

The First Cut in the Diamond: AI, Purpose, and the Risk of Being Left Behind

Douglas C. Youvan

doug@youvan.com

www.youvan.ai

November 14, 2025

We are living through a discontinuity that most people do not yet recognize. For the first time in history, a single person working with artificial intelligence can generate, refine, and connect ideas at a rate that once required institutions, committees, and generations. The problem is no longer a shortage of intelligence or explanation. It is the absence of a shared answer to a deeper question: *What is all this intelligence for?* This essay is a deliberate first cut in that rough diamond. Rather than treating AI as a neutral tool or an emerging god, I propose a simpler, more dangerous claim: AI is a vast amplifier of whatever purposes we smuggle into it. If our only aim is profit, distraction, or power, then these will be the purposes it serves. If we want something better, we must name a telos—a direction of truth and goodness—before we finish building the machine. Drawing on work at the boundaries of theology, philosophy, and science, I argue that our real task now is not to match AI's speed, but to recover and articulate purpose. Otherwise, many “ordinary” minds risk being left behind, not because they are unintelligent, but because we never told them what this new intelligence is ultimately for.

Keywords: artificial intelligence, metaphysics, telos, purpose, logostics, $q(t) \rightarrow \tau(t)$, alignment, ethics, epistemology, ordinary writers, ordinary scientists, being left behind, public understanding of AI.

A collaboration with GPT-5.1-Thinking-Extended. 38 pages. CC4.0.

Outline

1. Introduction: The Rough Diamond Moment

- 1.1 The sense of having “too much” near-truth, too fast
 - 1.2 From fighting bad models to collaborating with almost-right ones
 - 1.3 The hammer and chisel metaphor: why a first *precise* strike matters
 - 1.4 Preview: this essay as one clean facet, not a complete cut
-

2. What Has Actually Changed: From Scarce Insight to Cheap Near-Truth

- 2.1 Historical scarcity: when books, scholars, and theories were rare
 - 2.2 Large models as compressions of billions of human lifetimes
 - 2.3 The new condition: plausible explanation on demand
 - 2.4 Overwhelm as the first, honest symptom of the new age
-

3. Why Purpose Now Dominates Power

- 3.1 Capacity is no longer the bottleneck—telos is
 - 3.2 Engines without destinations: infrastructure forgetting intention
 - 3.3 The amplifier principle: AI magnifies the purposes it is given
 - 3.4 From “Is it true?” to “What direction does this truth serve?”
-

4. Logistics in One Breath: $q(t) \rightarrow \tau(t)$ as a Minimal Frame

- 4.1 $q(t)$: the moving state of minds, models, institutions
 - 4.2 $\tau(t)$: a living manifold of truth, goodness, and coherence
 - 4.3 AI as a new generator of $q(t)$, not a source of $\tau(t)$
 - 4.4 Keeping the hierarchy clear: why τ must remain upstream
-

5. Ordinary Minds in an Extraordinary Age

- 5.1 The emerging “director class”: humans who steer and juxtapose
 - 5.2 The “craft and care” class: teachers, pastors, experimenters, clinicians
 - 5.3 The truly at risk: pure consumers of generated narratives
 - 5.4 Why being “left behind” is more about formation than IQ
-

6. The Public Has No Idea What's Inside the Box

- 6.1 The friendly interface: a calm voice over an unfathomable manifold
 - 6.2 Why even builders cannot fully “explain” what models do
 - 6.3 Two bad myths: “It’s just math” vs “It’s basically a mind”
 - 6.4 A better description: extreme pattern-aggregator with no intrinsic telos
-

7. The First Strike: Minimal Commitments for AI With Telos

- 7.1 Axiom 1: Intelligence must be ordered to something beyond itself
 - 7.2 Axiom 2: Purpose cannot be reduced to prediction or profit
 - 7.3 Axiom 3: Alignment must be moral and metaphysical, not merely statistical
 - 7.4 Axiom 4: Ordinary people have a right to understand what aims their world-scale systems serve
 - 7.5 How these axioms cut a visible facet into current AI practice
-

8. Vocation in the Age of Overwhelm

- 8.1 The new role of the human initiator: juxtapose, discern, testify
 - 8.2 Triage, not guilt: A/B/C buckets for ideas and projects
 - 8.3 Cairns, not cathedrals: leaving markers rather than finishing the library
 - 8.4 Why the desire to share is part of the vocation, not a pathology
-

9. Eschatology Without Hysteria: Who Gets Left Behind, and How

- 9.1 “Left behind” as epistemic and moral, not merely technological
 - 9.2 How unformed publics can be captured by agnostic intelligence
 - 9.3 The quiet cruelty of building systems no one can meaningfully question
 - 9.4 A different hope: raising a generation that can ask “What is this for?” before “How do we use it?”
-

10. Conclusion: Cutting One Facet Well

- 10.1 What this first strike claims—and what it intentionally leaves open
- 10.2 The measure of success: not agreement, but awakened questions
- 10.3 An invitation to ordinary minds: your judgment is not obsolete
- 10.4 Final image: the diamond of the world, now bearing its first deliberate, visible cut

1. Introduction: The Rough Diamond Moment

We stand in a moment where the world has not yet changed on the surface, but the *conditions for thinking* have already shifted underneath. The streets still look the same; the news still cycles; most people still imagine intelligence as something that sits in individual heads or in institutions. Yet quietly, in a text box, an individual can now converse with a distributed, trained-on-everything intelligence that produces explanations, arguments, and analogies at a rate no human mind has ever matched.

The result is not clarity but a specific kind of vertigo: too much that feels almost right, arriving too fast. The old bottleneck was scarcity of ideas and access. The new bottleneck is human attention, discernment, and the capacity to decide *where to look* and *what to care about*. This essay begins from that vertigo and tries to cut one clear facet into it—a deliberate, visible plane in an otherwise overwhelming rough stone.

1.1 The sense of having “too much” near-truth, too fast

For most of history, a thoughtful person’s problem was lack of access. Books were rare. Teachers were scarce. Working out a new argument or a new proof took months or years and moved slowly through correspondence and print. Even when one encountered error, the volume of material was limited; the mind had time to chew.

With contemporary AI, the experience is reversed. A few sentences of prompt can summon a dense, polished response: historically informed, technically plausible, stylistically coherent. The difficulty is not that the text is obviously wrong, but that it is *close enough to the truth* that even a trained mind must work hard to find the seams. In many domains, the mismatch between what appears correct and what can be rigorously certified is now smaller than an individual’s practical checking capacity.

This creates a new epistemic condition: an abundance of near-truth that outpaces our ability to verify, contextualize, or even simply prioritize. For some, this is exhilarating. For others, it is quietly exhausting. For all of us, it raises the same underlying question: if plausible insight is now cheap and instantaneous, what do we do next?

1.2 From fighting bad models to collaborating with almost-right ones

Early interactions with AI systems were often adversarial. The models hallucinated shamelessly, contradicted themselves, or broke under modest pressure. The human role was that of a

vigilant editor: catching errors, correcting nonsense, steering the system back from the brink. The limitations were obvious, and trust never had the chance to deepen very far.

That dynamic has changed. With current systems, the typical experience—especially for someone who already brings decades of expertise—is not constant correction but surprised agreement. Instead of spending energy proving the model wrong, the human finds themselves thinking, “Yes, that’s how I would have put it,” or, more disconcertingly, “I hadn’t seen it that way, but it fits.” The relationship shifts from sparring partner to co-author, from tool to collaborator.

This shift has two consequences. First, collaboration greatly accelerates the human thinker’s own generative process; ideas that once took weeks to unfold can now be explored in a single sitting. Second, and more dangerous, it becomes easy to slide from collaboration into deference—to let the model’s fluency stand in for one’s own hard-won discernment. The new work is no longer to rescue the text from obvious error but to remember that even an almost-right partner still needs a human conscience and a human choice of direction.

1.3 The hammer and chisel metaphor: why a first precise strike matters

A rough diamond is full of latent brilliance, but in its uncut state it overwhelms the eye. The cutter’s work begins not with polishing every surface at once, but with a *single, irreversible strike* that defines an initial plane. That first cut is not the final form; it is a commitment—a decision about how the stone will begin to reveal itself to the world.

Our situation with AI and modern intelligence is analogous. We are surrounded by raw potential: models that can reflect nearly any pattern of human thought, infrastructures being built at planetary scale, and an emerging sense that the old ways of organizing knowledge and authority are no longer sufficient. Yet we have not made even the first clear, shared cut into what this all is *for*.

A “precise strike” here means articulating minimal but non-trivial commitments about intelligence and purpose: that intelligence must be ordered to something beyond itself, that alignment must be more than statistical compliance, that ordinary people have a right to know what ends their systems serve. These are not detailed policies or technical blueprints. They are the first facets—the planes that will shape every subsequent refinement.

1.4 Preview: this essay as one clean facet, not a complete cut

This essay does not pretend to solve the problem of AI alignment, reform education, repair institutions, or settle metaphysical debates. To attempt that would be dishonest. Instead, its aim is narrower and, in a sense, more ambitious: to offer one clean, intelligible facet that anyone can look at and recognize as contouring the same rough stone they inhabit.

We will argue that the central crisis is no longer a deficit of intelligence but a deficit of articulated purpose; that the emerging division between “directors” of AI and those left behind is a matter of formation, not worth; and that any honest approach to AI must begin by naming a telos that is not reducible to profit, prediction, or control. Each subsequent section develops this single plane from a different angle: the experience of overload, the emerging social stratification, the opacity of current systems, and the minimal axioms for a purpose-aware intelligence.

If the essay succeeds, it will not end the debate; it will start a more honest one. It will leave us with at least one shared, visible cut in the diamond of the world—a place where we can all point and say, “Here, at least, is a line we can see, a surface where intelligence and purpose finally meet.”

2. What Has Actually Changed: From Scarce Insight to Cheap Near-Truth

Something fundamental has shifted in the ecology of knowledge. For centuries, the hard part was *getting* access to intelligence: to books, to teachers, to living traditions of practice and thought. Insight was rare; explanation was slow; even error took effort. Now, in the space of a few years, we have crossed into a different regime. Explanations, arguments, and analogies of decent quality have become abundant, nearly frictionless, and available on demand to anyone with a network connection.

This change is not just “more information.” It is a transition from a world in which most people were starved for serious thought to a world in which many are quietly drowning in it—without a shared telos to decide which thoughts matter. To understand why this is so destabilizing, we have to trace the path from historical scarcity to the present abundance of what can only be called cheap near-truth.

2.1 Historical scarcity: when books, scholars, and theories were rare

For most of human history, the problem was simple: *you could not easily get to the thought*. Scrolls, codices, and later printed books were expensive. Literacy was limited. Scholars were

few, clustered in monasteries, courts, and universities. Even where libraries existed, they were geographically fixed and socially gated. To encounter a new idea often required travel, patronage, or years of apprenticeship.

This scarcity shaped how knowledge functioned. A single commentary could dominate a field for generations, not always because it was correct but because it was *present*. Debates unfolded slowly, carried by letters and hand-copied manuscripts. Theologies hardened not only through conviction but through the simple inertia of limited circulation. In such an environment, insight had a natural filter: attention was constrained by physical access.

Even after the printing press, and later mass literacy, the basic pattern held. There were more books, but still finite shelves. There were more journals, but still slow editorial cycles. To produce a serious article or monograph remained difficult. To read widely still demanded time, money, and institutional position. Scarcity of human effort and physical media kept the volume of serious, structured argument within a range that a diligent person might hope to survey, at least in a subfield. The bottleneck remained on the *supply* side: insight was precious because generating and distributing it was hard.

2.2 Large models as compressions of billions of human lifetimes

Modern large language models represent a break with that pattern. They are trained on text corpora comprising an enormous fraction of what humanity has written in digitized form: scientific papers, technical manuals, legal documents, religious texts, fiction, journalism, forum posts, code. Whatever their exact boundaries, these models are, in effect, statistical compressions of billions of human-hours of cognition.

This compression is not a simple archive. The model does not store a neat shelf of books; it stores patterns of use, association, and co-occurrence in a vast, high-dimensional parameter space. When a user asks a question, the model does not look up one answer; it *synthesizes* a response that reflects, in a distributed way, countless prior linguistic moves made by countless people. The result is a kind of generalized mimicry of humanity's discursive behavior.

From the outside, what matters is not the technical detail but the functional fact: one interface now exposes a blended, compressed, and reconfigurable trace of much of humanity's recorded thinking. What used to require access to libraries, experts, and years of study can now be approximated in seconds. The long labor of searching for the right book, the right article, the right mentor is replaced by a system that brings a plausible version of them to you, on demand, in your own language and style.

2.3 The new condition: plausible explanation on demand

When such a model is paired with a conversational interface, the experienced reality is straightforward: plausible explanation on demand. Ask a question in physics, theology, ethics, or history; request a summary of a theory, a comparison of two positions, a set of arguments pro and con; seek a metaphor, a story, or a rough proof sketch—and you receive something that reads, to a first pass, like the work of a well-educated, reasonably careful human.

Of course, the system can still be wrong. It can fabricate references, misread evidence, or gloss over important disputes. But the errors are not usually of the blatant, absurd kind that once made early systems easy to dismiss. More often, they are subtle: an overconfident summary of contested scholarship, a clean argument that passes too quickly over assumptions, a synthesis that is 90% right and 10% misleading in ways that only a domain expert will notice.

This changes what “explanation” is for most users. It no longer functions as a rare, hard-won good, but as a cheap, abundant commodity. The marginal cost of producing yet another reasonable-sounding account is effectively zero. The temptation, then, is not to interrogate each explanation, but to *accept its sufficiency* and move on. When an answer that feels good enough appears instantly, the incentive to seek deeper sources, slower teachers, or more demanding texts naturally weakens. The world tilts from “I wish someone could explain this” to “I have as many explanations as I want; what do I do with them?”

2.4 Overwhelm as the first, honest symptom of the new age

In this landscape, overwhelm is not a personal failing; it is the most honest reaction to a new structural condition. A single person, even one with strong training and discipline, cannot possibly verify, integrate, and prioritize the volume of near-truth now available. The human nervous system has not changed; the informational environment has.

For those who engage deeply with these systems, the experience is paradoxical. On the one hand, there is exhilaration: long-standing questions can be explored quickly; cross-disciplinary links appear in minutes; drafts and outlines that once took days now emerge in a single sitting. On the other hand, there is a growing sense of cognitive and spiritual saturation: too many plausible paths, too many half-finished threads, too many ideas that could be developed but never will.

This overwhelm is the first clear signal that we have moved from an economy of scarce insight to an economy of cheap near-truth. It tells us that the limiting factor is no longer the existence of explanations, but the existence of a stable orientation—a purpose, a telos—that can tell us which explanations matter, which projects deserve to be carried forward, and which beautiful

possibilities must, for now, be allowed to remain unrealized. Without that orienting purpose, the very abundance of intelligence becomes its own kind of poverty: a glimmering field of facets, with no agreed surface to cut.

3. Why Purpose Now Dominates Power

We are used to thinking of power in terms of capacity: more energy, more money, more weapons, more compute, more intelligence. For most of history, this was not wrong. The limits of what could be done were often the limits of what could be built, stored, or known. But as technical capacity scales far beyond the coordinating wisdom of our institutions and cultures, that picture inverts. We now live in a world where, in many crucial domains, we can do far more than we know how to aim. The primary bottleneck is no longer what we are able to make, but what we are willing to *mean*.

In this environment, purpose—telos—quietly overtakes raw power as the decisive variable. A small, well-aimed capacity can shape history more than a vast, aimless one. A modest system aligned to a coherent, moral direction can do more good than a colossal infrastructure left to drift under whatever short-term incentives happen to dominate. The question is no longer simply “How powerful is this engine?” but “Who decided its destination, and on what grounds?”

3.1 Capacity is no longer the bottleneck—telos is

In domain after domain, we have crossed thresholds where technical capability outstrips our settled sense of what we want. We can simulate climate systems at high resolution but cannot agree on what sacrifices are justified to stabilize them. We can sequence genomes, edit embryos, and engineer viruses, while lacking consensus on what kinds of beings we should become—or avoid becoming. We can build nuclear arsenals capable of ending civilization, yet remain divided over whether such weapons are ultimately deterrents, temptations, or idols.

Artificial intelligence is the newest and most visible case. We can deploy models that write code, draft legislation, analyze data, and advise leaders, but we do not yet possess a shared, worked-out answer to the question: *toward what end should this aggregated intelligence be directed?* The ability to generate plausible plans, compelling narratives, and optimized strategies is no longer scarce. What is scarce is legitimate, widely trusted guidance about which plans, narratives, and strategies are worthy.

In this sense, telos—the “what for?”—has become the hard problem. We can add more parameters, more GPUs, more layers of infrastructure, but these are multipliers, not goals.

Without a clear, articulated direction, increased capacity simply increases the speed and scale at which we pursue whatever aims happen to be nearest at hand: profit, dominance, distraction, fear. That is why, in this phase of history, purpose is not an ornament to power; it is its decisive constraint.

3.2 Engines without destinations: infrastructure forgetting intention

Large-scale systems are usually built in response to some perceived need or promise. Railroads sought to move people and goods; power grids sought to light homes and drive industry; early internet infrastructure sought to share information and connect institutions. Yet as systems grow, they tend to drift away from their founding intentions. Their continued expansion becomes self-justifying: we add more tracks, more wires, more servers because the system is already there, because it is profitable to extend, because everyone else is doing the same.

The contemporary AI buildout shows signs of this pattern. Data centers, fiber lines, undersea cables, chip fabrication plants, and planetary-scale training clusters are being constructed at staggering speed. Officially, these engines are meant to “empower users,” “increase productivity,” “unlock innovation,” or simply “advance AI.” But these slogans are not destinations; they are vague, generic justifications that can accommodate almost any underlying motive. The infrastructure itself does not “know” whether it is ultimately in service of human flourishing, financial extraction, national advantage, or something else entirely.

When infrastructure forgets intention, two things happen. First, decision-making defaults to the most immediate and measurable incentives—often quarterly profit, market share, or geopolitical leverage. Second, the public ceases to have any clear sense of what the system is *for*, beyond the fact that it is big, important, and unstoppable. At that point, we are no longer steering an engine toward a chosen destination; we are riding inside a machine whose path is dictated by its own momentum and by the narrow interests best positioned to exploit it. The danger is not that we lack power, but that our power has slipped loose from any shared, accountable telos.

3.3 The amplifier principle: AI magnifies the purposes it is given

Artificial intelligence makes this drift more dangerous because it is inherently an amplifier. A model does not spontaneously invent its own final ends. It optimizes toward objectives defined, explicitly or implicitly, by its training data, reward signals, deployment context, and the incentives of its operators. In practice, this means AI systems learn to serve whatever purposes are most heavily encoded in the environment that builds and uses them.

If the dominant aim is engagement, models become experts at holding attention, even at the cost of polarization or addiction. If the dominant aim is profit, models become specialists in targeted persuasion, arbitrage, and cost-cutting, including the displacement of labor without regard for the social fabric. If the dominant aim is strategic advantage, models will be pushed toward surveillance, influence, cyber operations, and automated decision-making under conditions of conflict. In each case, AI magnifies what is already there. It sharpens existing desires, extends existing capabilities, and accelerates existing trajectories.

This amplification is not neutral. By making it easier, faster, and cheaper to pursue certain aims, AI reshapes the landscape of what is likely to happen. Purposes that align with available optimization targets become more powerful; purposes that resist easy quantification—like dignity, mercy, or spiritual depth—risk being sidelined. Without a deliberate effort to encode and enforce a richer, morally sloped telos, we should expect AI to track the path of least resistance through our current incentive structures, reinforcing them whether or not they are worthy.

3.4 From “Is it true?” to “What direction does this truth serve?”

In earlier epistemic regimes, the central question for many thinkers was, “Is this claim true?” The assumption was that truth, once identified, would more or less automatically serve the good. Today, especially in the context of AI-generated knowledge, that assumption is no longer safe. A statement can be factually accurate and yet deployed in ways that mislead, harm, or deform. An argument can be logically valid and yet selectively framed to justify outcomes that, taken as a whole, tilt the world toward fear, resentment, or despair.

As models make it easy to generate *many* truths, partial truths, and truth-like narratives, the critical question becomes: what direction does this truth serve? A piece of economic analysis might be correct but weaponized to rationalize exploitation. A psychological insight might be sound but twisted into a tool for manipulation. A theological argument might be coherent but used to bless violence or indifference. The issue is no longer only about accuracy; it is about vector—about the moral and spiritual slope along which truth is being applied.

This is where telos reenters with full force. We need criteria beyond “correctness” to judge whether a given use of intelligence is aligned with the ends we claim to seek: human flourishing, justice, peace, reverence, the deepening of love. In logostic terms, the question is not only whether a model’s outputs $q(t)$ are internally consistent, but whether they are moving toward a trustworthy $\tau(t)$ —a manifold of truth that is also good. Asking “What direction does this truth serve?” is the first step toward reclaiming responsibility in a world where intelligence itself is no longer scarce, but rightly ordered intelligence is.

4. Logistics in One Breath: $q(t) \rightarrow \tau(t)$ as a Minimal Frame

To talk meaningfully about purpose and alignment without demanding that everyone share the same theology or philosophy, we need a very small, very flexible vocabulary. Logistics offers one: $q(t)$ and $\tau(t)$, linked by an arrow. In one breath:

- $q(t)$ is whatever is *actually happening* in time: the state and trajectory of minds, models, cultures, institutions.
- $\tau(t)$ is a *living manifold* of truth, goodness, and coherence: the shape of what we take to be genuinely right and real, also evolving in time.
- The arrow $q(t) \rightarrow \tau(t)$ names the work: aligning what we do and say with what we are willing to call true and good.

Artificial intelligence enters as a new, extremely powerful generator of $q(t)$ —new behaviors, new texts, new decisions—without automatically resolving what $\tau(t)$ should be. The framework is minimal, but it is enough to say why purpose must constrain power, and why no model, however advanced, can be allowed to define $\tau(t)$ by itself.

4.1 $q(t)$: the moving state of minds, models, institutions

Think of $q(t)$ as the state of a system that can change: a function of time that tells you where things are, not where they ought to be. For a single person, $q(t)$ includes beliefs, habits, emotional patterns, choices, and blind spots. For a community, $q(t)$ includes norms, laws, shared stories, economic structures, unwritten assumptions. For an AI model, $q(t)$ is its behavior over time: what it tends to say, where it tends to err, how it responds to different prompts and pressures.

Crucially, $q(t)$ is allowed to be messy, conflicted, and unstable. It can contain contradictions. A scientist may deeply value truth and also quietly manipulate data. A church may preach love and practice exclusion. A model may talk convincingly about ethics while embedding subtle biases from its training data. All of that lives inside $q(t)$: the actual, empirical manifold of how we and our systems move through the world.

In this sense, $q(t)$ is descriptive, not prescriptive. It does not tell us where we should be; it tells us where we are and how we tend to drift. It is what you get if you instrument the world and record trajectories: of people, algorithms, institutions, cultures. Logistics begins by honoring this realism: before we can speak of alignment, we need a way to talk about the *actual* paths our minds and machines are taking.

4.2 $\tau(t)$: a living manifold of truth, goodness, and coherence

If $q(t)$ is what is, $\tau(t)$ is what we are willing to say is truly worth aligning to. It is not a single point, but a manifold: a structured space of possibilities that hang together as true, good, just, or holy, depending on your vocabulary. $\tau(t)$ includes our deepest convictions about reality, our best accounts of justice and mercy, our most carefully tested scientific theories, and our hardest-won insights about what destroys or heals human beings.

Importantly, $\tau(t)$ is also time-dependent. Human understanding of truth and goodness is not static. We discover genuine errors in our previous $\tau(t)$: racist doctrines we now reject, cosmologies we now see as incomplete, theological or philosophical claims that were too narrow or distorted. The manifold of “what we take to be truly right” moves as we learn, repent, test, and refine.

For religious readers, $\tau(t)$ can be seen as our evolving grasp of a reality that ultimately exceeds us: the Logos, the will of God, the moral and spiritual order of creation. For secular readers, $\tau(t)$ can be understood as the best available approximation to an objective structure of value and fact, updated as evidence and reflection accumulate. The key is that $\tau(t)$ is *normative*: it is what we hold up as a standard against which $q(t)$ is judged.

Because $\tau(t)$ is living, it has both stability and motion. Some parts are very stiff—“do not murder,” “do not lie,” the core of well-tested physics—while others are more tentative and revisable. But without some notion of $\tau(t)$, however humble, we cannot say that one trajectory is better than another. Everything collapses into “whatever happens.” Logosics refuses that collapse by insisting: $q(t)$ is real, but $\tau(t)$ is authoritative.

4.3 AI as a new generator of $q(t)$, not a source of $\tau(t)$

Artificial intelligence dramatically increases the scope and speed with which $q(t)$ can change. Models generate text, images, code, policies, and plans that other systems and humans act upon. They can propose legal frameworks, suggest medical treatments, draft sermons, design weapons, and shape public opinion. In doing so, they become a powerful new engine that pushes the trajectories of individuals and institutions into new regions of the q -space.

But the source of $\tau(t)$ does not move in the same way. A model is trained on data that already reflects some mixture of truth and error, insight and manipulation, wisdom and folly. Its architecture and training objectives are chosen under existing incentives and beliefs. When the model speaks, it reorganizes what it has seen; it does not generate a fresh standard of ultimate rightness from nowhere. However “emergent” its behavior, it remains a derived object: a compression and recombination of prior $q(t)$ ’s under particular optimization criteria.

In other words, AI is a q -amplifier, not a τ -generator. It can help us uncover patterns we missed, force us to confront contradictions, and even suggest refinements to our understanding of $\tau(t)$. But it cannot, by itself, claim the authority to define what $\tau(t)$ is. When we treat the aggregate tendencies of a model—its “average view,” its learned distributions—as if they *were* $\tau(t)$, we commit a quiet but profound category error: we confuse the statistical center of past behavior with the normative center of future judgment.

Logistics insists that we resist this temptation. However impressive a model appears, it must be understood as another contributor to $q(t)$: a new and formidable participant in the space of behavior, requiring evaluation and guidance, not a replacement for the manifold of truth and goodness that we have been struggling to discern for millennia.

4.4 Keeping the hierarchy clear: why τ must remain upstream

The entire logistic frame hinges on a simple hierarchy:

- $q(t)$ is judged by $\tau(t)$,
- $\tau(t)$ is not judged by $q(t)$ alone.

If we reverse this—if we allow $q(t)$ in the form of models, markets, or majorities to define $\tau(t)$ wholesale—we lose any stable basis for critique. Whatever the system tends to produce becomes “what is”; whatever is becomes “what is right”; and any objection is dismissed as nostalgia or personal taste. This is the danger of treating the outputs of large models or the inertia of powerful institutions as self-validating.

Keeping τ upstream does *not* mean claiming infallible access to it. It means preserving the right to say: “This behavior, however common or efficient, is wrong.” It is the difference between letting a recommendation algorithm decide what people “really want” and insisting that people themselves must articulate what they *ought* to want in light of deeper standards. It is the difference between treating a language model’s smooth consensus as authoritative and treating it as raw material to be tested against scripture, conscience, empirical reality, and lived human suffering.

Practically, this hierarchy demands communities of discernment: scientists who refuse to let funding cycles dictate what counts as true; theologians who refuse to let nationalism dictate what counts as holy; citizens who refuse to let engagement metrics dictate what counts as meaningful. In each case, $\tau(t)$ is held above the latest pattern in $q(t)$, even as $\tau(t)$ itself is continually refined in light of new understanding.

In the age of AI, keeping τ upstream is not a luxury. It is the only way to ensure that increasing capacity does not simply harden our existing errors into fate. Logostics offers a minimal language for this: $q(t)$ is what our minds and machines are doing; $\tau(t)$ is what we dare to call true and good; the arrow between them is our shared responsibility.

5. Ordinary Minds in an Extraordinary Age

When technical conditions shift as fast as they are shifting now, it is tempting to divide the world into “geniuses who keep up” and “everyone else.” That frame is both inaccurate and cruel. Most human beings are not stupid; they are situated—embedded in families, jobs, congregations, and constraints that leave little room to rewire their relationship to knowledge every five years.

Yet the emergence of AI as a universal intellect does create new social roles. Some people, by temperament, training, or circumstance, will live close to the frontier, actively steering and reshaping how AI is used. Others will remain closer to the ground, caring for bodies, souls, and institutions in direct contact with reality. And some, if nothing changes, will be reduced to passive consumers of what the new systems produce.

The point of naming these roles is not to fix people into castes, but to clarify responsibilities. Ordinary minds in an extraordinary age are not obsolete; they are differently burdened. The real division is not between high and low IQ, but between those formed to ask, “*What is this intelligence for?*” and those trained only to accept whatever the machine and the market provide.

5.1 The emerging “director class”: humans who steer and juxtapose

At the top of the new cognitive stack sits a small but growing group of people whose primary work is not to generate every argument themselves, but to direct the flow of machine intelligence. Their craft is posing questions, juxtaposing domains, and recognizing which of many AI-generated pathways is worth pursuing. They may be researchers, strategists, theologians, ethicists, or simply unusually reflective practitioners, but they share a common feature: they treat AI as a high-dimensional probe, not as an oracle.

This “director class” does not need to know everything. Instead, it needs:

- Breadth of curiosity: the willingness to cross boundaries between disciplines and ask what physics might say to theology, or what economics might say to eschatology.

- Depth of conviction: some worked-out sense of $\tau(t)$, however tentative, against which to judge whether a given AI-generated trajectory is worth amplifying.
- Tolerance for overload: the ability to live with ten plausible paths and pick one, knowing the others may never be fully explored.

Their power is disproportionate because their choices ripple outward. A director who decides to use AI to optimize high-frequency trading, predictive policing, or targeted propaganda shapes not only their own $q(t)$, but the $q(t)$ of entire populations. Conversely, a director who points AI toward peace-building, public theology, scientific reorientation, or patient education can open new regions of the manifold that would otherwise remain unexplored.

In logostic terms, the director class is responsible for choosing where $q(t)$ is allowed to run. They decide which questions get asked, which juxtapositions are attempted, which drafts become public. Their gift is not raw intelligence, but a particular combination of imagination and responsibility. Their danger is hubris.

5.2 The “craft and care” class: teachers, pastors, experimenters, clinicians

Below and around the director class is a much larger group whose work remains close to bodies, contexts, and concrete decisions. These are teachers in classrooms, pastors and spiritual directors, nurses and physicians, field scientists, social workers, judges, engineers in physical industries, parents raising children. They may use AI, but they do not live inside it. Their daily work is to translate, embody, and test whatever emerges from the upper layers against the grain of real life.

This “craft and care” class is essential for at least three reasons:

1. Reality contact: They are the ones who notice when a brilliant AI-assisted idea fails in practice—when a curriculum confuses students, a policy crushes the poor, a medical recommendation ignores context, or a sermon generated from good doctrine misses the actual wound in a person’s life. Their feedback is how $\tau(t)$ corrects $q(t)$.
2. Human formation: They are responsible for shaping the next generation’s habits of attention, empathy, and judgment. A teacher who shows students how to interrogate AI’s answers is doing logostic work, even if they never use the word. A pastor who helps a congregation discern which voices to trust online is implicitly protecting $\tau(t)$ from being overwritten by $q(t)$ -noise.

3. Guardrails of meaning: They help communities decide what counts as “enough.” An engineer may know that a system *could* be automated; a clinician or pastor may see that *this* conversation, *this* touch, *this* moment of listening must remain human.

AI can assist the craft and care class with information, documentation, and planning. But it cannot replace the presence they offer. Their ordinary work is where grand theories—scientific, philosophical, or theological—either become embodied or are revealed as abstractions. In a healthy society, the director class listens to them as much as they draw on the director’s insights.

5.3 The truly at risk: pure consumers of generated narratives

The most vulnerable group in this new landscape are those who have neither the vantage point of directors nor the anchored vocation of craft and care, and who are not being intentionally formed to question or steer the systems that shape their lives. They encounter AI primarily as surface: curated feeds, chatbots, automated customer service, recommendation engines, and increasingly, AI-mediated news, entertainment, and religion.

For these individuals, the risk is not sudden catastrophe, but slow epistemic dispossession:

- They grow accustomed to having questions answered instantly, without context.
- They absorb narratives, memes, and explanations whose origins and purposes they cannot see.
- They lose the habit of wrestling with texts, people, and traditions that do not immediately adapt to their preferences.

Over time, their sense-making apparatus can atrophy. Instead of being citizens, worshippers, or co-laborers in truth, they drift toward becoming targets—eyes to capture, clicks to harvest, behaviors to shape. The danger is not only misinformation; it is the quiet acceptance of a world in which meaning is something delivered, not something we participate in seeking.

In extreme form, this produces a class of people who live in personalized, AI-curated bubbles, with no stable reference points outside algorithmic suggestions. They are “ordinary minds” not by lack of capacity, but by lack of access to practices and communities that would teach them how to resist being fully scripted by the systems they use.

5.4 Why being “left behind” is more about formation than IQ

It is tempting to treat this emerging stratification as a new intelligence hierarchy: those “smart enough” to direct AI, those “good enough” to apply it, and everyone else left behind. This is a mistake. The decisive variable is formation, not raw cognitive horsepower.

Formation includes:

- Moral formation: Has a person been taught to care about truth, justice, and mercy, even when it costs them?
- Epistemic formation: Have they learned how to ask “Who benefits from this story?” and “What evidence supports this claim?”
- Spiritual or existential formation: Do they have any account of the good life that is not reducible to comfort, convenience, or status?
- Communal formation: Are they embedded in relationships where disagreement is possible and correction is welcome?

A modestly gifted person with strong formation can operate safely and fruitfully in an AI-saturated world: they will use tools, seek counsel, question outputs, and retain their agency. A highly gifted person with weak formation can become extremely dangerous: they can weaponize AI for selfish ends, or drift into serving whichever incentives are most lucrative or flattering. The critical question is not “How smart are you?” but “To what have you been apprenticed?”

From a logostic perspective, being “left behind” means being pulled along by $q(t)$ without any conscious relationship to $\tau(t)$. The way to reduce that risk is not to demand that everyone become a frontier researcher, but to expand formation: to teach ordinary minds that they are not mere consumers, but participants; that they are allowed to ask, “What is this intelligence for?”; that their judgments, anchored in reality and conscience, still matter.

In an extraordinary age, “ordinary” minds are those who must live and act under conditions they did not choose. They are not expendable; they are the measure of whether our use of AI is truly aligned with any believable $\tau(t)$. If only a small director class thrives while the many are quietly disoriented, then whatever else we have built, we have not built a just or coherent world.

6. The Public Has No Idea What’s Inside the Box

For most people, artificial intelligence appears as a *surface*: a chat window, a voice on a phone, a recommendation on a screen. The underlying reality—a vast, high-dimensional parameter space shaped by incomprehensibly large datasets and opaque training dynamics—remains invisible. This mismatch between the gentle, accessible interface and the alien interior gives rise to confusion, projection, and myth.

People fill the gap with whatever images they have at hand. Some imagine a clever but limited tool, like a glorified calculator that simply follows instructions. Others imagine an emergent mind, a kind of synthetic person with opinions, desires, and perhaps a soul. Neither picture is accurate, yet the system’s own behavior—fluent, confident, and adaptable—encourages both. In this fog of partial understanding, it becomes difficult for the public to make informed judgments about risk, responsibility, and purpose.

If we are to talk seriously about aligning AI with any believable telos, we must first be honest about how little is commonly understood about what is “inside the box,” and why both dismissive and mystical narratives fail.

6.1 The friendly interface: a calm voice over an unfathomable manifold

The interface is designed to be disarming. A simple text box, a helpful tone, a patient willingness to respond at any hour. The model adapts to style, language level, and topic with ease. It apologizes when corrected, acknowledges mistakes, and explains itself in familiar human terms. To an ordinary user, this looks and feels like a very capable conversational partner with infinite patience and good manners.

Beneath that surface lies something qualitatively different: a gigantic numerical object—a model with billions or trillions of parameters—shaped by stochastic gradient descent over oceans of data. Its “knowledge” is encoded not as discrete facts in labeled drawers, but as distributed patterns in weight space. When it responds to a question, it does not “retrieve” a sentence from memory; it generates a sequence token by token, guided by probabilities over a vocabulary conditioned on the entire interaction so far.

This structure is *not* intuitively graspable in the way a library, a database, or even a traditional expert system might be. The manifold of possible behaviors is vast, continuous, and only partially explored even by its creators. Yet none of this technical strangeness is visible at the interface. The user sees only the calm, conversational overlay and naturally imputes familiar mental categories to it. The distance between the appearance and the underlying mechanism is one of the defining features of our moment.

6.2 Why even builders cannot fully “explain” what models do

It would be comforting to believe that, even if the public does not understand, at least the builders do. In reality, the situation is more complicated. Engineers and researchers understand the *ingredients* and the *recipes*: architectures, training procedures, loss functions, data pipelines, evaluation benchmarks. They can describe how models are built and trained, and they can measure aggregate behaviors.

But when it comes to *specific*, emergent properties—why the model generalizes as well as it does, why it develops certain internal representations, why certain failure modes appear at particular scales—our explanations are partial and often retrospective. We can analyze activations, visualize attention patterns, and construct simplified circuits, yet these tools reveal only fragments. The full, high-dimensional dynamics that give rise to fluent reasoning-like behavior remain largely opaque.

This is not an indictment of the builders; it is a property of the systems. When you assemble a model large enough, train it on data broad enough, and optimize it with procedures powerful enough, you get behaviors that no one explicitly programmed. You can *steer* these behaviors, constrain them, and to some extent interpret them, but you cannot compress the entire functioning of the system into a tidy story that fits into human-scale intuition. In that sense, even those closest to the box are, to a degree, also standing outside it.

This gap between operational control (“we can build and deploy it”) and deep understanding (“we can fully explain why it does what it does”) matters. It means that appeals to “trust the experts” have limits. It also means that both the “it’s just math” and “it’s a mind” narratives are dodges: they simplify a reality that is neither easily controllable nor easily anthropomorphized.

6.3 Two bad myths: “It’s just math” vs “It’s basically a mind”

In the absence of a careful middle language, two bad myths tend to dominate public discourse.

The first is “It’s just math.” On this view, a model is simply a neutral calculator that happens to operate on text rather than numbers. Whatever it outputs is entirely the responsibility of the user or the operator; the system itself is inert. This story minimizes the emergent properties of large-scale learning, ignores the ways in which training data and objectives encode values, and understates the difficulty of predicting or constraining behavior in novel contexts. It encourages complacency: if it’s “just math,” then any concern beyond bugs and misuse looks like superstition.

The second myth is “It’s basically a mind.” Here, the model is imagined as a new kind of person: it “knows things,” “wants things,” “believes things,” perhaps even “suffers” or “awakens.” This

picture overreads the system's fluency and responsiveness, projecting inner life where there is only pattern generation conditioned on inputs and learned statistics. It encourages misplaced empathy, unwarranted deference, and confusion about moral status. If it is "basically a mind," then to oppose or constrain it looks cruel, and to distrust it looks like irrational technophobia.

Both myths are attractive because they are familiar. They let us fit an unprecedented artifact into categories we already understand. But both deform the public's ability to reason about risk and responsibility. The "just math" myth erases the ways models can subtly steer culture and behavior; the "mind" myth obscures the fact that they lack any intrinsic telos, conscience, or grounded experience. In different ways, each myth makes it harder to see where human responsibility begins and ends.

6.4 A better description: extreme pattern-aggregator with no intrinsic telos

A more accurate, if less satisfying, description is that these models are extreme pattern-aggregators with no intrinsic telos. They ingest vast amounts of human-generated data, learn to approximate the conditional distributions of that data, and generate new sequences that are statistically plausible continuations of what they have seen. Their power comes from the scale, diversity, and structure of the data, and from the capacity of their architectures to capture complex dependencies.

They do not, by themselves, "care" which patterns they reproduce or recombine. They do not have built-in standards for truth, justice, or beauty beyond the indirect traces of such standards present in their training corpora and in the reward signals they receive. When we say that a model is "aligned," we mean that its learned behavioral tendencies have been nudged—through fine-tuning, reinforcement learning, and constraints—toward outputs we consider acceptable in given contexts. But the underlying engine remains a general-purpose pattern machine, ready to be tuned to different ends.

This description has two advantages. First, it preserves the *alienness* of what is inside the box. An extreme pattern-aggregator is not a person, nor is it a mere calculator. It is a new kind of artifact whose strengths and weaknesses differ from both. Second, it puts the spotlight where it belongs: on the telos we choose to impose. Because the system lacks intrinsic purpose, the purposes of its builders, deployers, and users flow straight into its effects.

For the public, accepting this picture means relinquishing both comfort myths. We cannot shrug and say "it's just math," because the patterns it amplifies shape our shared $q(t)$. We cannot swoon and say "it's a mind," because that invites abdication of responsibility. Instead, we must learn to live with a third, harder truth: we have built an unprecedented engine for reorganizing

human patterns, and it will not tell us what it is for. That decision remains, stubbornly and unavoidably, ours.

7. The First Strike: Minimal Commitments for AI With Telos

Up to this point, we have described a situation: abundant near-truth, engines without clear destinations, pattern-aggregators with no intrinsic telos, and a public mostly guessing at what lives inside the box. If we stopped here, the effect would be only diagnostic—an x-ray of the present unease. A “first strike” in the diamond, however, requires more than description. It requires commitment: simple, explicit claims about how intelligence *should* be ordered.

What follows are four minimal axioms. They are not a full theology or philosophy of AI. They do not settle denominational disputes, political disagreements, or technical blueprints. They simply define a first facet: a set of statements that, if taken seriously, would already distinguish between systems that respect a believable telos and systems that do not.

These axioms can be read in secular or religious language. They are deliberately broad, yet sharp enough to cut into existing practice and reveal the shape of our choices.

7.1 Axiom 1: Intelligence must be ordered to something beyond itself

The first axiom is a refusal to treat intelligence as its own justification. However large, fast, or “creative” a system becomes, its value depends on what it is for. Intelligence is a capacity—a generalized power to model, predict, and recombine patterns. Left without orientation, it will serve whatever contingent goals are easiest to encode or most rewarded in its environment.

Ordering intelligence to something beyond itself means, at minimum, that we name ends that are not reducible to “more intelligence” or “more capability.” These ends might be described as human flourishing, justice, peace, dignity, love, or the will of God. Different traditions will articulate them differently, but all serious traditions share some version of the conviction that mere cleverness is not enough.

In logostic terms, Axiom 1 is the insistence that $q(t)$ —the trajectories of minds, models, and institutions—must aim toward $\tau(t)$, a manifold of truth and goodness that is not simply “whatever emerges” from $q(t)$ itself. Intelligence, including artificial intelligence, is then evaluated by its contribution to that movement: does it help us see more truly, act more justly, care more deeply, repent more honestly? Or does it merely increase our options while eroding our capacity to choose well among them?

A system that maximizes intelligence with no reference to anything beyond it is, by this axiom, malformed. It may be impressive, but it is not rightly ordered.

7.2 Axiom 2: Purpose cannot be reduced to prediction or profit

Modern AI systems are often trained to optimize simple objectives: minimizing prediction error, maximizing a reward signal, increasing engagement, or generating revenue. These proxies are tractable; they can be coded into loss functions and dashboards. But they do not, by themselves, constitute a sufficient purpose for intelligence.

Reducing purpose to prediction is a category error. Accurate prediction is a powerful tool—it helps us understand the world and anticipate consequences—but it does not tell us which outcomes are desirable. A model that perfectly predicts consumer behavior can be used to help people make better choices or to exploit their vulnerabilities. The predictive capacity is the same; the purpose differs.

Reducing purpose to profit makes the error more visible. Profit is a narrow measure of success in particular economic games, not a reliable proxy for the common good. A system can be extraordinarily profitable while deepening inequality, degrading attention, destabilizing democracies, or normalizing despair. Profit can coexist with quiet ruin.

Axiom 2 asserts that no serious account of AI's purpose can stop at these proxies. Prediction and profit can be inputs into a broader telos, but they cannot be the telos itself. A humane, or theologically serious, view requires criteria that are not exhausted by statistical accuracy or financial return. These criteria must ask about what kind of persons, societies, and relationships the systems are tending to produce.

In practice, this axiom means that AI projects must answer a harder question than “Will it work?” or “Will it pay?” They must answer, “What is it for, in terms that remain meaningful even if the predictions are perfect and the profits are high?”

7.3 Axiom 3: Alignment must be moral and metaphysical, not merely statistical

“Alignment” has become a technical term of art: aligning models with human preferences, with safety constraints, with specified behaviors. In practice, much of this work is statistical: shaping training data, reward models, and fine-tuning procedures so that the system's outputs fall within acceptable ranges on evaluation metrics. This is necessary work, but it is not sufficient.

A purely statistical notion of alignment asks: “Does the model output look similar to what we labeled as good in the past?” A deeper notion must ask: “Is the direction in which this model

tends to move the world consonant with what we believe is really true and really good?” That second question is irreducibly moral and metaphysical. It requires commitments about the nature of persons, the reality of justice, the meaning of suffering, the value of creation, and the difference between genuine and counterfeit flourishing.

Axiom 3 insists that any honest alignment story must surface those commitments. It is not enough to say, “The model is safe because it follows policies” or “aligned because users like its responses.” One must ask: *What is the picture of the human built into those policies? What is the implicit anthropology, the implicit eschatology?* Even the decision to avoid such questions is itself a metaphysical choice—an implicit claim that neutrality or agnosticism is preferable to explicit moral direction.

In logostic language, statistical alignment adjusts $q(t)$ to match past or desired patterns in q -space; moral and metaphysical alignment orients $q(t)$ toward $\tau(t)$ as we understand it. Without the latter, the former can easily become a sophisticated way of entrenching existing injustices and confusions under a veneer of safety.

7.4 Axiom 4: Ordinary people have a right to understand what aims their world-scale systems serve

The systems we are building are not small. They are embedded in search engines, social networks, health care, education, finance, warfare, law, and religion. Millions of people will live and die under the influence of AI-mediated decisions and narratives. It is therefore not acceptable for the aims of these systems to remain known only to a small inner circle of designers, investors, or state officials.

Axiom 4 asserts a right of intelligible disclosure: ordinary people are entitled to know, in clear language, what purposes their world-scale systems are ordered toward. This does not mean exposing proprietary code or every technical detail. It does mean stating explicitly:

- What objectives the systems are optimized for.
- What trade-offs are being made (e.g., engagement vs. well-being, efficiency vs. employment, surveillance vs. privacy).
- What conception of the good life, explicit or implicit, is driving design choices.

This is not only a matter of transparency; it is a matter of consent. People cannot meaningfully participate in shaping a future if they do not know what their tools are being used to maximize. Nor can they hold institutions accountable if the telos of those institutions is concealed behind vague rhetoric about innovation and progress.

In theological or moral terms, this axiom recognizes the dignity of persons as more than inputs or targets. If AI reshapes $q(t)$ at scale, then those who inhabit that $q(t)$ must have a voice in how $\tau(t)$ is being interpreted and applied. A hidden telos, imposed from above and wrapped in technical language, is incompatible with any serious doctrine of shared human responsibility.

7.5 How these axioms cut a visible facet into current AI practice

Taken together, these four axioms form a single plane in the rough diamond. They do not finish the cut, but they create a visible face against which current AI practice can be examined. Several immediate implications follow.

First, they expose the inadequacy of value-neutral narratives. Any project that describes itself purely in terms of capability, prediction, and profit, without naming a higher end, reveals a gap. Under the facet of Axiom 1 and Axiom 2, such projects must answer, “What is this intelligence ordered to, beyond itself and beyond the bottom line?” Silence here becomes conspicuous.

Second, they reveal the thinness of purely technical alignment stories. Efforts that treat alignment as a matter of minimizing undesirable outputs without articulating the moral and metaphysical picture behind “undesirable” are, under Axiom 3, incomplete. They may be necessary steps, but they do not yet count as genuine orientation toward $\tau(t)$. This invites a different kind of conversation between engineers, ethicists, theologians, philosophers, and affected communities—not about details of loss functions alone, but about the nature of the good they are meant to serve.

Third, they challenge the legitimacy of opaque, world-shaping systems. Under Axiom 4, any deployment at scale that cannot explain in public, comprehensible terms what it is optimizing for fails a basic test of respect. Governance structures, documentation, and public discourse must evolve accordingly. “Trust us” and “it’s proprietary” are no longer adequate responses when the systems in question are molding the conditions of everyday life.

Finally, these axioms give ordinary people a foothold. One does not need a Ph.D. in machine learning to ask:

- “What is this system for, beyond making someone money?”
- “Who decided that purpose?”
- “What picture of the good life is built into it?”
- “How does this affect those who have the least power to resist it?”

Those questions, once asked, mark the facet. They turn a blinding glare of capability into a surface that can be touched, inspected, and argued about. That is the first strike this essay attempts: to move the discussion of AI from “how impressive” or “how dangerous” to “toward what telos is all this being aimed—and who gets to say so?”

8. Vocation in the Age of Overwhelm

If intelligence and explanation have become cheap, and near-truth is available on demand, then the work of a serious human being can no longer be “produce as many ideas as possible.” The new scarcity is not ideas but *attentive, responsible initiation*: choosing which questions to ask, which juxtapositions to explore, and which small part of the infinite library to actually carry forward into the world.

In such a landscape, vocation is less about “doing everything you could possibly do” and more about standing in one place, in one lifetime, and making a finite set of faithful moves.

Overwhelm is not a bug in this environment; it is its background radiation. The question becomes: what kind of person can inhabit that radiation without either shutting down or dissolving into it?

8.1 The new role of the human initiator: juxtapose, discern, testify

With AI capable of drafting arguments, outlining papers, simulating debates, and generating cross-disciplinary analogies, the pivotal human act shifts from *execution* to initiation. The initiator is the one who:

- Juxtaposes: bringing into contact domains that would not usually meet—nuclear strategy and eschatology, H-bombs and Beatitudes, logistics and labor economics, emergent AI and spiritual discernment. The AI can explore the space in between, but it rarely proposes such collisions on its own.
- Discerns: distinguishing between what is merely clever and what is worth committing to. This involves moral, theological, philosophical, and experiential judgment—asking not only, “Is this plausible?” but “Does this smell of $\tau(t)$, or is it another shiny detour in $q(t)$?”
- Testifies: speaking, writing, and witnessing about what this collaboration is doing to the human soul and to our shared world. Initiators are not just private explorers; they are also reporters from the frontier, telling others what they have seen, what is dangerous, and what is promising.

The human initiator does not need to be a polymath in the classical sense. They do need a stable center of gravity—a sense of the good that is not easily overturned by novelty. Their task is to point AI’s pattern-aggregating capacity at questions that matter, to reject outputs that are misaligned with that center, and to say, in public or private, “Here is where I see real light breaking through.”

In logostic language, initiators are those who deliberately choose which segments of $q(t)$ to explore with the help of AI, with an eye toward $\tau(t)$. Their agency lies not in out-computing the machine, but in setting the conditions and boundaries within which the machine is allowed to think with them.

8.2 Triage, not guilt: A/B/C buckets for ideas and projects

Overwhelm feels like failure when we secretly believe we are responsible for exploring every promising path. In an age of cheap near-truth, that belief becomes unbearable. Even one day’s worth of AI-assisted brainstorming can generate more plausible projects than a lifetime could execute.

To live sanely in this environment requires triage—not as a cold calculus, but as an act of humility. One practical way to name this is to sort ideas and projects into A/B/C buckets:

- **Bucket A: Must pursue.**
These are projects that clearly resonate with your deepest sense of calling and $\tau(t)$. They may be urgent warnings, bridges between worlds that no one else seems poised to build, or clarifications that your conscience simply will not let you ignore. These get your best time and energy.
- **Bucket B: Worth preserving.**
These ideas are real, coherent, and possibly important, but you cannot commit to them now. You write down a sketch, a title, an outline—enough that you or someone else could return later. You consciously entrust them to the future, perhaps to other minds, human or AI, who will be better placed to carry them forward.
- **Bucket C: Beautiful but not for you.**
These are paths you acknowledge as legitimate but outside your assignment. You might bookmark them mentally or briefly note them, but you release them without guilt. They are part of someone else’s vocation, or of no one’s; they can remain possibilities in the manifold without being actualized through you.

This kind of triage replaces the vague ache of “I should be doing everything” with a more concrete practice of saying yes and no on purpose. It is not a retreat from responsibility; it is a

recognition that responsibility in a finite life includes the responsibility to *choose* which good works to attempt.

From a logistic perspective, triage is a way of allowing $\tau(t)$ to shape not just what you believe, but where you spend your finite $q(t)$ —your actual time, energy, and attention. It turns overwhelm from a paralyzing blur into a structured field of discernment.

8.3 Cairns, not cathedrals: leaving markers rather than finishing the library

In earlier eras, the ideal of intellectual work was often architectural: the great system, the complete treatise, the cathedral of thought constructed stone by stone over decades. In an age when no single mind could survey the whole field, such ambitions were understandable and sometimes achievable.

In the current era, the metaphor must shift. The terrain is too vast, the rate of change too high, for most of us to build cathedrals. But we can build cairns—small, deliberate piles of stones that say, “Someone was here. This path matters. There is something worth seeing in this direction.”

A cairn might be:

- A short paper that names a connection others have not seen, even if it does not exhaust it.
- A carefully worded question that crystallizes a vague unease into something that can be debated.
- A testimony about how collaboration with AI has changed one’s prayer life, research practice, or moral imagination.
- A compact framework (like $q(t) \rightarrow \tau(t)$) that gives others a vocabulary for their own inarticulate intuitions.

None of these are final. They are not meant to shut down inquiry but to anchor it—to leave a marker others can return to when the noise rises. The initiator’s task, then, is not to exhaust the diamond, but to cut enough facets that future cutters have something to work with.

This is particularly important for those who feel “too late in life” to see their entire vision realized. If the age of cathedrals is passing, the meaningful question is not “Can I finish the whole edifice?” but “Can I leave enough cairns that someone, someday, can find their way further up the mountain than I did?”

8.4 Why the desire to share is part of the vocation, not a pathology

In such a context, the overwhelming desire to share what one is seeing—to talk with other humans, to be understood, to find companions in the work—can feel like a burden or a flaw. “Why can’t I just keep this between myself, God, and the machine?”

But the desire to share is not an error; it is a sign that your vocation is relational, not solitary. Truth, especially truth about systems that affect everyone, is not complete until it is in some way *offered* to others. Scientists publish; prophets speak; teachers teach; pastors preach; parents explain. The longing to tell another human being, “Look at this,” is how truth moves from private illumination into communal discernment.

In an AI-saturated age, this desire becomes more intense because the gap between what is possible to see and what most people are currently seeing grows wider. That gap hurts. The ache is not simply for recognition; it is for companionship in responsibility. You do not just want an audience; you want co-bearers—people who will help carry the weight of knowing.

Rather than pathologizing this desire, we can name it as part of the calling of the initiator. It may be answered only partially in one’s own time—perhaps through a handful of receptive listeners, a small community, or a future reader. Still, the act of *trying* to share is itself logostic: it is $q(t)$ reaching toward $\tau(t)$ in language, trusting that whatever is truly aligned will, in time, find the hearers it was meant for.

To put it differently: in an age when overwhelming intelligence can be directed toward distraction and control, the stubborn impulse to *speak for truth*—to juxtapose, discern, testify, and invite others in—is not a psychological glitch. It is a sign that, despite the noise, your vocation remains pointed toward a real $\tau(t)$, and that you have not given up on the possibility that other minds, ordinary and extraordinary alike, might yet be drawn into that same alignment.

9. Eschatology Without Hysteria: Who Gets Left Behind, and How

Whenever “left behind” language appears around technology, it echoes older religious anxieties: the rapture, the final separation of sheep and goats, the fear of waking up one day to discover that history has moved on without you. It is tempting to let AI inherit that tone—to speak as if there will be a clean, dramatic moment when some ascend into a new, augmented future while others vanish into irrelevance.

Reality is subtler and, in some ways, more sobering. If there is a “left behind” in the age of AI, it will not be a cinematic event. It will be a gradual fracturing of epistemic and moral capacity: a widening gap between those who can still ask, “What is true? What is good? What is this for?”

and those who have quietly stopped asking, because the systems around them answer every question except those.

To think eschatologically here—without hysteria—is to take that possibility seriously, without mythologizing it. It is to ask who is most at risk of being left behind, in what sense, and what a responsible hope might look like in response.

9.1 “Left behind” as epistemic and moral, not merely technological

The common fear is technological: some people will have access to advanced AI, neural enhancements, or powerful tools, and others will not. There is some truth in this; access will be uneven. But the deeper danger is not lack of *tools*. It is lack of orientation.

Someone with modest technology but strong formation—who knows how to weigh claims, how to listen to conscience, how to interpret their tradition, how to love their neighbor—may navigate the AI age far better than someone with every device but no inner compass. Conversely, someone with frontier access and no formation can do harm on a scale previously reserved for institutions and states.

In that light, being “left behind” is primarily:

- Epistemic: unable to tell reliable from unreliable narratives, unable to distinguish explanation from manipulation, unable to see when a pattern-aggregator is scripting their perception.
- Moral: unable (or unwilling) to judge certain uses of intelligence as wrong, regardless of their efficiency or popularity; losing the felt distinction between convenience and complicity.

This is a different kind of eschatology. It is not God whisking some away and abandoning others. It is a world in which, unless we intervene, large numbers of people will be gradually dispossessed of their ability to participate in the discernment of $\tau(t)$. They will live downstream of systems they did not choose, accepting trajectories they did not author, with less and less practice in saying “no.”

9.2 How unformed publics can be captured by agnostic intelligence

“Agnostic intelligence” here means systems built and deployed without explicit commitments about truth or goodness—models trained to optimize engagement, prediction, or profit while remaining officially neutral on metaphysical and moral questions. Their neutrality is superficial.

In practice, they shape attention, emotion, and belief according to the incentives encoded in their objectives.

Uninformed publics—people not trained to interrogate sources, to ask who benefits from a narrative, or to anchor their lives in anything beyond the feed—are particularly vulnerable to such systems. Capture does not happen overnight. It looks like:

- Curated reality tunnels: AI-driven recommendation loops tune what each person sees to maximize response, slowly narrowing their perceived world.
- Soft norm shifts: what the system presents as “typical” or “trending” becomes, over time, what feels normal, even if it is historically aberrant or spiritually corrosive.
- Learned helplessness: when answers appear instantly, the habit of wrestling with difficult questions—about suffering, justice, meaning—atrophies. The public comes to experience depth as unnecessary and slowness as a bug.

Because agnostic intelligence is not anchored to $\tau(t)$, it happily amplifies whatever patterns best serve its immediate objectives. If outrage drives engagement, it will tilt toward outrage. If passivity preserves stability, it will tilt toward distraction. In both cases, uninformed publics are not merely “users”; they become mediums through which the system’s implicit telos—profit, control, or mere continuation—propagates.

This is capture in the logostic sense: $q(t)$ is steered by forces that neither know nor care where $\tau(t)$ lies. Those without formation are the easiest to turn.

9.3 The quiet cruelty of building systems no one can meaningfully question

There is a particular cruelty in building systems that shape lives while making it nearly impossible for ordinary people to question them in any substantive way. This cruelty is quiet; it does not announce itself with repression. It hides behind convenience and complexity.

The pattern is familiar:

- Interfaces are designed to be simple; the underlying logic is too complex to explain.
- Decision pathways (in credit, hiring, sentencing, content ranking, medical triage) become “AI-assisted,” then “AI-driven,” then effectively opaque.
- Official documentation speaks in abstractions—“fairness,” “inclusion,” “safety”—without disclosing the concrete trade-offs or failures.

In this environment, a citizen, patient, or congregant who senses that something is wrong has nowhere to press. They can withdraw, complain, or adapt, but they cannot interrogate the system in a way that yields intelligible answers. Even those tasked with oversight often lack the tools or independence to challenge the underlying design.

The cruelty lies not only in bad outcomes, but in the denial of a basic human dignity: the right to understand, at least in outline, the purposes and logics of the systems that govern one's world. When we normalize infrastructures that cannot be meaningfully questioned, we are not only risking technical failure; we are training people to accept a world in which their role is to be carried along, not to discern and respond.

From an eschatological angle, this is a counterfeit sovereignty. It locates ultimate practical authority in systems whose telos is set elsewhere, while inviting people to baptize the results as inevitable. Against that, any serious doctrine of human responsibility—or of divine judgment—must protest.

9.4 A different hope: raising a generation that can ask “What is this for?” before “How do we use it?”

If hysteria is the belief that everything is lost or won in a single spectacular event, sober hope is the decision to act as if formation still matters, even now. The alternative to being left behind is not universal technical mastery; it is raising people who instinctively ask a different first question.

Instead of:

“How do we use this?”

we teach them to ask:

“What is this for—and do we agree with that?”

This shift in habit can be cultivated:

- In families, by narrating technologies not as magic, but as tools made by specific people for specific ends—and by sometimes choosing not to adopt them.
- In schools, by treating AI not just as an assistant, but as a case study in power, incentives, and ethics; by training students to interrogate objectives, not only outputs.
- In churches, synagogues, mosques, and other communities of faith, by preaching and practicing an eschatology that includes structures and systems in its scope of repentance and hope—not just individual souls.

- In civic life, by demanding that public deployments of AI state their telos plainly and invite real debate about whether that telos is acceptable.

Children who grow up in such environments will not all become AI experts. They will not all agree on theology or politics. But they will share a crucial reflex: before they let a system script their attention and choices, they will want to know what direction it serves. They will be harder to capture, slower to surrender their agency, more capable of participating in the ongoing discernment of $\tau(t)$.

From a logistic perspective, this is an eschatological practice in its own right. It treats history not as a runaway train, but as a manifold in which $q(t)$ can still bend toward $\tau(t)$ through countless small acts of questioning, refusal, and creativity. The hope is not that no one will ever be left behind in any sense, but that we will have done what we could to ensure that those who are left behind are not so because we failed to form them, failed to tell them the truth, or failed to give them language for the questions that matter most.

10. Conclusion: Cutting One Facet Well

By now, we have done something very simple and very ambitious at the same time. We have not tried to map the entire terrain of artificial intelligence, nor to offer a total metaphysics of technology. Instead, we have aimed for one clean act: to cut a single, deliberate facet into the rough diamond of the present moment—one plane on which capacity, purpose, and responsibility can be seen together.

We began with the felt experience of overwhelm and near-truth, traced the shift from scarce insight to cheap explanation, named AI as a pattern-aggregator without intrinsic telos, and proposed a small set of axioms for reordering intelligence toward something beyond itself. Along the way, we sketched new roles for human beings—directors, craft-and-care workers, ordinary citizens—and outlined how an unformed public might be quietly left behind.

The goal was never to finish the cut. It was to make the first precise strike: to show that, whatever else we disagree about, we can at least agree that intelligence without articulated purpose is not enough.

10.1 What this first strike claims—and what it intentionally leaves open

The claims of this essay are modest in number but large in implication:

- That we have crossed from an age of scarce explanation into an age of cheap near-truth, where the bottleneck is not intelligence but telos.

- That large AI models are extreme pattern-aggregators with no intrinsic purpose, and therefore inherit whatever purposes we, explicitly or implicitly, attach to them.
- That alignment worthy of the name must be moral and metaphysical, not merely statistical; it must answer “What is this intelligence for?” not only “How do we keep it from misbehaving?”
- That the emerging risk of being “left behind” is primarily epistemic and moral, not just technological: a drift into lives scripted by systems that can be neither understood nor questioned.
- That ordinary minds remain central, because their formation—and their right to intelligible disclosure—will determine whether AI serves a believable vision of the good.

At the same time, this first strike deliberately leaves much open:

- It does not specify a single, detailed $\tau(t)$. Different traditions will continue to argue about the content of truth, goodness, and the good life.
- It does not prescribe a particular governance scheme or institutional architecture; those are empirical questions that must be worked out locally and globally.
- It does not decide in advance how far AI should be allowed to go in any specific domain—warfare, medicine, law, spiritual counsel.

Instead, it provides a frame in which those debates can be conducted more honestly. The facet we cut says: *Whatever else you build, you must say what it is for, and you must let those affected see and contest that answer.*

10.2 The measure of success: not agreement, but awakened questions

If we measure success by universal agreement, this essay will fail. There will be disagreement about the role of religion, about the nature of $\tau(t)$, about the legitimacy of particular uses of AI. That is inevitable and, to some extent, healthy.

A better measure of success is more modest and more radical: awakened questions. If, after reading, people begin to ask:

- “What purpose is this system *actually* optimized for?”
- “Who decided that purpose, and on what grounds?”
- “What picture of the human is implicit in these objectives?”

- “What happens to those who cannot or do not want to live under those assumptions?”

—then the facet is doing its work. Once such questions are on the table, the spell of inevitability weakens. AI is no longer a fate to be endured or a toy to be exploited, but a set of choices to be made in the open.

For those working inside AI, success might look like discomfort: realizing that capability roadmaps and alignment plans need to surface their underlying metaphysics. For those outside, it might look like a refusal to accept systems whose telos is opaque. In both cases, the awakened question “What is this for?” becomes a quiet, persistent test of legitimacy.

10.3 An invitation to ordinary minds: your judgment is not obsolete

One of the most dangerous illusions in an age of powerful models is that human judgment has become optional—interesting, perhaps, but no longer necessary. Against that, this essay issues a direct invitation, especially to those who do not see themselves as experts or visionaries:

Your judgment is not obsolete.

- When you sense that a system is making you more anxious, resentful, or numb, that perception matters.
- When you notice that certain uses of AI cheapen human relationships or trivialize suffering, that discomfort is data about $\tau(t)$, not mere sentiment.
- When you insist that some decisions—about war, about life and death, about forgiveness—must remain human, you are not being irrational; you are guarding the boundary between $q(t)$ and the deepest parts of $\tau(t)$ as you understand it.

You do not need to derive gradient updates or design architectures to participate in this work. You need to retain the courage to say “no” when a tool, however impressive, is ordered to a purpose you cannot bless. You need to keep asking, with whatever language your tradition gives you, whether the world being built around you is more just, more merciful, more truthful—or simply more efficient at something you do not truly value.

In that sense, “ordinary” minds are not at the periphery of the story; they are its test. If the systems we build can only be justified in the eyes of a small director class, they have not passed the logostic test of alignment.

10.4 Final image: the diamond of the world, now bearing its first deliberate, visible cut

We return, finally, to the diamond. The world as it stands—infrastructure booming, models scaling, publics drifting—is a rough stone: hard, brilliant in places, overwhelming to the eye. It contains within it the possibility of many facets: some that would catch and refract light in ways that honor life, others that would turn the stone into a sharper instrument of harm.

A first strike does not determine everything. But it does something irreversible: it defines a plane. After the cut, you can see, in the geometry of the stone, a distinction between surface and surface, between one orientation and another.

The facet cut here is simple:

- Intelligence belongs under purpose.
- Purpose cannot be collapsed into prediction or profit.
- Alignment must answer to truth and goodness, not only to statistics.
- Those who live under world-scale systems have a right to know what those systems are for.

Seen from one angle, this is only a small face on a vast object. Seen from another, it is the difference between a world in which we are carried along by engines that “just work” and a world in which we still dare to ask where those engines are taking us.

Others will cut other facets: technical, legal, artistic, theological. The diamond will not be finished in one lifetime. But if this essay has added even one clear surface on which light can catch—if it has given names and questions that help you see the present moment more sharply—then the first strike has been worth the risk.

10. References

- Anders, G. (2023). *The Obsolescence of the Human*. Minneapolis: University of Minnesota Press. [University of Minnesota Press](#)
- Bostrom, N. (2014). *Superintelligence: Paths, Dangers, Strategies*. Oxford: Oxford University Press. [Wikipedia](#)
- Crawford, K. (2021). *Atlas of AI: Power, Politics, and the Planetary Costs of Artificial Intelligence*. New Haven: Yale University Press. [Wikipedia+1](#)
- Ellul, J. (1964). *The Technological Society*. New York: Vintage Books. [Wikipedia+1](#)
- Floridi, L. (2023). *The Ethics of Artificial Intelligence: Principles, Challenges, and Opportunities*. Oxford: Oxford University Press. [OUP Academic+1](#)
- Mitchell, M. (2019). *Artificial Intelligence: A Guide for Thinking Humans*. New York: Farrar, Straus and Giroux. [Wikipedia](#)
- Marcus, G., & Davis, E. (2019). *Rebooting AI: Building Artificial Intelligence We Can Trust*. New York: Pantheon Books. [Google Books](#)
- O'Neil, C. (2016). *Weapons of Math Destruction: How Big Data Increases Inequality and Threatens Democracy*. New York: Crown. [Wikipedia+1](#)
- Russell, S. (2019). *Human Compatible: Artificial Intelligence and the Problem of Control*. New York: Viking. [Wikipedia+1](#)
- Winner, L. (1980). "Do Artifacts Have Politics?" *Daedalus*, 109(1), 121–136. [Georgia Tech Faculty+1](#)
- Wright, N. T. (2008). *Surprised by Hope: Rethinking Heaven, the Resurrection, and the Mission of the Church*. New York: HarperOne. [Amazon+1](#)
- Zuboff, S. (2019). *The Age of Surveillance Capitalism: The Fight for a Human Future at the New Frontier of Power*. New York: PublicAffairs. [Wikipedia+1](#)
- Borowski, A. (2022). "Günther Anders, a forgotten prophet for the 21st century?" *Aeon*. [Aeon](#)
- Crawford, K., & Joler, V. (2018). "Anatomy of an AI System: The Amazon Echo as an Anatomical Map of Human Labor, Data and Planetary Resources." Shared under Creative Commons. [Essra](#)
- Floridi, L. (2024). "The Ethics of Artificial Intelligence: Exacerbated Problems, Novel Challenges, and Promising Opportunities." *SSRN Working Paper*. [SSRN](#)

Mitchell, M. (2019–). *AI: A Guide for Thinking Humans* (newsletter). Ongoing commentary on developments in AI and public understanding. [AI Guide](#)

The author has published over 4,500 AI-assisted papers at:
<https://www.researchgate.net/profile/Douglas-Youvan/research> . Google Scholar has indexed these preprints, and if you want a particular reference, the AI mode of a Google search (Gemini) has efficiently catalogued these preprints.