

Open Source This Month: a Cheaper DeepSeek, an Open Code Corpus, and Audits That Check Themselves

Kabui, Charles

2026-09-22

[Read at ToKnow.ai](#)

Open Source This Month: a Cheaper DeepSeek, an Open Code Corpus

DeepSeek, OpenBMB, OpenResearch and Cloudflare on cost and verification

890

bytes of cache per token,
a quarter of V4-Flash

\$0.27

cost per Artificial Analysis
Intelligence Index task

550B

tokens of open code data:
400B selected, 150B made

September 22, 2026

ToKnow.ai

Open-source releases this month were about cost and infrastructure, not capability jumps. DeepSeek's [V4.1-Flash](#) is a 552-billion-parameter multimodal model that activates 16 billion per token and fits a one-million-token context into a key-value cache of 890 bytes per token,

a quarter of its predecessor's. [Artificial Analysis](#) rates it 39 on its Intelligence Index at 226 tokens per second and \$0.27 per task. [OpenBMB](#) released UltraData, which archives 192 million public GitHub repositories and ships about 400 billion tokens of selected code plus 150 billion tokens of generated exercises. [OpenResearch](#) turns Claude Code, Codex, Cursor and OpenCode into research agents with git-tracked experiments, and Cloudflare's [security-audit-skill](#) runs six-phase code audits whose findings are independently checked.

The effect is a lower floor for building on open models. A team can fine-tune on a curated code mixture instead of scraping GitHub, or run a fast multimodal model locally. DeepSeek trails the leader on Artificial Analysis' index by ten points but costs a fraction per task, because its design saves memory rather than chasing accuracy. OpenBMB's own numbers show selection helps by itself: continued training on the selected code tier beat the unselected tier by 7.8% on EvalPlus.

Openness still needs qualifiers. UltraData's licence is Apache-2.0, but its content comes from public repositories, and the Cloudflare skill earns trust by separating confirmed findings from ones needing validation. A survey of AI in games ([arXiv 2609.16679](#)) makes a related point: results rarely carry between engines, so evidence has to be rebuilt for each setting.

Read More: [the previous open-source radar](#), [on forecasting models and torrent-distributed weights](#)

Sources:

- [DeepSeek-V4.1-Flash model card](#) (Hugging Face)
- [DeepSeek V4.1 Flash: performance and price analysis](#) (Artificial Analysis)
- [UltraData-Code dataset](#) (Hugging Face)
- [OpenResearch: a local-first workspace for research agents](#) (GitHub)
- [Cloudflare security-audit-skill](#) (GitHub)
- [AI for Games in the Foundation Model Era](#) (arXiv)

***Disclaimer:** For information only. Accuracy or completeness not guaranteed. Illegal use prohibited. Not professional advice or solicitation. **Read more:** [/terms-of-service](#)*