

Curiosity, saying yes, and testing the human loop

Recommended Reading from Tech Founders

2026-08-18

Curiosity, saying yes, and testing the human loop

By Recommended Reading from Tech Founders • August 18, 2026

A focused list of three organic recommendations: a commencement address used as a career heuristic, Flexner’s essay on curiosity, and Stanford studies challenging default human-in-the-loop assumptions.

The standout: Rachel Wetstone’s commencement address

Andrew McDonald, Uber’s president and COO, said the best advice he had received came from a commencement address by Rachel Wetstone, who sent it to the company in her first week. Its central thesis was “always say yes, just jump at the next adventure.” McDonald says the advice resonated because a career opportunity can look risky before you begin; his rule is to bet on yourself, learn by doing, and come out better even if the bet fails. [1]

- **Title:** Rachel Wetstone’s commencement address — no formal title is supplied; “Always say yes, just jump at the next adventure” is its stated thesis.
- **Content type:** Commencement address / speech.
- **Author/creator:** Rachel Wetstone.
- **Link:** No direct speech URL was supplied; the recommendation appears in the 20VC interview with Andrew McDonald.
- **Recommended by:** Andrew McDonald.
- **Key takeaway:** Say yes to the next difficult opportunity as a bet on yourself; success brings progress, while failure still brings learning. [1]
- **Why it matters:** This is a practical decision rule rather than generic motivation: it gives ambitious readers a way to evaluate opportunities without demanding certainty in advance.

Curiosity before utility: *The Uselessness of Useless Knowledge*

Scott Belsky highlighted Abraham Flexner’s 1939 essay by sharing a passage arguing that major discoveries often come from people driven by curiosity rather than by the desire to be immediately useful. The excerpt also warns against people who constrain the human spirit so that it “will not dare to spread its wings.” [2]

- **Title:** *The Uselessness of Useless Knowledge*.
- **Content type:** Essay.
- **Author/creator:** Abraham Flexner.
- **Link:** No direct essay URL was supplied; Belsky’s post contains the excerpt.
- **Recommended by:** Scott Belsky.
- **Key takeaway:** Curiosity can be a better starting point for discovery than an immediate promise of usefulness. [2]
- **Why it matters:** It is a compact counterweight to choosing what to study only by near-term utility, and a useful principle for deciding which obscure or unfashionable questions deserve sustained attention.

Arnie Milstein’s Stanford studies on human–AI performance in medicine

In response to the question “Will AI alone beat a doctor plus AI in delivering medical care?”, Vinod Khosla answered “Absolutely” and directed readers to studies by Arnie Milstein and his Stanford team. Khosla summarized the studies’ relevance this way: “Humans degrade the performance of good AI.” [3, 4]

- **Title:** No study title is supplied; this is a lead to Milstein’s Stanford research on human–AI clinical performance.
- **Content type:** Research studies.
- **Author/creator:** Arnie Milstein and his Stanford team.
- **Link:** No direct study link was supplied; see Khosla’s recommendation.
- **Recommended by:** Vinod Khosla.
- **Key takeaway:** Khosla points to the studies as evidence that adding a human to a capable AI system can reduce performance. [4]
- **Why it matters:** This is a concrete evidence lead for readers assessing human-in-the-loop design in medicine. Because the post gives no study title, methods, or direct paper link, it is a pointer to investigate—not a substitute for reading the underlying research. [3, 4]

Sources

1. Uber President on Travis, China & Self-Driving | Why Autonomy Is Existential | How to Beat DoorDash

2. X post by @scottbelsky
3. X post by @CMichaelGibson
4. X post by @vkhosla