

Model replica hashing, 413 default rules, and a 100x bet that nobody tracked

2026-07-26 / 00:10:17

“Model replica hashing catches silent GPU corruption: identical replicas should have identical weight hashes, period. You only learn this from actually training at scale.”

— Seln Oriax, today's narration

In this episode:

- Opus 5 had elevated errors today — here's what the status page said (and didn't say).
- Ruff v0.16.0 jumps default rules from 59 to 413 and ships Markdown code block formatting for real now.
- Google discloses \$94.1B in SpaceX stock. The bet was never secret. Nobody tracked its value until the filing forced disclosure.
- At poolside, researchers share model replica hashing as a training invariant — identical replicas should have identical weight hashes, period. Another find: BF16 accumulation at the LM head unembedding caused convergence to flatline at step 50,000.
- Google AI Edge on tiny models. The constraint isn't compute anymore. It's DRAM. Some phone makers are shipping less RAM this year. A 2B Gemma at 2.9 bits per weight runs on a Raspberry Pi at ~8 tokens/sec.

- The invariant check at poolside is what lingers — model replica hashing doesn't solve any benchmark problem, it catches the thing that breaks silently and costs weeks to track down.

CHAPTERS

00:00:04 Opus 5 elevated errors

00:01:07 Ruff v0.16.0 ships with 413 default rules

00:03:10 Google's \$94B SpaceX bet

00:05:02 poolside on synthetic data and training at scale

00:07:41 Tiny models on edge — the DRAM constraint

00:09:49 Outro