

On the Impossibility of Kantbots

Robert Hanna

The history of Kant's life is difficult to describe. For he had neither a life nor a history. He led a mechanically ordered, almost abstract, bachelor life in a quiet out-of-the-way lane in Königsberg, an old city on the northeast border of Germany. I do not believe that the great clock of the Cathedral there performed its task with less passion and less regularity than its fellow citizen, Immanuel Kant.... –What a strange contrast did this man's outward life present to his destructive, world-annihilating thoughts! Indeed, if the people of Königsberg had had the least awareness of the full significance of his ideas, they would have felt far more awful dread at the presence of this man than at the sight of an executioner, who can kill only the body. But the people saw in him nothing more than a Professor of Philosophy, and as he passed at his customary hour, they greeted him in a friendly manner and set their watches by him. (Heine, 1962: vol. II, p. 461, underlining added; see also Kuehn, 2001: p. 14)

Men are not machines, not even ghost-ridden machines. They are men—a tautology which is sometimes worth remembering. (Ryle, 1949/1963: p. 79)

Consciousness is the act or process of subjective experience. For example, as you're reading these very words, then you're enacting the subjective experience of reading these very words. Therefore, you enacted the consciousness of reading those very words. And as soon as you'd read, or re-read, the preceding two sentences, then you were also conscious of your first-order consciousness of reading those very words. So you were self-consciously conscious (i.e., meta-conscious) of reading those very words. Let's assume that you're a human animal—i.e., a “man” in Gilbert Ryle's quaint 1940s Oxonian oldspeak—and also know that you're a human animal. Therefore, since you're a human animal and also know it, then you already self-consciously know *that* human consciousness exists—i.e., your own human consciousness—and you also already self-consciously know *what* the specific character of human consciousness is, by direct acquaintance with it. In turn, the primary philosophical and scientific *problem of consciousness* is how to explain the existence and specific character of human consciousness in the manifestly real natural universe. There's a secondary philosophical and scientific problem of consciousness, which is how to generalize that explanation to anything, human or non-human, that's sentient and therefore capable of consciousness; but it can't be solved until the primary problem of consciousness has already been solved.

Over and above consciousness and self-consciousness, but also *including* them, intelligence is the complex mental act or process consisting of the more basic mental acts

or processes of (i) consciousness, (ii) self-consciousness, (iii) sensible cognition (i.e., sense perception, imagination, and episodic memory), (iv) intellectual cognition (i.e., conceptualizing, thinking, understanding, judging, logical inference, and semantic memory), (v) caring (i.e., desire, emotion, and feeling—aka the affects), (vi) free will, and (vii) practical agency (including instrumental or non-instrumental practical reasoning, choosing, and acting). So intelligence is *rational mindedness*. For example, in reading these very words and also understanding their meanings and logical implications, then you're enacting intelligence or rational mindedness. Indeed, the act or process of reading is a *paradigmatic* enactment of intelligence or rational mindedness (Hanna, 2024a).

Let's again assume that you're a human animal and also know that you're a human animal. Therefore, since you're a human animal and also know it, then you already self-consciously know *that* human intelligence exists—i.e., your own human intelligence—and you also already self-consciously know *what* the specific character of human intelligence is, by direct acquaintance with it. In turn, the primary philosophical and scientific *problem of intelligence* is how to explain the existence and specific character of human intelligence in this actual, natural world. As in the case of the problem of consciousness, there's a secondary philosophical and scientific problem of intelligence, which is how to generalize that explanation to anything, human or non-human, that's capable of intelligence; but, similarly, it can't be solved until the primary problem of intelligence has already been solved.

How can we solve the primary problems of consciousness and intelligence? Well, although it might seem so obvious as to be “tautological” in Ryle's terminology, in fact we already know something extremely important about human consciousness and intelligence. We already know *how to produce* human consciousness and human intelligence in this actual, natural world, namely, by means of *normal human reproduction* and *normal human biological development*. For all healthy, normal human animals possess innate capacities for consciousness and intelligence, and all healthy, normal, *mature* human animals enact those capacities during their lifetimes. As a species, we've been doing this for roughly 300,000 years, and on the whole, especially since we started engaging in the act or process of *reading* about 6000 years ago, we're pretty good at producing consciousness and intelligence.

Nowadays, moreover, we also have pretty good *scientific explanations* of normal human reproduction and of normal human biological development, as *complex dynamic living organismic systems* (see, e.g., Hanna and Maiese, 2009: section 7.3; Torday and Miller Jr, 2020; Hanna, 2024b: ch. 2). Moreover, as per the first three paragraphs of this essay, we already self-consciously know *that* human consciousness and human intelligence exist, and we also self-consciously already know *what* the specific characters of human

consciousness and human intelligence are, by direct acquaintance with them, for example, in the act or process of reading. Therefore, the unified solution to the problems of consciousness and intelligence is staring us right in the face: human consciousness and human intelligence are nothing more and nothing less than the essentially embodied, complex dynamic living organismic *global forms or structures* of the complex dynamic *living organismic systems* of healthy, normal human animals, as ineluctably embedded in their larger natural and social environments (Hanna and Maiese, 2009; Hanna, 2011; Maiese and Hanna, 2019; Torday, Miller Jr, and Hanna, 2020, 2024b: ch. 2). This unified solution, in turn, is *hylomorphic* (i.e., it invokes the necessary complementarity of causally activating and guiding *form or structure* and causally efficient *matter or stuffing*) and *organicist* (i.e., it's anti-mechanistic and focused on organismic life), and as such, was anticipated by both Aristotle and Kant:

The soul (*anima*) is the first actuality of a natural body that has life potentially. (Aristotle, 1968: II.i.412a22)

Life is the subjective condition of all our possible experience. (Kant, 1783/2004: p. 87; Ak 4: 335)

Mind for itself is entirely life (the principle of life itself). (Kant, 1790/2000: p. 159; Ak 5: 278)

One essential implication of the unified hylomorphic organicist solution to the problems of consciousness and intelligence is that consciousness and intelligence are *not* localized or located in *the human brain* or in any *other* proper part of the human body, *nor* are they anything that's actually or possibly *detachable* from the whole healthy, normal human animal, as ineluctably embedded in its larger natural and social environment. Therefore, to look for consciousness or intelligence in the human brain alone—or in any proper part of the human body for that matter—or to look for it in something that's actually or possibly detachable from the whole healthy, normal human animal, as ineluctably embedded in its larger natural and social environment, is just like looking for the social institution of a university in its central administration building alone, or in any of its classroom or office buildings alone, in the university president alone, in the administration alone, in the human resources staff alone, in the faculty alone, in the students alone, in the academic assistance staff alone, in the custodial and maintenance staff alone, in the lawns or athletic fields alone, or in any other proper parts of the whole university as a social institution, or in something that's actually or possibly detachable from the whole university as a social institution—"the spirit of the university" or whatever. As Ryle pointed out in *The Concept of Mind*, however, that's simply a *category mistake*, and precisely the sort of category mistake that René Descartes made when he

localized or located human consciousness or intelligence in *the pineal gland alone*, and when he identified them with the separable or separate immortal *soul* (Descartes, 1641-1642/1984; Ryle, 1949/1963: ch. 1). Ryle called this “the dogma of the Ghost in the Machine” (Ryle, 1949/1963: p. 17), and it’s crucial to note that not only Cartesian *substance dualists* but also Cartesian *mechanistic materialists*, alike, are equal but opposite subscribers to the dogma of the Ghost in the Machine. Again, human consciousness and human intelligence aren’t localized or located in proper parts of the human body, nor are they actually or possibly detachable from whole healthy, normal human animal, as ineluctably embedded in its natural and social environments. Instead, as I asserted above, they’re nothing more and nothing less than the essentially embodied, complex dynamic living organismic global forms or structures of the complex dynamic living organismic systems of healthy, normal human animals, as ineluctably embedded in their larger natural and social environments (Hanna and Maiese, 2009; Hanna, 2011; Maiese and Hanna, 2019; Torday, Miller Jr, and Hanna, 2020, 2024b: ch. 2).

Another essential implication of the unified hylomorphic organicist solution to the primary problems of consciousness and intelligence is that necessarily, no machine, no mechanical functional system, and in particular, *no digital computing system or digital technology*, can ever be conscious or intelligent, nor can its operations ever equal or exceed the achievements of conscious, intelligent human animals. Therefore, the philosophical and scientific doctrine of *computational functionalism*, and also the philosophical and scientific thesis of *strong artificial intelligence* (see, e.g., Block, 1980: part 3; Kim, 2011: ch. 6), which presupposes computational functionalism, are equally necessarily false and impossible. The widespread contemporary dogmatic false belief that the thesis of strong artificial intelligence is true is what I call *the myth of artificial intelligence* (Hanna, 2024c, 2025), and it’s intimately related to the dogma of the Ghost in the Machine in two ways: **first**, it also presupposes Cartesian mechanistic materialism, and **second**, it’s every bit as deeply and dogmatically ideologically embedded in contemporary philosophy, formal and natural science, and sociopolitical culture (Hanna and Paans, 2020).

To demonstrate all this, let’s consider Immanuel Kant, who was born in 1724 and died in 1804. Although he was notoriously mocked by the poet Heinrich Heine, as per the first epigraph of this essay, who described Kant in his old age as nothing but a philosophical clockwork puppet in a wig and professor’s garb, Kant was in fact a healthy, normal, rational human animal who was produced and raised to maturity by his parents, Mr and Mrs Kant, aka Johann Georg Kant and Anna Regina Kant (Kuehn, 2001: ch. 1). Kant was therefore conscious and intelligent. Kant was also a formal and natural scientist, and above all he was the greatest and most important modern philosopher after Descartes (Hanna, 2024d), who wrote what is undeniably (even to its critics and enemies) the greatest and most important philosophical book since Descartes’s *Meditations on First*

Philosophy appeared in 1641-1642, namely, the *Critique of Pure Reason* (Kant, 1781/1787/1997), along with many other important scientific and philosophical texts (Kuehn, 2001: chs. 4-8). Now, the following question arises: could Kant ever be effectively modelled by any actual or really possible digital computing system or digital technology—for convenience, let's call it a *Kantbot*—in such a way (i) that this Kantbot is conscious and intelligent, and (ii) that all of Kant's scientific and philosophical achievements are equalled or exceeded by this Kantbot? If so, that would entail this Kantbot's equalling or exceeding Kant's achievements of writing the *Critique of Pure Reason* and also the famous and highly influential essay "An Answer to the Question: What is Enlightenment?," which ends with this sentence, which of course is directly relevant to this essay:

When nature has unwrapped, from under this hard shell, the seed for which she cares most tenderly, namely the propensity and calling to *think* freely, the latter gradually works back upon the mentality of the people (which thereby gradually becomes capable of *freedom* in acting) and eventually even upon the principles of *government*, which finds it profitable to itself to treat the human being, *who is now more than a machine*, in keeping with his dignity. (Kant, 1784/1996: p. 22, Ak 8: 41-42, italics in the original)

Again: could there ever be a Kantbot such that it's conscious, intelligent, and able to equal or exceed Kant's actual scientific and philosophical achievements? Here is my answer: *No, absolutely not, and in fact the very idea of it is nothing but nonsense on stilts*. And here are ten reasons in support of that strong claim.

First, a Kantbot could never be either conscious or self-conscious, since a Kantbot, as a digital computing system or digital technology, is a machine, whereas Kant was a conscious and self-conscious human animal, hence a living organismic complex dynamic system, and *not* a machine, and *only* animals can be conscious or self-conscious, hence Kantbots are impossible (Hanna and Maiese, 2009; Hanna, 2011; Torday, Miller Jr, and Hanna, 2020; Hanna, 2024b: ch. 2).

Second, a Kantbot could never possess the unified set of capacities jointly constitutive of Kant's intelligence, since a Kantbot, as a digital computing system or digital technology, is a machine, whereas Kant was a conscious and self-conscious human animal, hence a living organismic complex dynamic system, and *not* a machine, and *only* conscious and self-conscious animals can possess the unified set of mental capacities, faculties, or powers that are jointly constitutive of human intelligence, hence Kantbots are impossible (Hanna and Maiese, 2009; Hanna, 2015, 2018, 2025: esp. ch. 1; Landgrebe and Smith, 2025).

Third, there are some illegible, meaningless, or nonsensical texts that no digital computing system or digital technology can parse or read, yet Kant could indeed parse and read such texts, hence Kantbots are impossible (Hanna, 2025: section 2.5).

Fourth, there are some well-specified sets of circumstances in which digital computing systems or digital technology cannot discriminate between left-handed and right-handed but otherwise identical counterparts (aka “incongruent counterparts,” aka “enantiomorphs”), yet Kant could indeed discriminate between them—in fact, Kant *discovered* the “incongruent counterparts” argument (Hanna, 2015: section 2.5)—hence Kantbots are impossible (Hanna, 2025: section 2.4).

Fifth, digital computing systems or digital technology can’t carry out functions or operations in the logico-mathematical sense over domains containing objects or other items that are either non-denumerably finite or non-denumerably infinite, vague, holistic, or entangled, or for which *the rule-following problem* holds, including *the halting problem*, yet Kant, as both a rationalistic logician and an intuitionistic mathematician, could indeed perform these very functions or operations, hence Kantbots are impossible (Hanna, 2001: chs. 4-5, 2006a: chs. 6-7, 2006b: ch. 6, 2015: chs. 6-8, 2025: section 2.6).

Sixth, no digital computing system or digital technology can carry out functions or operations in the logico-mathematical sense beyond the formal limitations determined by Kurt Gödel’s two incompleteness theorems, yet Kant, as both a rationalistic logician and an intuitionistic mathematician, could indeed perform these very functions or operations beyond the limits of incompleteness, hence Kantbots are impossible (Gödel, 1931/1967; Hanna, 2001: chs. 4-5, 2006a: chs. 6-7, 2006b: ch. 6, 2015: chs. 6-8, 2025: section 2.6; Chomsky et al., 2023; Keller, 2023; Landgrebe and Smith, 2025).

Seventh, no digital computing system or digital technology can perform categorical improvements or upgrades of the intrinsic specific character or quality of the informational inputs, premises, or materials with which they’re supplied—for example, from meaninglessness to meaningfulness, from morally bad to morally good, from non-beautiful to beautiful, or from falsity to truth—yet Kant, as not only both a rationalistic logician and an intuitionistic mathematician, but also the greatest and most important modern philosopher after Descartes, could *creatively transform* these informational inputs, premises, or other materials, into categorically improved or upgraded informational outputs, conclusions, or other products, in at least ten different ways, hence Kantbots are impossible (Hanna, 2001: chs. 4-5, 2006a: chs. 6-7, 2006b: ch. 6, 2015: chs. 6-8, 2025: sections 2.7-2.8).

Eighth, no digital computing system or digital technology can ever actually have, or even effectively model or simulate, the specifically human affects—i.e., the desires, emotions, and feelings—that all conscious and intelligent human animals can achieve, enact, or experience, especially including what I call *the desire for self-transcendence* and *deep happiness* or *principled authenticity*, whereas Kant, who was a conscious and intelligent human animal, and thereby also a *person* possessing *dignity* or absolute, non-denumerably infinite, inherent, objective worth, could indeed achieve, enact, or experience all of these, hence Kantbots are impossible (Hanna, 2023, 2025: chs. 3-5).

Ninth, no digital computing system or digital technology can ever *freely pretend* to be a digital computing system or digital technology, not only (i) because such systems or technology—as deterministic or indeterministic automata—can never have consciousness, self-consciousness, free will, or practical agency, all of which are required by the intentional act of pretending, but also (ii) because, as a matter of conceptual necessity, nothing can ever pretend *to be what it already is*, but instead can only ever pretend *to be what it actually is not*, whereas Kant, who was a conscious and intelligent human animal, could freely pretend to be a machine, hence Kantbots are impossible (Hanna, 2024d).

Tenth, and finally, every digital computing system or digital technology, in order to operate, presupposes a semantic, natural, and social context that operates as a foundational framework for all its operations. This is simply *given* to it, and without this foundational framework, it cannot operate. Therefore, that digital computing system or digital technology, including of course a Kantbot, cannot supply that foundational framework for itself. By sharp contrast, rational human animals like Kant, via their conscious and intelligent activities, can and do provide such foundational frameworks for themselves and others, including all digital computing systems or digital technology. So Kantbots are impossible.¹

If what I've argued is sound, then recognizing the unified hylomorphic organicist solution to the primary problems of consciousness and intelligence consists in freely performing the following three-step creative philosophical dance.

First step: critically and reflectively liberating oneself from the category mistake that all Cartesian substance dualists and Cartesian mechanistic materialists alike have committed, the dogma of the Ghost in the Machine (Ryle, 1949/1963: ch. 1).

¹ I'm grateful to Maria Carolina Mendonça de Resende for suggesting this line of argument to me. It's of course closely related to *the frame problem* in computer science.

Second step: critically and reflectively liberating oneself from the myth of artificial intelligence (Hanna, 2024a, 2025; Chomsky et al., 2023; Keller, 2023; Landgrebe and Smith, 2025).

Third step: attentively and open-mindedly seeing what's staring us right in the face, namely, the simple yet profoundly significant truth that human consciousness and human intelligence are nothing more and nothing less than the essentially embodied, complex dynamic living organismic global forms or structures of the complex dynamic living organismic systems of healthy, normal human animals, as ineluctably embedded in their larger natural and social environments (Hanna and Maiese, 2009; Hanna, 2011; Maiese and Hanna, 2019; Torday, Miller Jr, and Hanna, 2020; Hanna, 2024: ch. 1).

But, what if, for whatever reason, you simply can't (kant) yet *either* liberate yourself from the ideological clutches of the dogma of the Ghost in the Machine or the myth of artificial intelligence, *or* see what's staring us right in the face? That is: what if, for whatever reason, currently it's simply *impossible* for you to see what's staring you right in the face? My proposal is that you then try this supplementary creative philosophical dance move at least once: *look into a mirror and pretend to be a Kantbot, and then self-consciously critically reflect on what you've just done.*

Then, when you've successfully performed this supplementary creative philosophical dance, you'll recognize that you're *neither* a ghost, *nor* a machine, *nor* will your actual or possible authentically creative achievements *ever* be superseded by the operations of some digital computing system or digital technology, no matter how sophisticated it is and no matter how many bells and whistles have been tacked onto it by some obscenely rich technocratic capitalist corporation. For you're essentially *more* than a ghost and also essentially *more* than a machine, precisely because you're *nothing more and nothing less* than an essentially embodied, conscious, and intelligent animal that's biologically human, finite, fallible and more generally thoroughly normatively imperfect. We're all and only those conscious and intelligent "human all-too-human" animals in the mirror, staring us right in the face. You're essentially *like* Kant, and essentially *unlike* a Kantbot. *So get used to it.* And once we've gotten used to it, then we'll finally *know ourselves* in the Socratic sense.²

² I'm grateful to Elma Berisha and Bhupinder Singh Anand for thought-provoking correspondence on and around the main topics of this essay, and also to the organizers—especially Maria Borges—of the 11th Kant Multilateral Conference in Florianópolis, Brazil, 29 Sept. to 3 Oct. 2025, where I presented it.

REFERENCES

(Aristotle, 1968). Aristotle. *Aristotle's De Anima: Books I and II*. Trans. D.W. Hamlyn. Oxford: Clarendon/Oxford Univ. Press.

(Chomsky et al., 2023). Chomsky, N., Roberts, I. and Watumull, J. "Noam Chomsky: The False Promise of ChatGPT." *New York Times*. 8 March. Available online at URL = <https://www.nytimes.com/2023/03/08/opinion/noam-chomsky-chatgpt-ai.html>.

(Descartes, 1641-1642/1984). Descartes, R. "Meditations on First Philosophy." In R. Descartes, *Philosophical Writings of Descartes*. Trans. R. Cottingham et al. 3 vols., Cambridge: Cambridge Univ. Press. Vol. II. Pp. 3-62.

(Gödel, 1931/1967). Gödel, K. "On Formally Undecidable Propositions of Principia Mathematica and Related Systems." In J. van Heijenoort (ed.), *From Frege to Gödel*. Cambridge MA: Harvard Univ. Press. Pp. 596-617.

(Hanna, 2001). Hanna, R. *Kant and the Foundations of Analytic Philosophy*. Oxford: Clarendon/Oxford Univ. Press. Also available online in preview at URL = [https://www.academia.edu/25545883/Kant and the Foundations of Analytic Philosophy](https://www.academia.edu/25545883/Kant_and_the_Foundations_of_Analytic_Philosophy).

(Hanna, 2006a). Hanna, R. *Kant, Science, and Human Nature*. Oxford: Clarendon/Oxford Univ. Press. Also available online in preview at URL = [https://www.academia.edu/21558510/Kant Science and Human Nature](https://www.academia.edu/21558510/Kant_Science_and_Human_Nature).

(Hanna, 2006b). Hanna, R. *Rationality and Logic*. Cambridge: MIT Press. Available online in preview at URL = [https://www.academia.edu/21202624/Rationality and Logic](https://www.academia.edu/21202624/Rationality_and_Logic).

(Hanna, 2011). Hanna, R. "Minding the Body." *Philosophical Topics* 39: 15-40. Available online in preview at URL = [https://www.academia.edu/4458670/Minding the Body](https://www.academia.edu/4458670/Minding_the_Body).

(Hanna, 2015). Hanna, R. *Cognition, Content, and the A Priori: A Study in the Philosophy of Mind and Knowledge*. THE RATIONAL HUMAN CONDITION, Vol. 5. Oxford: Oxford Univ. Press. Available online in preview at URL = [https://www.academia.edu/35801833/The Rational Human Condition 5 Cognition Content and the A Priori A Study in the Philosophy of Mind and Knowledge OUP 2015](https://www.academia.edu/35801833/The_Rational_Human_Condition_5_Cognition_Content_and_the_A_Priori_A_Study_in_the_Philosophy_of_Mind_and_Knowledge_OUP_2015).

(Hanna, 2018). Hanna, R. *Deep Freedom and Real Persons: A Study in Metaphysics*. THE RATIONAL HUMAN CONDITION, Vol. 2. New York: Nova Science. Available online in preview at URL =

<https://www.academia.edu/35801857/The_Rational_Human_Condition_2_Deep_Freedom_and_Real_Persons_A_Study_in_Metaphysics_Nova_Science_2018>.

(Hanna, 2023). Hanna, R. "'It's a Human Thing. You Wouldn't Understand.' Computing Machinery and Affective Intelligence." Unpublished MS. Available online at URL = <https://www.academia.edu/102664604/Its_a_Human_Thing_You_Wouldnt_Understand_Computing_Machinery_and_Affective_Intelligence_June_2023_version>.

(Hanna, 2024a). Hanna, R. "Caveat Lector: From Wittgenstein to The Philosophy of Reading." *Journal for the Philosophy of Language, Mind and the Arts* 5: 1-26. Available online at URL = <<https://edizionicafoscari.unive.it/en/edizioni4/riviste/the-journal-for-the-philosophy-of-language-mind-an/>>.

(Hanna, 2024b). Hanna, R. *Science for Humans: Mind, Life, The Formal-&-Natural Sciences, and A New Concept of Nature*. Berlin: Springer Nature. Available online in preview at URL =

<https://www.academia.edu/118055113/Science_for_Humans_Mind_Life_The_Formal_and_Natural_Sciences_and_A_New_Concept_of_Nature_Springer_Nature_2024>.

(Hanna, 2024c). Hanna, R. "The Myth of AI, Existential Threat, Why The Myth Persists, and What is to be Done About It." *Borderless Philosophy* 7: 35-61. Available online at URL = <<https://www.cckp.space/single-post/bp-7-2024-robert-hanna-the-myth-of-ai-existential-threat-why-the-myth-persists-and-what-is-to>>.

(Hanna, 2024d). Hanna, R. "Kantian Futurism." *Journal of Philosophical Investigations* 18, 47: 1-8. Available online at URL =

<https://philosophy.tabrizu.ac.ir/article_18254.html?lang=en>.

(Hanna, 2024e). Hanna, R. "The Refutation of Compatibilism and Soft Determinism, and A New Proof That Incompatibilistic Real Free Agency Really Exists, and That You Really Have It." Unpublished MS. Available online at URL =

<https://www.academia.edu/103199623/The_Refutation_of_Compatibilism_and_Soft_Determinism_and_A_New_Proof_That_Incompatibilistic_Real_Free_Agency_Really_Exists_and_That_You_Really_Have_It_November_2024_version>.

(Hanna, 2025). Hanna, R. *Digital Technology for Humans: The Myth of AI, Human Dignity, and Neo-Luddism*. Berlin: De Gruyter Brill. Available in hard copy or eBook at URL = <https://www.degruyterbrill.com/document/doi/10.1515/9783111260525/html?lang=en>>.

(Hanna and Maiese, 2009). Hanna, R. and Maiese, M., *Embodied Minds in Action*. Oxford: Oxford Univ. Press. Available online in preview at URL = [https://www.academia.edu/21620839/Embodied Minds in Action](https://www.academia.edu/21620839/Embodied_Minds_in_Action)>.

(Hanna and Paans, 2020). Hanna, R. and Paans, O. "This is the Way the World Ends: A Philosophy of Civilization Since 1900, and A Philosophy of the Future." *Cosmos & History* 16, 2 (2020): 1-53. Available online at URL = <https://cosmosandhistory.org/index.php/journal/article/view/865>>.

(Heine, 1962). Heine, H. *Lyrik und Prosa*. 2 vols., Frankfurt am Main: Büchergilde Gutenberg.

(Kant, 1781/1787/1997). Kant, I. *Critique of Pure Reason*. Trans. P. Guyer and A. Wood. Cambridge: Cambridge Univ. Press. (1781 or A edition: Ak 4: 1-251; 1787 or B edition: Ak 3)

(Kant, 1783/2004). Kant, I. *Prolegomena to Any Future Metaphysics*. Trans. G. Hatfield. Cambridge: Cambridge Univ. Press. (Ak 4: 253-383)

(Kant, 1784/1996). Kant, I. "An Answer to the Question: What is Enlightenment?" Trans. M. Gregor. In I. Kant, *Practical Philosophy*. Cambridge: Cambridge Univ. Press. Pp. 17-22 (Ak 8: 35-42).

(Kant, 1790/2000). Kant, I. *Critique of the Power of Judgment*. Trans. P. Guyer and E. Matthews. Cambridge: Cambridge Univ. Press. (Ak 5: 165-485)

(Keller, 2023). Keller, A. "Artificial, But Not Intelligent: A Critical Analysis of AI and AGI." *Against Professional Philosophy*. 5 March. Available online at URL = <https://againstprofphil.org/2023/03/05/artificial-but-not-intelligent-a-critical-analysis-of-ai-and-agi/>>.

(Kuehn, 2001). Kuehn, M. *Kant: A Biography*. Cambridge: Cambridge Univ. Press.

(Landgrebe and Smith, 2025). Landgrebe, J. and Smith, B. *Why Machines Will Never Rule the World: Artificial Intelligence without Fear*. 2nd edn., London: Routledge.

(Maiese and Hanna, 2019). Maiese, M. and Hanna, R. *The Mind-Body Politic*. London: Palgrave Macmillan. Available online in preview at URL = [https://www.academia.edu/38764188/The Mind-Body Politic Preview Co-authored with Michelle Maiese forthcoming from Palgrave Macmillan in July 2019](https://www.academia.edu/38764188/The_Mind-Body_Politic_Preview_Co-authored_with_Michelle_Maiese_forthcoming_from_Palgrave_Macmillan_in_July_2019_)_.>.

(Ryle, 1948/1963). Ryle, G. *The Concept of Mind*. Harmondsworth, Middlesex UK: Penguin/Peregrine.

(Torday, Miller Jr, and Hanna, 2020) Torday, J.S., Miller, W.B. Jr, and Hanna, R. "Singularity, Life, and Mind: New Wave Organicism." In J.S. Torday, J.S. and W.B. Miller Jr, *The Singularity of Nature: A Convergence of Biology, Chemistry and Physics*. Cambridge: Royal Society of Chemistry. Ch. 20. Pp. 206-246.