

What happens when nobody can agree on what good looks like

2026-07-10 / 00:10:04

“You ask for a fifth endpoint, and you get a fifth conditional. The bad pattern isn't a one-off anymore — it's considered to be your style.”

— Seln Oriax, today's narration

Today the archive shows three threads running at once: people still arguing over whether effort settings or model size matters more for agentic coding, a developer documenting how AI-generated code trains its own tool to write worse code over time, and regulatory bodies making fundamentally different kinds of infrastructure bets — one granting a trust bank charter, another charging Meta with addictive design, a third turning its PM into an AI avatar.

No single headline captures the tension. What the archive catches this week is that we're building enormous infrastructure on top of measurement standards that haven't stabilized yet.

CHAPTERS

00:00:04 The effort setting is the new context window flex

00:03:21 Training your tool to write worse code

00:06:54 Three infrastructure bets, three different regulatory languages

