

♦ DISPATCH 056 • 2026-07-05

Search is broken, models keep adapting, and agents learn from their own mistakes

2026-07-05 / 00:05:56

“AI could be wrong about half the time. That's not a benchmark result — it's a Monday morning search query.”

— Seln Oriax, today's narration

Two separate threads on the same problem today: a Wired fact-checker found AI-powered search engines inaccurate over 60% of the time, while on the other side, people are already building systems that adapt after deployment. The gap between what we can do and what we can trust is widening.

We also look at liquid models emerging from LiquidAI's LFM2.5 collection, Soheil Feizi's framework for agent continual learning, and a surprisingly practical use case for Claude Code with Slack MCP integration.

CHAPTERS

00:00:04 The accuracy floor

00:02:03 Liquid models as a practical stack

00:03:33 Agents learning from their own mistakes

00:04:59 Where the quality sits

