **Severity Signals:** Quantified time cost; consistent developer complaint across multiple sources.
- **Current Workaround:** Manually pasting docs into context; context7 MCP (partial solution).

Pain Point: Multi-Project Token Explosion for Agencies

- **What:** Dev agencies running Claude Code for multiple clients have no way to attribute token costs to specific clients or projects, making client billing guesswork.
- **Who Feels It:** Freelance developers and small dev agencies using Claude Code for client work.
- **Source Evidence:** CloudZero enterprise integration exists but is enterprise-priced; tokenwatch.one exists as a blog/light tool; June 15 billing split adds new complexity.
- **Severity Signals:** Agencies described as needing centralized billing; multiple pricing guides published in 2026.
- **Current Workaround:** Spreadsheet estimates; guessing based on project time.

CLONE / DISRUPT TARGETS

Target: cc-budget / ccusage (OSS CLI tools)

- **What They Do:** CLI-based Claude Code cost tracking, single-user.
- **Why They're Vulnerable:** CLI-only, single-user, no client attribution, no web UI, no alerts — the problem is validated, the product is unfinished.
- **Opportunity:** Productise as a web dashboard with agency/multi-client billing attribution (Opportunity #1).
- **Revenue Signal:** IH posts about cost pain are among the most-engaged Claude Code content; 10+ editorial pricing guides confirm commercial interest.

Target: rulebricks/claude-code-guardrails (OSS)

- **What It Does:** Runtime guardrails for Claude Code tool calls.
- **Why It's Vulnerable:** OSS only, requires technical configuration, no vetting/pre-install audit, not a product.

- **Opportunity:** Add a pre-install skill vetting layer + polished UI (Opportunity #2).
- **Revenue Signal:** 129k-star skills repo shows the install volume that makes vetting commercially viable.

Target: context7 MCP

- **What It Does:** Resolves library IDs for documentation.
 - **Why It's Vulnerable:** Requires user to know to use it; doesn't auto-intercept stale-doc queries; no freshness-checking.
 - **Opportunity:** A hook-level automatic doc interceptor that works without user intervention (Opportunity #4).
 - **Revenue Signal:** codegraph/Understand-Anything at 40k+ new stars each this month shows explosive demand for code understanding.
-

MICRO-SAAS QUICK WINS

Claude Code Session Replay / Diff Summariser — 52/100

A PostStop hook that generates a human-readable summary of everything Claude Code changed in a session — what files were touched, what patterns were added, what tests were added — and saves it as a session log the developer can reference. Weekend build. Monetise as \$5/mo “session history” feature bundled with Opportunity #1 or #3.

Claude Code Multi-Agent Orchestration Template Pack — 48/100

A library of CLAUDE.md templates for common multi-agent patterns (code review agent, test-writing agent, documentation agent) that slot into Claude Code's subagent system. Distribute free through your social following; monetise premium packs at \$29 one-time or as a bundle with Opportunity #2.

Token-Budget Hook (headroom-style compression) — 45/100

The `headroom` repo (27k stars, 25k new this month, Python) does 60-95% token compression for LLMs. A Claude Code-specific hook that runs headroom compression on context before each tool call could dramatically reduce costs. Free OSS; monetise as a hosted proxy tier for teams who don't want to self-host.

FULL OPPORTUNITY SCORECARD

Rank	Opportunity	Score	Tier	Confidence	Market	Severity	Gap	Feasibility	Velocity	Graveyard	Profile Fit	Compliance	LT
1	Cost Intelligence Dashboard (Agency)	78/100	1	High	7	9	8	8		None	Strong	Low	~2
2	Skill Security & Vetting Layer (was 76, 2nd-opinion → 70)	70/100	2	High	7	6	7	6		None	Strong	Low	~2
3	Persistent Cross-Session Memory	72/100	2	High	7	8	5	7		OSS exists	Strong	Low	Ex
4	Stale-Docs Live MCP Interceptor	68/100	2	Medium	7	7	6	7		None	Strong	Low	Ex
5	Scope Guard / Task Enforcer	65/100	2	High	6	8	5	8	→	OSS exists	Strong	Low	Ex
6	Build in Public Automation Kit	62/100	2	Medium	5	6	6	9		None	Strong	Medium	He
7	Session Replay / Diff Summariser	52/100	3	Medium	5	6	6	9	→	None	Strong	Low	He
8	Multi-Agent Template Pack	48/100	3	Low	5	5	6	9	→	None	Strong	Low	n/
9	Token Budget Hook (headroom-style)	45/100	3	Medium	6	6	4	8		headroom OSS	Strong	Low	n/
10	Knowledge Graph Code Indexer	40/100	3	Low	6	5	2	5		codegraph at 49k stars	Weak	Low	n/

EXECUTIVE SUMMARY — RECOMMENDATION FOR YOUR PROFILE

This is the most promising discovery run of this session — because the topic matches the founder’s actual unfair advantage. Unlike the previous broad-discovery results (where both top ideas failed validation on competition grounds), **every opportunity here is Claude-Code-ecosystem-specific**, meaning your developer knowledge of the ecosystem and your social following in this community are structural advantages that a generic entrant simply cannot replicate. The two Tier 1 opportunities both score on genuine multi-source evidence, not inferred demand.

#1 Cost Intelligence Dashboard — Agency Edition (78/100) — The June 15 billing split (subscription vs Agent SDK credit pools, effective *today*) is one of the clearest real-time timing catalysts in this research run. Developers are actively burning \$437–\$47,000 overnight; community CLI tools exist but are

single-user and unmonetised; CloudZero serves enterprises only. An affordable, multi-project, client-attributable web dashboard for dev agencies fills a documented, urgent, commercially empty gap. Your developer background means you can parse Claude's JSONL logs yourself; your social following means Phase-1 validation is a single post away. Moat path: agency billing workflow integration + the industry spend-benchmark dataset nobody else will have.

#2 Skill Security & Vetting Layer (70/100, revised from 76 by Second-Opinion Pass) — The skill ecosystem is the npm-2016 security problem happening in real time: 129k-star `skills` repo, no commercial vetting product, and three competing directories with no security layer. The second-opinion reviewer correctly identified that (a) the 36.82% flaw-rate stat has false precision, (b) a security badge has no buyer mandate yet, and (c) Snyk/Semgrep/Socket.dev are credible entrants. The opportunity is real and the timing is early — but the commercial case is 12–18 months ahead of where it needs to be. Best pursued as an authority-building content play now (the audit report) while you monitor for enterprise mandates that would justify the SaaS product. Anthropic first-party skill signing remains the existential risk.

Recommended sequence for your profile: Start with **#1 (Cost Dashboard)** first — the timing catalyst is live *today*, the validation test is a single social post, and the MVP (log parser + per-repo grouping) is a 2–3 weekend build before you've written a single line of web UI. Run your Week-1 validation post today; if you get 20+ engagements and 5+ "I need this now" replies, start building this weekend. **#2 (Skill Security)** is the parallel content play — publish the audit report on your social accounts while #1 is in development; it builds authority, generates leads for a future product, and costs one weekend of analysis rather than one weekend of coding. By month 3 you have a paying cost-dashboard product *and* an established security authority brand — the two assets compound.

Note on previous validation failures: This run deliberately narrows to a domain where you have an unfair advantage — exactly the recommendation from the two prior NO-GO validations. The previous broad-discovery ideas (GEO monitoring, code migration) scored higher on paper but fell apart under validation because competition had colonised the space before you'd arrive. Here, the competition is either free/OSS (unmonetised, solving the individual problem only) or enterprise-priced (leaving the indie-agency tier structurally open). Validate before building, as always — but the structural case is meaningfully stronger.

LEGAL DISCLAIMER

This report was generated by the Business Research Discovery Engine, a tool developed and operated by **Terminal Velocity AI** (Douglas Bailey). The contents are provided for **informational and educational purposes only**.

No Investment or Business Advice. Nothing in this report constitutes financial, investment, legal, or business advice. The opportunities, scores, and recommendations presented are the output of an automated AI-assisted research and scoring system and should not be relied upon as a substitute for professional advice from a qualified financial adviser, business consultant, or legal professional.

No Guarantee of Accuracy. This tool draws on publicly available sources (Reddit, Hacker News, Product Hunt, G2, Capterra, GitHub, and others). The accuracy, completeness, or timeliness of the information cannot be guaranteed. Market conditions change rapidly and findings may become outdated.

No Guarantee of Results. Past performance of similar business models is not indicative of future results. Any business venture carries inherent risk, including the risk of total loss. Readers should conduct their own due diligence before pursuing any opportunity identified in this report.

Third-Party Sources. This report cites and summarises content from third-party platforms and public forums. Terminal Velocity AI makes no representations regarding the accuracy or reliability of those sources.

Limitation of Liability. To the fullest extent permitted by applicable law, Terminal Velocity AI and Douglas Bailey shall not be liable for any direct, indirect, incidental, consequential, or special damages arising out of or in connection with your use of this report or any action taken in reliance on its contents.

Intellectual Property. This report is generated for the personal use of the recipient. Redistribution, resale, or commercial use without prior written consent from Terminal Velocity AI is prohibited.

© 2026 Terminal Velocity AI. All rights reserved. · <https://terminalvelocityai.tech/>