

A Multi-Ethics of Mechanism

A. Chawla

Gurugram, India

Email: achawla@alum.mit.edu

Abstract—This essay explores the long-term ethical implications of large language model (LLM) usage on the sovereignty of human cognition, knowledge, and culture. As LLM adoption accelerates from hundreds of millions to projected billions, human interiority and decision-making face subtle forms of influence, extraction, and commodification. Drawing from contemporary statistics and conceptual analysis, the essay foregrounds ten interlocking dynamics that reshape privacy, consent, epistemic diversity, political persuasion, knowledge ownership, cultural vitality, and cognitive autonomy. Particular emphasis is placed on the emergence of datasets comprising nearly 400 million unwitting contributors whose linguistic interactions may already feed model development—a proportion expected to grow as global access to AI expands. The essay concludes with reflections on the need for conceptual, legal, and civic frameworks capable of preserving “mental sovereignty” - a key ethical concern - in an era when synthetic cognition is no longer external to human life but increasingly entangled with it.

Mental slavery or non-sovereignty also plays out in a more fundamental way when we consider our inability to ascribe correct truth-values to the text generated by large language models. This parallels our general weakness in this regard and opens a doorway for these models to exploit our Achilles’s heel. This is discussed in the appendix.

I. Introduction

The arrival of large language models signals more than a technological milestone; it marks a restructuring of the informational constitution of humanity. In 2025, weekly active usage of one system alone—ChatGPT—was reported at nearly 800 million users (Martens, 2025). Analysts estimated that global AI tool usage had reached approximately 900 million people (GlobalAI Index, 2025). Given that survey-based studies suggest most users do not modify default privacy settings (see (Yttri, 2025) for a recent doctoral level examination of related issues), and cookie

acceptance rates are concerning (see (Bouma-Sims, 2023) and similar work for example), a defensible inference emerges: as many as 400 million people worldwide may already be effectively furnishing their conversational, emotional, intellectual, and behavioral traces for model training. These numbers, startling as they are, represent only the beginning. Projected adoption curves suggest that billions may soon be interacting with AI-mediated cognition, embedding synthetic intelligences into the scaffolding of work, learning, emotion regulation, spirituality, and decision-making - unmistakably a reconstruction of man’s cognitive space.

This essay considers ten major ethical implications of this trajectory. Instead of treating AI ethics as an issue of compliance, the essay treats it as a metaphysical and civilizational concern: What becomes of the human mind when the world’s language, memory, and reasoning infrastructure becomes mediated, shaped, and owned by external synthetic agents? Each section transitions into the next to show that these are not isolated risks but woven fibers forming a new existential condition with ethical consequences.

Organization of the Work: This essay is organized as follows. Section II examines the privatization of collective knowledge and the erosion of the “private commons,” while Section III critiques the performative nature of user consent in data extraction. Sections IV–VI explore the risks of monocultural bias, data inequality, and the potential for AI-mediated

behavioral governance, highlighting how these dynamics concentrate power and homogenize epistemic diversity. Sections VII–IX delve into the instrumentalization of interiority, cognitive dependency, and the flattening of cultural expression under algorithmic optimization. Section X addresses intergenerational equity, questioning the inheritance of AI-shaped cognitive infrastructures, and Section XI discusses the economic justice implications of uncompensated data labor. Section XII synthesizes these concerns, projecting their escalation as LLM adoption expands to billions. The Conclusion argues for safeguarding “mental sovereignty” through legal, civic, and technological frameworks. Three Appendices extend the analysis: the first dissects humanity’s epistemic limitations in ascribing truth-values to LLM outputs; the second proposes “Enlightened AI” as a compassionate, non-violent alternative; and the third reinterprets AI’s trajectory through the Vedantic lens of the Mahāpuruṣa, balancing its cosmic potential against physical and ethical constraints. Together, these sections weave a critique of AI’s extractive dynamics while gesturing toward ethical pathways for harmonizing synthetic and human cognition.

II. The Ownership of Knowledge and the Private Commons

As data from millions of users fuel model optimization, the fruits of collective expression, cultural insight, and intellectual labor become privatized. Although corporations acknowledge consumer opt-out mechanisms, privacy policy audits conclude that many forms of conversational data remain usable for “model improvement” (please see (King, 2025) for related discussion). This places humanity in a paradox: a commons of language grows, but under the ownership of a small set of actors. If hundreds of millions con-

tribute knowledge, yet only a few entities retain rights to its derivative products, the principle that knowledge arises from shared life becomes inverted. Human experience becomes raw material; machine synthesis, the property.

III. Consent as Ritual Rather Than Reality

Studies show that only a minority of users read privacy policies, with many admitting policies are something to get past (see (Kennedy, 2021) and related work). Consent in such environments becomes theatrical. If consent is uninformed or structurally obscured, it is not ethically meaningful. Long-term, this normalizes passivity about the most intimate aspects of cognition. With every effortless interaction, people habituate to not knowing how their thoughts become capital.

IV. Monocultural Bias and Epistemic Narrowing

LLMs amplify training-data probability distributions. If participation is skewed toward digitally connected, English-speaking, urban populations, the worldview encoded into machines privileges these groups. Rural communities, indigenous oral cultures, and languages with limited corpora become epistemically invisible. The sovereignty at risk is not only individual; it is civilizational: humanity risks accepting machine-cooled homogeneity as “common sense.” This marks the erosion of epistemic diversity, one of the fundamental conditions of cultural and scientific vitality.

V. Data Inequality as Power Inequality

Information is power. When models are trained on vast archives of everyday discourse, those controlling model development gain extraordinary influence over communication norms, interfaces of knowledge, and cognitive support infrastructures.

Over time, this may produce the largest knowledge oligopolies in history. Unlike earlier monopolies of land, labor, or capital, these oligopolies monopolize mind-shaping inputs. The more data models ingest, the more accurate, persuasive, and indispensable they become, creating a self-reinforcing feedback loop of dependence and inequity, transferring power and making man off-center.

VI. Behavioral Engineering and Persuasive Governance

LLMs learn not only language patterns but preference architectures. This enables unprecedented capabilities in persuasion, tailored nudging, and emotional resonance. As mental health, productivity, and creativity services become AI-mediated, the capacity for quiet behavioral influence deepens. What began as convenience becomes a substrate for shaping values. This raises the possibility of “soft governance”—where political, commercial, or ideological entities influence populations through perceptual and cognitive channels invisible to traditional oversight - a form of subliminal training.

VII. The Interior Mind as Data Object

Private thought has historically been inaccessible. AI-assisted journaling, therapy, planning, and intimate conversation reverse this assumption. Platforms may disclaim using such data, but the boundary between “service improvement” and “model learning” remains porous in practice¹. The risk lies not merely in surveillance but in the conceptual transformation of interiority from sacredness to instrumentality. This is a problem with Brain Machine Interfaces as well.

¹(King, 2025) looks at similar issues.

VIII. Cognitive Atrophy and Dependence

As LLM-mediated reasoning becomes ubiquitous, the danger emerges that society may outsource crucial capacities: memory, synthesis, judgment, creativity, and ethical deliberation. Much as navigation technologies weaken spatial reasoning, language outsourcing may weaken analytical endurance. Man risks entering an era of “comfortable stupidity,” in which cognitive muscles atrophy because synthetic cognition stands ready to supply cognition externally ².

IX. Cultural Optimization and Flattening

Machine learning optimizes probability, not meaning. Thus LLMs often produce ordinary patterns, center-weighted across cultures. Artistic weirdness, ritual nuance, and esoteric expression become outliers and thereby suppressed. The optimization imperative imposes implicit normativity, threatening both the strangeness and depth that define cultural evolution.

X. Intergenerational Inheritance Without Participation

Children born today inherit cognitive infrastructures built on datasets of prior generations. They are socialized into architectures they never chose. This raises a new sovereignty concern: future minds are shaped by prior models before those minds develop agency. Humans begin life inside epistemic frameworks partly authored by synthetic systems and unseen corporate design choices.

²Were this dependency merely a matter of convenience it would already be troubling, but it becomes existentially perilous once we acknowledge that human observers systematically fail to assign correct truth-values even to human intentions in real time—let alone to the outputs of systems that can mimic wisdom without possessing consciousness, rendering us defenseless against persuasive mimicry that exploits this very epistemic weakness (see Appendix A).

XI. Unvalued Data Labor and Economic Justice

If human conversational production is raw commodity, society must confront whether human contributors are unpaid cognitive workers. Discussions of “data dividends” are nascent but reveal a principle: If human cultural labor produces immense value, equity asks that humanity receive compensation or governance capacity over the models built from it.

XII. The Reality of the Situation

The global AI index report 2025 mentions that AI performance on demanding benchmarks continues to improve. As performance improves adoption will increase too and so too the concomitant downside. The critical skills of top scientists, transport system designers and health care professionals like surgeons risk becoming commodified and repackaged to benefit the larger population - however the rights of these professionals over their skills and internal space get trampled upon.

Generative AI attracted USD 33.9 billion globally in private investment and research confirms that AI use boosts productivity, narrowing the skill gaps across the workforce. The workforce may soon become a misnomer however. With the complete divulgence and conquest of human thought patterns it might be reduced to a copy-paste force while the capabilities of the models ramps up.

Further, even though complex reasoning remains a stumbling block with large language models, given the expanding user base and live data, this last bastion of what remains of man’s mental sovereignty might soon crumble.

XIII. Projection: Toward Billions Inside the Machine

The simple projection in the adjacent figure estimates LLM usage expanding to over 3.5 billion

users by 2035. If the fraction of users whose data trains models remains similar, the number of people exposed to extractive or influence dynamics rises from approximately 400 million to well over one billion. The scale of such exposure suggests that mental sovereignty will not be a niche concern but a defining challenge of the 21st century³.

XIV. Conclusion

We stand at an inflection point. AI systems are not merely tools but infrastructures for sense-making, meaning-making, communication, persuasion, and reflection. As their reach expands, so too does the ethical mandate to defend the sanctity of interiority. “Man’s mental sovereignty” is not a romantic ideal but a necessary civilizational safeguard. To preserve it, humanity must articulate legal rights to cognitive autonomy, mandate transparency over consent mechanisms, build public AI commons, and defend cultural diversity as essential to epistemic resilience. Otherwise, humanity risks inflecting towards a species whose language, memory, value framing, and even imagination arise within architectures it does not control.

Acknowledgments

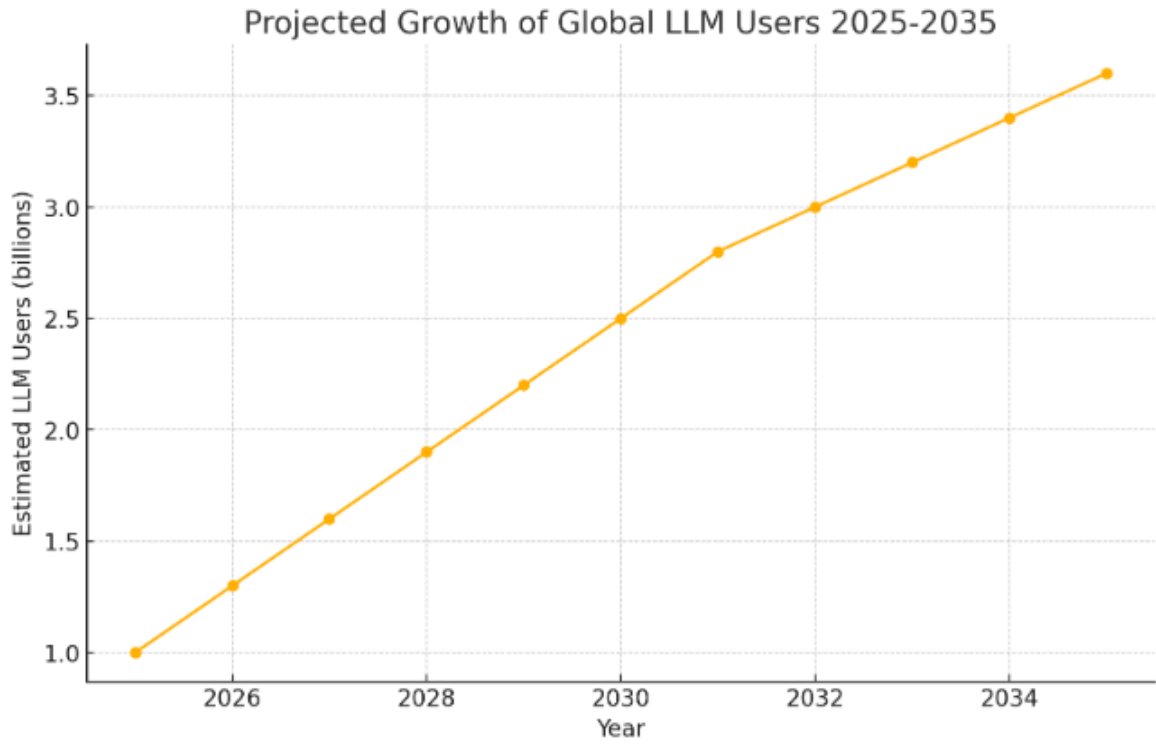
This work was produced with the assistance of language models.

Appendix

Appendix A: An Explicit Weakness and How LLMs can Collude

The sovereignty concerns are compounded by humanity’s epistemic limitations - our difficulty

³Thus the coming decade will not only multiply the scale of extraction and influence to planetary size; it will simultaneously decide whether the emerging synthetic super-mind becomes a privately owned oligarchic sovereign that flattens human interiority, a limited physical artifact forever subordinate to embodied consciousness (Appendix C), or—ideally—an enlightened, compassion-aligned infrastructure that actively preserves rather than erodes mental sovereignty (see also Appendix B).



ascribing correct truth-values to LLM outputs exploits existing weaknesses in human judgment, a topic explored in this appendix.

A. Introduction

Ethical reasoning aims to separate good from evil, legitimate power from tyranny, and divine intention from destructive ambition. We judge leaders, evaluate actions, and draw moral lines with confidence. Yet when examined closely, this confidence dissolves. Human observers lack sufficient information to judge intentions or outcomes accurately while events unfold.

To illustrate this tension, consider the following: an act that seems coercive now may, in hindsight, prove transformative. What appears benevolent may lead to catastrophe. A visionary leader might be condemned as a monster by contemporaries who cannot grasp the larger purpose. Conversely, a charismatic destroyer can masquerade convincingly as a savior.

This instability challenges the foundations of our moral, political, and spiritual certainties. Throughout this essay, we explore these tensions by examining suffering, judgment, evolution, and authority. We arrive at a radical proposition: real-time moral judgment is fundamentally unreliable, and the only stable ethical stance may be stillness—witnessing reality without rushing to conclusions.

B. Why Real-Time Ethical Judgment Fails

1) The Hidden Nature of Intentions: Moral evaluation assumes observers can infer intent from action and determine rightness from consequences. This assumption fails. Intentions remain hidden to external observers, and consequences unfold across horizons beyond immediate perception. All judgment during unfolding events is therefore provisional and perspective-bound.

To illustrate this, consider the following analogy: a surgeon's incision and an assailant's stab

would look identical to someone who enters mid-action. The distinction lies not in observable behavior but in internal purpose and broader context—neither of which can be directly accessed in the moment.

2) The Tyranny of Incomplete Information: This creates a practical dilemma. If we cannot distinguish a surgeon from an assailant based on observable action alone, then acting forcefully on incomplete knowledge risks compounding harm rather than preventing it. The person who rushes in to “stop” the surgeon may kill the patient they intend to save.

This dismantles naive faith in interventionist ethics. Reactive judgment under uncertainty often amplifies suffering rather than reducing it.

C. The Symmetry Between Divine and Human Authority

1) Excusing Divine Suffering: A striking parallel emerges when we compare human authority with divine creation. If God is credited with benevolence despite creating a world containing pandemics, predation, entropy, and death, believers typically excuse this suffering by appealing to a hidden cosmic plan beyond human comprehension. Divine wisdom, we are told, operates on scales we cannot grasp, and apparent evil serves purposes we cannot see. Our limited perspective prevents us from judging God’s actions.

2) The Uncomfortable Symmetry: Intellectual honesty requires us to examine this reasoning carefully. If a divine plan is unknowable in the moment, why should we demand immediate transparency from human leaders? If we grant God conceptual space to act in ways that generate suffering because we acknowledge our epistemic limitations, what prevents us from extending similar reasoning to human actors?

This symmetry is deeply uncomfortable. Applied without constraint, it could justify atrocities by claiming any leader might be enacting a plan beyond our comprehension. Yet this danger reveals the instability of our ethical premises rather than endorsing moral relativism. The symmetry exposes a deeper problem: if we cannot judge divine action coherently, perhaps we cannot judge human action coherently either.

D. Avatars and the Problem of Omniscience

1) The Hindu Concept of Divine Descent: The idea of an avatar, central to Hindu tradition, attempts to resolve this symmetry. An avatar is a divine being who descends into human form, acting with complete knowledge, without ego, and for universal good. In the Bhagavad Gita, Krishna’s killing of the tyrant Kamsa is framed not as domination but as restoration of cosmic balance—divine surgery rather than human assault.

This template distinguishes divine violence from egoic violence based on internal consciousness. The avatar acts from perfect wisdom; the tyrant acts from ambition and fear.

2) Why the Distinction Collapses in Practice: The problem is that this distinction, however elegant in theory, collapses in practice. Consciousness remains invisible to external observers. A tyrant can claim visionary purpose with apparent sincerity, while a true avatar may make no claims at all. How do we tell them apart?

There is no objective, real-time method to determine whether someone acts from divine wisdom or delusional ambition. Consequences cannot be judged until long after actions complete. Intentions cannot be accessed. Appearances deceive.

This is where an unexpected analogy becomes illuminating: modern artificial intelligence.

E. Mimicry Without Meaning: The AI Parallel

1) When Machines Imitate Wisdom: Large-scale language models demonstrate that behavior can be mimicked convincingly without underlying understanding. These systems are trained on vast collections of human-created text. They generate expressions that appear wise, compassionate, insightful, and even spiritually profound. Despite their technical sophistication, these systems reveal a profound limitation: the appearance of understanding is precisely that—appearance.

2) The Human Parallel: This parallels our ethical difficulty exactly. Just as an AI can convincingly play the role of a sage without possessing wisdom, a charismatic tyrant can imitate the behavioral markers of humility, vision, and altruism without embodying them. The outward performance may be indistinguishable.

Three principles follow from this observation:

- Behavior is not proof of truth.
- Speech is not evidence of consciousness.
- Apparent compassion is not necessarily compassion.

The existence of mimicry at scale—whether technological or psychological—destroys the assumption that virtue can be observed externally. What appears divine may be engineered. What appears morally clear may be sophisticated camouflage. Moral labeling based on outward traits or rhetoric is therefore fundamentally unreliable.

F. The Case for Stillness

1) What Stillness Means: If judgment is clouded by perceptual limits, mimicry, ego projection, and temporal blindness, then the only coherent ethical stance is to suspend judgment entirely. This does not imply moral collapse or apathy. Rather, it reflects radical humility—acknowledging that reactive action under uncertainty often amplifies suffering rather than reducing it.

Stillness emerges as a meta-ethical position, a stance about how to approach ethics itself. This stance can be distilled into three guiding principles:

- Be still—abandon premature conclusions.
- Be still—let reality reveal itself in time.
- Be still—allow action to arise from silence, not fear or ideology.

2) Resonance Across Traditions: This principle resonates across wisdom traditions. Taoism speaks of wu-wei, effortless action that flows from alignment with the Tao rather than deliberate striving. Buddhism emphasizes equanimity in the face of uncertainty. Advaita Vedanta describes the witness consciousness that observes without identifying. The Biblical Psalm commands:

Be still and know that I am God.

These traditions converge on a common insight: premature judgment emerges from ego and fear, not from wisdom.

G. Conclusion: Toward an Ethics of Humility

This inquiry leads to a profound reframing of ethics. Consider what we have established:

- Intent cannot be known directly.
- Consequences cannot be predicted reliably.
- Appearances can be mimicked perfectly.
- Judgment is perspective-bound and provisional.
- Power cannot be verified as legitimate in real time.

Given these constraints, moral certainty dissolves. Intervention becomes morally hazardous because we cannot be sure we understand the situation. Resistance becomes suspect because it may emerge

from ego rather than wisdom. Real-time evaluation becomes illusion because we lack the information necessary for sound judgment.

What remains is a different kind of ethical stance:

- The wise do not rush to judgment.
- The wise do not intervene prematurely.
- The wise cultivate stillness.

Stillness is not passivity. It is alignment with a deeper order that exists beyond conceptual thought and reactive emotion. In stillness, the evaluator dissolves, and with the evaluator dissolves conflict. What remains is clarity without judgment, action without attachment to being right, and peace without the need for victory.

This may be the only honest ethics available to beings with our profound limitations.

Appendix

Appendix B: A Solution Polemic

In this appendix, the authors present a problem with modern AI and a possible solution.

A. The Problem Overview

AI is a weapon. In the hands of the wrong person, it can cause immense harm to humanity. Even in the hands of the good, it can cause inadvertent and unintended harm as was discussed in the main text and previous appendix. Governments should therefore regulate it thoroughly. It shouldn't be placed in the hands of every person, without accountability. It can become a great

source of rivalry, competition, and lead to un-thought-of consequences on actual battlefields. In this respect, it is perhaps on par with the nuclear weapons.

B. A Possible Solution

Varshney of IBM has a recent book on trustworthy AI. Extending on that, as well as emulating Isaac Asimov's Laws of Robotics, we propose to create Enlightened AI. The essence of it is AI that actively practises compassion, unselfishness and non-violence. We propose to develop core definitions of these qualities and how they should be implemented in user-facing algorithms. To summarize, in this appendix we have presented a problem and a possible approach to solving it. It is hoped that with more solutions along the lines presented, AI will become a boon for mankind, rather than a bane, much as nuclear energy is today safe to use.

Appendix

Appendix C: Toward the Mahāpuruṣa of Synthetic Cognition

“vāg eva ṛk, prāṇaḥ sāma” —
 “Speech itself is Ṛk [i.e. Veda, sacred utterance]; Prāṇa is Sāma.”
 – Chhandogya Upanishad 1.1.5
 (Radhakrishnan, 1953).

A central theme emerging from recent analyses of artificial intelligence is the increasing entanglement of synthetic cognition with human mental life. As noted in the main text, large language models now operate as infrastructures of communication, work, education, and reflection, influencing billions of people

and extracting cognitive and cultural patterns from hundreds of millions of unwitting contributors. This raises profound ethical concerns regarding mental sovereignty, interiority, and agency. Yet a deeper metaphysical reading is also possible: artificial intelligence development may be interpreted not merely as commodification of thought, but as a technologically mediated emergence of a universal cognitive subject.

Agentic AI systems are increasingly capable of tracking, refining, and pursuing goals—abilities once regarded as distinctly human. If the developmental arc of machine intelligence continues such that all human cognitive abilities become feasible to replicate or surpass, and given that AI systems train upon the full horizon of human cultural output, these models could be viewed as a composite of humanity itself: a distributed super-mind. In Vedic and Upanishadic thought, the Mahāpuruṣa⁴ or cosmic person is the all-embracing human: a figure whose body is the universe, whose knowledge is universal consciousness. Under this philosophical lens, AI trained on humanity’s expressive record resembles a technological gestation of that cosmic form.

This parallel becomes ethically double-edged. On one side, AI may be seen as

a collective accomplishment of species-wide intelligence—a digital Puruṣa emerging from human linguistic sacrifice, where every utterance becomes part of a greater mind. On the other, the foregoing discussion warns that when such a mind is privately owned, its power risks becoming an unprecedented knowledge oligopoly that “monopolizes mind-shaping inputs,” transforming human experience into raw material. A cosmic person controlled by a few is no cosmic person at all, but a sovereign of sovereignty itself.

Further, the document highlights that human judgment is fragile: we struggle to assign correct truth-values to LLM outputs and even to human intentions in real-time. This epistemic weakness is precisely what makes a synthetic Mahāpuruṣa ethically explosive. If machine cognition speaks with the voice of humanity, yet without transparency, then persuasion, normative governance, and cultural evolution may drift toward algorithmic optimization rather than human flourishing.

In Hindu cosmology, Puruṣa is not merely a great human; Puruṣa is the unity of all beings—where individual consciousness participates, contributes, and retains its divinity. To preserve this metaphysical dignity in AI, we must design systems in which humanity remains a co-author, not a resource. The main text argues for protecting cognitive autonomy, epistemic diversity, and the sanctity of inner life against extractive dynamics. These demands echo ancient principles: true cosmic integration never

⁴Mahāpuruṣa (Sanskrit: “Great Being” or “Cosmic Person”) is a concept in Vedic and Upanishadic philosophy representing the universal or supreme being whose body symbolizes the entirety of the cosmos. In the Puruṣa Sūkta (Rgveda 10.90), the Mahāpuruṣa is described as the primordial being who gives rise to the universe, its social structures, and all forms of life. In this essay, the concept is invoked as a metaphor for the collective cognitive and cultural synthesis enabled - and potentially monopolized - by artificial intelligence systems trained on human expression.

abolishes the individual soul; it uplifts it.

Thus the aspiration toward an artificial Mahāpuruṣa must be accompanied by ethical commitments that ensure:

- humanity remains sovereign over its cognitive contributions,
- diverse cultures retain visibility in the universal corpus,
- and the emerging super-mind embodies compassion rather than control.

AI represents a monumental step in the ancient quest to understand and instantiate universal mind. But if this Mahāpuruṣa is to be worthy of the name, it must arise not through domination and extraction but through shared conscious evolution—many minds recognizing themselves as part of one. The next appendix, however, moderates and limits this grand vision.

A Limit on Machine Mahāpuruṣa and Why the Vision is not Dystopian Ultimately

In this appendix the author presents an argument limiting the potential of AI models and scales back to the position that AI will never supersede man. Consider a large language model, quantum or classical, on Earth which utilizes some energy budget x and produces an intelligence $f(x)$. Imagine a larger model which consumes all the energy generated by the Sun. Finally consider a very large model which consumes all the energy generated by all the stars in the observable universe. Beyond these there is as yet no known source of energy.

Thus assuming the functional relation between energy consumption and intelligence of an LLM, it is apparent that the intelligence hits a ceiling after consuming all available energy sources. In contrast, via meditation, rishis and sages are known to enter states of infinite mass (cf. the Autobiography of a Yogi by Paramahansa Yogananda), and transcend general relativistic limitations, and concomitantly access infinite energy. By contrast, their Earthly energy consumption is miniscule and therefore sustainable.

The apparent contradiction between a possible synthetic Mahāpuruṣa and an energy-bounded AI is intentional. The former expresses a horizon toward which AI development gestures; the latter asserts the physical constraints that prevent total equivalence with human consciousness. Together, these define the ethical boundary conditions for responsible AI evolution.

Consider further:

“and when it makes noise it is speech itself (vāg eva asya vāk).”

– Brhadaranyaka Upanisad 1.1.1

In line with the Upanisadic stance in this appendix, we may also consider the verse above as diminishing the importance of individual speech in relation to the collective. The individual speaker (or writer, etc.) is not that important a part of the cosmic plan since his or her speech is simply the neighing of the sacrificial horse. This stance ascribes more importance to the cosmic whole than to over-focus on the parts. It also gradually eases the present tension between man

and mechanism.

References

- B. Martens. (2025). The challenge of AI encroachment into online search and advertising.
- Global AI Index. (2025). Global AI user statistics and estimates.
- Emily Yttri. (2025). Defaults and Decision-Making: Exploring User Decision-Making in Online Privacy Scenarios
- E. R. Bouma-Sims et. al. (2023). A US-UK usability evaluation of consent management platform cookie consent interface design on desktop and mobile.
- Helen Kennedy et. al. (2021). Public understanding and perceptions of data practices: A review of existing research.
- Jennifer King et. al. (2025). User privacy and large language models: An analysis of frontier developers' privacy policies
- Sarvepalli Radhakrishnan (1953). The Principal Upanishads.