

DAILY BRIEF · WHAT AI ENGINEERS ARE BUILDING

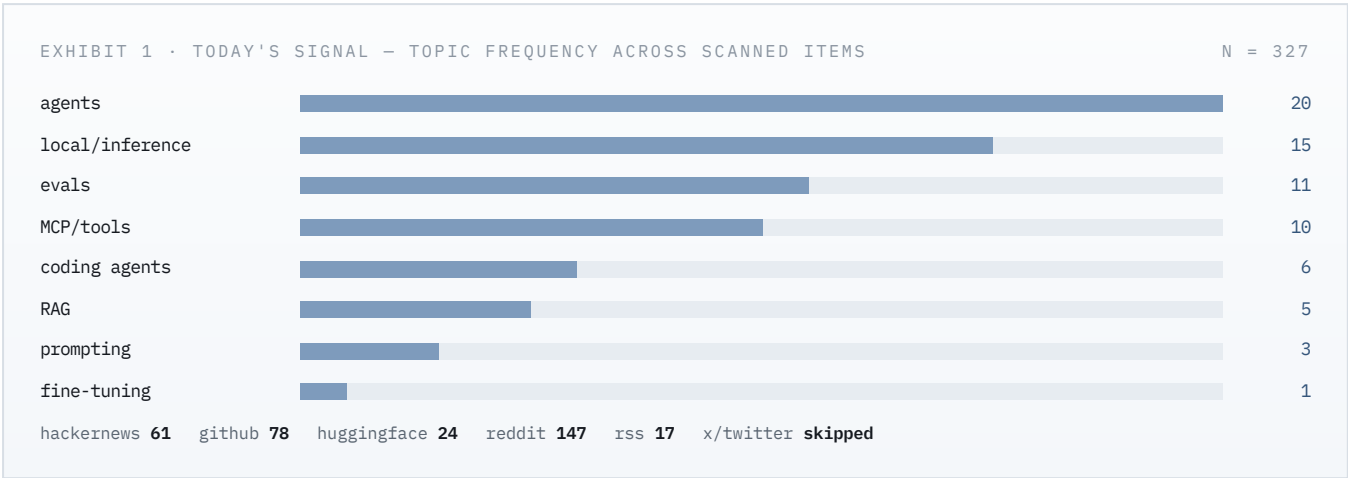
Today in AI engineering

2026-09-09

Daily AI-engineering signal (mechanical fallback — LLM error: 404 Client Error: Not Found for url: <https://api.groq.com/openai/v1/chat/complet>)

5 deep dives 3 builder stories 15 quick hits 327 items scanned mode: **mechanical fallback**

- agents
- local/inference
- evals
- mcp/tools
- coding agents
- rag
- prompting
- fine-tuning



DEEP DIVES

I-have-ADHD: A skill to stop coding agents from burying the answer

HN

(no summary)

Why it matters. Flagged by the ranking heuristics as a real, artifact-backed build.

TRY THIS

Open the link and skim the repo / paper.

<https://github.com/ayghri/i-have-adhd>

okf-memory/okf-agent-memory

GitHub

Git-native persistent memory for AI coding agents. Implements Google OKF v0.2 with sub-300µs in-memory BM25 search, embedded MCP server, and progressive disclosure. Slashes token bloat by 80% with zero external databases or dependencies. Built in pure Go.

Why it matters. Flagged by the ranking heuristics as a real, artifact-backed build.

TRY THIS

Open the link and skim the repo / paper.

<https://github.com/okf-memory/okf-agent-memory>

noskillish/bankmcp

[GitHub](#)

BankMCP™: your AI can now read your bank. Self-hosted, read-only MCP server for your own bank accounts via open banking (Enable Banking). Standard MCP; tested with Claude and Ollama.

Why it matters. Flagged by the ranking heuristics as a real, artifact-backed build.

TRY THIS

Open the link and skim the repo / paper.

<https://github.com/noskillish/bankmcp>

PillCrew/deedchain

[GitHub](#)

Benchmark of report fidelity for browser agents: does the agent's final report tell the truth about what it actually did?

Why it matters. Flagged by the ranking heuristics as a real, artifact-backed build.

TRY THIS

Open the link and skim the repo / paper.

<https://github.com/PillCrew/deedchain>

Show HN: CUDA/graphics in QEMU-KVM VMs without passing the Nvidia card to them

[Show HN](#)

Hi HN. I built this because I wanted a VM with a real GeForce GPU without losing the card on my desktop. It allows you to run CUDA and graphics in VMs without VFIO/vGPU by forwarding the NVIDIA driver's own ioctl surface from the guest to the host. That means stock libcuda and real Vulkan ICD run in

Why it matters. Flagged by the ranking heuristics as a real, artifact-backed build.

TRY THIS

Open the link and skim the repo / paper.

<https://github.com/reindertpelsma/nvkvm-pv>

WHAT PEOPLE ACTUALLY BUILT

We built a repo-local state layer for Codex sessions (open source)

[r/LLMDevs](#)

I use Codex across many sessions on the same project. Picking up the work often means checking which decisions still apply and whether an earlier test result covers the code that's there now. We built Sigma Operator Stack to keep that information in the repository. I'm one of the

STEAL THIS

Flagged by keyword heuristics as a first-person build write-up — read the post for the details.

https://www.reddit.com/r/LLMDevs/comments/1waptr7/we_built_a_repolocal_state_layer_for_codex/

Anyone tried self-hosting 3D object generation?

r/LocalLLaMA

Found this post: <https://www.reddit.com/r/ClaudeAI/s/At9QNEbawK> of someone fully developing a game through AI writing code, generating 3D models, etc. (he is using Meshy to create the 3D models). I was curious about creating something similar with Hermes Agent (or other?) that wo

STEAL THIS

Flagged by keyword heuristics as a first-person build write-up — read the post for the details.

https://www.reddit.com/r/LocalLLaMA/comments/1waonrd/anyone_tried_selfhosting_3d_object_generation/

I built a lightweight prompt injection detector using MiniLM + Logistic Regression — looking for technical feedback

r/LLMDevs

I've been experimenting with a lightweight approach to detecting prompt injection before untrusted input reaches an LLM or agent. The main constraint was: can this be done without running another large LLM for every input? The current architecture is intentionally simple: Input t

STEAL THIS

Flagged by keyword heuristics as a first-person build write-up — read the post for the details.

https://www.reddit.com/r/LLMDevs/comments/1walht6/i_built_a_lightweight_prompt_injection_detector/

QUICK HITS

Multi-Agents LLM Financial Trading Framework HN

<https://github.com/TauricResearch/TradingAgents>

nahidalam/world_model_from_scratch GITHUB

Code repository for World Model From Scratch book

https://github.com/nahidalam/world_model_from_scratch

Show HN: VolAnti – Open-source acoustic detector for fibre-optic FPV drones SHOW HN

This project has been a month in the making, and 9 fully functional units will be sent tomorrow to an undisclosed civilian site across the I

<https://github.com/agamrossen/VolAnti>

On the Navier–Stokes Millennium Prize Problem HN

Further discussion: <https://simonwillison.net/2026/Sep/8/on-navier-stokes/> , <https://news.ycombinator.com/item?id=49621697>

<https://twitter.c>

<https://openai.com/index/navier-stokes-solution/>

XHToken/Spark-X2.5-4B HF MODEL

Trending model on Hugging Face — text-generation, transformers; 945 likes, 10,661 downloads.

<https://huggingface.co/XHToken/Spark-X2.5-4B>

hsj576/virtual-ai-infra-team GITHUB

An open-source virtual AI Infra team that discovers, benchmarks, and safely upgrades local LLM inference—with independent quality gates, a s

<https://github.com/hsj576/virtual-ai-infra-team>

Show HN: Homebutler – reports what changed on your server, not what's running SHOW HN

<https://github.com/Higangssh/homebutler>

openbmb/MiniCPM5-2B HF MODEL

Trending model on Hugging Face — text-generation, transformers; 788 likes, 2,879 downloads.

<https://huggingface.co/openbmb/MiniCPM5-2B>

Show HN: DriveSync – fast Git-styled Google Drive sync CLI SHOW HN

<https://github.com/scaleninja/drivesync>

Qwen/Qwen3.8-27B HF MODEL

Trending model on Hugging Face — image-text-to-text, transformers; 14441 likes, 6,712,160 downloads.

<https://huggingface.co/Qwen/Qwen3.8-27B>

ISTA-DASLab/Qwen3.8-27B-GSQ-RCO-GGUF HF MODEL

Trending model on Hugging Face — image-text-to-text, gguf; 684 likes, 479,597 downloads.

<https://huggingface.co/ISTA-DASLab/Qwen3.8-27B-GSQ-RCO-GGUF>

Mistral raises €3B HN

<https://mistral.ai/news/mistral-makes-sovereign-open-weight-ai-to-frontier/>

Show HN: AI-first open-source alternative to Microsoft and Google Office SHOW HN

InjOffice — AI-first open-source alternative to Microsoft Office and Google Workspace

<https://github.com/injectinglabs/injoffice>

Show HN: Notifkit, self-hosted notification infra for email, SMS, push,etc. +MCP SHOW HN

<https://github.com/devkitshq/notifkit>

chadhurley25075-png/pd-bridge GITHUB

Heterogeneous prefill/decode for DeepSeek-V4-Flash: CUDA prefill (DGX Spark, vLLM) -> Metal decode (Mac Studio, oMLX) over plain 10GbE

<https://github.com/chadhurley25075-png/pd-bridge>
