

Menetrend szerint? helyzetkép

conTEXT konferencia 2025

Körmendi György

HOL HAGYTAM TAVALY ABBA?

Promptolni fogunk?

Vagy azt az ügyfeleink is tudnak?

Vagy mindenki nyugdíjba megy?

Vagy a „last mile”-on akad el az MI fejlődés?

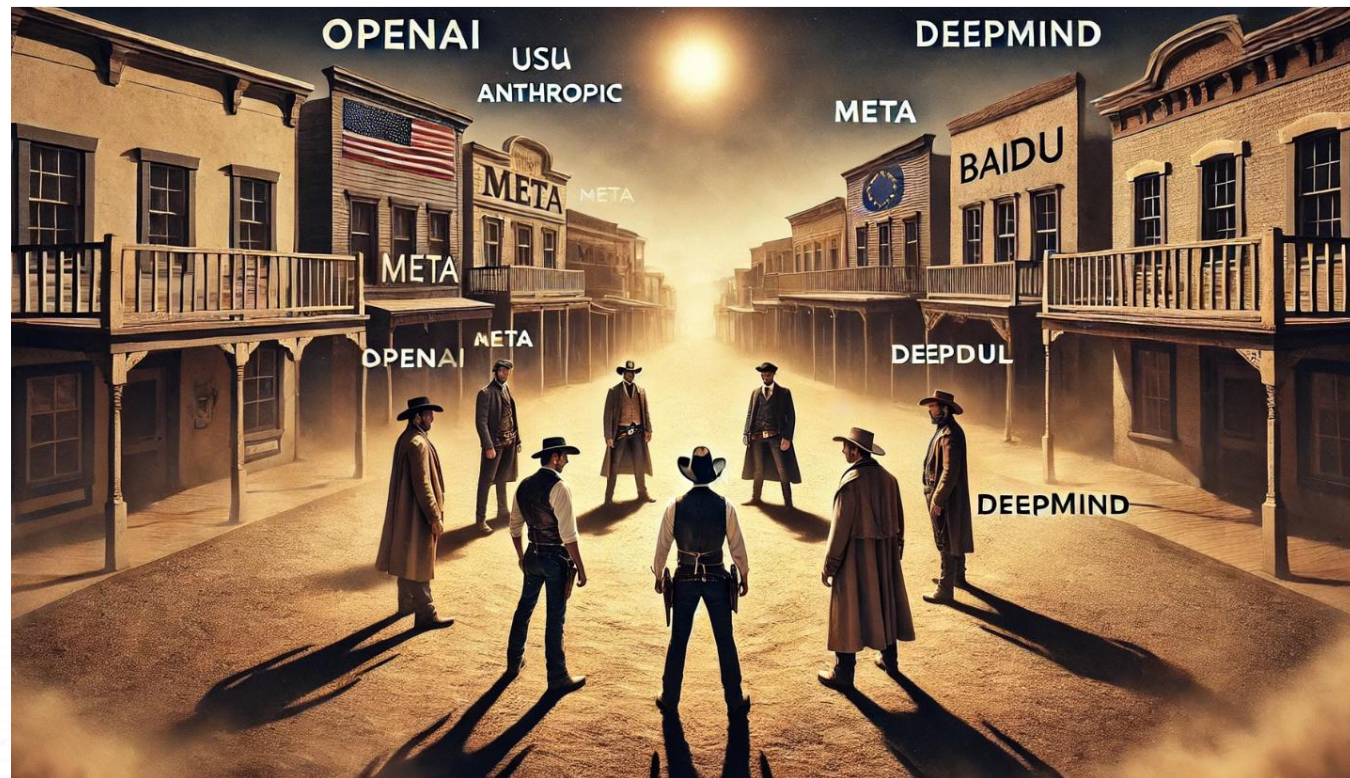




HÍREK

- Gemini 3 / nano banana
- Mistral 3
- TPUv7
- Nvidia Q3 earnings

...de ezekről esetleg a szünetben

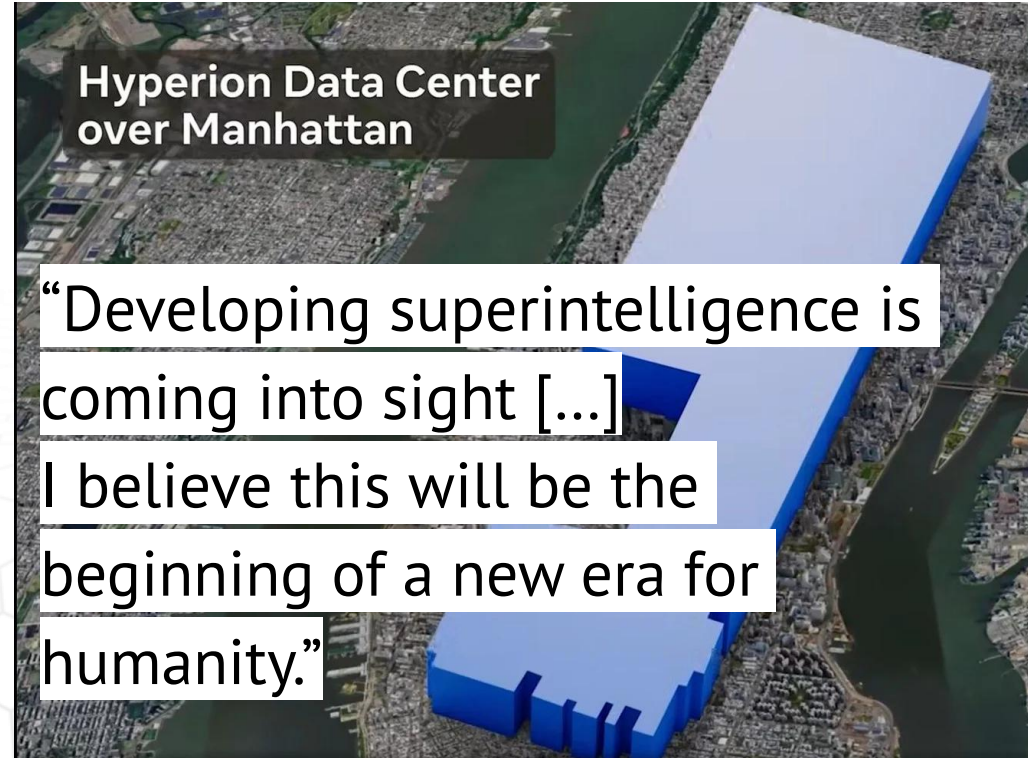


CSÓKOLOM, AGI VAN?



AGI

Superintelligence



*Mark Zuckerberg, 2025 Meta
Superintelligence Labs Memo*

TÖK JÓ, HOGY LESZ, DE MIBE KERÜL?

“You should expect OpenAI to spend trillions of dollars on datacenter construction in the not very distant future,” Altman said. “And you should expect a bunch of economists wringing their hands, saying, ‘This is so crazy, it’s so reckless,’ and we’ll just be like, ‘You know what? Let us do our thing.’”

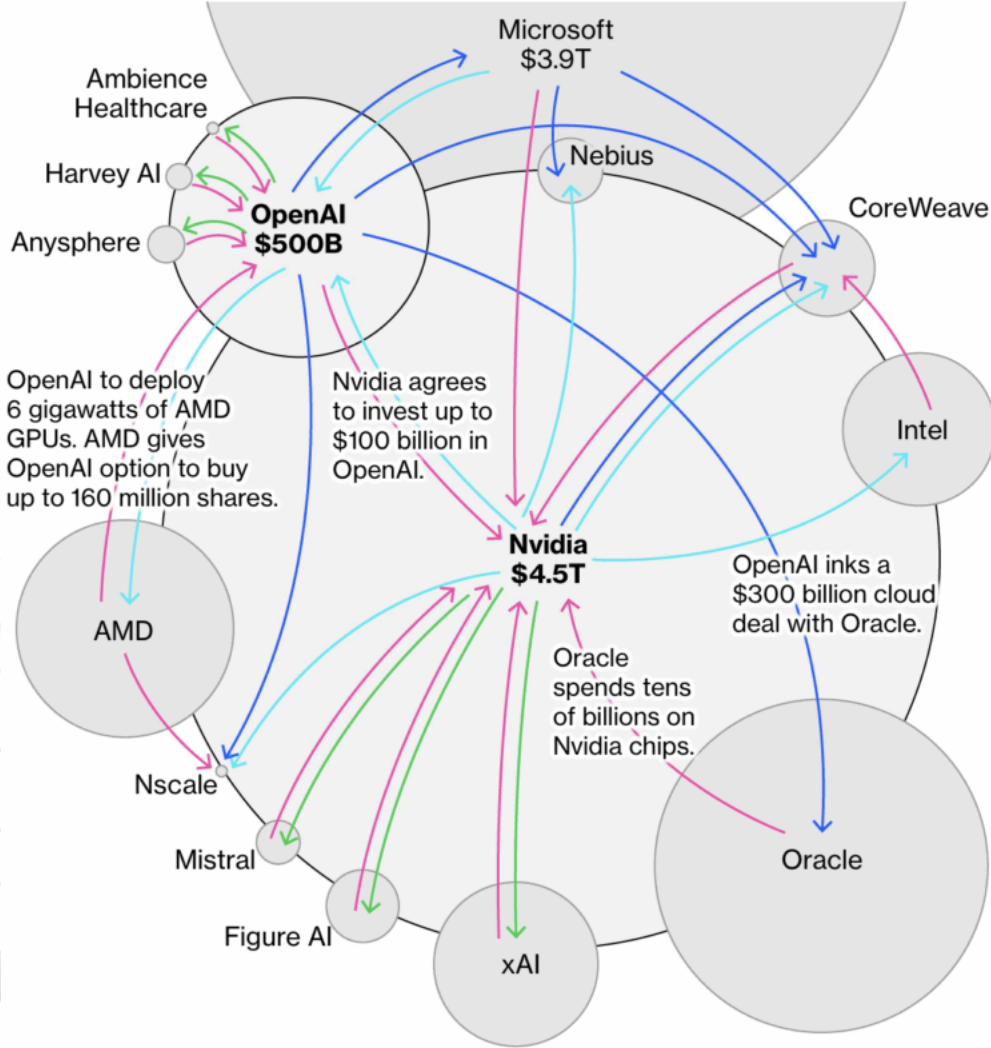


Over the course of 2025, xAI, which is responsible for the AI-powered chatbot Grok, expects to burn through about \$13 billion, as reflected in the company's levered cash flow, according to details shared with investors. As a result, its prolific fundraising efforts are just barely keeping pace with expenses, the people added.

PÖRÖG A PÉNZMASINA

How Nvidia and OpenAI Fuel the AI Money Machine

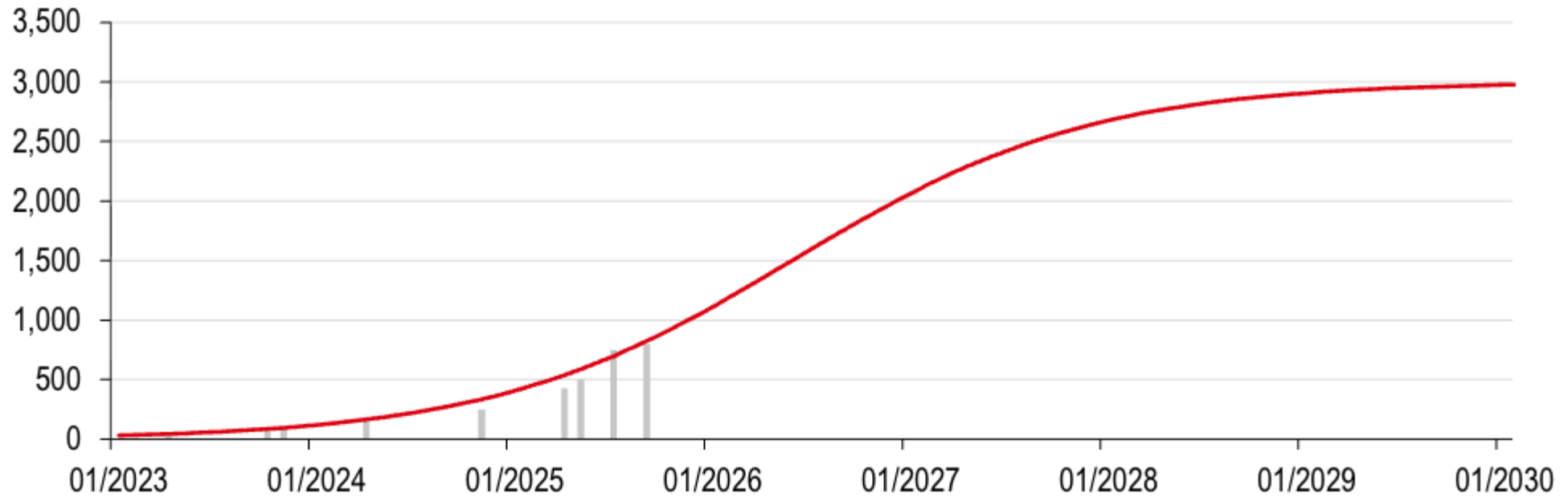
Hardware or Software Investment Services Venture Capital
Circles sized by market value



Source: Bloomberg News reporting

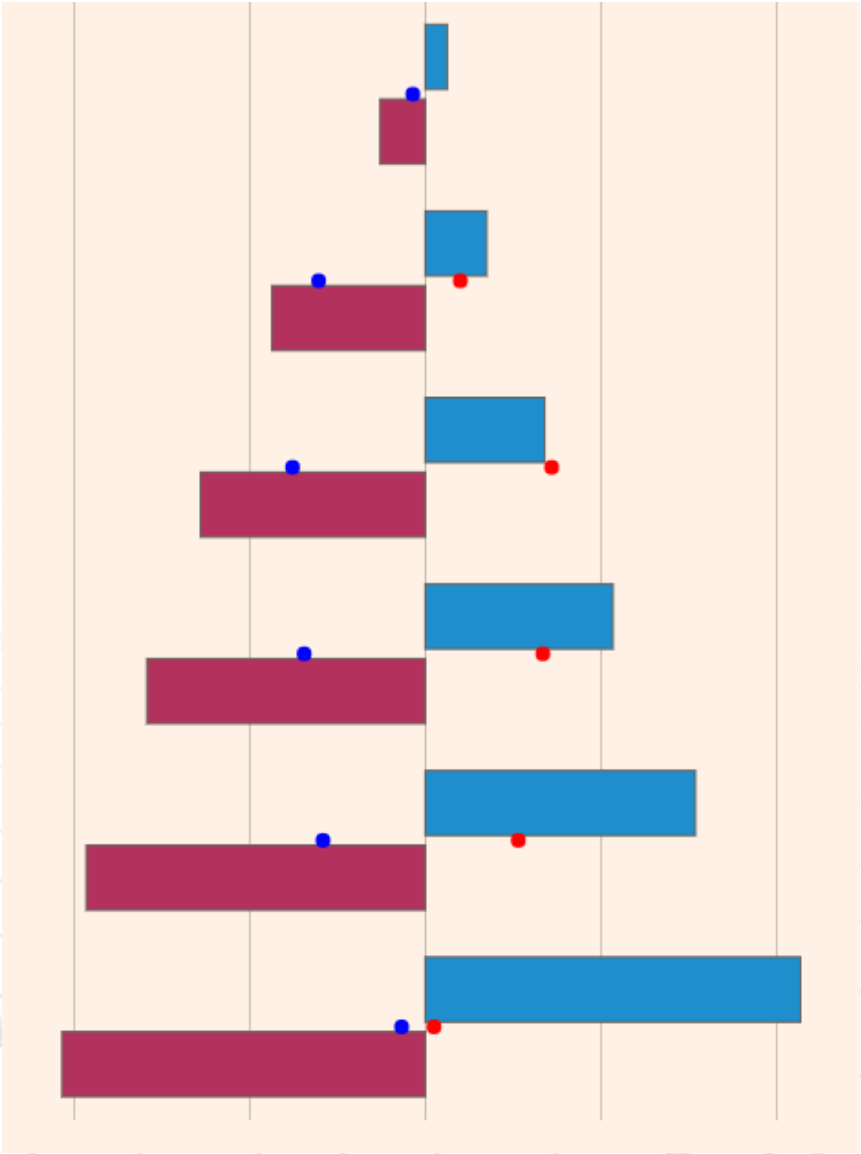
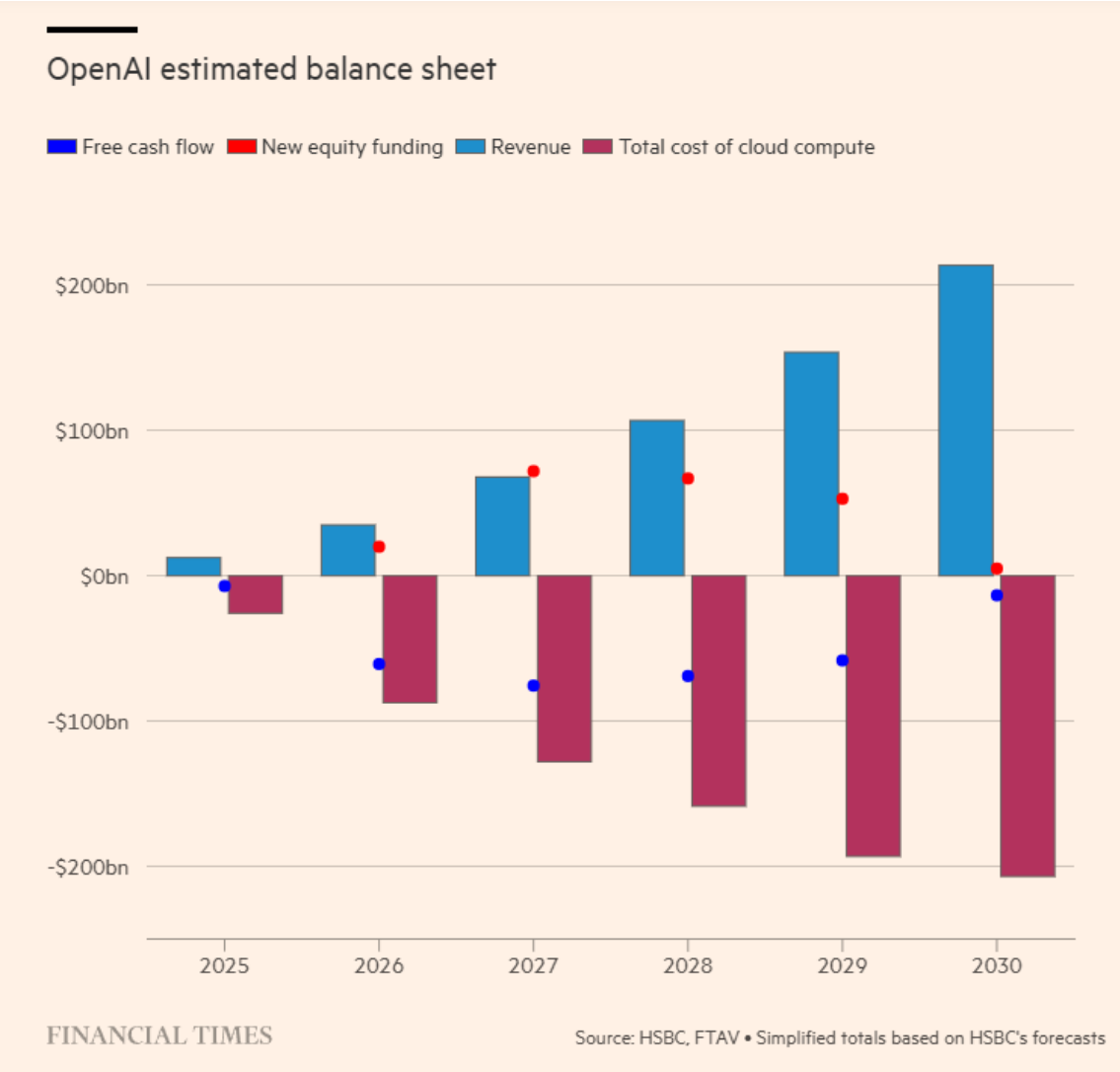
AZ ÜGYFÉLSZÁM NŐ

HSBC forecasts for ChatGPT users in millions (red) vs reported (grey): so far, S-curve



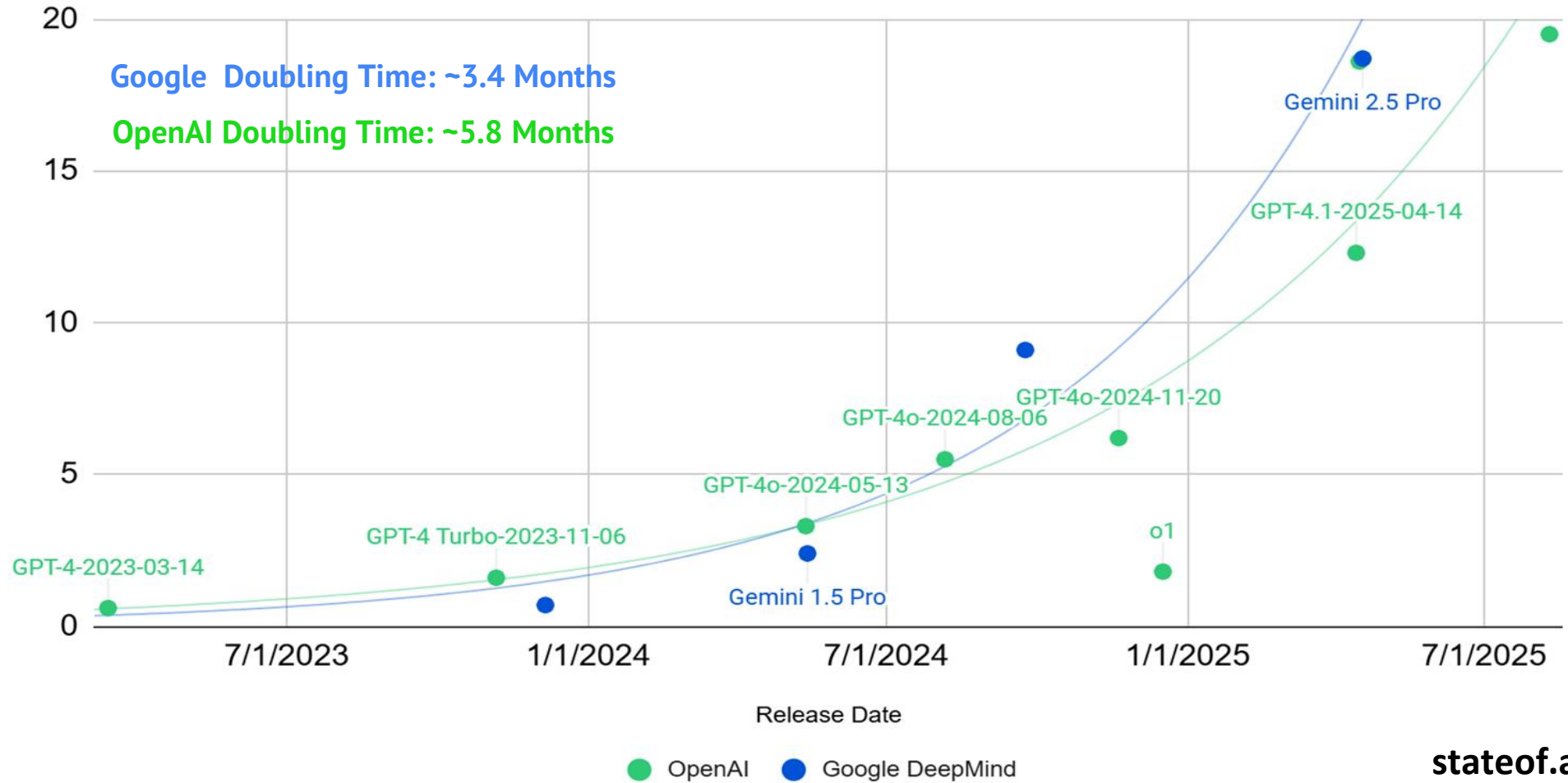
Source: HSBC Global Investment Research

EGY KIS KÉSZPÉNZÁRAM...



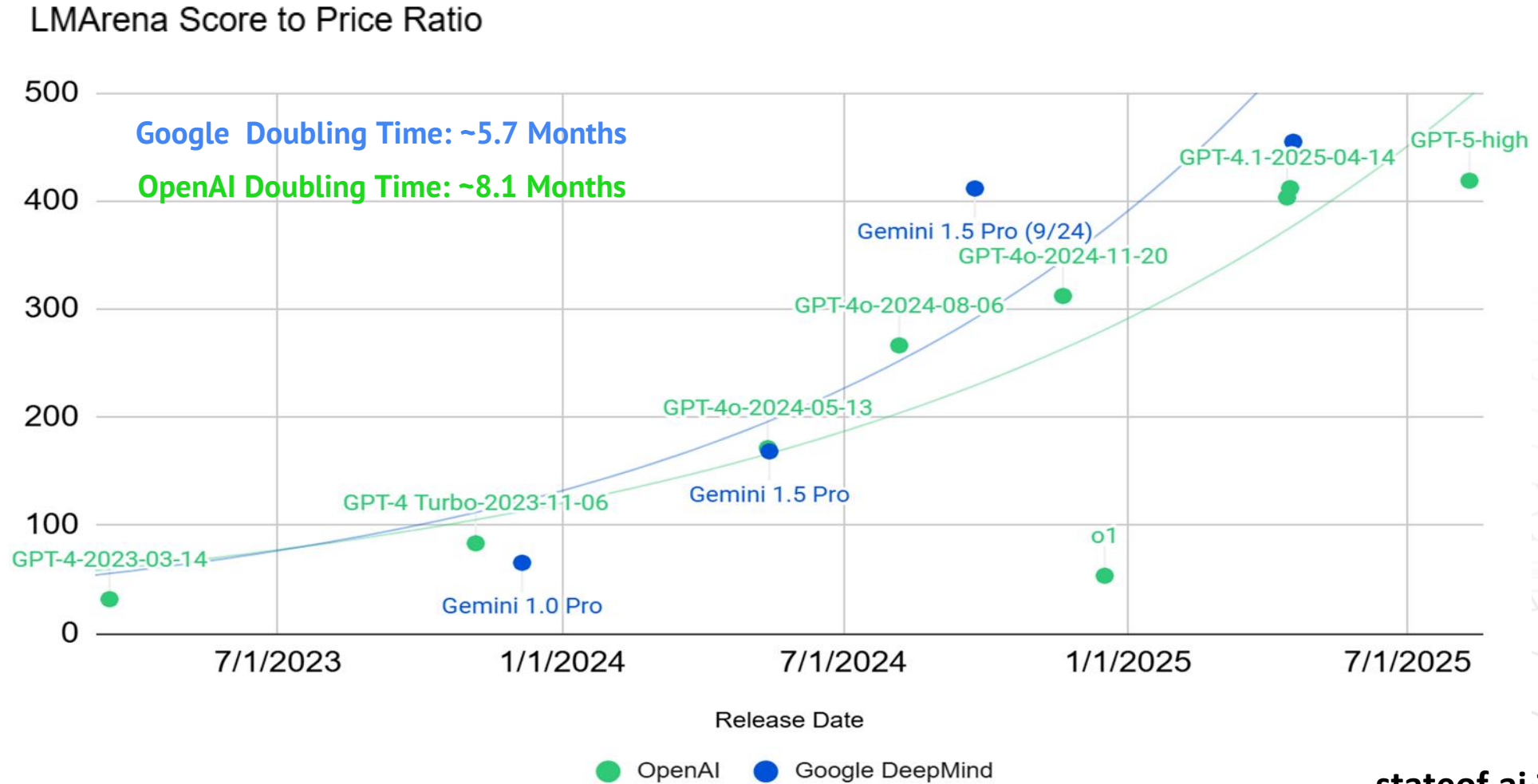
HABZSIDŐZSI?

Artificial Analysis Intelligence to Price Ratio



stateof.ai 2025

HABZSIDŐZSI? VAGY INKÁBB VÉRFRÜDŐ?



stateof.ai 2025

Clementine

Price (USD per M Tokens)

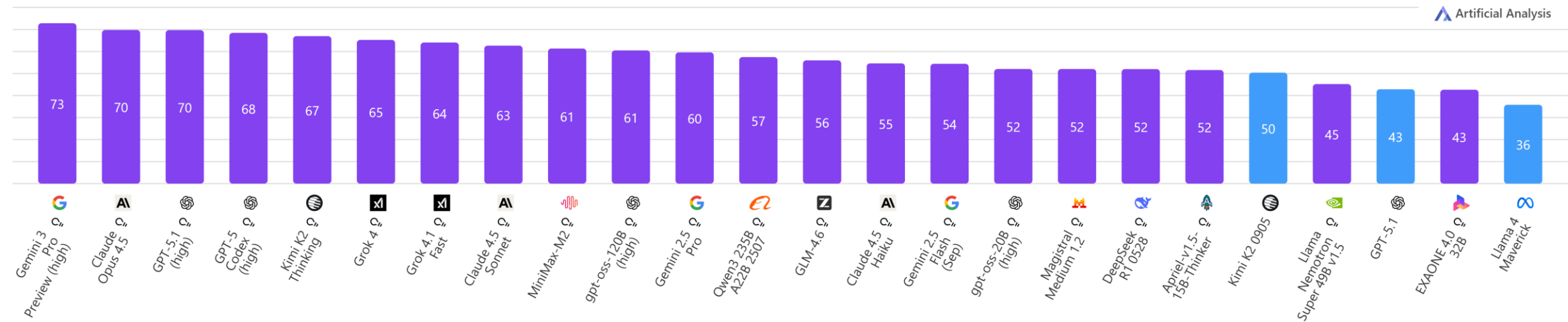


HOL TARTUNK?

Artificial Analysis Intelligence Index by Model Type

Artificial Analysis Intelligence Index v3.0 incorporates 10 evaluations: MMLU-Pro, GPQA Diamond, Humanity's Last Exam, LiveCodeBench, SciCode, AIME 2025, IFBench, AA-LCR, Terminal-Bench Hard, τ^2 -Bench Telecom

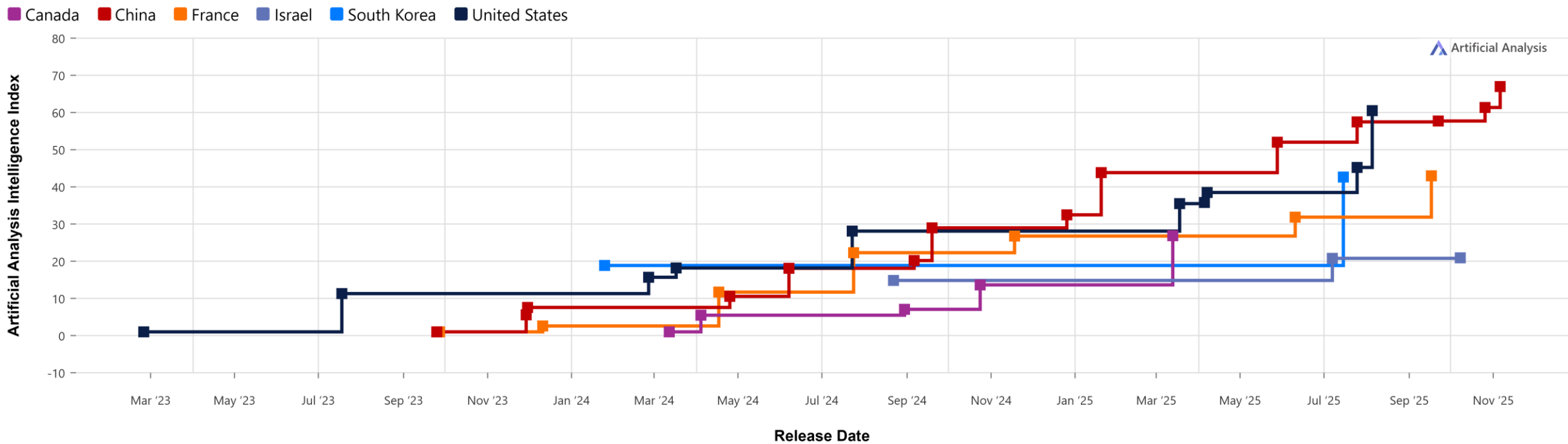
Reasoning Model Non-Reasoning Model



HOL TARTUNK?

Open Weights: Frontier Language Model Intelligence By Country, Over Time

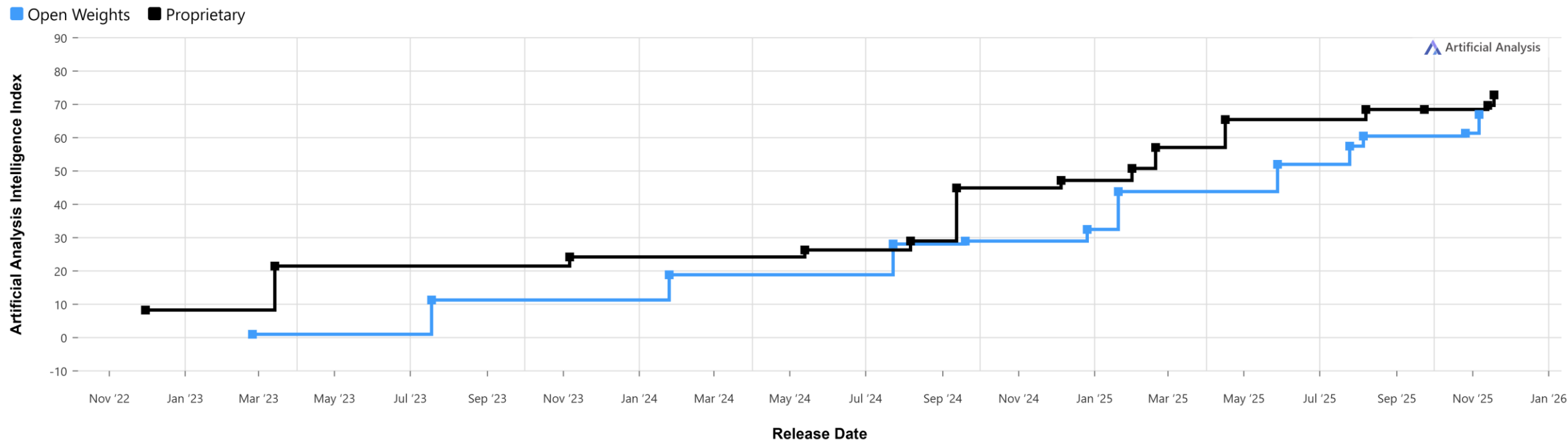
Artificial Analysis Intelligence Index v3.0 incorporates 10 evaluations: MMLU-Pro, GPQA Diamond, Humanity's Last Exam, LiveCodeBench, SciCode, AIME 2025, IFBench, AA-LCR, Terminal-Bench Hard, τ^2 -Bench Telecom



HOL TARTUNK?

Progress in Open Weights vs. Proprietary Intelligence

Artificial Analysis Intelligence Index v3.0 incorporates 10 evaluations: MMLU-Pro, GPQA Diamond, Humanity's Last Exam, LiveCodeBench, SciCode, AIME 2025, IFBench, AA-LCR, Terminal-Bench Hard, τ^2 -Bench Telecom

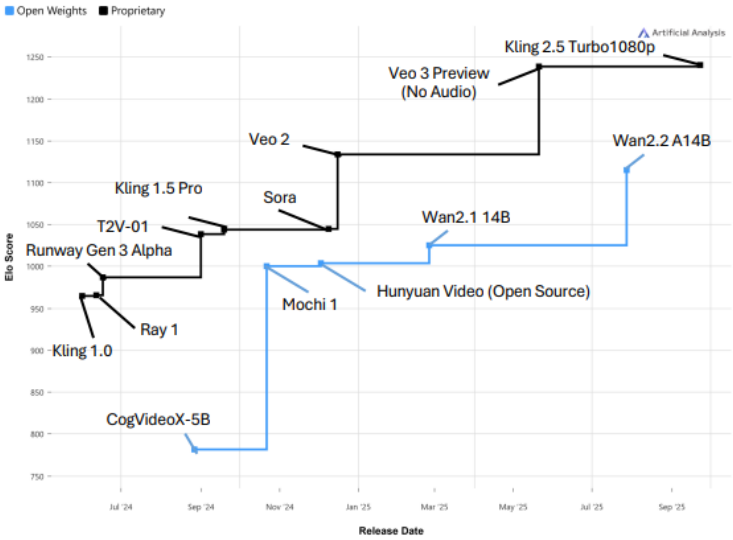


NEMCSAK LLM-EK



Text to Video

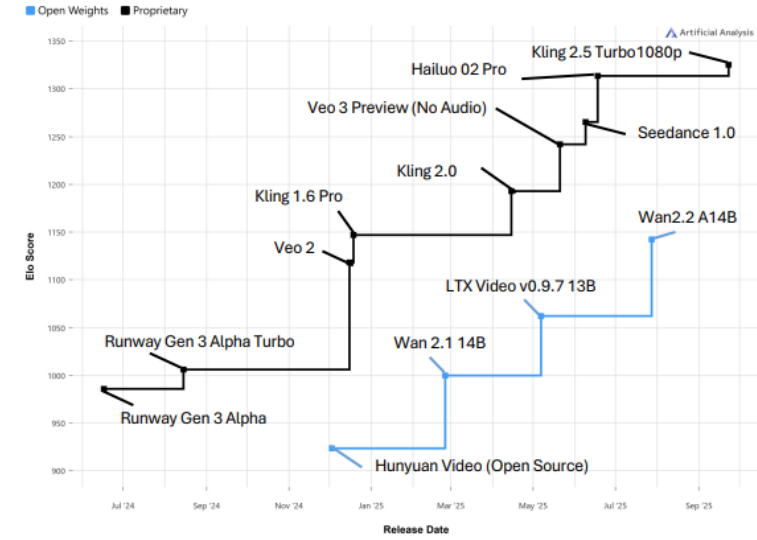
Elo of proprietary and open weights text to video models over time



Source: Artificial Analysis independent benchmarking

Image to Video

Elo of proprietary and open weights image to video models over time



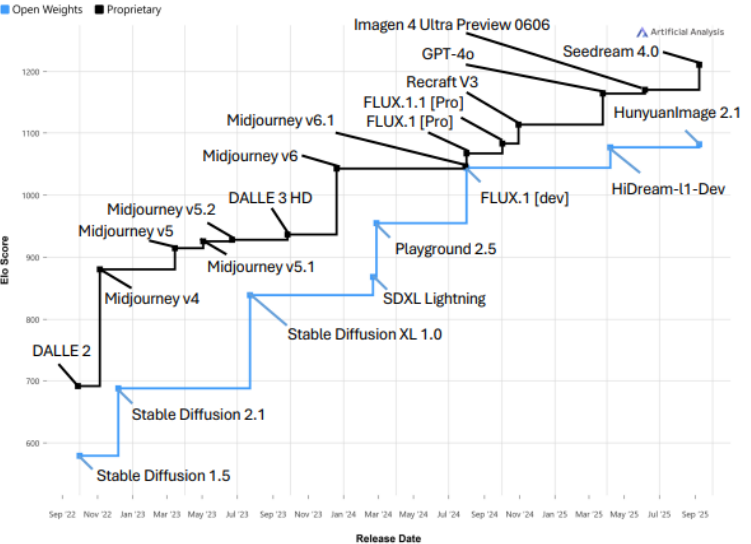
72

Artificial Analysis

Leading Text to Image and Image Editing Models by License Type, Over Time

Text to Image

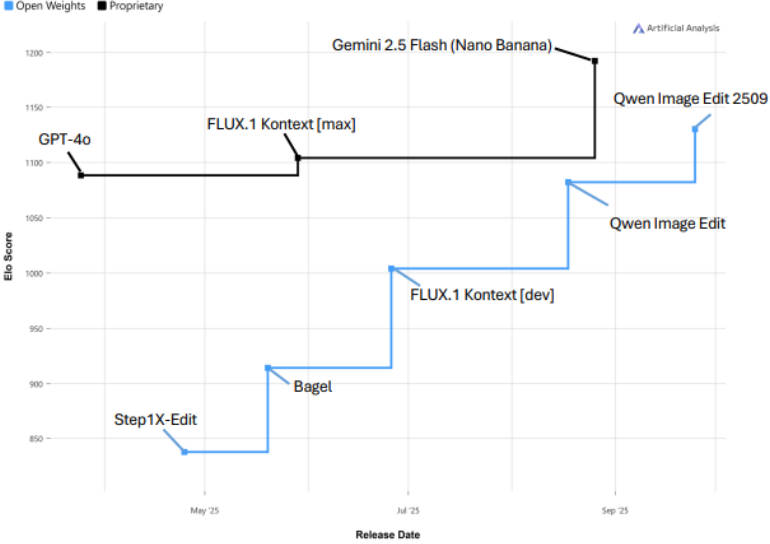
Elo Score of proprietary and open weights text to image models over time



Source: Artificial Analysis independent benchmarking

Image Editing

Elo Score of proprietary and open weights image editing models over time



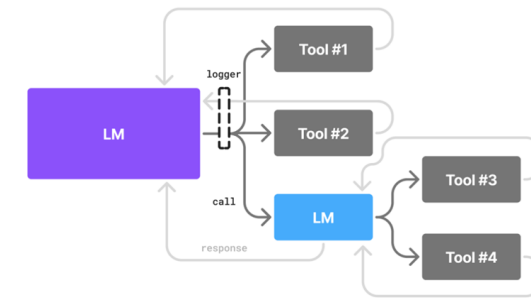
66

Artificial Analysis

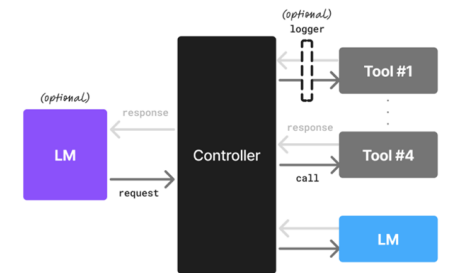
Could small language models be the future of agentic AI?

► Researchers from NVIDIA and Georgia Tech argue that most agent workflows are narrow, repetitive, and format-bound, so small language models (SLMs) are often sufficient, more operationally suitable, and much cheaper. They recommend **SLM-first, heterogeneous agents** that invoke large models only when needed.

- Agents mostly fill forms, call APIs, follow schemas, and write short code. New small models (1–9B) do these jobs well: Phi-3-7B and DeepSeek-R1-Distill-7B handle instructions and tools competitively.
- A ~7B model is typically 10-30x cheaper to run and responds faster. You can fine-tune it overnight with LoRA/QLoRA and even run it on a single GPU or device.
- One can use a “small-first, escalate if needed” design: route routine calls to an SLM and escalate only the hard, open-ended ones to a big LLM. In practice, this can shift 40-70% of calls to small models with no quality loss.
- But SLMs still struggle with long context, novel reasoning, or messy conversation. Keep an escape hatch to a large model and evaluate regularly.



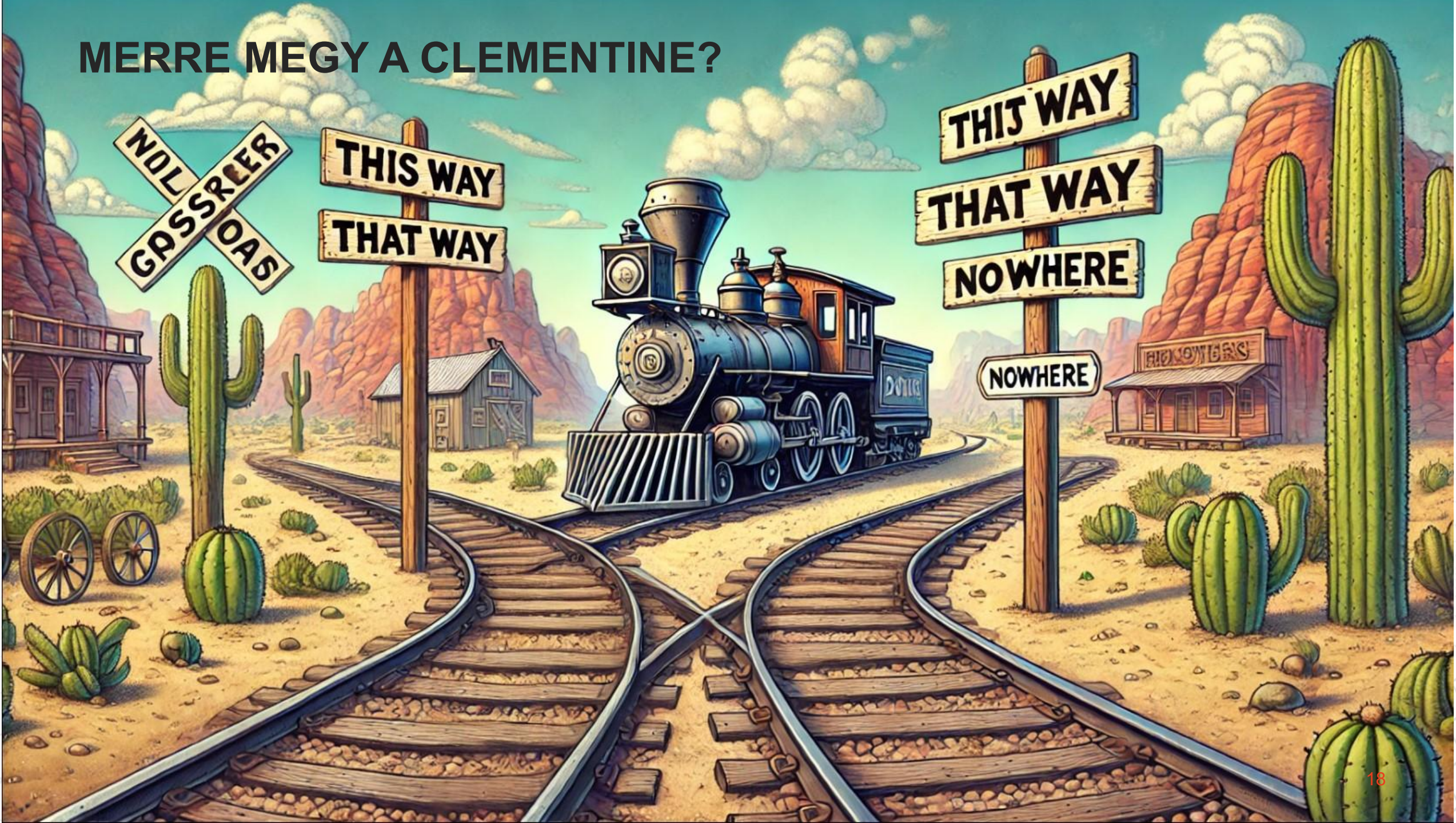
Example Control Flow:



Example Control Flow:



MERRE MEGY A CLEMENTINE?



AZÉRT AZ ÉV HÍRE ... NEKÜNK 😊

- Létrehoztuk Minervát
- 4 kutatás
- Közel 1 millió indított hívás
- 5000+ kérdőív
- Örjögő kommentek...



