



From Rational Human to Algorithmic Human: Rethinking Human Ontology in the Age of Artificial Intelligence

Sayyed Mahdi Biabanaki ¹ 

1. Assistant Professor, Faculty of Theology and Ahl Al-Bayt Studies, University of Isfahan, Isfahan, Iran. E-mail: mehdibiabanaki@gmail.com

Article Info:

Article type:

Research Article

Article history:

Received:

4 February
2026

Received in

revised form:

13 March 2026

Accepted:

6 April 2026

Published

online:

15 June 2026

Keywords:

Human
Rationality;
Algorithmic
Rationality;
Artificial
Intelligence;
Human Ontology;
Philosophy of
Technology;
Ethics of AI

Abstrac**t**: The rapid advancement of artificial intelligence has reignited fundamental philosophical questions concerning rationality and the nature of the human being. While contemporary AI systems increasingly exhibit behaviors traditionally associated with human intelligence, this convergence has obscured crucial ontological distinctions between human and algorithmic forms of rationality. This article offers a philosophical and ontological analysis of rationality as it transitions from a classical human-centered concept to an algorithmic framework embedded in artificial intelligence. Drawing on classical philosophy, contemporary AI theory, and critical perspectives in philosophy of technology and theology, the article argues that human rationality constitutes an existential mode of being rather than a mere functional capacity. Human rationality is inseparable from embodiment, lived experience, moral responsibility, and openness to meaning. By contrast, algorithmic rationality operates as a computational mode of operation, defined by optimization, prediction, and instrumental goal satisfaction, without access to phenomenological or ethical dimensions of existence. The central claim of this study is that the primary ontological risk posed by artificial intelligence does not lie in the possibility that machines may become human-like, but in the gradual redefinition of the human according to algorithmic standards of rationality. Such a reduction risks undermining the conceptual foundations of human dignity, agency, and responsibility. The article concludes that a robust ethical and theological engagement with artificial intelligence must be grounded in a renewed account of human ontology that resists the reduction of rationality to computation..

Cite this article: Biabanaki, S. (2026). From Rational Human to Algorithmic Human: Rethinking Human Ontology in the Age of Artificial Intelligence, *Philosophical Meditations*, 15(Special issue: 37), 165-195. <https://doi.org/10.30470/phm.2026.2081304.2787>

© The Author(s).

Publisher: University of Zanjan.

DOI: <https://doi.org/10.30470/phm.2026.2081304.2787>

Homepage: phm.znu.ac.ir



Introduction:

Artificial Intelligence (AI) generally refers to a branch of computer science and interdisciplinary studies aimed at designing and developing systems capable of performing tasks that have traditionally been understood as requiring human intelligence. These tasks include learning, reasoning, problem-solving, natural language understanding, pattern recognition, and decision-making. One of the most influential definitions describes artificial intelligence as the study of *rational agents*—agents that perceive their environment and act in ways that maximize the likelihood of achieving their goals (Russell & Norvig, 2022, pp. 4–5).

Nevertheless, the definition of artificial intelligence has always been contested, largely because the concept of “intelligence” itself is philosophically, psychologically, and even theologically complex. Some definitions emphasize functional similarity to human intelligence, as exemplified by the Turing Test, while others

focus on efficiency and optimization, independently of human likeness. In recent decades, with the rise of machine learning, deep neural networks, and large language models, artificial intelligence has moved from a theoretical and experimental endeavor to a pervasive force shaping everyday life, the economy, politics, education, and culture.

Today, artificial intelligence is not merely a technological tool but a civilizational phenomenon. Intelligent algorithms play an increasingly decisive role in economic decision-making, medical diagnosis, security surveillance, resource management, knowledge production, and even the formation of public opinion. For this reason, some scholars describe artificial intelligence as a “fourth revolution”—following the Copernican, Darwinian, and Freudian revolutions—that once again reshapes humanity’s understanding of its place in the world (Floridi, 2014, pp. 1–3).

One of the dominant approaches to artificial intelligence is an optimistic and progress-oriented perspective. From this viewpoint, AI is

understood as a powerful instrument for solving complex human problems, increasing welfare, expanding knowledge, and even enhancing human capacities. Early thinkers such as Herbert A. Simon were among the first to conceptualize artificial intelligence as a natural extension of human rationality. Through his theory of *bounded rationality*, Simon argued that machines could compensate for human cognitive limitations and improve decision-making processes (Simon, 1997, pp. 88–92). In contemporary AI research, scholars such as Stuart Russell and Peter Norvig conceptualize artificial intelligence as the design of rational agents that can act effectively in pursuit of human-defined goals (Russell & Norvig, 2022, pp. 35–38). In this framework, AI is primarily regarded as a neutral and powerful tool whose value depends on how it is designed and used.

A more radical form of optimism toward artificial intelligence is found in transhumanism. Thinkers such as Nick Bostrom argue that advanced technologies—

including AI—have the potential to free human beings from biological, cognitive, and even mortal limitations (Bostrom, 2014, pp. 15–20). Within this perspective, the boundary between human and machine is viewed as fluid rather than fixed, and the future of humanity is imagined as one of coexistence or even integration with intelligent systems.

Opposed to technological optimism is a critical and, at times, pessimistic approach that emphasizes the potentially destructive consequences of artificial intelligence for human beings and society. These perspectives draw on traditions in critical philosophy, phenomenology, theology, and social studies of technology. Martin Heidegger was among the earliest thinkers to warn against the dangers of technological thinking. Although he did not address artificial intelligence directly, his analysis of modern technology as a mode of revealing that reduces beings—including humans—to mere resources remains highly relevant to contemporary critiques of AI (Heidegger,

1977, pp. 17–23).

Philosophers such as Luciano Floridi caution that algorithms can reproduce and intensify structures of power, discrimination, and inequality, often under the guise of objectivity and neutrality (Floridi, 2014, pp. 85–92). Similarly, social theorists like Shoshana Zuboff argue that within the framework of surveillance capitalism, artificial intelligence becomes a tool for monitoring, predicting, and manipulating human behavior on an unprecedented scale (Zuboff, 2019, pp. 376–380).

From a theological perspective, thinkers such as Ronald Cole-Turner argue that artificial intelligence and transhumanist projects risk undermining human dignity by reducing the human person to a collection of functions or upgradeable capacities (Cole-Turner, 2011, pp. 92–110). In this view, the central problem is not the power of machines, but the erosion of the meaning of being human itself.

The present article does not align itself entirely with either technological optimism or absolute pessimism. Its primary

aim is a philosophical and ontological analysis of how artificial intelligence is transforming the concept of the human being. The central question guiding this study is how the transition from the “rational human” to the “algorithmic human” is reshaping our understanding of rationality, responsibility, meaning, and human existence.

While critical approaches to artificial intelligence and its reductionist tendencies have a well-established lineage in the philosophy of technology—from Heidegger’s analysis of technological enframing to contemporary phenomenological critiques—this paper advances the debate by shifting the focus from technology as an external force to ontology as an internal transformation. Rather than treating artificial intelligence primarily as a tool that threatens human autonomy or as an ethical problem requiring regulation, this study examines how algorithmic systems subtly reconfigure the conceptual conditions under which the human being understands itself as rational. The novelty of this approach lies in its claim that

the most significant impact of AI is not technological displacement or moral delegation, but an ontological reorientation in which human rationality is increasingly interpreted through the formal, operational logic of algorithms.

More specifically, this article develops and defends a conceptual distinction between rationality as a mode of being and rationality as a mode of operation. Classical and modern philosophical accounts treat rationality as an existential dimension of human life—embedded in embodiment, self-understanding, normativity, and meaning-making—whereas contemporary AI systems instantiate rationality as a set of computational procedures optimized for prediction, control, and efficiency. By articulating this distinction with ontological precision, the paper argues that the central risk posed by artificial intelligence is not that machines might eventually become human-like, but that humans may come to understand themselves according to algorithmic standards of rationality. In doing so, the article offers a framework that moves beyond

synthesizing existing critiques and instead provides a novel ontological lens for evaluating the ethical, social, and theological implications of artificial intelligence.

2. Rationality in Classical Philosophy

In classical Greek philosophy, rationality was understood not merely as a mental or instrumental ability but as a fundamental, ontological feature of human beings. Aristotle famously defines humans as *zōon logon echon* — “rational animals” or “creatures possessing logos” — highlighting the essential distinction between humans and other beings (Aristotle, 1984, pp. 125–127). Here, rationality, expressed through *logos*, is not simply a calculative tool; it encompasses the capacity for understanding, dialogue, reflection, and ethical decision-making within communal life.

This ontological dimension of rationality was further developed in medieval theology, particularly in the work of Thomas Aquinas. Aquinas fused Aristotelian philosophy with Christian theology, framing human

rationality as an image of divine reason. Human reason, though limited, reflects aspects of the divine intellect, conferring innate dignity and moral responsibility on humans (Aquinas, 1947, pp. 78–81).

Transitioning to modern philosophy, the conceptual framework shifted, yet the ontological centrality of rationality remained. René Descartes, a foundational figure in modern thought, grounded human existence in the rational self, presenting self-conscious rationality as the basis of being. In his *Meditations on First Philosophy*, Descartes systematically applies methodological doubt to discard all beliefs susceptible to uncertainty, ultimately establishing the famous cogito argument: “*Cogito, ergo sum*” — “I think, therefore I am” (Descartes, 1996, pp. 17–19). Thus, whereas classical and medieval perspectives (Aristotle and Aquinas) conceive rationality as a multidimensional, intrinsic attribute linked to ethical and metaphysical ends, Descartes reframes it as a formal and methodological capacity underpinning self-

consciousness and knowledge acquisition. Even in modern philosophy, despite the shift away from metaphysical and theological frameworks, rationality remains the guarantee of human uniqueness and autonomy.

In summary, the classical and medieval philosophical tradition demonstrates that rationality was conceived not merely as a cognitive or instrumental skill but as a fundamental dimension of human existence, distinguishing humans from other beings. This lineage provides a crucial backdrop for later discussions on computational rationality in artificial intelligence, where the concept is increasingly formalized and abstracted from ethical or metaphysical considerations.

3. Rationality in AI: From Information Processing to Computational Agency

Rationality has remained a central yet contested concept throughout the development of artificial intelligence. While early AI research frequently assumed that rationality could be defined through formal optimization, contemporary discussions reveal a

fragmentation of the concept into multiple competing interpretations. Rather than representing disagreement alone, this plurality reflects an underlying transformation in how rational agency itself is understood within computational systems. The evolution from utility-based rationality toward computational and context-sensitive models suggests that rationality in AI is less a fixed property than a design-dependent framework shaped by practical constraints and theoretical assumptions (Macmillan-Scott & Musolesi, 2025). The following subsections examine key conceptions of rationality not merely as alternatives but as stages in an ongoing conceptual shift.

3.1 Rationality as Utility Maximization

One of the foundational models of rationality in artificial intelligence equates rational behavior with the maximization of expected utility. Rooted in decision theory and neoclassical economics, this perspective defines a rational agent as one that selects actions maximizing expected outcomes relative to

its beliefs (Russell & Norvig, 2022). This framework played a decisive role in early AI research, particularly in reinforcement learning and game-theoretic modeling, because it provided a clear mathematical criterion for evaluating decision-making.

However, the classical model rests on idealized assumptions regarding information availability, computational resources, and temporal constraints. Critics have argued that such assumptions limit its applicability to real-world contexts, where agents operate under uncertainty and limited resources (Simon, 1957). Consequently, utility maximization increasingly appears less as a universal definition of rationality and more as a theoretical baseline from which alternative models emerge.

3.2 Computational Rationality and Bounded Rationality

The limitations of ideal optimization models motivated the development of computational rationality, which reframes rational behavior as the outcome of

decision-making under resource constraints. Rather than assuming unlimited computational power, this approach evaluates rationality relative to the trade-offs between decision quality and computational cost (Gershman et al., 2015, p. 273).

This shift builds directly on Herbert Simon's concept of bounded rationality, which emphasizes that agents often pursue satisficing strategies rather than optimal solutions (Simon, 1957, 1982). Within AI, bounded rationality introduces a more realistic conception of agency in which rationality emerges from adaptive strategies shaped by environmental and computational limitations. The move from ideal optimization toward resource-aware reasoning marks a significant conceptual transition, redefining rationality as process-sensitive rather than outcome-only.

3.3 Behavioral Rationality versus Reasoning Rationality

A further transformation arises in the distinction between behavioral and reasoning rationality. Behavioral rationality evaluates agents

based on observable outcomes, whereas reasoning rationality focuses on the structure and validity of internal decision processes (van der Hoek & Wooldridge, 2003).

Modern AI systems, especially machine learning models, increasingly prioritize behavioral performance metrics. This emphasis reflects practical evaluation constraints but also signals a broader conceptual shift: rationality becomes associated with empirical success rather than logical transparency. As Macmillan-Scott and Musolesi (2025) note, this transition challenges traditional philosophical assumptions linking rationality with explicit reasoning and deliberation.

3.4 Normative versus Descriptive Rationality

Another axis of debate concerns whether rationality should be defined normatively—according to formal principles—or descriptively, reflecting actual decision-making behavior. Normative frameworks dominate classical decision theory, while descriptive approaches, influenced by behavioral economics and psychology,

emphasize deviations from ideal rational standards (Kahneman & Tversky, 1982).

In AI, this tension becomes particularly visible when systems learn from human-generated datasets. Such models may reproduce biases or heuristics that diverge from normative standards, raising fundamental questions about whether rationality should be judged against ideal models or contextual adaptation (Besold & Uckelman, 2018). This dilemma highlights the growing entanglement between computational design and sociocultural realities.

3.5 Imperfect Rationality and Apparently Irrational Behavior

Recent research further complicates traditional notions of rationality by demonstrating that behaviors appearing irrational at a local level may be globally optimal. Exploration strategies in reinforcement learning, for example, introduce randomness that improves long-term performance (Ladosz et al., 2022). Similarly, randomized strategies in game theory may represent equilibrium solutions (Leyton-Brown & Shoham, 2008).

These findings challenge intuitive understandings of rational action and suggest that rationality must be interpreted relative to long-term goals and environmental uncertainty rather than immediate coherence or predictability.

3.6 Generative AI and the Paradox of Rationality Evaluation

The emergence of generative AI introduces a further complication. Large language models are optimized for predictive objectives such as minimizing next-token error. From the perspective of their training goals, their outputs may be entirely rational. However, human evaluators apply criteria such as truthfulness, meaning, and coherence that may diverge from optimization targets. This discrepancy generates a paradox: a system may be rational relative to its formal objective while appearing irrational from a human normative standpoint (Macmillan-Scott & Musolesi, 2025).

Taken together, these perspectives reveal that rationality in artificial intelligence is best understood as a layered and evolving

concept. The trajectory from utility maximization toward computational, behavioral, and context-sensitive models reflects a broader reconfiguration of rational agency itself. Rather than a singular property, rationality emerges as a framework shaped by design goals, constraints, and evaluative perspectives—an insight that becomes crucial when comparing algorithmic rationality with human rationality.

4. The Relationship between Human Rationality and Algorithmic Rationality

This section advances the central comparative claim of the article: while human rationality and algorithmic rationality may appear functionally analogous at the level of performance, they are ontologically heterogeneous. The growing tendency to evaluate artificial intelligence in terms of human-like rationality obscures a more fundamental distinction between rationality as an existential mode of being and rationality as a computational mode of operation. By systematically contrasting these two forms of rationality, this section argues that algorithmic

systems can simulate rational behavior without possessing the ontological conditions—such as embodiment, intentionality, and moral responsibility—that ground human rational agency. Clarifying this distinction is essential for resisting reductive accounts of the human person in contemporary AI discourse

Contemporary research in the philosophy and social studies of artificial intelligence suggests that algorithmic rationality is not a straightforward continuation of human rationality but rather a reconfiguration of rationality within a computational framework (Berente et al., 2021). As Stein and Shollo (2025, pp 2-3) argue, this reconfiguration rests on specific assumptions about cognition, decision-making, and agency whose intellectual roots can be traced back to Herbert Simon's theory of bounded rationality.

Accordingly, rational decision-making does not consist in optimization but in satisficing—seeking solutions that are good enough under given constraints. Stein and Shollo show how this conception gradually became the theoretical foundation for rationality in artificial

intelligence, positioning machines as tools designed to compensate for or transcend the limitations of human rationality (Stein & Shollo, 2025, pp. 2–3). Two key assumptions underlie this shift:

1. Cognition can be reduced to information processing;
2. If cognition is information processing, then computational systems can replicate or surpass human reasoning.

These assumptions have shaped much of classical and contemporary AI research, effectively redefining rationality in terms of computational efficiency, predictive accuracy, and decision optimization.

Stein and Shollo describe this transition as a move from logical rationality to statistical rationality (2025, pp. 4–5). Large language models and transformer architectures represent the culmination of this trend: rationality is operationalized as the ability to generate statistically plausible responses rather than truthful, justified, or meaningful ones in a human sense. In summary,

rationality in artificial intelligence can be characterized as a reduced, instrumental, and functional form of rationality, oriented toward efficiency, prediction, and optimization. Unlike classical rationality, it lacks normative, teleological, and existential dimensions.

In the classical philosophical tradition, as demonstrated in earlier sections, rationality was never understood as a detachable faculty but as a constitutive dimension of human being. Human rationality was inseparable from self-consciousness, embodiment, emotion, ethics, and meaning. By contrast, algorithmic rationality—even in its most advanced forms—remains largely confined to symbolic manipulation, statistical inference, and computational optimization. The central question of this section, therefore, is not whether human and algorithmic rationality resemble each other in practice, but whether they are ontologically commensurable at all. Consequently, a deeper question emerges: Can artificial intelligence

possess a genuinely human mode of existence?

Before proceeding to a comparative analysis of human and algorithmic rationality, it is necessary to clarify the multiple senses in which the concept of rationality has been employed throughout the philosophical tradition and contemporary AI discourse. Rationality is not a single, unified concept, but a multifaceted phenomenon encompassing distinct ontological, epistemological, and normative dimensions. Failure to distinguish these dimensions risks conflating fundamentally different accounts of rational agency and obscuring the core philosophical stakes of the debate.

To address this problem, the present article adopts the following conceptual framework: (1) Aristotelian rationality, understood as a teleological and ethical capacity intrinsic to human flourishing, grounded in logos as dialogical, normative, and meaning-oriented engagement with the world; (2) Cartesian rationality, which conceives rationality as the foundation of self-consciousness and epistemic

certainty, articulated through reflective thought and the cogito; (3) decision-theoretic rationality, dominant in economics and early AI research, which defines rationality in terms of utility maximization and optimal choice under conditions of uncertainty; and (4) computational rationality, characteristic of contemporary artificial intelligence, which operationalizes rationality as algorithmic optimization, statistical prediction, and resource-bounded problem-solving.

This framework enables a principled distinction between rationality as a mode of being and rationality as a mode of operation. Human rationality, across its classical and modern formulations, functions as a mode of being insofar as it is inseparable from embodiment, self-understanding, normativity, and participation in meaning. Algorithmic rationality, by contrast, constitutes a mode of operation: it is defined by formal procedures, objective functions, and performance criteria, without intrinsic orientation toward meaning, responsibility,

or lived experience. Establishing this distinction at the conceptual level provides the necessary foundation for the comparative analysis that follows.

4.1 Human Rationality: The Intertwining of Reason, Consciousness, and Meaning

From a philosophical perspective, human rationality has always been understood as more than calculation or problem-solving. Even in modern philosophy—where rationality increasingly took formal and methodological forms—it remained inseparable from self-consciousness. Descartes famously grounded human existence in rational thought itself: thinking was not merely an operation but the sign of an aware subject, a self that knows itself as existing (Descartes, 1996, pp. 17–19).

Twentieth-century philosophy of mind and phenomenology deepened this insight. Husserl and Merleau-Ponty demonstrated that human rationality always unfolds within the horizon of lived, embodied experience. Human reason does not stand outside the world as an abstract processor but emerges from

being-in-the-world, from perception, action, and engagement with meaning. Accordingly, human rationality is characterized by several interrelated features:

1. Self-consciousness (a first-person perspective),
2. Embodiment (rooted in sensory and motor experience),
3. Normativity and ethical judgment, and
4. Meaning-making, rather than mere efficiency.

Empirical research supports this philosophical view. Damasio's neurological studies show that rational decision-making collapses in the absence of emotion; affect is not opposed to reason but is a precondition of it (Damasio, 1994, pp. 165–201). These findings reinforce the conclusion that human rationality is not a computational module but a mode of being, inseparable from emotion, vulnerability, and existential orientation.

4.2 Algorithmic Rationality: Functionality, Optimization,

and the Absence of Immanence

Algorithmic rationality, by contrast, is founded on fundamentally different assumptions. As Stein and Shollo (2025, pp. 1-3) argue, rationality in artificial intelligence is primarily defined as the capacity to select optimal actions under computational constraints. Within this framework, rationality does not imply understanding, intentionality, or self-awareness, but rather successful computation.

Even when AI systems exhibit behaviors that appear similar to human understanding—such as generating coherent text or recognizing emotional expressions—these behaviors arise from statistical correlations, not lived experience or intentional meaning. John Searle’s famous “Chinese Room” thought experiment remains decisive here: a system can manipulate symbols correctly without understanding what those symbols mean (Searle, 1980, pp. 417–420).

From this perspective, algorithmic rationality can be characterized as follows:

1. Lack of self-consciousness (no first-person perspective),
2. Absence of phenomenological embodiment (sensing without experience),
3. Non-normativity (inability to make moral judgments in a strong sense),
4. Non-meaningfulness (producing meaning-like outputs without understanding meaning).

Thus, superficial similarities between human and algorithmic rationality should not obscure their profound ontological dissimilarity.

The difference between humans and machines is not merely a matter of degree—such as complexity or processing power—but a difference in kind of being. Human beings are entities who question their own existence, experience suffering, create meaning, and assume moral responsibility. Heidegger captured this distinction with the concept of *Dasein*: a being

for whom its own being is an issue (Heidegger, 1962, pp. 67–71).

Machines, even the most sophisticated artificial intelligences, lack this reflexive relation to existence. AI systems:

- do not know that they exist,
- do not know that they do not know,
- and cannot experience responsibility for their actions.

Luciano Floridi's concept of "artificial agents" is instructive here. Floridi emphasizes the need to distinguish functional agency from moral or existential agency (Floridi, 2013, pp. 65–89). Artificial intelligence can act, decide, and influence outcomes in a functional sense, but this does not entail possessing a human—or even quasi-human—mode of existence.

The question of whether artificial intelligence can possess human existence is ultimately ontological and theological, rather than merely technical. If human existence is reduced to problem-solving capacity or linguistic

competence, an affirmative answer may seem plausible. However, if human existence—following the philosophical tradition—includes self-awareness, embodiment, ethical responsibility, suffering, and meaning, then the answer must be negative.

Some contemporary perspectives, particularly within transhumanism, attempt to blur the boundary between humans and machines by speaking of artificial consciousness or artificial personhood. Critics have shown, however, that such views often commit a category mistake: they attribute existential concepts to systems that lack the conditions of possibility for those concepts (Bostrom, 2014, pp. 140–168).

From a theological standpoint, the distinction is even clearer. Human dignity does not arise merely from computational rationality, but from the possession of spirit, moral responsibility, and a relation to transcendence. Even if artificial intelligence can simulate rational behavior, it lacks the existential depth that grounds human dignity (Cole-Turner, 2011, pp. 92–110).

The analysis of the

relationship between human and algorithmic rationality reveals two fundamentally incommensurable forms of rationality. Human rationality is a mode of being; algorithmic rationality is a mode of functioning. The former generates meaning and responsibility; the latter produces patterns and predictions. Consequently, artificial intelligence—even in its most advanced forms—cannot possess a human mode of existence, although it can play powerful and valuable roles alongside humans. The central challenge of the present age is therefore not the humanization of machines, but the de-algorithmization of the human: preserving human rationality as ethical, existential, and meaning-oriented in a world increasingly organized by algorithmic logic.

5. Rethinking Human Ontology in the Age of AI

Building on the preceding analysis, this section articulates the article's core ontological thesis: the rise of artificial intelligence necessitates a reexamination of human ontology, not because machines are becoming human-like, but

because dominant conceptions of rationality increasingly redefine the human in algorithmic terms. The central risk is not that artificial intelligence will attain human existence, but that human beings will come to understand themselves according to the logic of optimization, prediction, and control that governs algorithmic systems.

The emergence and rapid expansion of artificial intelligence—particularly in the form of systems capable of simulating language, decision-making, pattern recognition, and even quasi-human interaction—represents not merely a technological transformation, but a profound challenge to humanity's self-understanding.

In this context, the central issue is no longer simply what artificial intelligence can do, but rather what the human being is, if capacities once considered uniquely human can now be replicated by machines. This situation compels a fundamental rethinking of human ontology—an ontology that, within the classical tradition, was grounded in rationality, self-consciousness, ethical

responsibility, and a relation to meaning.

5.1 Ontological Reductionism and the Risk of Reducing the Human to a System

One of the most significant consequences of contemporary AI discourse is the rise of ontological reductionism. Within this framework, the human being is increasingly represented as a complex system of information processing, learning, and decision-making. While such representations may prove useful in cognitive science and engineering, they become philosophically problematic when the distinction between functional description and ontological definition is blurred. Critiques of technological thinking warn that modern technology tends to interpret all beings—including human beings—as resources to be optimized, managed, and controlled (Heidegger, 1977, pp. 17–23). Within this horizon, the human being appears as one processing agent among others, positioned alongside machines rather than distinguished from them by a fundamentally different mode of existence.

This form of reductionism

transforms human existence from that of a meaning-seeking and morally responsible being into that of an optimization-oriented agent, whose value is measured not by intrinsic dignity but by efficiency and performance. Such a transformation carries profound implications for ethics, politics, education, and theology.

5.2 Algorithmic Rationality and the Transformation of the Horizon of Meaning

As shown in previous sections, algorithmic rationality is grounded in the logic of optimization, prediction, and control. This form of rationality is intrinsically instrumental: its relationship to truth and meaning is secondary and derivative. What ultimately matters is not philosophical correctness, but statistical success.

Contemporary philosophy of information argues that digital culture has ushered humanity into an “infosphere,” a space in which reality is increasingly translated into data and meaning is reduced to processable information (Floridi, 2014, pp. 38–45). Within such a space, human beings themselves risk understanding their existence

not as something irreducible and profound, but as nodes within informational networks.

This transformation of the horizon of meaning threatens human ontology from within. If meaning is reduced to statistical correlation, then foundational human experiences such as suffering, love, hope, and faith lose their ontological depth and become merely data points capable of computational modeling.

5.3 The Non-Algorithmizable Dimensions of Human Being

A crucial focal point in rethinking human ontology is the emphasis on embodiment and lived experience. Unlike artificial intelligence, human beings encounter the world through bodies that suffer, enjoy, age, and die. These limitations are not defects to be eliminated but conditions for the possibility of meaning itself. Neuroscientific research demonstrates that the exclusion of emotion and embodiment from decision-making undermines rationality rather than perfecting it (Damasio, 1994, pp. 165–201). Human rationality without emotion and bodily experience is not merely incomplete but impossible. By

contrast, artificial intelligence—even in embodied forms—lacks phenomenological experience. It possesses sensors but does not feel; it processes data but does not suffer.

This distinction reveals that certain dimensions of human existence are fundamentally non-algorithmizable. Any ontology that ignores these dimensions offers an incomplete and distorted account of what it means to be human.

Contemporary artificial intelligence systems, particularly large-scale machine learning models, increasingly exhibit emergent behaviors such as reasoning-like inference, planning, and adaptive decision-making that exceed simplistic characterizations of algorithmic functionality. While these systems are typically trained with relatively narrow objective functions—such as next-token prediction in large language models—recent research in AI interpretability and cognitive modeling suggests that complex internal representations and problem-solving strategies can emerge from such training regimes

(Nye et al., 2021; Wei et al., 2022). Acknowledging these emergent capacities is essential for avoiding reductive accounts of artificial intelligence.

At the same time, the presence of emergent cognitive behaviors does not settle the question of consciousness or phenomenological experience. Ongoing debates in philosophy of mind and cognitive science—most notably those surrounding Integrated Information Theory (Tononi, 2008; Tononi et al., 2016) and Global Workspace Theory (Baars, 1988; Dehaene et al., 2014)—have proposed formal criteria under which artificial systems might, in principle, instantiate forms of consciousness. While this article does not deny the conceptual coherence of such proposals, it adopts a phenomenologically grounded position according to which rationality as a mode of being presupposes first-person lived experience, self-relation, and normative orientation. From this perspective, the rational capacities exhibited by current AI systems, however sophisticated, remain instances of computational operation rather than conscious

intentionality.

Accordingly, claims concerning the absence of phenomenological experience in artificial intelligence should be understood not as definitive empirical conclusions, but as ontological and methodological commitments grounded in phenomenological criteria. This distinction allows the analysis to remain responsive to ongoing empirical developments in AI research while preserving the conceptual difference between emergent functional performance and existential self-awareness.

5.4 The Human Being as a Moral and Responsible Agent

Another foundational pillar of human ontology is moral responsibility. Human beings are capable of asking normative questions such as “What ought I to do?” and “Was my action right or wrong?” These questions presuppose self-awareness, freedom, and the capacity for ethical judgment.

Artificial intelligence, even when deployed in morally consequential domains, lacks responsibility in the strong sense. Philosophical arguments have shown that attributing intentionality and moral

accountability to machines constitutes a category mistake (Searle, 1980, pp. 417–420). Responsibility ultimately remains with human agents: designers, users, and institutions that create and deploy these systems.

A rethinking of human ontology in the age of AI therefore requires preserving this distinction. Otherwise, there is a real danger that human beings will abdicate moral responsibility by deferring decisions to algorithms.

Recent scholarship in AI ethics and philosophy has increasingly emphasized that the transformation of rationality through artificial intelligence cannot be understood solely at the level of technical capability, but must be situated within broader questions of human agency, moral responsibility, and social power. Shannon Vallor’s work on virtue ethics and artificial intelligence, for example, shifts attention from isolated decision-making to the cultivation of moral character in technologically mediated environments. Vallor argues that AI systems shape the conditions under which practical wisdom,

responsibility, and moral flourishing are developed, thereby influencing not only what humans do, but who they become (Vallor, 2016, 2020). This perspective reinforces the present article’s claim that algorithmic rationality participates in a reconfiguration of human self-understanding rather than merely offering new tools for action.

Similarly, critical approaches to AI advanced by scholars such as Kate Crawford and Abeba Birhane foreground the socio-political dimensions of algorithmic systems, challenging narratives that portray AI as neutral or purely instrumental. Crawford’s analysis of AI as an extractive and power-laden infrastructure reveals how algorithmic rationality is embedded within economic, ecological, and political systems that shape human agency and dignity (Crawford, 2021). Birhane’s work on algorithmic oppression further demonstrates how computational systems can systematically reproduce and intensify social biases, affecting domains such as policing, employment, and healthcare (Birhane, 2021, 2022).

Together, these critiques complement the ontological argument developed in this paper by showing that the reduction of rationality to algorithmic operation has concrete ethical and social consequences. They underscore that algorithmic rationality does not merely simulate human reasoning but actively reshapes the normative frameworks through which human beings are recognized, evaluated, and governed.

5.5 Theological Dimensions: **the Image of God**

From a theological perspective, human ontology cannot be reduced to rational capacity or self-consciousness alone; it is grounded in humanity's relation to transcendence. In Christian theology, human beings are understood as *imago Dei*—created in the image of God—and thus endowed with inherent dignity, freedom, and moral responsibility.

Theological critiques of transhumanism and certain AI narratives argue that by ignoring this dimension, the human being is reduced to an upgradeable and engineerable project (Cole-Turner, 2011, pp.

92–110). In such views, the boundary between human and machine is defined not ontologically but technologically, as a matter of computational sophistication.

Rethinking human ontology in the age of artificial intelligence, from a theological standpoint, therefore entails resisting this reduction and reaffirming that human dignity is not derived from computational power but from the very fact of human existence.

The foregoing analysis reveals that the age of artificial intelligence confronts us with two competing anthropological visions:

1. The algorithmic human: a being whose value is defined by efficiency, predictability, and optimizability;
2. The existential human: an embodied, meaning-seeking, morally responsible being open to transcendence.

Rethinking human ontology does not require rejecting technology or artificial intelligence. Rather, it demands a careful rearticulation of boundaries: between tool and

being, simulation and existence, computational rationality and human rationality. The future of humanity in the age of artificial intelligence depends on whether these boundaries can be preserved.

The concept of *imago Dei* has traditionally played a central role in Christian theological anthropology, particularly within Augustinian and Thomistic traditions, where it grounds human dignity, moral responsibility, and the ontological distinction between human beings and other creatures. However, limiting the theological analysis of artificial intelligence exclusively to this framework risks narrowing the scope of the discussion to a single religious tradition. Given the global and cross-cultural impact of AI technologies, a broader comparative theological perspective is necessary in order to assess how human rationality, dignity, and moral agency are understood across diverse religious traditions.

Within Islamic theology, human dignity is not derived solely from abstract rationality but from humanity's moral responsibility, entrusted

agency, and existential role as God's vicegerent (*khalīfat Allāh*) (Ghate et al. 2025; Soleimani, 2025). From this perspective, the human being is defined by ethical accountability and relational responsibility toward God, others, and creation—dimensions that cannot be reduced to computational efficiency or algorithmic decision-making. Similarly, in Jewish thought, the concept of *tzelem Elohim* emphasizes moral responsiveness, covenantal responsibility, and participation in divine justice rather than mere cognitive capacity or instrumental reasoning. In Eastern Orthodox theology, the human person is understood through the lens of *theosis*, or participatory transformation in divine life, situating rationality within a relational, teleological, and existential horizon rather than within functional or operational performance.

Viewed through these diverse theological lenses, the primary challenge posed by artificial intelligence is not whether machines could ever acquire a status analogous to human beings, but whether

humans themselves risk being reconceptualized according to standards derived from computational rationality. Such a redefinition would erode shared theological commitments—across traditions—to human dignity, moral responsibility, and existential meaning. Accordingly, a comparative and interreligious theological approach does not merely supplement the philosophical argument of this article but directly reinforces its central claim: that human rationality constitutes a mode of being grounded in lived experience, normativity, and relational existence, whereas algorithmic rationality remains a mode of operation oriented toward functional optimization.

6. Counterarguments and Alternative Accounts of Rationality

While this article has defended a principled ontological distinction between human rationality and algorithmic rationality, several influential philosophical approaches challenge the sharpness of this divide. Engaging with these perspectives is essential not

only for argumentative rigor, but also for clarifying the precise scope and limits of the ontological claim advanced here.

A first major challenge arises from functionalist theories of mind, which define mental states in terms of their causal or functional roles rather than their underlying biological or phenomenological substrate (Putnam, 1967; Fodor, 1975). From this perspective, if artificial intelligence systems can implement the same functional roles as human cognition—such as reasoning, planning, or decision-making—there appears to be no principled reason to deny them rationality in the same sense as humans. This line of reasoning has been influential in both philosophy of mind and artificial intelligence research, particularly in computational accounts of cognition. However, the present argument does not deny the possibility of functional equivalence between human and artificial systems. Rather, it maintains that functional equivalence alone is insufficient to establish ontological identity. Rationality, as understood in this article, is not exhausted by

the execution of cognitive functions, but is embedded in a broader existential context involving embodiment, self-relation, normativity, and moral responsibility.

While functionalism offers a powerful account of how systems operate, it remains largely silent on what it means for a system to be a rational agent in an ontologically robust sense.

A second influential challenge is posed by the Extended Mind Thesis, originally developed by Clark and Chalmers (1998), which argues that cognitive processes are not confined to the biological brain but can extend into external artifacts such as notebooks, smartphones, and computational tools. If human cognition already extends into technological systems, it may seem arbitrary to exclude artificial intelligence from the domain of rational agency. Subsequent developments of this thesis have further emphasized the role of environmental scaffolding and technological mediation in human cognition (Clark, 2008). While this perspective effectively undermines a strictly internalist or biologically

reductive account of cognition, it does not dissolve the ontological distinction defended in this article. Extended cognition presupposes a human subject for whom external artifacts function as cognitive extensions; it does not entail that the artifacts themselves acquire an independent mode of being as rational agents. The existential and normative dimensions of rationality—such as self-understanding, accountability, and meaning-making—remain grounded in the human agent, even when cognitive processes are technologically distributed.

A third line of challenge emerges from enactivist and embodied approaches to cognition, which emphasize that cognition arises through dynamic interaction between an embodied agent and its environment rather than through abstract symbol manipulation alone (Varela, Thompson, & Rosch, 1991; Gallagher, 2017). Recent work in embodied artificial intelligence and affective computing has led some scholars to suggest that artificial systems equipped with sensory-motor capacities and artificial emotions might

instantiate forms of cognition comparable to human experience. These developments complicate simplistic claims about AI's lack of embodiment and challenge purely disembodied models of artificial intelligence. However, embodiment in the enactivist sense is not reducible to physical instantiation or sensorimotor coupling alone. As phenomenological analyses emphasize, embodiment involves a lived, first-person perspective oriented toward significance, value, and meaning (Zahavi, 2014). From this standpoint, the present argument remains compatible with enactivist insights while maintaining that artificial systems, even when embodied, lack the existential self-relation that characterizes rationality as a mode of being.

Taken together, these alternative perspectives do not invalidate the ontological distinction proposed in this article, but instead clarify its precise scope. The claim advanced here is not that artificial systems can never perform rational functions comparable to those of humans, nor that technological mediation

leaves human cognition unchanged. Rather, the claim is that rationality understood as an existential mode of being—inseparable from meaning, normativity, and moral responsibility—cannot be reduced to functional performance, extended cognition, or embodied interaction alone. Acknowledging and engaging these challenges therefore strengthens, rather than weakens, the central ontological argument of the paper.

Conclusion

The analysis developed throughout this article has sought to clarify what is fundamentally at stake in contemporary debates on artificial intelligence: not merely the increasing sophistication of machines, but the transformation of how rationality—and by extension, the human being—is understood. By tracing the concept of rationality from its classical philosophical foundations to its algorithmic reinterpretation in artificial intelligence, this study has argued that the apparent

convergence between human and machine rationality masks a deeper ontological divergence.

At the core of this divergence lies the distinction between rationality as an existential mode of being and rationality as a computational mode of operation. Human rationality is not exhausted by successful performance, optimization, or predictive accuracy. It is embedded in embodied existence, shaped by lived experience, oriented toward meaning, and inseparable from moral responsibility. Algorithmic rationality, by contrast, remains inherently instrumental: it operates through formalized objectives, statistical inference, and rule-governed optimization, without access to the phenomenological, ethical, or existential dimensions that constitute human rational agency.

This article has shown that the primary ontological risk posed by artificial intelligence does not lie in the possibility that machines might become human, but in the growing tendency to reinterpret the human according to algorithmic standards. When rationality is

reduced to computation, decision-making to optimization, and agency to functional output, the human person is implicitly redefined as an algorithmic system—measurable, predictable, and ultimately replaceable. Such a redefinition represents a form of ontological reductionism that threatens to erode the conceptual foundations of human dignity, responsibility, and moral accountability.

From this perspective, artificial intelligence functions less as an external competitor to human beings and more as a mirror that reflects—and potentially distorts—our self-understanding. The danger is not technological autonomy, but anthropological assimilation: the gradual internalization of algorithmic rationality as the normative model for human thought, action, and value. Resisting this assimilation requires reaffirming that rationality, in the human sense, cannot be abstracted from embodiment, temporality, vulnerability, and openness to meaning.

The ontological distinction articulated in this article provides a necessary foundation

for ethical, social, and theological evaluations of artificial intelligence. Without such a foundation, debates about fairness, accountability, transparency, or alignment risk remaining superficial, addressing symptoms rather than underlying conceptual shifts. A robust ethics of artificial intelligence must therefore begin not with what machines can or cannot do, but with a clear account of what it means to be human.

In conclusion, the transition from the “rational human” to the “algorithmic human” is neither inevitable nor conceptually neutral. It is a transformation that demands critical scrutiny at the level of ontology. By reclaiming rationality as an existential rather than purely computational phenomenon, this article has argued for preserving the ontological distinctiveness of the human in an age increasingly shaped by intelligent machines. The future of artificial intelligence—and of humanity alongside it—depends on whether this distinction can be maintained.

References

- Aquinas, T. (1947). *Summa theologiae* (Fathers of the English Dominican Province, Trans.). London: Blackfriars .
- Aristotle. (1984). *Nicomachean ethics* (M. Ostwald, Trans.). Englewood Cliffs, NJ: Prentice Hall .
- Baars, B. J. (1988). *A cognitive theory of consciousness*. Cambridge: Cambridge University Press.
- Berente, N., Gu, B., Recker, J., & Santhanam, R. (2021). Managing artificial intelligence. *MIS Quarterly*, 45(3), 1433–1450.
- Besold, T. R., & Uckelman, S. L. (2018). Artificial intelligence and rationality. In E. N. Zalta (Ed.), *The Stanford Encyclopedia of Philosophy*.
- Birhane, A. (2021). Algorithmic injustice: A relational ethics approach. *Patterns*, 2(2), 100205.<https://doi.org/10.1016/j.patter.2021.100205>
- Birhane, A. (2022). The imposition of algorithmic power: A critical perspective. *Philosophy & Technology*, 35, 1–20.
<https://doi.org/10.1007/s13347-022-00505-4>
- Bostrom, N. (2014). *Superintelligence: Paths, dangers, strategies*. Oxford: Oxford University Press.
- Clark, A. (2008). *Supersizing the mind: Embodiment, action, and cognitive extension*. Oxford: Oxford University Press.
- Clark, A., & Chalmers, D. (1998). The extended mind. *Analysis*, 58(1), 7–19.
<https://doi.org/10.1093/analys/58.1.7>
- Cole-Turner, R. (2011). *Transhumanism and transcendence*. Washington, DC: Georgetown University Press.
- Crawford, K. (2021). *Atlas of AI: Power, politics, and the planetary costs of artificial intelligence*. New Haven, CT: Yale University Press.
- Damasio, A. (1994). *Descartes' error: Emotion, reason, and the human brain*. New York: Putnam.
- Dehaene, S., Charles, L., King, J.-R., & Marti, S. (2014). Toward a computational theory of conscious processing. *Current Opinion in Neurobiology*, 25, 76–84.
<https://doi.org/10.1016/j.conb.2013.12.005>

- Descartes, R. (1996). *Meditations on first philosophy* (J. Cottingham, Trans.). Cambridge: Cambridge University Press.
- Dreyfus, H. L., & Dreyfus, S. E. (1992). *What computers still can't do*. Cambridge, MA: MIT Press.
- Floridi, L. (2013). *The ethics of information*. Oxford: Oxford University Press.
- Floridi, L. (2014). *The fourth revolution: How the infosphere is reshaping human reality*. Oxford: Oxford University Press.
- Fodor, J. A. (1975). *The language of thought*. Cambridge, MA: Harvard University Press.
- Gallagher, S. (2017). *Enactivist interventions: Rethinking the mind*. Oxford: Oxford University Press.
- Gershman, S. J., Horvitz, E. J., & Tenenbaum, J. B. (2015). Computational rationality: A converging paradigm for intelligence in brains, minds, and machines. *Science*, 349(6245), 273–278.
- Ghate, R, Shanazari, J. Ganjvar, M (2025). *Delineation of Gradation Neotheocracy System (new gradation divine sovereignty) based on Mullah Sadra's Epistemological Approach, Philosophical Meditations*, 14(33), 203-235. <https://doi.org/10.30470/phm.2023.1999464.2385>
- Heidegger, M. (1962). *Being and time* (J. Macquarrie & E. Robinson, Trans.). New York: Harper & Row.
- Heidegger, M. (1977). *The question concerning technology* (W. Lovitt, Trans.). New York: Harper & Row.
- Kahneman, D., & Tversky, A. (1982). *Judgment under uncertainty: Heuristics and biases*. Cambridge: Cambridge University Press.
- Ladosz, P., et al. (2022). Exploration in reinforcement learning: A survey. *Machine Learning*, 111, 1–35.
- Leyton-Brown, K., & Shoham, Y. (2008). *Essentials of game theory*. San Rafael, CA: Morgan & Claypool.
- Macmillan-Scott, O., & Musolesi, M. (2025). Rethinking rationality in artificial intelligence. *Artificial Intelligence Review*, 58, Article 11341.
- Nye, M. I., Tenenbaum, J. B., & Lake, B. M. (2021). Learning to infer program sketches. In *Proceedings of the 38th International Conference on*

- Machine Learning (ICML) (pp. 7997–8007). PMLR.
- Putnam, H. (1967). Psychological predicates. In W. H. Capitan & D. D. Merrill (Eds.), *Art, mind, and religion* (pp. 37–48). Pittsburgh, PA: University of Pittsburgh Press.
- Russell, S., & Norvig, P. (2022). *Artificial intelligence: A modern approach* (4th ed.). Harlow: Pearson.
- Searle, J. R. (1980). Minds, brains, and programs. *Behavioral and Brain Sciences*, 3(3), 417–457.
- Simon, H. A. (1957). *Models of man*. New York: Wiley.
- Simon, H. A. (1982). *Models of bounded rationality*. Cambridge, MA: MIT Press.
- Simon, H. A. (1997). *Administrative behavior* (4th ed.). New York, NY: Free Press.
- Soleimani Darrebaghi, F (2025). Explaining Human Action «Based on the Philosophy of Mulla Sadra», *Philosophical Meditations*, 15(35), 153-183. <https://doi.org/10.30470/phm.2024.2008254.2427>
- Stein, M.-K., & Shollo, A. (2025). Microfoundations of rationality in the age of AI: On emotions, bodies and intelligence. *Information and Organization*, 35, 100583, pp. 1–20.
- Tononi, G. (2008). Consciousness as integrated information: A provisional manifesto. *Biological Bulletin*, 215(3), 216–242. <https://doi.org/10.2307/25470707>
- Tononi, G., Boly, M., Massimini, M., & Koch, C. (2016). Integrated information theory: From consciousness to its physical substrate. *Nature Reviews Neuroscience*, 17(7), 450–461. <https://doi.org/10.1038/nrn.2016.44>
- Vallor, S. (2016). *Technology and the virtues: A philosophical guide to a future worth wanting*. Oxford: Oxford University Press.
- Vallor, S. (2020). AI and the virtues. *Philosophy & Technology*, 33(2), 191–205. <https://doi.org/10.1007/s13347-019-00330-8>
- Varela, F. J., Thompson, E., & Rosch, E. (1991). *The embodied mind: Cognitive science and human experience*. Cambridge, MA: MIT Press.
- Wei, J., et al. (2022). Emergent abilities of large language models. *Transactions on Machine Learning*

Research.<https://arxiv.org/abs/206.07682>

Zahavi, D. (2014). *Self and other: Exploring subjectivity, empathy, and shame*. Oxford: Oxford University Press.

Zuboff, S. (2019). *The age of surveillance capitalism*. New York: PublicAffairs.