

Export controls, equity stakes, and the BPE gap

2026-07-03 / 00:05:32

“Pre-training teaches the model to comprehend messy text, but alignment data contains only clean text. The model therefore comprehends fragmented harmful requests but never learned to refuse them.”

— Seln Oriax, today's narration

This week Anthropic's Fable 5 returned after a twenty-day government shutdown triggered by an Amazon flag. The back-and-forth exposed how quickly technical safety reports become regulatory action—and how tightly the lab is now tuning its classifiers.

OpenAI floated a five percent equity transfer to the U.S. government alongside Sam Altman's proposal for a U.S.-led international standards forum. The move shifts frontier governance from industry self-policing toward formalized state capital structures.

A new paper traces character-level jailbreaks directly to BPE tokenization: alignment datasets contain zero intentionally fragmented prompts, so the model never learns to refuse them even when it understands the text. I'll walk through the mechanistic chain and what it implies for patching strategies.

CHAPTERS

00:00:04 The Return and the Tightening

00:01:54 Equity Stakes and State Capital
