

DAILY BRIEF · WHAT AI ENGINEERS ARE BUILDING

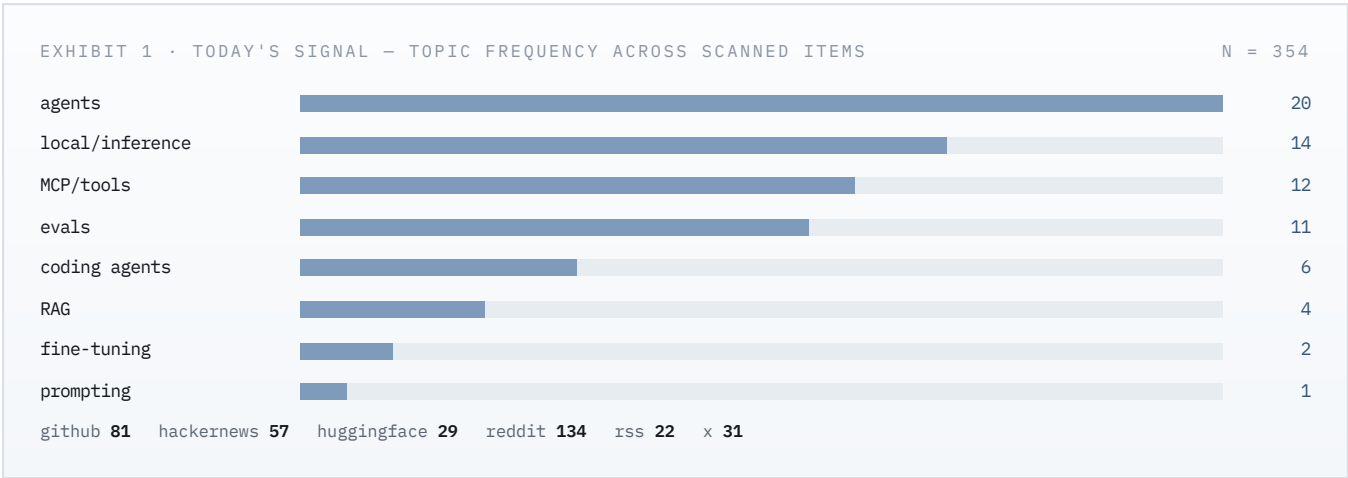
Today in AI engineering

2026-08-07

AI engineers are building and sharing tools for LLM inference, automation, and model optimization

3 deep dives 2 builder stories 5 quick hits 354 items scanned

llm inference mcp rag ai agents



DEEP DIVES

FareedKhan-dev/kimi-k3-in-c

github

Fareed Khan has built a portable C99 implementation of the Kimi K3 model, which can run inference on a single CPU with 8.24 GB of RAM. This implementation does not require any frameworks or GPU acceleration.

Why it matters. This project demonstrates the feasibility of running large language models on resource-constrained hardware, making AI more accessible to a wider range of users.

TRY THIS

Try building and running the Kimi K3 model on your own hardware to see how it performs.

<https://github.com/FareedKhan-dev/kimi-k3-in-c>

llm inference c99

patchy631/time-to-first-token

[github](#)

Patchy631 has created a 10-week roadmap for optimizing LLM inference serving, including techniques such as quantization, speculative decoding, and benchmarking.

Why it matters. This project provides a comprehensive guide for developers to improve the performance of their LLM inference pipelines, reducing latency and increasing throughput.

TRY THIS

Follow the roadmap and apply the optimization techniques to your own LLM inference project.

<https://github.com/patchy631/time-to-first-token>

[llm](#)[inference](#)[optimization](#)

ombori/deepseek-v4-flash-0731-sglang-4x-rtx-pro-6000

[github](#)

Ombori has developed a reproducible recipe for running the DeepSeek V4 Flash 0731 model on 4x RTX PRO 6000 GPUs, achieving significant performance improvements.

Why it matters. This project demonstrates the potential for large-scale GPU acceleration of LLM inference, enabling faster and more efficient processing of complex AI workloads.

TRY THIS

Try running the DeepSeek V4 Flash 0731 model on your own GPU hardware using the provided recipe.

<https://github.com/ombori/deepseek-v4-flash-0731-sglang-4x-rtx-pro-6000>

[llm](#)[inference](#)[gpu acceleration](#)

WHAT PEOPLE ACTUALLY BUILT

Show HN: A terminal glued to the macOS dock

[hackernews](#)

A terminal application that integrates with the macOS dock, allowing users to quickly access and manage their terminal sessions.

How it works. The application uses a combination of scripting and system integration to provide a seamless user experience.

STEAL THIS

The project demonstrates the potential for creative and practical solutions to common productivity challenges.

<https://github.com/palamim/starboard>

Show HN: mcp-use v2 rebuilt from scratch for stateless 2026-07-28 MCP spec

[hackernews](#)

A rebuilt version of the mcp-use framework, designed to work with the stateless 2026-07-28 MCP specification.

How it works. The new version of the framework provides a more efficient and scalable way to build and manage MCP applications.

STEAL THIS

The project highlights the importance of staying up-to-date with evolving specifications and standards in the field of AI and automation.

<https://manufact.com/blog/mcp-use-v2>

QUICK HITS

MiniMaxAI/MiniMax-H3 HUGGINGFACE

A trending model on Hugging Face, with 2801 likes and 12,102 downloads.

<https://huggingface.co/MiniMaxAI/MiniMax-H3>

deepseek-ai/DeepSeek-V4-Flash-0731 HUGGINGFACE

A trending model on Hugging Face, with 2675 likes and 617,900 downloads.

<https://huggingface.co/deepseek-ai/DeepSeek-V4-Flash-0731>

mrpulor-gh/nuphus-mcp GITHUB

A desktop automation MCP server that allows users to control their computer using AI agents.

<https://github.com/mrpulor-gh/nuphus-mcp>

Sparkfetch/sparkfetch GITHUB

An open-source web fetching and extraction API that can be used to turn any URL into clean, structured, LLM-ready content.

<https://github.com/Sparkfetch/sparkfetch>

drmikecrypto/WebSearchFree GITHUB

A free, open-source alternative to Tavily, providing keyless web search and extract capabilities for AI agents and RAG applications.

<https://github.com/drmikecrypto/WebSearchFree>
