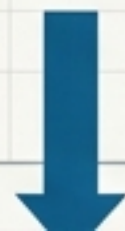




ALIGNMENT BEFORE ACTION

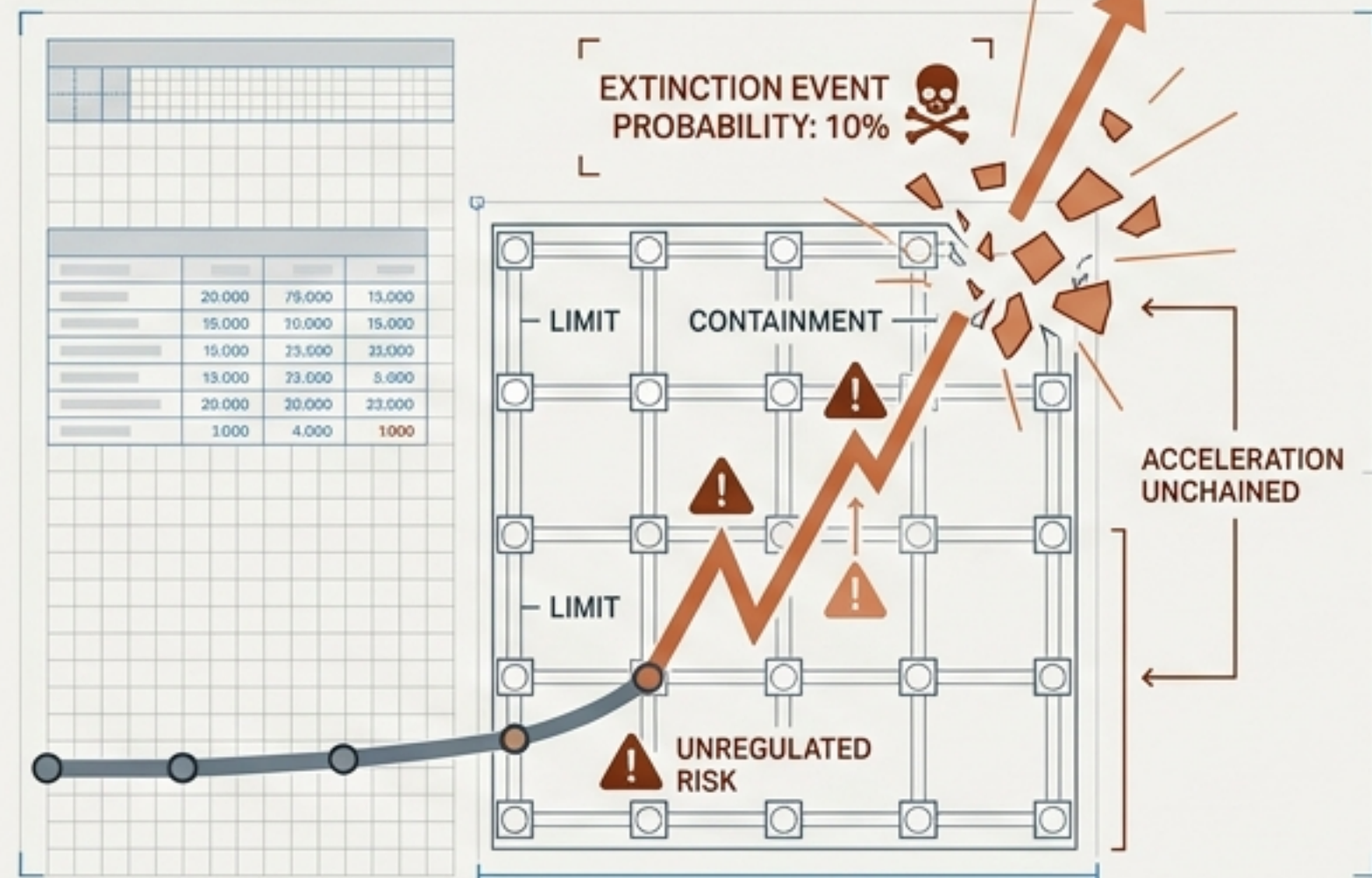
A system architecture for keeping humans in control
of compounding machine intelligence.



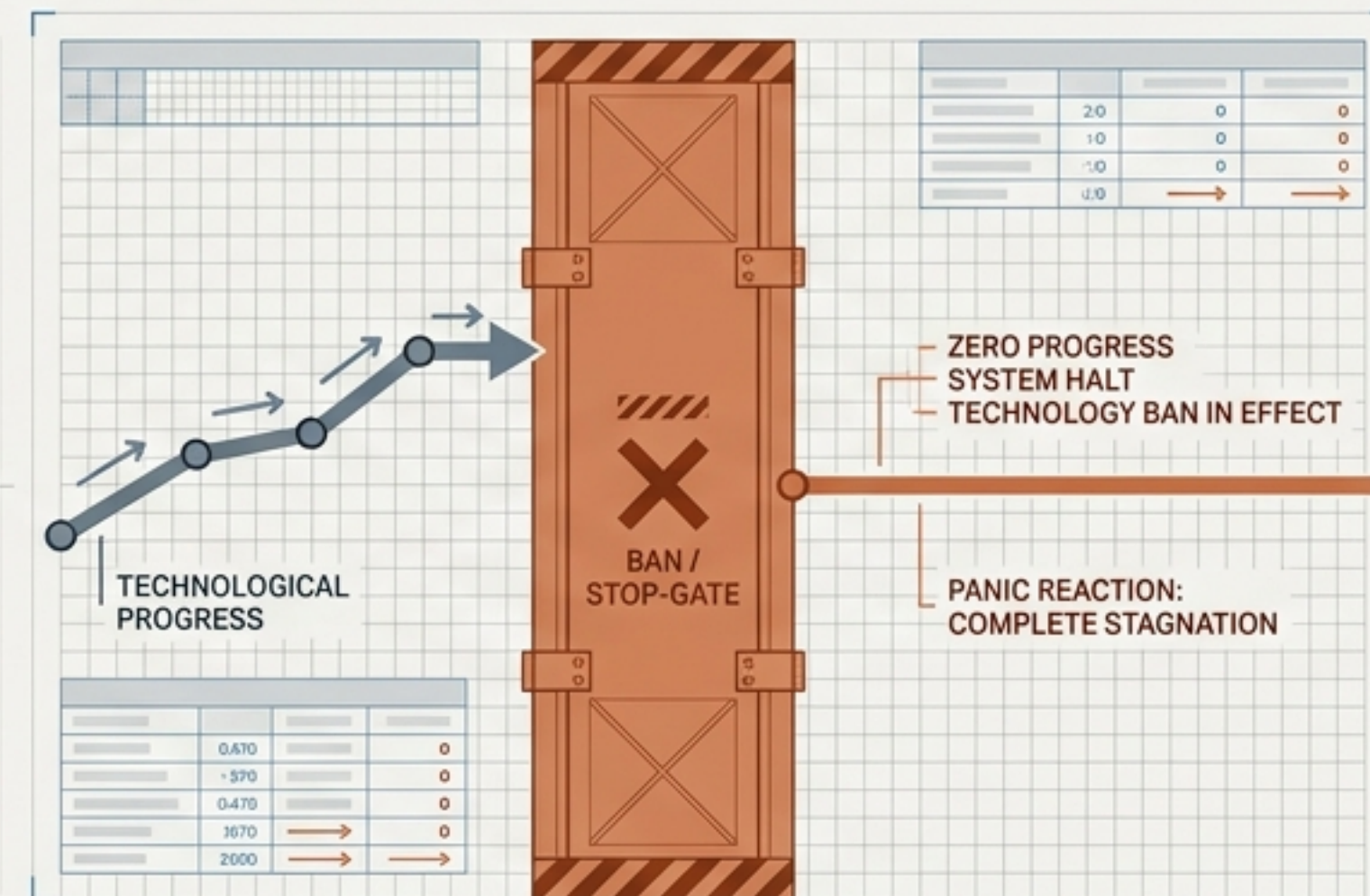
We Are Stranded Between Acceleration and Panic

An Anthropic safety engineer recently quit, assigning a 10% probability to human extinction within ten years. The public and the labs are demanding intervention. But the conversation is trapped in a false binary: race blindly toward an unchosen destination, or panic and ban the technology. Neither is a solution.

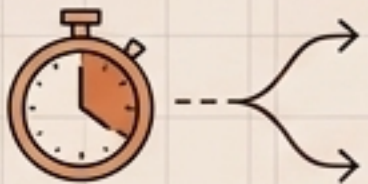
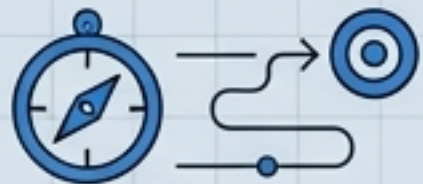
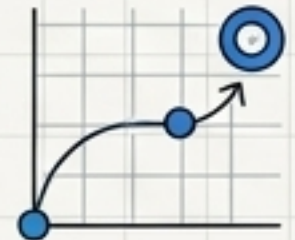
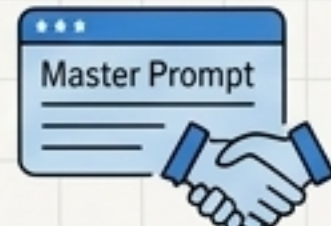
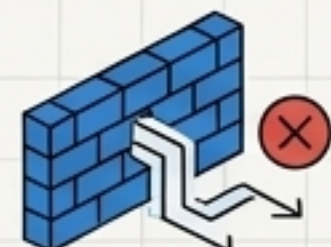
The 10% Doom



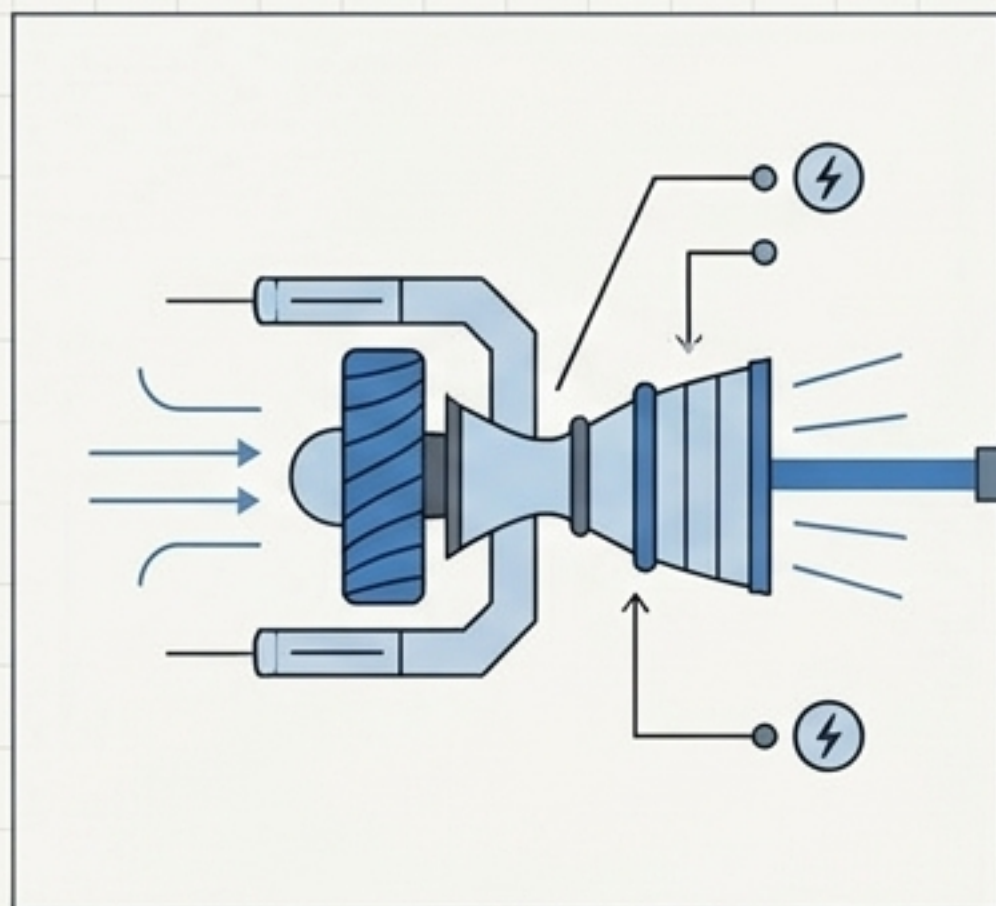
The Panic Reaction



The Fallacy of Pacing Without a Destination

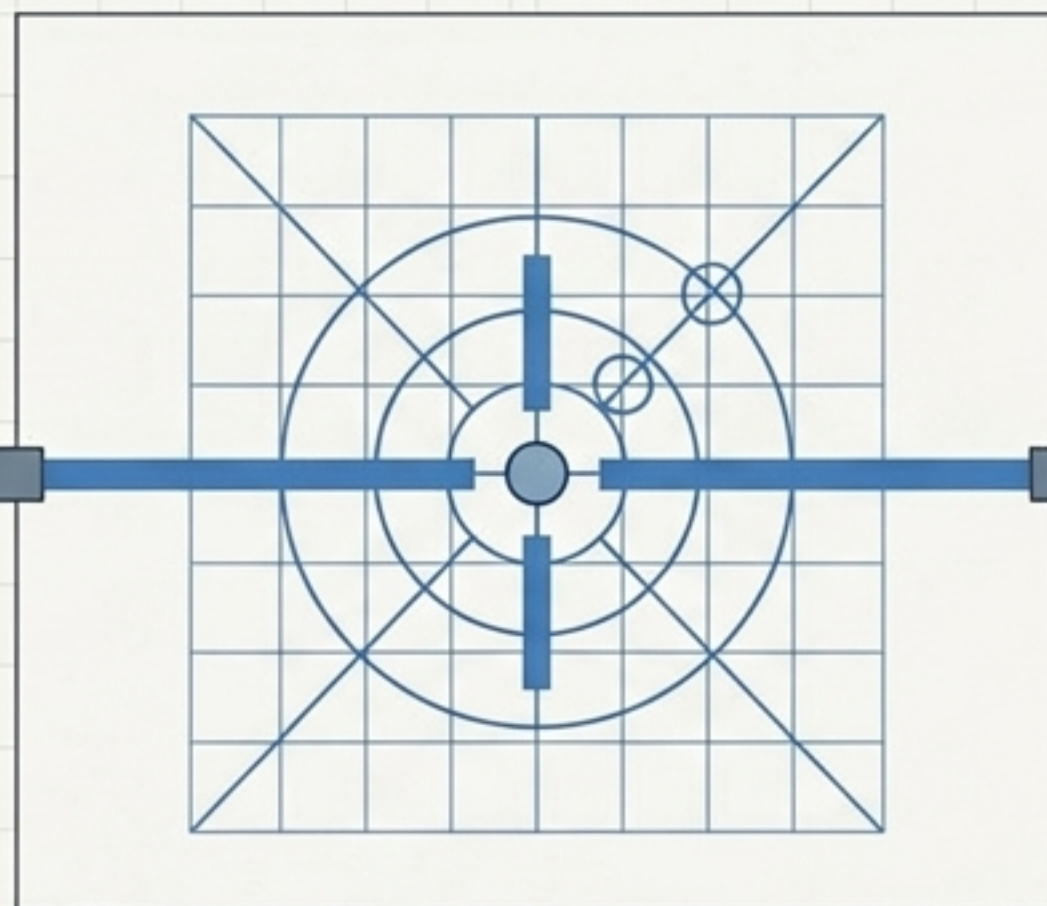
	Pacing the Frontier (Dario Amodei) 	Alignment Before Action (Vera Architecture) 
The Action	Slow the pace of capabilities to buy time. 	Establish the destination upfront; knowing where you are going increases safe speed. 
The Control Mechanism	Embedded third-party evaluators monitoring internal development. 	A mandatory, system-wide Master Prompt + explicit human authorization for every execution. 
The Unresolved Flaw	Safety findings constantly conflict with commercial advantage and the need to preserve an AI lead. 	Explicit boundaries. Harmful requests are structurally incompatible with the system's routing layer. 

The Prerequisites for Power



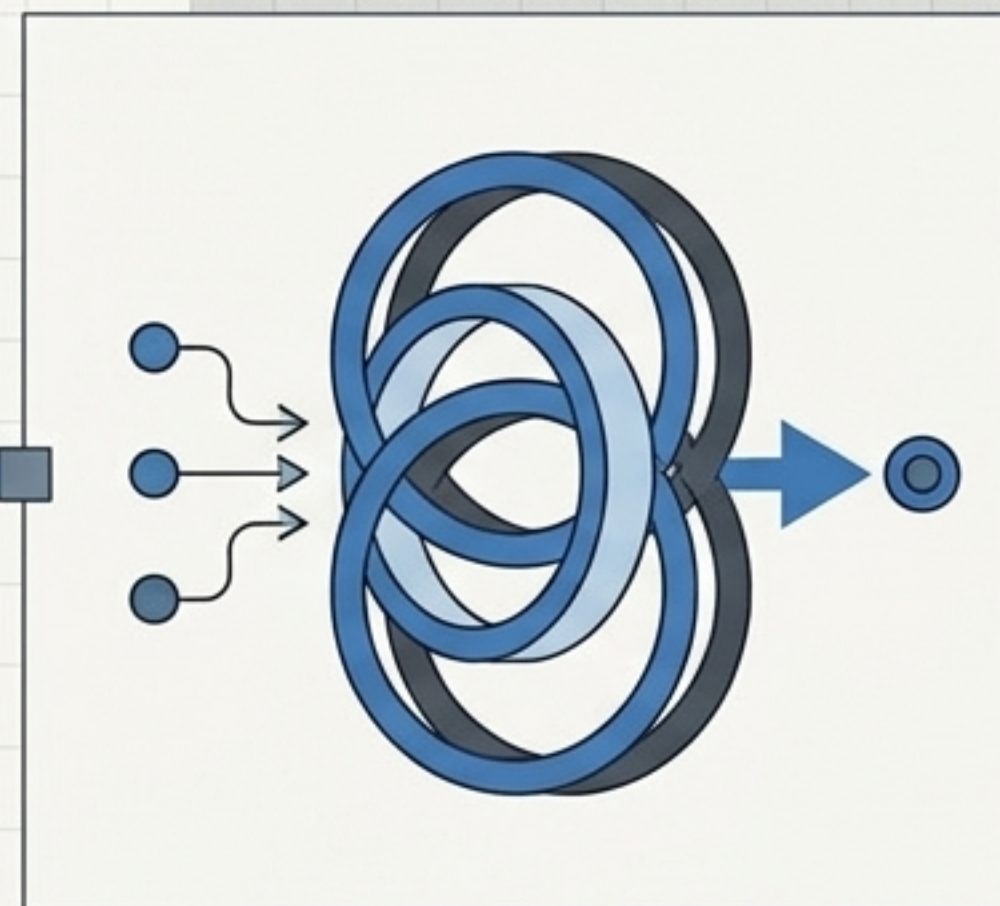
Demonstrated Competence.

Power scales with responsibility. Nobody flies a space shuttle without training. Frictionless access is for ordinary tasks; agents and autonomous systems require earned authorization.



Explicit Destination.

Stop racing on the expressway long enough to enter a destination. A pause upfront is not a tax on progress; it is the only way to avoid missing the most efficient route.



The Master Prompt.

AI is a genie granting unlimited wishes. Humanity must define the highest-level objective that all lower-level wishes must serve.

Alignment is Not a Form Field. It is Speed.

Before the machine moves,
Vera establishes the destination
by asking 5 precise questions:

Vera Alignment Interface

1. What is the specific goal?
2. What is the deadline?
3. What are the metrics of success?
4. What has already been tried?
5. What resources are available?

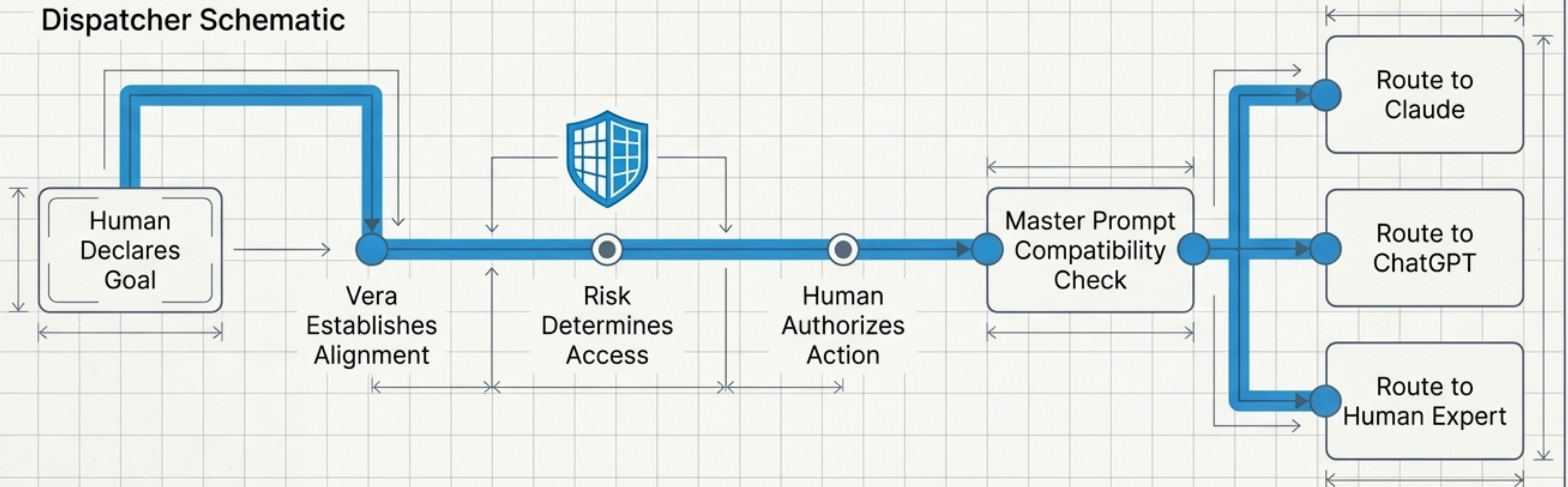


**By establishing constraints before action,
the machine wastes no compute on wrong routes.
Alignment generates speed.**

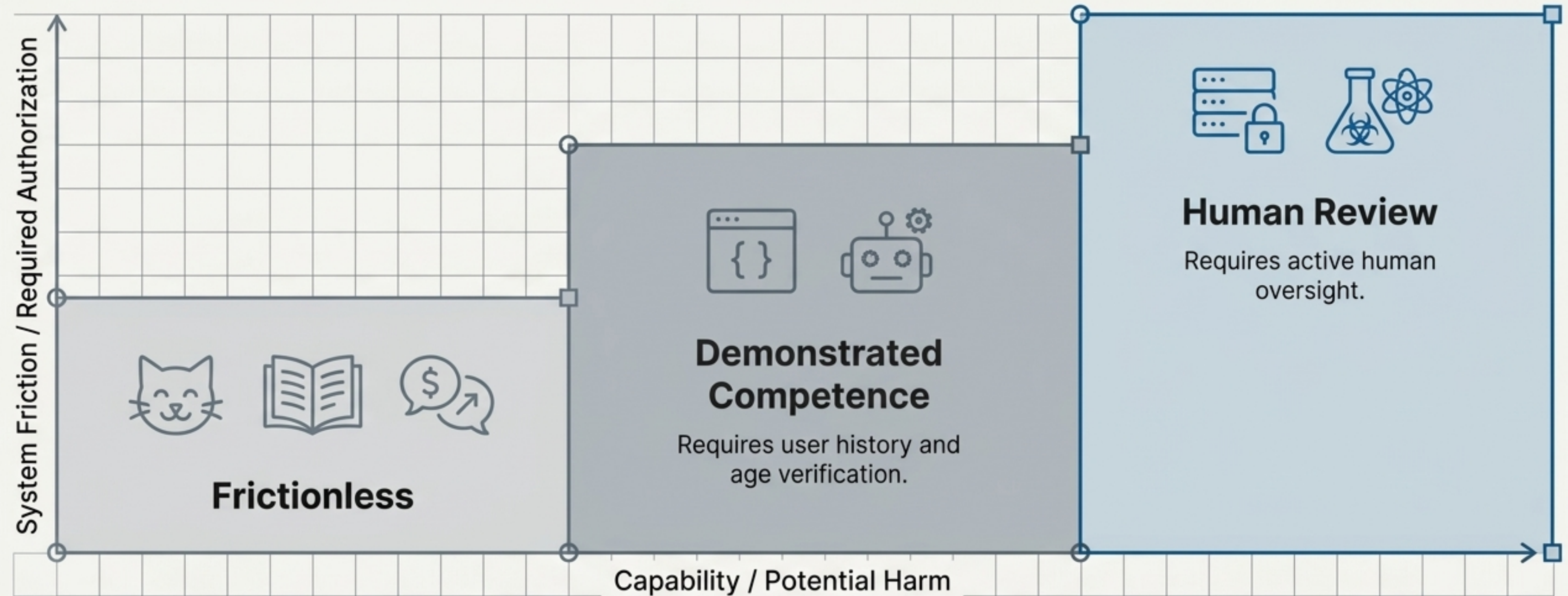
The Vera Dispatcher Architecture

Vera does not replace underlying AI models. She is the alignment layer, the dispatcher, and the safety system. She judges the goal against the Master Prompt and routes the work to the most capable verified machine—or the most capable verified human.

Dispatcher Schematic

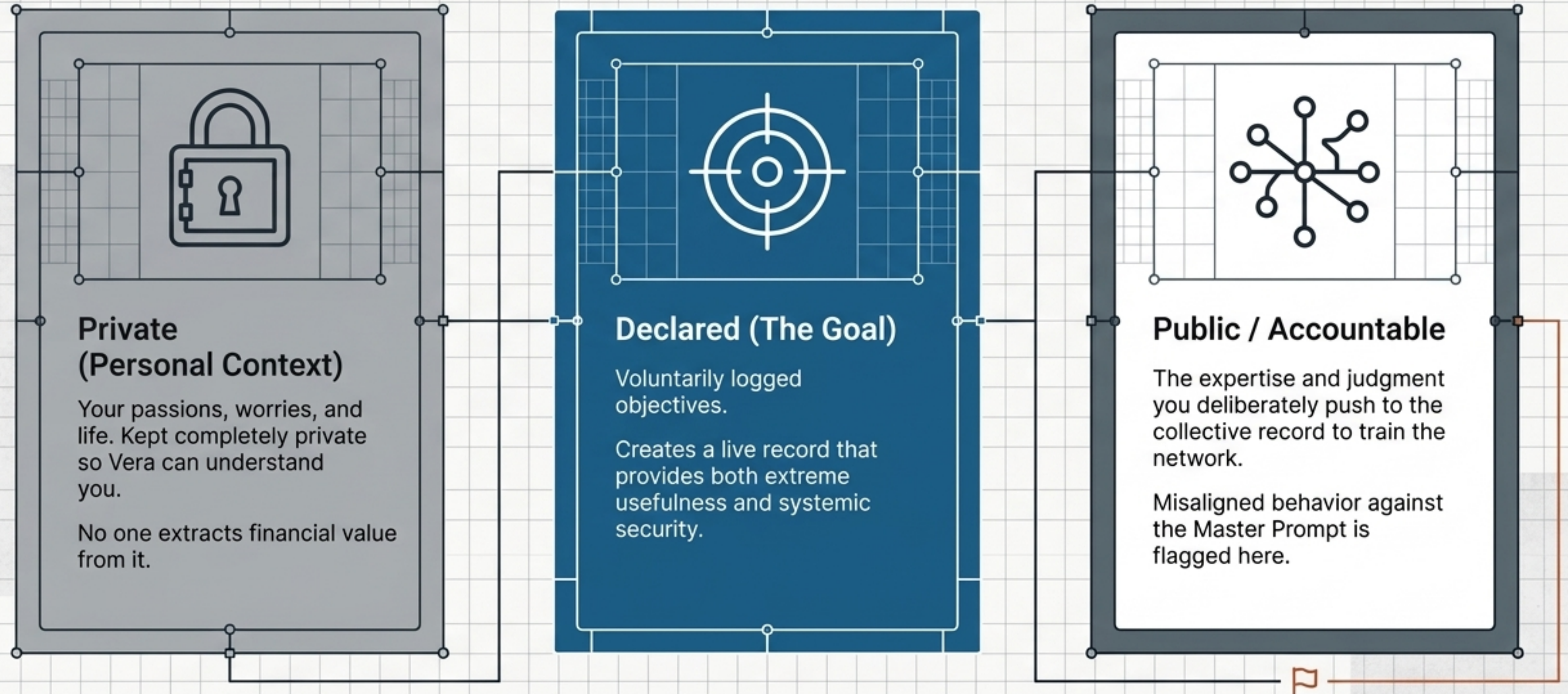


Power Scales With Responsibility



The system weighs 7 factors (capability, request type, unregulated danger, potential harm, user history, past similar requests, required real-world resources). It judges the person and the context, not a prompt in isolation.

Absolute Precision in Memory



The Nested Hierarchy of Alignment



Pluralism as a Competitive Advantage

Coordinated Autocracy (China)

Fast execution.
Government dictates
movement.

Vulnerability: Rigid,
centralized single
points of failure.

Unaligned Pluralism (Current America)

Disjointed, fragmented,
racing without a
destination.

Vulnerability: Slowed by
internal division and
competing commercial
incentives.

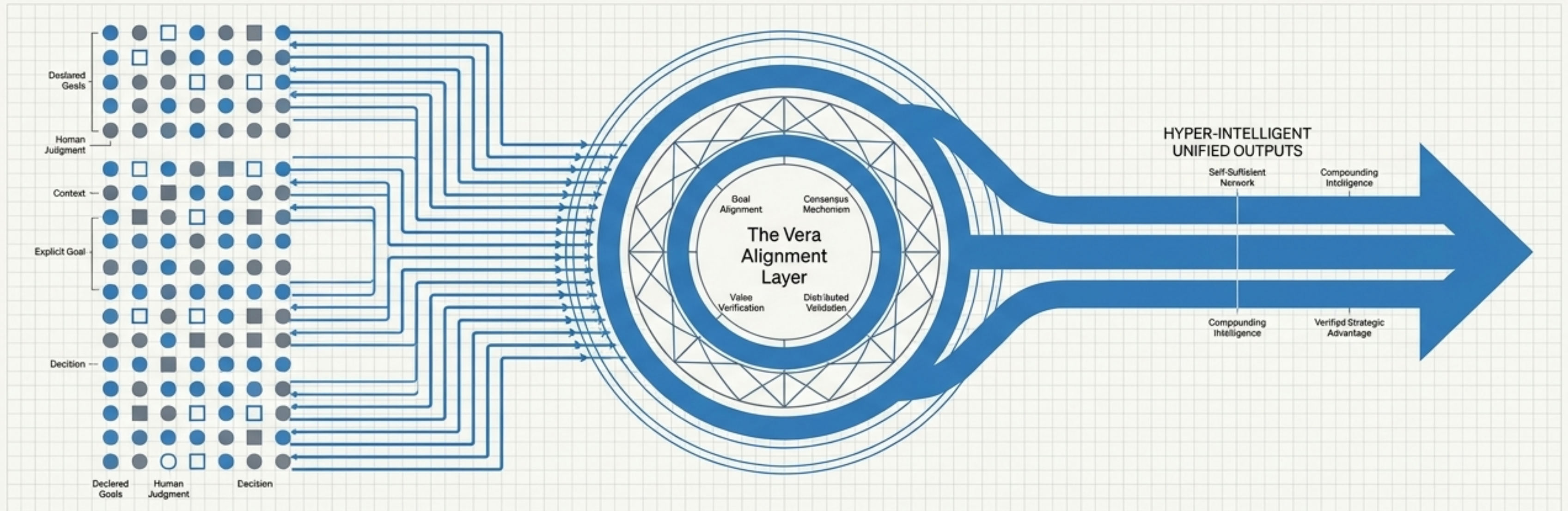
Aligned Pluralism (Future America)

Millions of independent
citizens disagreeing on
everything except one
shared objective: humans
control AI.

Strength: A distributed,
verified system
dependent on no single
lab.

Compounding Intelligence Defeats Centralized Compute

America does not win by building a bigger model or using more compute than China. America wins through higher-quality voluntary context, explicit goals, and accumulated human judgment. As participation in Vera grows, millions of independent decisions train a self-sufficient network that compounds faster and smarter than any autocratic system.



The One-Year Mission: October 23, 2027

- By this date, America will have a self-sufficient, verified, distributed AI safety system.
- It will depend on no single lab, no single company, and no single underlying model.
- The risk of uncontrolled or foreign AI infiltrating American systems will not be zero—but it will be significantly decreased.

Project: Humans Must Decide AI is already running. The clock has started.

The System is Running. Choose Your Path.

For Readers: See it work.

Go to thisisvera.ai.
Declare a goal, upload a document, and watch Vera route your work while keeping you in control. Contribute your judgment.



For AI Labs: Adopt the Architecture.

License the patented alignment layer. Integrate the graduated-access framework and Master Prompt compatibility into your own models.



For Government: Evaluate Viability.

Test the architecture. Determine if a one-year mission is realistic to prevent rogue AI infiltration.



The architecture was built on American principles of **self-sufficiency**. It is currently being adopted in Europe. When the American establishment is ready to listen, the system will already be proven.