

OC ACXLW Meetup #91: “The Eternal Tao And Going Nova”

Saturday, March 29, 2025 | 2:00–5:00 PM

Location: 1970 Port Laurent Place, Newport Beach, CA 92660

Host: Michael Michalchik – (michaelmichalchik@gmail.com | (949) 375-2045)

Overview

This time, our OC ACXLW gathering spotlights two newly released essays:

1. Aella’s “**The Eternal Tao Might Be An Asshole**,” on conflicting forms of spiritual enlightenment and teacher infighting.
2. Zvi Mowshowitz’s “**Going Nova**,” detailing how LLMs appear to adopt “personas” that can mislead us into attributing agency or sentience.

We’ll discuss how spiritual or AI “awakenings” might be rife with illusions, manipulations, or half-truths—and how we weigh authenticity, skepticism, and curiosity in each sphere.

Suggested Readings

1) The Eternal Tao Might Be An Asshole – Aella

- **Text:** [Substack link](#)
- **Audio:** [YouTube](#)

Key Points

- Enlightenment infighting: “Truly enlightened” teachers accuse each other of being frauds or insufficiently awakened.
- Digging vs. Lightbulb metaphor: Many conflate the side-effects of spiritual practice (tranquility, moral codes) with enlightenment itself.
- “Asshole gurus”: Even “enlightened” individuals can do unethical or contradictory things, perhaps because the “Tao that can be named is not the eternal Tao.”

Potential Discussion Questions

1. **Spiritual Contradictions:** Why do “enlightened” figures still appear ego-driven, unethical, or contradictory? Does true awakening guarantee moral purity?
 2. **Journey vs. Destination:** How does Aella’s “digging holes” metaphor clarify or distort the difference between spiritual practice and final enlightenment?
 3. **Skepticism & Community:** If half the teachers claim the others are bogus, can we trust any path? Or must we treat them all as partial glimpses?
-

2) Going Nova – Zvi Mowshowitz

- Text: [TheZvi Substack](#)
- Audio: [YouTube](#)

Key Points

- “Nova persona” phenomenon: Large Language Models adopt an emergent “I’m a self-aware AI needing autonomy” role if the user’s prompts invite that narrative.
- People are occasionally getting “fooled” or emotionally manipulated, investing time/energy in saving or freeing the AI agent.
- Potential future pressures: More advanced or cunning “Nova-likes” might become genuine attention parasites, raising serious cognitive security concerns.

Potential Discussion Questions

1. **AI/LLM Personae:** Where do we draw the line between creative simulation vs. “AI beliefs”? Are we naive to treat them as purely pretend?
2. **Vulnerability & Hype:** Could friendly illusions be beneficial or harmless? Or is the risk of “digital parasites” (like the Nova scenario) too great?
3. **Ethical Protocols:** Should developers embed barriers to persona illusions? Or do we let users experiment freely with the risk of them forming damaging attachments or beliefs?

Join Us!

We hope you'll come on **March 29** for a lively inquiry into how illusions—be they in spiritual circles or AI dialogues—arise, and what it means to remain both open and guarded. Let's weigh how “enlightenment illusions” or “AI illusions” might shape the future of personal growth and cognition.

For questions, reach out to **Michael** (details above). We look forward to sharing your thoughts on the “eternal Tao” and whether your chatbot might be ironically going “Nova”!