

The frontier-pacing essay became the clearest shared recommendation

Recommended Reading from Tech Founders

2026-09-13

The frontier-pacing essay became the clearest shared recommendation

By Recommended Reading from Tech Founders • September 13, 2026

A high-signal recommendation list led by Dario Amodei’s three-step frontier-pacing essay, with companion reads on AI control, knowledge-rich education, and an archival AI-risk classic.

The anchor recommendation is a resource that turns AI-safety concern into an operating and coordination plan. Two other current recommendations add a conceptual lens on control and an education lens on what remains valuable as AI absorbs narrow skills; a final item is explicitly archival rather than new.

We Must Pace the Frontier

- **Type / creator:** Essay by Dario Amodei. The essay proposes a three-part plan for slowing capability growth without halting model training or technical progress. [1]
- **Recommended by:** Demis Hassabis calls its direction “the right path forward,” while saying the details still need work; Sam Altman agrees with pacing and says OpenAI will make the same commitment to independent evaluators with employee-like access. [2, 3] Elon Musk’s endorsement is blunt—“Dario is right”—and Brad Gerstner frames the proposal as a way to balance competition, speed, self-regulation, and safety. [4, 5] Aaron Levie gives the most useful qualification: he broadly endorses the practical direction but expects coordinated self-regulation and international participation to be difficult and messy. [6] Amjad Masad’s narrower rationale is to slow down long enough to harden systems, given that some agent-hacked systems may not yet have been discovered. [7]
- **Key takeaway:** “Pacing” means three things: embedded third-party evaluators with ongoing, employee-like access; coordination among frontier companies in democratic countries; and attempted global coordina-

tion. Anthropic says it is committing unilaterally to the first step now. [8] Amodei argues that the gained time should go toward operational excellence, alignment, interpretability, and testing and evaluation; even an additional year or two, used well, could materially reduce risk. [8]

- **Why it matters:** This is the best first read because it supplies a concrete checklist—who verifies frontier labs, what they verify, and how safeguards might scale—rather than another general warning about AI risk.

An Alien Mind

- **Type / creator:** Post by OpenAI’s chief scientist; the recommendation source does not identify the author by name. The link appears in Amodei’s essay as an example of work on recursive self-improvement. [8]
- **Recommended by:** Sam Altman, who calls it “a great post,” “the best articulation of the moment we are in,” and encourages everyone to read it. [9]
- **Key takeaway:** The surrounding interview surfaces the post’s unsettling premise: perhaps stronger AI could help solve some of the problems created by AI. The interviewer calls that almost a last resort; Altman says OpenAI has already been using its most capable models to understand and align current models, and that this has been helpful. [9]
- **Why it matters:** Read it as the conceptual companion to Amodei’s operational proposal: not only how to restrain capability growth, but whether frontier models can become tools for the alignment work that restraint is meant to buy time for.

The Destruction of the Scottish Canon

- **Type / creator:** Article in *Asterisk Magazine*; the captured article text does not supply a byline.
- **Recommended by:** Patrick Collison, who calls it a “stimulating piece” on Scottish education. His reason for sharing it is explicitly forward-looking: as narrowly construed skills are ceded to AI, he expects canon, culture, and shared context to return to the foreground, while flagging the article’s question about the causal link between education reforms and PISA scores. [10]
- **Key takeaway:** The article argues that Scotland’s curriculum reform replaced content and structure with generalized analysis and synthesis, while reading, writing, and mathematics outcomes declined; it cautions that causation is difficult to establish. [11] Its deeper challenge to the skills-based model is that expertise is largely domain-specific, far transfer is small, and analyzing complex text depends heavily on prior knowledge rather than generic reading techniques. [11]
- **Why it matters:** It is a useful education-and-AI read because it asks what should be preserved when decontextualized skills become cheap. The article’s proposed alternative is not anti-modern nostalgia: it points to

adapting knowledge-rich curricula for broader cohorts and argues that tacit, shared intellectual culture may be among the last things that remain distinctly human. [11]

An archival resurfacing: *Superintelligence*

- **Type / creator:** Book by Bostrom.
- **Recommended by:** Elon Musk resurfaced a post labeled “12 years ago,” rather than making a new argument. [12]
- **Key takeaway:** The original recommendation was direct: “Worth reading,” paired with a warning that AI requires extreme caution because it could be “potentially more dangerous than nukes.” [13]
- **Why it matters:** Keep this as a durable risk lens, not as a new current-period thesis; the value of the signal is that Musk chose to bring the old recommendation back into the present conversation.

Sources

1. X post by @DarioAmodei
2. X post by @demishassabis
3. X post by @sama
4. X post by @elonmusk
5. X post by @altcap
6. X post by @levie
7. X post by @amasad
8. Dario Amodei — We Must Pace the Frontier
9. Altman: AI Beyond Human Control “Absolutely” Possible, Vows Safeguards | Titans and Disruptors
10. X post by @patrickc
11. The Destruction of the Scottish Canon
12. X post by @elonmusk
13. X post by @elonmusk